

CONVERSATION TWO

3 Ps on a Pod: AI Disclosure on Substack and in Academia

Patrick Adler, Parth Goyal and Peter Matanle on disclosure scores, the stigma attached to partial AI use, authorship and credit, and whether transparency survives contact with the incentives around it.

Recorded 25 August 2026 · Runtime 1:26

Speakers Patrick Adler, University of Hong Kong · Parth Goyal, Parth's Data Stack · Peter Matanle

Transcript generated by Zoom, edited by Claude Opus 5

Learn more about [Parth](#) and [Peter](#) here.



Generated by ChatGPT Images 2.0 (GPT Image 2)

CONTENTS

01:37	The Slop Detective study: what Pangram measures, and which categories score highest
05:46	Why a fifteen or twenty per cent score carries stigma
09:09	The humanities' objections: inauthenticity, plagiarism-by-training, energy and water
15:54	The doping analogy: human achievement contest, or problem-solving enterprise?
18:25	Nearly everything in academia rests on predecessors who go unsighted
21:29	Graduating into a post-AI workforce, and worrying the skills never formed
24:09	Authorship and credit, where citations are the only currency academia has
25:44	The hackathon: agents making solo work cheaper than collaboration
28:57	AI slop as boilerplate with a new engine
33:13	Telephone numbers, and what the freed capacity is actually for
34:40	Writing a blog to stop the slipping
36:44	Picasso at Lascaux: we have learned nothing
40:00	Bridging the gap when the lower-order jobs are gone
42:32	Fields without benchmarks, and losing a coding moat that was paid for
45:57	Returns to expertise: does prior skill make you a better supervisor of agents?
48:51	The BS detector, and whether it can be built alongside AI
52:58	Handwritten essays plus designed AI tasks, joined by post-task discussion
59:02	Tacit knowledge, cookbooks, and why chefs still charge a premium
1:01:27	Methodology as the fulcrum of a good study
1:02:55	Rigour against enrolment: competition as the thing that drives slop
1:06:22	A modest proposal for a doping journal
1:08:31	Whether disclosure survives its own incentives
1:10:16	Body cameras for academics, mostly tongue-in-cheek
1:14:06	Standing by every word of a text that is half generated
1:16:06	Disclose when you benefit from perceived ownership
1:16:17	Journal editing as a precedent for crediting curation
1:17:37	Stigma as gatekeeping, and how enclosure feeds anti-intellectualism
1:21:29	Research that can't be explained to an intelligent outsider
1:23:28	Learn to Write Badly, and deliberate impenetrability in social science

PATRICK ADLER 00:01

We're live, welcome to the second conversation of Disembedded. Last night I was joined by, three friends, and tonight I'm joined by two folks I'm meeting for the first time. Even happier to talk to these folks.

Because it's nice to meet new people, and it's a good. A group of people, because we have someone who comes from outside of academic scholarship. And it has different norms, I think, when it comes to AI use. And also has done an analysis that I think is very illuminating, for us academics.

And then we have another academic, Peter, who works, you know, in a kind of cognate area, and, you know, area studies, you know, we sort of both work on depopulation a little bit. So, slight connection there, but still definitely not in my, milieu here at HKU, and, not solidly in my discipline. So, without further ado, we'll start with Parth. Parth, you have done, and we'll post the link, you've done this wonderful analysis, of disclosure on Substack.

So maybe you could start by just talking about the context for the study. You're a Substacker, you can give some context for that, and then the study that you did.

PARTH GOYAL 01:37

Yeah, thank you. I mean, first of all, thank you for posting this, I think it's super cool that you're recording it pretty online, and I'm glad you find my study, wonderful. So yeah, on the Substack itself, it's just something I do on the side, alongside my job, like, I like to think in terms of, data, and, like, there's so many problems, I think that we can solve that way, and just, like, pursue, like, whatever I'm curious about, effectively. So on the study itself, [Slop Detective](#), the idea was I recently found out about Pangram, the AI detection software, when Substack, announced that they were going to integrate it from, I think, mid-July onwards, where you could start scanning people's posts to see if it's human-detected, AI detected, all of that.

And then, as someone as anyone who ever scrolls through LinkedIn, you can see, like, how annoying it is when you just see people, like, posting slop, right? It's just, like, the most annoying thing ever, and then people and then you see all this engagement it gets, and I'm just I'm just curious, like. Okay, now that we actually have a reliable way to detect, AI software, can we actually see, how are users engaging with it? And I use Substack as, like, my choice of platform, partly because it's one of my favorite platforms, and the fact that they introduced the feature, I was like, okay, might as well give this a shot.

And then for Pangram, at least like the earlier models, you need, like, a minimum of 50 to 100 words of text, so Substack tends to run itself good for that kind of analysis. So I was just scanning a bunch of Substack posts, putting it through Pangram API, seeing how much was AI detected, how much was human detected, and then scraping, like, engagement metrics and stuff, and the TLDR of, like, what I found out is A is, like, the people who chose to disable the feature, like, not allowing their readers to scan their subside post, on average had almost 3 times the share of AI use, right? Which is quite interesting, because, like, if you kept up with the initial discourse around when the feature came out, a lot of people were angry. They're like, oh my god, like, why are you doing this?

You know, like, we have, like, authentic connection with our readers, leads to witch hunts and whatnot. And it's just, like, this is just an adverse selection problem, like, regardless of what people say, this is kind of, like, what you expect the behavior to be. And then another kind of intuitive finding was, like, smaller publications, on average, tend to use more AI. I think this is reflective of, people's writing ability, because obviously, like, if you have someone who writes extremely well, you would expect them to have more subscribers, more engagement.

I think there's also a part you can say that they might not have as many resources to do, like, full-time writing or something, right? They need to, like, use AI tools to speed them up. So that's interesting. And then, in Substack, people can self-tag their posts in terms of, like, what category it is.

You have, like, finance-type business, whatever, right? And then I found that, on average, like, financial writing and tech, adjacent writing tends to have higher share. Of AI usage, you would think that this is because, like, they are heavy adopters, but I feel like by the time you're in 2026, because all these posts are based on just 2026, I feel like everyone has kind of started using these tools quite heavily, so that didn't feel like a good enough, explanation. One thing that Pangram did on their own, like, their own independent study, is they found out, like, how much AI usage is there in, like, each of the different platforms.

And you can see, like, for example, LinkedIn or X, and some of the other platforms we expect, like, people are using them for, like, commercial intent purposes, right? Because, like, I feel like a lot of the time, like, people are selling something there. They tend to have much higher AI proliferation. And just using that as an analogy, I thought, like, that's kind of why I expect finance writing and tech writing had, like, at least on Substack, higher share of AI use.

Yeah, I mean, that's, like, the TLDR of, like, my post, happy to dive into anything you think is interesting.

PATRICK ADLER 05:46

I mean, I think it's great, and I think, as you point out in the post. The really useful aspect of this metric is it's continuous. So, you by opting in, you can, you know, theoretically be opting in to, like, 15-20% AI use. And what's interesting, because you talk about the initial, reaction to this, and this was very controversial, I'm most fascinated that people would be squeamish about disclosing their scores 15%, 20%, and, you know, granted, I don't exactly know what 15% means, but it doesn't sound very high, and I think given, you know, the widespread use of these tools to do, like, lower-order functions, like, I don't see why there'd be stigma at all.

Why do you think. There's stigma. Or some self-consciousness.

PARTH GOYAL 06:53

I have a pretty cynical view on this, to be honest. Like, a lot of publications that I used to follow, and then once the tool came out, and I also, like, started paying for, like, the lower-tier background, and started putting through posts, and I'm just like, this is kind of sad. A lot of people I follow, who I have

been, like, intuitively feeling like, oh, I'm just reading kind of, like, sloppy text, I'm just like, oh no, these authors are great. And I put their posts, and I'm like, oh, I increasingly see that once they have enough engagement, they're starting to kind of automate away some of the things that they have been previously doing.

I think like, in terms of, like, the stigma, I think, like, it might just be, like, any percentages of AI usage in your tests might signal that none of your work is original. It could be something like, oh, you know. I'm just speculating here, I don't know. I would think, like, if I see 20% score on someone's text, right, I'm like, oh, maybe, like, they generated it fully using AI, whatever the research question is, and then they just, like, edited the text to make it sound, like, human.

Right? So, I think, like, that's maybe it's more of, like, the signaling around it, but, like, any, like, you know, decent chunk of AI-ness in your text, might indicate to your readers that, like, this is not original work. Like, it's not worth their time. Which I disagree with.

I think, like, it's perfectly okay to use somewhat, right?

PATRICK ADLER 08:10

Yeah. Well, and maybe this is where, like, cultural norms come into it, and this is a good opportunity to bring Peter in, because Peter, in our area. You'll find journals that have, like, very, austere AI policies, such that 15%, 20%, that's, like, unpublishable. And, you know, last night we had a whole discussion about the journal Progress in Human Geography requesting what they call light AI usage, which I think is probably sub-15 or 20%.

So can you understand that desire to opt out, for people to opt out, you know, so they're they're using these tools, but their culture doesn't accept even that level of use, so they opt out. Do you think that's a kind of academic instinct, at least in the social sciences where we are.

PETER MATANLE 09:09

Yeah, I think in the social sciences, and even more so the humanities, there's an instinctive, and particularly if you like the softer side, or the more cultural side, or the more social side of the social sciences. As opposed to maybe economics, or I'm just speculating here. There's an instinctive aversion to the use of, AI tools. Because I think there's a there's a suspicion or an assumption that it's not really the author's work.

But I. I'm kind of very agnostic about this sort of thing, and I'm really here to try and learn about, where the ethical boundaries lie, because they seem to lie in different places, in different subjects. So we see, over on one side of the university, AI in common use, and for all kinds of different purposes. And, university academics advocating the use of AI because of its apparent reliability, its speed of operation, that it can save you a lot of time, it can perform complex calculations and tasks that would take a human being a long time to perform, and maybe they wouldn't be able to perform.

And can be used, for example, extremely useful as diagnostic tools, for example, in medicine and engineering and so on, which really saves a lot of people's lives. And yet, over on the over in the

humanities, we have a group of academics who are really, some of them, very vociferously against AI. I think that comes from a number of different directions. One is a stylistic.

Direction, in the sense that the style of writing, it's is apparently not produced by humans, so it's it's in some way inauthentic. But there's also other reasons for that, so, I think quite a lot of, people in social science and humanities are against the use of AI tools, particularly in writing, because of the environmental impacts of AI, of data centers, the consequences in terms of energy usage and water usage and building these places in remote places, which might be harmful for nature, and this sort of thing. So there's a lot of different things to unpack here, but personally, I'm really agnostic about this. I'm really wondering I.

I hesitate to say that some of those voices, perhaps, are not thinking through what real communication actually is, and what AI can be used for in humanities and social sciences, and who is really the author of anything? Because another aspect of the AI controversy, if you like, in social sciences and humanities is that apparently authors who use AI are plagiarizing other authors because of the way in which AI agents are trained. So this is a this is plagiarism going on. And so on.

But I'm, you know, I'm really rather agnostic about that, and the more I think it through, the more I think that not that those criticisms are misplaced, but that we need to have, actually, a deeper understanding and a deeper appreciation for what actually constitutes human communication, human-to-human communication. As well as what constitutes human to AI, or AI to human communication.

PATRICK ADLER 13:26

Yeah, I mean, because I guess, citing a corpus of theory and facts, which is just how we're trained, is I mean, it's not plagiarism because you're attributing, but, I mean, there are like, you're attributing kind of, like, the end-of-pipe part of that. There is a kind of common knowledge built up before that is very analogous to what's in this black box AI system. You're not citing the entire thing, and maybe that's that's kind of the that's a flavor of what you're getting at there, right? Like.

PETER MATANLE 14:05

I'm really struggling to try and describe it and articulate it, because I don't at the moment, I don't really know what I think about this, if you see what I mean. I haven't come to any kinds of conclusions or decisions, and really, I'm not an AI expert, but all of us in academia, and I suppose, you know, I'm speaking for a lot of people here who. Who are faced with the conundrum of, AI is here to stay. The public, the general public is not going to accept a ban on AI use.

Because of its, for example, its applications in medicine and engineering and so on that I already, mentioned. So, and most employers of our the graduates that graduate from my universe, you know, universities in the UK, are going to be compelled by their employ most of the graduates are going to be compelled by their employers to use AI tools, so that if so if we just say in the university, no AI tools, then we're not preparing our students for graduation, and we're not making them employable, and it would be a nonsense to do that. So, you know, there are lots and lots of ways in which I'm

thinking, well, we have to use this thing, and not just it's there for our use, and we could use it if we wanted to, or and so on. It's we have to use it!

And the consequence of that is, where do we use it? When do we use it? Where does the consensus lie in terms of its potential uses in academia and in training university students to become employable, for example.

PATRICK ADLER 15:54

Yeah, I think, like, part of why this conversation is so good is because we are part of this really weird culture, and I think it's a culture that's gotten weirder and weirder over the past decade or so, or a couple decades, and, like, Parth represents someone, like, outside that culture, like, presumably Parth, like, you went to university, but, like you know, you're kind of standing on the shore more than we are, and, you know, like, bringing you back in, one thing I was struck by this conversation we had last night, and a few conversations I've had lately, is, like. AI is revealing of this fact that, like, we seem to be doing different things, like, even in the same department. So, I think we're in the business of clarifying social issues, understanding social structures better. And if that's what we're in the business of doing, then the more social structure we can uncover, and the faster we can do it, the better.

Like, that's what I see myself doing. But then I talk to people, and I read, like, this Progress in the Human Geography editorial, and, you know, they're talking about cognitive shortcuts, like, it's a bad thing to take cognitive shortcuts. Like, presumably, if you use, like, a calculator, that's that's a cognitive shortcut that's, like, not so good. And it's almost like this model and I'm not trying to disparage the model, I think it's just a different model, is one of, like, human achievement, where scholars are supposed to, like, take their brains and apply them to the highest and best use, and everyone else is supposed to kind of marvel at that and revel in that, and this is, like, more of a model that's, like.

Competitive cycling, like, non-doping version, or, like, chess, non-computer version, and I and, you know, maybe I'll throw it to Peter, and then we could get Parth's, like, take on all this. I see myself as not in the human achievement business. I see myself more in the, like. I'm trying to figure stuff out on behalf of society, and that's why I think I'm more into these tools.

And I, you know, I dope, you know, this is a whole project that's, like, I'm doping right now, I'm a 20% AI user, or maybe even higher. I mean, I still maintain a lot of autonomy over this process, but, like, you know, some people look at this and they're completely scandalized. Peter, does that kind of resonate with you, that formulation?

PETER MATANLE 18:25

Oh, absolutely. I'm very much on your side in the pragmatic applications of what we're doing in academia. That, for me, there are very few of us that really should reasonably dignify themselves with the assumption that they're sort of some kind of a, you know, a blue skies thinker, an unusually talented and creative blue skies thinker who's creating really very new and original ideas. Nearly everything that we do in academia is it relies upon people previously who go unsighted.

And, you know, I've read hundreds and hundreds of articles and books, and yet I write an article which is really built upon that record of scholarship and that accumulated learning in a very significant way. And I write an article, or even a book, and I cite. You know, barely 5% of that, or even or less. And yet, apparently, AI is persona, in inverted commas, non grata.

So, but my principal aim in academia, obviously, is to try and resolve problems, and to try and help, you know, identify problems, and maybe potentially critique different approaches to those problems. And help practitioners who are actually facing those problems themselves, resolve them so that, you know, we can improve the conditions of life for humans and other animals and species. So, if AI is going to enhance that process, speed it up, and we really do need it speeded up when we consider that so little action has already so far been taken on, really monumental human crises, such as climate change, the biodiversity loss, pollution, war, and social and political fragmentation, and so on. We really do need to speed these processes up, and quite frankly.

Traditional methods of academia are not quick enough. They've proved they've proven to be not quick enough. We need something else to help us, I think.

PATRICK ADLER 21:01

So, obviously Peter and I are pretty aligned. Parth, what's your take, like, looking over the fence, or, like, this conversation that we've had, or, like, do you in this, like, Substack community, or, like, in other communities you might be part of, do you also recognize people who. Had this more, like, human achievement-oriented version of, like, how content should be created.

PARTH GOYAL 21:29

Yeah, I mean, as I'm hearing you guys speak about this as well, like, in my head, they're, like, two they're, like, these two distinct strands with this particular issue. Like, one is, like, the fact that I mean, this is true for all technology, that it enables you to do, like, be more productive, do better work, right? Like, if you just take coding, for example, like, previously, you might have needed to know, like, low-level language or something. Like, now I'm glad I can just, like, almost talk to it and get it done.

For most of us, coding is a means to an end, right? Like, I don't think anyone can disagree, like. That's, like, a good thing, and it makes you more productive, and all these things you guys said, like, you know, speeding up the process. I think the issue that a lot of people have, especially, like, me, because I recently graduated university, and then I ended up being in the workforce in, like, a post-AI world, where I'm just like, I don't think I've learned how to do these things in, like, a hard way, right?

I'm already offloading so much, on the tool. And that's where it comes to this idea of, like. Private property, or, like, authenticity or ownership, right? It previously, like, obviously, you write papers, it's a process of, like, reading stuff, and then connecting random ideas, and you come up with some good idea that you want to publish, right?

But then, before, I think you could easily say, or, like, assign, you know, taking plagiarism outside, you can assign that the value of that output to the researcher, being like, oh, this is the person who's made

it, like, you can, like, give them credit. But, like, now it's increasingly blurry, where you can't really say, like, oh, does this person deserve credit? Like, is this their original work? Apart from the quality of the work itself, I think that's, like, one of the things that everyone's having an issue with and doesn't really know how to answer.

And then second is, obviously, you have cases where, like, sometimes you read an, like, paper or a draft of a paper, and you're just like, okay, this is, like, extremely difficult to, like, read or understand. And I think it becomes annoying when, like, a reader also realizes, like, oh, the author themselves also don't understand this. They've unloaded this burden on me to understand. So I think, like, there are these two distinct problems, like, one is, like, the progress and the pragmatic of, like, oh, this is good to push science forward, and then the other is, like, at the same time, how do you attribute these things to people?

Because, like, people derive whatever, like, the identity and social connections and all of that to this other piece. That's, like, my kind of, like, take on it. I don't have an answer, that's why I was, like, curious, because, like, you guys have had, like, some experience, like, working without these tools, right? Like, you've developed some of these skills over time, then obviously now you're using these tools.

Do you guys feel this kind of weird estrangement from, like, your actual work? And how do you guys see yourself as, like, where am I adding, like, value in this process?

PATRICK ADLER 24:09

Maybe to be fair to us and our colleagues, I really like how you frame things, Parth. There's probably no domain where individual ownership, authorship over output, intellectual output, is so consequential, like, that's everything. You know, citation count is, like, the currency in academia, so that is the extent to which you personally have been credited for influential ideas. And so I think when we like, obviously AI is going to introduce regimes that eat away at authorship a lot, because you've got these, like, decentralized way of authoring things, or you're just making it easier for some machine to author things.

And, you know, I think for us. It's just gonna be very hard to expect for us to even be middle adopters of this technology. Like, we're probably going to be very late adopters because it's so like, because we just don't have the institutions that are equipped for it, but, like, in the real world Parth. Do you think, like, there's different institutions or different arrangements that are developing around, like, group authorship, or, like, co-authorship with AI?

Like, are there models that you could tell us about that, like, maybe would be influential within the academy?

PARTH GOYAL 25:44

I wish I had a good answer for you over this. I mean, like, in, like, in obviously, like, commercial settings, like, an employee doesn't really get authorship, right? Like, it's really much all about the output. And I guess, like, the revenue you can generate with it and whatnot.

And in those settings, I guess, like, these considerations are not, like, a big issue. It's very much, like, a result-oriented, like, oh, you know. I think this is developing good models, I think, like, this is, like, spinning up good intelligence or whatnot, which I can use for other tasks, and that's, like, perfectly okay. I don't have this issue of ownership.

Because I'm not that's not what, as a business, we are selling. Right? And I think the other content creation bit, like, that you see this, like, which is why, like, people slowly have been kind of this distaste for, like, whenever they identify anything, like, sloppy online, and they're just, like, not very happy with it. Because, like, you know, they are people are online because they are interacting with another person they can identify, and, like, when you read something like that, you just cannot.

And it's just like, what am I doing here? Like, why am I online? Yeah, I mean, I'd love to find out if there's, like, a nice model on, actually, group authorship that I recently discovered. I went to a hackathon, like, last weekend, right?

And then it was, like, a simple you're given a prompt, like, a challenge, you have, like, a few hours to develop something, right? I talked to the organizers after, and I was like, oh, like, you know, you guys had given us the offer of, like, forming teams, but surprisingly, a lot of us ended up doing solo projects. Is this, like, a common thing? But now, like, people say, oh, increasingly we're seeing as, like, these coding tools are improving, people actually want to collaborate a lot less at these events, because it's, like, such a friction to get your agent to talk to some other person's agent, when I can just, like, keep spinning off stuff on my own.

And which is, like, quite surprising to me, because I was like, oh, I would have imagined, like, human collaborations are hard, because you have to communicate so many things, but, like, apparently this is also, like, one thing where, like. A solo person has kind of enabled so much more that they might seem like, this is such a friction, like, telling someone, like, what my setup is like, and who's gonna collaborate is gonna slow me down. So yeah, I don't know if you guys are slowly experiencing this while, like, researching and stuff, but I totally see it sometimes, and it's like, oh, it is actually, like, I might as well just do this on my own, like, why do I need to share things?

PATRICK ADLER 28:05

I mean, arguably, that's the premise for my project, right? Like, the whole thing is, like, I've there's this so, like, in my field, there's this whole idea that, like, these, like, knowledge spillovers between people locally, especially through face-to-face interaction, are generating a lot of benefit, and so I'm saying, oh, AI lets me. And, like, this is verboten in my field. AI really lets me, like, forego all that, so I'm, like, I'm like your colleagues at the hackathon.

I'm, like, opting into that. But I don't know, Peter, you find that I mean, I know you're kind of an agnostic on the big picture, and you know, maybe you're, like, a light or moderate user of AI, but do you find yourself becoming more cloistered as a result of it.

PETER MATANLE 28:57

Yeah, I do, and I mean, I use AI for mundane tasks that would take me much longer to perform if I hadn't used AI. I before AI, I used software to do that. And before software, I used a typewriter, or whatever.

I mean, you know, so I think that a lot of this is sort of, discussions really around not all of it, I mean, it's really discussions around who really is the author of this, and what and its authenticity. But, you know, when we read AI slop on LinkedIn, which is I mean, Parth is great in describing this. It's it's so funny when you read this stuff. It's so obvious, and it's actually really boring to read.

The point about this is we've had AI slop ever since people started writing.

It's just it's just that this time, it's a machine producing yes. Because in the past, I mean, you know, in the last hundred years, we used to call AI slop boilerplate. Which was basically, you know, when corporations, you know, or what did we call it? Managerialism, or, you know, all these kinds of expressions we use for really boring, obfuscating language where people don't really mean what they're saying, and they're not really thinking about what they're writing, and so on.

We've had all this before, and you know, this goes way back, you know, since, you know. Time immemorial that leaders have stood up on podiums and talked to us as if we're children, and used the same language, you know, from one place to another, patronizing people with this nonsensical language that even a fool can really, you know, really start to take apart and critique and then throw it back. So this time it's just the machine that's doing the same thing, really. And people should be careful of it.

I, you know, yes, I do use AI, and particularly for unimportant tasks, for playing around. For doing things that are rather mundane and boring, and I'd rather spend my time doing something else. And it occasionally enhances what I've produced. It gives, you know, also the idea, and I know this is probably going to be quite a controversial thing to say, but at least for an adult, and I'm not talking about children here.

Or younger people, even. But at least for an adult, or someone at my age, who's, you know, whose brain is way past that sort of early developmental stage. The idea that AI would impair my cognitive functions, I think, is probably contrary to the evidence that I see from my own behaviors, because actually, it frees me. It does all these mundane, boring tasks that really don't need much thinking about, and frees me up, and gives me more time to do the real thinking, so I can go out on a walk.

And really think through the consequences of what I've just produced. And come back, and enhance it still further. So actually, for me, I find that AI actually improves my cognitive functions. If someone tested me for that, they might come to a different conclusion.

You know, that's how I feel it works for me. That's the evidence that I've seen in my own behaviors and in my own outputs, but, is that really the case? I really don't know, because I haven't got inside my head with a bunch of wires and you know, detection machines.

PATRICK ADLER 33:13

I mean, some people used to say, or maybe would still say, that, like, having a senior moment means you forget telephone numbers, but I personally, and this is a pre-AI technology, like, outsourcing thing. I personally know, like, 3 telephone numbers. My brain just refuses to memorize telephone numbers. And it's like, because I'm using technology as a crutch or a cognitive shortcut.

And, you know, I would like to think that brain power is being used for more productive things. But it does beg a very interesting question, because Parth, like, you as someone, you know, who really did a lot, like I mean, Peter's talking about, like, the sort of, like, most supple brain period. I think you're probably, like, you were probably in that period, and were in that period when you were in college. Do you find yourself grasping for?

Because people are talking about, like, are we going to lose some of these fundamental skills or knowledges? Do you find yourself grasping for knowledge that you do you feel like you should have learned, but you didn't because you came of, like, age in this AI period?

PARTH GOYAL 34:40

Yeah, I mean, like, full disclosure, this is partly why I, like, started writing a blog, because I was like, oh my god, like, I see myself, like. I don't know, like, something about, like, my brain is slipping, where I'm just, like, ever since, like, maybe, like, my for context, I graduated, like, 2-3 years ago, I'm 24, so I came out, like, 4 years ago, I was maybe, like, a sophomore, junior, whatever, like, in college, and ever since then, I've just been like, oh my god, I feel like I'm losing my ability to think, and this is why, like, I've picked up writing, because I'm like, this is, like, helping me in a way, right? Because I feel like I'm articulating, I'm thinking through things, and whatnot. So that's, like, definitely there, where, like, for a lot of people who are just getting started with their careers, it's, like.

Alienating, in the sense that you actually kind of feel like you've stopped your cognitive development if you don't take stock and actually try to put in the practice. So yeah, like, most definitely, that's, like, an extreme concern. And then that makes us think about our careers in the future, and being like, oh, like, then what then, if I just become, like, a vegetable, and I'm, like, not able to think for myself? Like, am I gonna be able to earn in the future?

Like, am I able to, like, climb up, like, the social ladder, right? Like, these are actually concerns that people have, like, around my age, and I think that's, like, why you're such, like. Strong opinions about it, online. I had, like, a question for you guys, actually, kind of on this thread, right?

Like, I don't know if you guys saw, like, about 2-3 weeks ago, you had this news about, like, how, some of the latest OpenAI, like, models started, like, doing some crazy math theorems, right? And a bunch of mathematicians were just like, oh my god, this is, like, either gonna put me out of work forever, and

I feel like so elated, and some of them are had the opposite, like, oh my god, this is the greatest thing ever, like, we can just do so much more math, right? Like, have you guys had, like, similar moments in your fields and professions, talking to colleagues, and like, if such a moment did come to your field, like, how would you feel about it?

PETER MATANLE 36:44

That question really brings me to what Picasso said when he saw the 15,000-year-old cave art in Lascaux in France. And, Picasso observed looked at all this cave art. And he said, we have learned nothing. I mean the point he was making is that, actually.

Humans don't change, and we do change. And we adapt to circumstances, and we adapt to the circumstances and the tools that we have around us. But fundamentally, actually, there's not much difference between you and I and the people who were living 10, 15,000 years ago, there's probably not that much difference between us and the people who were living 200,000 years ago on the plains of the Serengeti in Africa. In terms of the way we approach life, and risk, and opportunity, and those sorts of things.

I. I think at universities, we really have to integrate AI into what we do in a very serious way, and I don't think it's right just to say we can't we're not going to and we can't use it. The genie is out of the bottle. And we just have to deal with it.

And the way to do it is to actually openly integrated with traditional skills, and traditional skills building, and with much more open and frank discussions, amongst ourselves as academics, but also between academics and students, and amongst students. As to what these skills really are, and how we can honestly and openly integrate who we are as humans, and what we do as humans, with these machines, because these machines actually are produced by humans anyway. So in a sort of, you know, it's a very crude analogy, but I do think it's kind of true that Picasso said that, because, you know. It's not that different from sharpening a flint.

PARTH GOYAL 39:32

Well, I mean, sorry to push back on this. I think I get the point, like, just on the flint analogy, I think A lot of people are worried that this is one of those things where it, like, starts sharpening itself without the human in the loop. Which I think it feels just alienating, but it also could be one of those things that younger generation is living through it for the first time, so it feels like, oh my god, this is, like, the most novel thing ever. Maybe that's wrong.

Yeah,.

PETER MATANLE 40:00

On that matter, I watched it, actually, kind of in preparation for this talk, I watched a really interesting video, or film made by, BBC last night on streaming BBC. About the use of AI robots in, you know, to replace, taxis and humanoid robots and so on. I mean robotic cars and so on, as well as humanoid robots, and the discussion over that question that you ask, Parth, which is a really important question,

which is how do younger people bridge the gap now. That is appearing between graduating from university and then going into jobs that require higher-order human functions, where the intervening gap of lower-order human functions is increasingly being replaced by AI.

So people who go into the labor force at the age of 18, or 21, or 24, or whatever, they need a period of a few years to get used to being in the labour force, to learn really quite yeah, quite basic and mundane functions, and to learn to get around the labor force, learn to get around their workplaces, and build up basic knowledge and basic skills to be able to go into those, higher order, positions. And by the way, I mean, you know, people have been losing their minds, in their first and second jobs ever since first and second jobs ever existed. Because they, you know, they've been asked to, for example, dig a hole and fill it up again. You know, these sorts of things, when, you know, people really have been losing their minds and thinking, you know, why did I go to university if I actually have to if all I'm doing right now is data entry into an Excel spreadsheet?

So, again, I think we sometimes I get your point, but we sometimes do exaggerate the newness of AI, but at the same time, your point about the flint sharpening itself is a really important one. And, I'm I'm sorry, I don't have the answer to that one.

PATRICK ADLER 42:32

I mean, in terms of the math question, Parth, like. I don't think my discipline is structured I mean, like, it's said about chess, like, to use another chess analogy. That, like, chess players are miserable. Because there's only one chess player in the world who gets to say they're the best in the world.

This is such a sheer meritocracy. There can only be one person on top. You ask any of my colleagues, who the best human geographer in the world? They'll have a bunch of different answers, and, you know, half of them will say themselves, and this is a very comfortable place to be.

Because everyone gets to, like, if they don't think that they're the best geographer in the world, they get to say to themselves, what if I am the best, or what if I am one of the best? Or they get to set up little, like, small worlds where everyone, you know, like, 10 people are citing each other's articles, and, like, it's as though everyone's really famous, because they are in the small world. In other words, we're not like math. We don't have these, like, enduring problems, with clear benchmarks that AI gets to help solve.

And therefore, we don't have we don't have the integrity of having these, like, sobering moments. But we do in a kind of smaller way, like, and I'll give you, like, one small version of this, it's with coding. Parth, if you were to, like, see the sorts of coding that gets called coding in my neck of the woods, you would probably just laugh at laugh. I mean, like, we're not coders, we just, like, we figure out how to do little fun we have to figure out how to do the same functions over and over again.

Sometimes we use the same scripts that we got from our advisors. But nonetheless, we all had to learn to do that minimal amount of coding in graduate school, and because we weren't selected to graduate school with an ability to code, it was hard for us. And then AI comes along. And, you know, like, let's

just take someone who graduates from a PhD program in 2022, and someone who enters a PhD program in 2022.

That person who graduates had to, like, work so hard to get the coding ability that the person who's entering in 2022 can get within a month with an AI tutor. And in this specific case, it's just depression. It's just, like, real sadness that, like, I spent so much time, in some cases so much money, to do this thing. I thought it made me special, I thought it gave me a moat in the labor market, and now.

And in this way, I'm no longer protected. So I think, like, at the edges, you get those little moments, but then in the core of, I think, a lot of our disciplines, it's like, we're not even equipped, we don't even have the, like, benchmarks to find out that we should be terrified.

PARTH GOYAL 45:49

Just to follow up, actually, I don't know, like, how long we're going on for, so I'll just, like.

PATRICK ADLER 45:54

Just as long as we want, yeah.

PARTH GOYAL 45:57

Oh, awesome, okay, great. So just on this thread that you put, I think it's interesting, right? Like. You say, like, it is kind of frustrating that you spend all this time developing these skills that, like, you know, someone can just easily pick up, or, like, arguably don't even need to learn anymore.

But, like, on the returns to expertise, like, do you feel like, now that you are, at least, like, this week, that you will, like, kind of, like, you know, AI Max, and try to get out, like, a really good paper or whatnot. Like, you feel like you have that moat, or, the advantage of being able to prompt it, or write the ask questions, or manage the AI agents better, because you did all those things before. So you have an idea of being like, okay, I've done this before, I kind of know what I want my AI agent to do, I can supervise them better, versus, like, someone who did not do all those, like, who did not learn all those skills from, like, first principles, where they just kind of, like, take it for granted or whatever, like, the agent spits out. And, like, have you had conversations with your colleagues about this, where you actually do see, like, a steep returns to expertise kind of thing, or is it actually, like, no, this is actually more decentralizing technology, where, like, the gap is closer now between, like, kind of a junior researcher and, like, a senior researcher.

PATRICK ADLER 47:12

So this is fascinating. I'm reflecting on this for the very first time, like, right now with you, Parth, and you, Peter. I think, like, grad school and, like, working as an academic has trained me to be within a discipline. Or, sorry, be within, like, a topic area within a literature, understand where the really interesting parts of that literature are.

And then, you know, I do, like, research design in graduate school, and, you know, I get this, like, sort of sense of, like, how to structure research questions, so how to, like, fill the gaps in the literature. And those skills. Are really serving me well now. And then I just get to, like, it takes me way less time to code, so I'm just way faster filling those gaps.

But I think for me, my comparative advantage has always been with these things that I'm doing now, like, I'm just naturally better at understanding where areas are in the literature, I'm naturally better at doing, like, big picture design stuff and conceptualization than techniques and so, it's helping me better. If I was more of a like, the sort of person that I would, like, rely on formerly before AI for advice, you know, great people with, like, great coding skills, I might be a little bit more terrified about what was going on. I don't know how you feel, Peter, or what you're observing from others.

PETER MATANLE 48:51

I think I've benefited enormously from my academic training in terms of developing my critical analytical skills. In being able to basically, you know, have a pretty good BS detector in my head. When I'm reading academic work, but also to be constantly asking questions of what I'm reading. These are not necessarily critical as in, you know, negative criticism.

They're critical as in, trying to think through, is that really true, or what are the consequences of that statement? Or, you know, constantly asking myself these questions as I'm as I'm going through someone else's text, when I'm reading it, and I'm wondering if AI will how academics will develop those skills alongside or with AI, I'm sure there will be, but I think this goes back to my call, if you like, for a proper integration of AI And a proper conversation that really identifies what we do as academics and how we learn how to do it. And how we can integrate AI into that process to enhance it, so that it becomes a really useful tool. But it's still a tool.

That it doesn't govern us. I'm sure, you know, my daughter, she's 18, she relied on a calculator in her maths education, as did I. So you know, we you know, to do all the trigonometry, the cosine, the sine, and the tangent, and all that sort of thing, make all those kinds of calculations. You can't do it I mean, you can without a calculator, but it just takes a long time.

And you have to learn it, and so on, and learning how to do that, I'm I appreciate, does help to develop cognitive skills. But it's how to use AI and maintain those cognitive skills, and how to use AI and maintain and develop our critical functions. And critical analytical abilities, I think, is, it will be really is a kind of a thing I'm I've been thinking about.

PATRICK ADLER 51:42

Well, this reminds me of a conversation we were having last night, where, like, these critical skills you're describing, Peter, seem. Especially useful in the context of AI gives you a response, and it seems really credible. But is it? And, like, with critical skills, you seem like the best.

Partner that an agent can have. But at the same time, when I describe a future academic to current academics that is doing a lot of grading AIs, checking AIs, they just feel, like, really bummed out. About

that kind of future, and there's this metaphor of, like, a reverse minotaur or something. Where you're not even, like, telling the minute the machine where to go, like, the machine's telling you where to go, and you're just kind of checking their work.

Do where do you kind of, like, fall there? Like, is, like, that kind of critical thinking really a, acceptable future, as, like, a main academic function?

PETER MATANLE 52:58

With teaching and grading and reading other people's work – students' work, and we go right up to PhD level here. I honestly think that there's absolutely no harm in asking students. On the one hand, as part of their work, to sit down and write essays by hand. I think writing by hand is an incredibly useful cognitive development tool.

Both in terms of. Taking notes during lectures, as well as and organizing material, and understanding what's really important and what isn't, because you've only got a certain amount of time to write down what that guy just said before he goes on to the next point. So, you know, writing notes and so on, for example, during lectures, saying to students, during this lecture, no electronic equipment allowed at all. You're going to do it by hand.

You're gonna write your essay, which is 50% of the assessment for this module, by hand. But at the same time, you also ask them, students, to do well-designed, to perform well-designed tasks that require them to use AI. And you act there the students then and you then take students into post-task evaluative discussions. And you ask the students, well, you know, what did you get out of writing that essay by hand?

What were you thinking while you were doing it? or producing your notes during that lecture. What were the difficulties and challenges that you experienced? And sharing that amongst the students, and then having post-task evaluative discussions with AI tasks as well.

And asking the students and seeing how the students have got ideas about how to integrate the two different forms of different tasks into something that merges the two and integrates the two together. That includes oral discussion, and so on. I think we're gonna have to move towards those kinds of, evaluations and assessments of students, but actually, I think that would be incredibly interesting to do. For the lecturer, to actually go into the to actually design, and ask students to do it, and go through the post-task evaluations and so on, to go through that whole process of assessment, I think would be incredibly interesting, because even with essays that are written by hand, I can tell you 9 out of 10 are really crushingly boring.

And 9 out of 10 AI essays are going to be crushingly boring, too.

I would think that some kind of set of integrative tasks that involve all of those different ways of producing and interacting between machine and human, and between human and lecturer and student, and all of those sorts of things would be incre so much more interesting and stimulating and fulfilling for the lecturer. Not and I think it would be even more so for the student.

PATRICK ADLER 56:36

I mean, yeah, some of that is, like, brand new thinking for me. It's fascinating. I get intuitively that there's something about, like, when I was revising for exams when I was a kid, there's something about writing that just I don't know, I don't know if it's because it's tactile or something, but just, like. Weighs down the information in your brain a little bit more.

Parth, what do you think? As a recent, graduate.

PARTH GOYAL 57:11

Yeah, I mean, I must say, I don't want to speak for everyone else, in terms of the recent graduate cohort, because, very different views on this. But, yeah, I mean, like, I really like the experiment you laid out. I remember seeing I think the Economist had, like, a nice, like, article about it a couple of months ago, where, like, they measured, like, students' grades on, like, homework versus, like, you know, in-person exams, and then, like, it's like, these are, like, two different distributions. Like, in the homework, everyone's, like, a superstar, right?

And then it's like, when they come in and do an exam, it's like, oh my god, you're not even the same person. Yeah, so I mean, like, I think that's, like, an alarming thing, like, for educators to find out about, right? Because then you're just, like, really questioning, like, is a student really getting out, like, getting the most out of this? I mean, like it's one of those things where, like, I think.

The skill set that is valuable, like, at any given point of time in any field, like, rapidly evolves, right? Like, that's, like, pretty natural. And that's, like, I think understandable as well, but I think one of the things that I've been thinking a lot about, and a lot of people online have been saying about, like, and I guess this is what you're doing, Patrick, with, like, recording yourself, researching, is, like, there's these tacit, like, unspoken, un like, almost untaught, but something you pick up doing the things over time that, like, tend to be extremely val- like, underappreciated, but, like, extremely valuable. And I'm just, like, trying to understand, like, in research, for instance, in your fields, like, what are some of those things, if you had to articulate, like.

That are, like, that is, like, you know, unwritten, but it's, like, extremely important to deliver good outcomes.

PATRICK ADLER 59:02

I don't know if we know. I mean, like, this was the original impetus for this project, it's like so much of academic work you just do tacitly. And we don't understand it, and video is, like kind of gives a lot more clues into what's going on. But, like, that's the whole thing, like, you know, it's it's why. You know, a famous chef, a world-famous chef is, it gets to charge a premium.

Because and it's why they get to sell cookbooks, which seem to give everything away, and then still charge a premium for their wage, because of the tacit knowledge. With the study, though, that you mentioned, Parth, like. And as a recent graduate, I do want your perspective on this. I think people are misreading that study a little bit.

I think they listen to that study, and they're like, oh, everyone plagiarizes in homework, but then when they have to, like, prove it, they can't do it, and so they're just not learning anything. As opposed to, like you know, like, there are presumably some more progressive professors who are assigning homework that you can't just, like, plagiarize through AI to do well on. So, like, I do just imagine, and I see through some of my students, like, that there is a kind of co-piloting skill base, maybe something along the lines of what Peter's describing, that's emerging here. And maybe we're still, even when we were succeeding in those cases, falling short of creating knowledge that people can demonstrate without any aid, but then there's some people, like scott cunningham, who's a really influential economist, talked about AI stuff, and he was talking today, like, why do we give these problem sets that don't allow for the equivalent of calculators? when that's not, you know, it's not kind of externally valid, like, when is that when are people ever gonna need that? So I know, it seems pretty complicated, as.

Annoying a conclusion as that is.

PETER MATANLE 1:01:27

I think a lot of it is in the more I do this, the more I realize that the real. The central core of a really good study A really good piece of research is the methodology. And the way the methodology is articulated in the paper. Because so much of the reliability of the data itself the way it's the way it's used, and the way it's the way it's that the outcomes are interpreted relies upon the methodology.

And so, more and more, I look at the quality of the methodology, when I'm reading other people's work. And more focusing myself, on my own methodologies. I really yeah, for me, I think that really is the kind of the fulcrum around which a really good study revolves.

PATRICK ADLER 1:02:51

And that's the sort of thing you get to think about more when you're walking.

PETER MATANLE 1:02:55

Yeah, when you're out, yeah. But the other thing that we need to consider, of course, is the fact that universities also kind of cheat. I say kind of cheat, because it's really hard to so one of the problems with introducing, more rigorous, assessment, tasks, for example, handwritten assessment tasks. If you start asking students to write handwritten assessment tasks.

You're probably the university is probably going to lose revenue. Because next door, you know, you're gonna have two universities like we do in Sheffield, and 3 or 4 in manchester, and 3 or 4 in leeds, and so on, and nottingham. These were all kind of neighbouring cities. And if you've got only a small number of universities who are putting students through that kind of rigor, there will be students who are going to see that and think, actually, that's what I want.

Because I really want my uni- I really want my fees to work for me. In terms of improving who I am and what I can do. But there are gonna be a I would guess a greater number of students who are actually thinking, well, yeah, I want my student money I want it to work for me, but I also kind of want

to enjoy myself, and I want some time to do this and that. And to be honest, that sounds like it's really hard.

And, you know, I'm or I'm fearful that I'll actually show myself up, or all kinds of different thoughts going through their heads. And they'll decide, okay, well, I'm going to go to the university 20 miles down the road that doesn't ask me to do that. And so universities that have to compete against each other in a market-based system, are going to really struggle if they unfortunately, I think, if in a market-based system where. So much of university revenue derives from student fees.

It's gonna be really tough, and that's what you know, competitiveness as much as anything else, is what drives AI slop, I think.

PARTH GOYAL 1:05:19

I have an interesting take on, like, what you just said. It sounds, like, very interesting to me, like, I wish we had enough time to, like, stop AI capabilities and run this study right now. It's just like, oh, like, set of universities that are, like, having those these two distinct approaches, where one is, like, quite, I guess you could say, forward, or very you know, thing, and then one where it's just, like, it still forces students to do these kind of, like, core exams that have always been done, while also, like, giving them AI tools for other activities, but, like, still doing that, like, I wonder if you did that, and then you fast-forwarded and looked at, like, graduate outcomes, and the amount of wages that students earn from these two universities. You kind of worked out like an answer, right?

Maybe, like, these students who went to this supposedly much challenging university are actually more valuable to employers, and they actually should be commanding higher student fees. I mean, it'd be great to find out, but I feel like I don't know if you have the time to run that study.

PETER MATANLE 1:06:17

Yeah, if you had a time machine, you could do it.

That's neat.

PATRICK ADLER 1:06:22

There's a lot of local variation. I think there's not enough. In other words, like, with any discipline, there should be a, like. Like, to use the kind of doping analysis, or analogy, like, there should be, like, a doping journal.

That's basically, like, use as much AI as you want. Like, we pro- we support it. Just, like, let us know how much you're using. And then there should be other journals that are completely strict, and at the same time, like, there should be departments that are more or less strict, and then that would allow you to and I think the latter you do probably see, but the former you don't.

You have these, like, you have this hurting that I think is really bad, but I take your point, Parth, like. There's just so much we need to know about what works. And, we can't know everything to what you

and Peter were saying. I guess, you know, I think we probably could talk forever about this, which is a good sign.

But, yeah, I mean, maybe any last thoughts from you two on, maybe we can return to just the topic of, like, transparency, which Parth study really focuses on, like. Do you see hope for more transparency? Because I think we will agree that, like, the adverse selection problem that Parth highlights is not good, like, that's net negative for everyone. Do you see much hope for more transparency, at least around AI?

Regardless of other directions. More or less transparency. Maybe, Peter, we could say for you, like, more or less transparency around AI use in academia, and Parth, more or less transparency in a substack, so people adopting, like, disclosure more.

PARTH GOYAL 1:08:31

Oh, sorry, go ahead, Peter.

Oh. Okay, sure, thank you. Yeah, I mean, I think, like, as a consumer of, like, Substack, I would, like, I like the fact that they have the feature, and I hope that we keep, you know. Having transparency around it.

And I think, like, if you start seeing these signals where, like, authors that are or, like, you know, whatever publications are, like, more transparent about it, they see more engagement, like, their users are happy, you know, you have more paying users, then, like, that's automatically, like, a signal to be, like, actually, like, transparency is a good thing, people like it. Right, whereas, like, if you see cases where, like, people just don't care. They just want to consume what they want to consume, it doesn't matter if humans make it, AI make it, then you just kind of have the situation where it doesn't matter. That's one.

I hypothesize that people will appreciate the authenticity and that kind of thing. The second is, like, how long can we keep adding, like, transparency scores to AI outputs. Like, I think, yeah, Pangram has nice papers and details, like, why they think they can keep doing this. I'm sure, like, this is, like, maybe labs in the future.

I don't know if it's, like, completely in their incentive to allow outputs to be completely transparent, like, you know, I'm sure, like, they are also thinking, like, maybe we should, like, make this as close to as, like, human-looking as possible, so people can't, like, tell if it's, like, slop or not, right? So those are the two things that I have, like, on transparency itself. Generally, all for transparency, but, like, I hope we can keep continue doing that, you know.

PATRICK ADLER 1:10:16

Fascinating and kind of scary, too, actually, as, like, a pro-transparency person, although, like, I guess my the thing I'm experimenting with is just, like, record everything. My colleagues in academia, I think, almost universally support body cameras for police officers, but then probably at the same threshold would say, oh, no, we shouldn't wear body cams as academics. But, like, we could just all wear body cams, we could just all and there's AI-aided analysis of, like, how much AI you're using, this is

kind of where I see this, like, extreme vision is kind of where I see, but, like, I think on the scoring, that is kind of scary, right, Parth? Because, like, scores are more middle-range solution, they're a lot more, probably, acceptable to people, and maybe, to your point, it's not a sustainable one.

Peter, how about you? Transparency in academia, and around AI use?

PETER MATANLE 1:11:20

Well, I'd like to see a lot more transparency, among academics. At the same time, I'd like to see probably a lot more acceptance of AI. Because we know that, properly used, it does contribute, and contribute to and enhance knowledge, and particularly solutions, and can speed things along.

So I'd like to see I'd like to see more acceptance of it, but at the same time, you have to have more transparency. Now, that's the age-old problem, isn't it? You know, we've had plagiarisers in academia ever since academia existed, we're not gonna we're not gonna get rid of that. The question is did you really do it?

And we need, you know, did you really use the tool, and did you really write those prompts, and so on? And we need methods for kind of verifying the response. How do we do that? Body cams?

I don't know. Bodycams, I mean, I never heard anyone suggest that, but I mean, the police, the military, they all use them.

PATRICK ADLER 1:12:47

I say it kind of tongue-in-cheek, but not really. Not really, I mean, I don't know. I mean, like.

PETER MATANLE 1:12:55

No, I'm saying it tongue-in-cheek, Patrick. I mean, I'm not I'm not at all advocating for that.

I'm absolutely saying it, tongue-in-cheek, but at the same time. You know, we know that academics cheat. Well, some. Some academics cheat.

You know, that's happened. Since time immemorial, people try and gain an unfair advantage in a highly competitive field where they feel, perhaps, a bit desperate, and they feel inadequate, or whatever it might be. They want to take shortcuts.

PARTH GOYAL 1:13:38

Just one thing that you mentioned. I mean, this has kind of been, like, a little bit of, like, a negative AI conversation, but I think that what Peter mentioned was extremely important, was that like, people are accepting of the fact that you might have a journal that is, like, 40, 50, 60%, like, article that's, like, AI, and people are just like, oh my god, like, even that's, like, 60% AI, and that's, like, that's fucking awesome. That's, like, a great article. That also should be, like, an outcome that is accepted, in my opinion.

PETER MATANLE 1:14:06

Yeah, I agree with that, because the point about using AI, if you produce a text. That you're happy with, and you're pleased with, and you've edited it, and you'll stand by every word that's on that page. But it happens to be 50% produced by AI, because you prompted AI in some way to so you start with the prompts, AI produces something, and then you edit it at the end, and eventually you've got something that you're pleased with, and you think, I can stand by every word on that page. I can argue it, till I'm blue in the face, and I can argue with anyone that the validity, the veracity of the truth of this work, then I don't see why.

I don't see why not. So long as it's transparent and people know. But this goes for academic work. Yeah.

I mean, people are using AI in all kinds of ways to, you know, for example, I don't know, to enhance an email or whatever to their boss, because they're nervous about what the boss might say, so they want to get the wording right, and all of this sort of thing. People are doing it all the time. God, if you actually asked everyone. Every single time they used AI to enhance a text.

Of any kind whatsoever, it has to be formally declared.

You know, again, you're just going to get massive amounts of cheating. So, the point is with academic work, for example, that goes through peer review and all of that sort of thing that is added to the scientific canon. You really need transparency on all of that, time, you know, enhancing an email to a boss, enhancing a post on LinkedIn, or actually producing slop on LinkedIn, you know, do we really need. To be real, this is 60% AI.

PARTH GOYAL 1:16:06

I think a good rule of thumb is, like, as long as the author is benefiting from the perceived ownership of something, then they should be covered. Yeah.

PATRICK ADLER 1:16:17

I like that. And I mean, Peter, like, you gesture towards, like. Like, a way in which academics might be able to lead. In that, like, we have been giving massive amounts of credit, so this is a credit area Parth, to editors for a long time.

So you edit a journal, you're not producing anything, but you get a status that's commensurate with producing. Like, you're like, that's for some people, that's a career achievement, is editing a journal. And it's, you know, you could not write an entire thing. You could barely write an editorial the whole time.

So maybe we kind of embrace that aspect of our tradition, where editing is a real skill, and we just add to the mix, well. What you're editing is getting a little bit more diverse. Maybe that'll maybe that'll help normalize this a little bit more, because I think we're all in agreement that there's a weird kind of, stigma, I think, in academics, and then probably for some of these people in Substack who are opting

out. And they don't want to even admit to any AI use at all, so, like, it's not just an academic that's probably specialized there.

PETER MATANLE 1:17:37

Going back to the stigma, which you, correctly, sort of, mentioned, earlier on, I think it was towards the beginning of our discussion. I think this is a really important point, because one thing that I think is a little bit, A part of this stigma is academics and, if you like, talented writers, sort of reserving for themselves. A certain kind of privilege, and excluding others. From being a part of the creative production community.

I think some of that stigma that is being I mean, this is a kind of a risky thing to say, because I can't think for other people, but I get the sense that sometimes, and this is this really is a small number of times, but I think I get the sense that sometimes it's, you know, sort of putting up walls around academia and saying, this is us, and this is what we do, and we're, you know, this is our knowledge, this is our creativity, this is our style. And, don't pretend to be a part of it. And, I think that in the end, actually has the reverse impact than is intended, in the sense that it contributes to this idea that a lot of people have outside of academia, that academia, is it encloses itself in ivory towers, and it prevents access to knowledge and technology and ideas and so on that really are actually publicly owned, because the public pays for these institutions. And it privileges certain types of knowledge, and it privileges certain the ways in which knowledge is produced.

In a way that it's it's a little bit exclusive.

And I think we need to get away from that, and part of being transparent with AI is trying to get away from that and saying, look, these tools are for everyone to use. But we need to use them responsibly, and we need to integrate them with real human skills and attributes.

PARTH GOYAL 1:20:37

I know, Patrick, you were trying to wrap up, but I just had a thread I wanted to pull on this and wanted to ask you guys if that's okay. So there's this idea of, like, you know, transparency and academics, like, you know, I know you're saying you're usually guarded about, like, this ivory tower or whatnot, like, I feel like a lot of people talked about this trend where, like, be it Substack or any other long-form so I was saying, like, more and more academics are on these platforms and kind of sharing, like, shorter versions of their research that are more accessible to people, and stuff like that, like, even on Twitter, right? Like, I'm curious, like, what do you guys think of this? Like, I put, as a consumer, I'm just like, oh my god, this is great.

This is, like, a journal article, normally, that I would not bother reading, because it's 30 pages, and it's like, text I, like, have to sit and think about, but, like, this 30-minute article version of this, I actually enjoy it, and I like to think about it.

PATRICK ADLER 1:21:29

I mean, I think, and, you know, this is very resonant with what Peter was saying about boilerplate, you know, this tendency for AI to reveal, and I'm gonna get to your specific question part, but, like, AI reveals. Language that we kind of put up with, and we didn't like, for being, like, fraudulent in some way. Because it's just as, you know, as though an AI could have just made it, like, the slop. And I think, in a kind of opposite way, research that cannot be transmitted to you as a very intelligent, but uninitiated person in a field.

Like, stuff that just can't achieve that, like, cross that barrier, I think is in some sense fraudulent, if you're unable to and I and, you know, I'm not a, like, particle physicist, right? So perhaps there is something that, like, there is some kind of domain that's you just can't ever, like, explain it to someone else, but I actually doubt that. I think that our ability to, at least in theory, cross that divide, have shorter, form maybe more entertaining form stuff, is a certification that we are doing something that's useful to society. And that we can explain to other people.

And explain to people who have no kind of interest in accepting it, in with your when you're in this cloistered, small world that I was talking about earlier. Anyone you're, like, sharing something with is, like, incentivized to pretend like it's like, it makes sense when it doesn't.

PETER MATANLE 1:23:28

I'm absolutely with you. I just went to get a book off my shelf. This is a really fantastic book. I'll just show it here.

It's called [Learn to Write Badly: How to Succeed in the Social Sciences](#). So the point of this book is that, I mean, the book starts out by saying. Actually, a lot of people are successful in the social sciences because they have consciously learned how to write badly.

And they deliberately write badly, in order to exclude large numbers of people from accessing the text with terminology, convoluted, phraseology, and all sorts of other different things, in order to appear very erudite. And this I'm very much an empirical social scientist, and this really annoys me about theorists. The theorists write in this impenetrable language deliberately to be impenetrable.

And that really annoys me, because they should be trying to do the opposite, to make their work accessible, but they do it because they think it enhances their career, and it makes other people it impresses other people, and so that, you know, they might be a successful how to succeed in the social sciences. So, I absolutely I love it when people make their research, however complex, accessible to intelligent and reasonably well-educated ordinary people, and open-minded ordinary people who are willing to engage with it. I love it, I think it's fantastic. The point about being exclusive that I was talking about earlier, and, is that?

It actually fuels, the anti-intellectualism that we see in so much of particularly anglo-American society. That it actually fuels the doubt. That people have in the validity of and the accuracy and authority and honesty and truth of science. Because it distances the scientist, or it actually distances the theorist very often.

From, the citizens that, really, they should be trying to inform and connect with. And I think if AI can enhance that, and if we're transparent about that, then fantastic.

PATRICK ADLER 1:26:21

I agree. It's a difficult conversation, because I think we are kind of aligned in our values a lot, but, you know. Sometimes you can get to some cool places with that. And I think we did do that.

Any final words? I'll let you guys go, you guys have been so generous with your time, but any final words?

PARTH GOYAL 1:26:51

Yeah, I mean, I just wanted to first thank you both, like, I've had, like, a very fun time talking to y'all. Glad that you're doing this, and then also, I must say, Peter, like, I feel like you should do more of these. The fact that you had a book out, like, to actually prove your point, and you had, like, a very nice, articulate, like, ending to this, which, like, I think resonated with me at least a lot, I'm just like, that's awesome, that's awesome.

PETER MATANLE 1:27:15

That was entirely spontaneous.

PARTH GOYAL 1:27:18

Yeah, no, it was great!

PETER MATANLE 1:27:21

Yeah, I mean, I loved our discussion. Thank you, Patrick. I really learned a lot.

Really nice to hear from Parth and what you're doing. It's really been interesting for me, and I feel like I've just sort of you know, if you have a 12-lecture module at university, I've just I've just been through 3 lectures, and I've learned so much from just interacting with you both. Really enjoyed it. Thank you so much.

PATRICK ADLER 1:27:52

Well, the feeling's mutual, and who knows, maybe we'll do it again at some point. Let's all be in touch. Great meeting you both, and thanks again for your generosity.

PETER MATANLE 1:28:04

No, thank you. You got yours.

PARTH GOYAL 1:28:06

Best of luck.

Transcript generated by Zoom and edited by Claude Opus 5: backchannel interjections and broken fragments removed,
a small number of misrecognised proper nouns corrected. Wording is otherwise the speakers' own.
Part of Disembedded – an economic geography paper written live, with AI tools and no human collaborators. disembedded.com