

Hello, everyone! Well, it's been a while. I...

I'm now back. We're doing a livestream about Q&A, about future of science and technology. I was gone for nearly a month.

In Europe, and finally got back, and now, of course, I got slightly sick, so I'm a little bit under the weather, so may not be as bouncy as I sometimes am.

And, bye.

had a very intense trip to Europe. I hadn't really been there seriously for close to a decade, and many things have been saved up to do, and had a good chance to, meet up with people who I hadn't seen in a while. I think the record was somebody I hadn't seen in 53 years.

Although I had exchanged email within that period of time. But, I had an interesting time. But anyway, now I'm back and, able to do these livestreams again.

So, let me see, we have a question here from Anon. What do you think of the current AI bubble? well, that's kind of a leading question. It sort of assumes that the things going on in the world of AI are a bubble.

I think, the, you know, like many things, kind of the... the... What?

AI... kind of... the current round of AI technology has kind of shaken loose a bunch of ideas and methodologies and things that can be done that will have long-term importance.

whether the promise of what's happening right now will really be fulfilled is a lot less clear. I mean, I think

that, the way I see, kind of, what's happened with AI, and with LLMs in particular, you know, in 2022, a new kind of wild animal was discovered that is the LLM, and it does all kinds of amazing things.

And now the question is, how do we harness it best for, kind of, human technology?

And, what I think is sort of increasingly clear

is that the harness for particular kinds of domains is important. The wild animal is what it is. You know, people have said the wild animal's going to get much, much better, stronger, faster, smarter, etc. There's not a lot of evidence for that.

The things that are happening, in terms of, you know, there have been a lot of kinds of, oh, it'll be fine-tuning, it'll be chain-of-thought reasoning, it'll be something or another.

None of those will be a mixture of experts. None of those things have really panned out in a big, big way. They're all incremental engineering improvements, but they didn't hit anything out of the park.

I think the, sort of the question is, well, what can the wild animal ultimately do? And, the,

And you know, how far can it get? How much can we actually evolve the wild animal?

I think... The thing that, kind of in our world.

Kind of, we've spent lots of effort making, sort of, precise computation work

which is really a different kind of thing from what's being done with LLMs and neural nets and so on. LLMs and neural nets are, I think, providing a reasonable approximation of kind of what humans do, kind of, right off the top of their head.

It's something that can be done faster, cheaper, you know, with more data, and so on, but fundamentally, the kinds of tasks are ones that humans can sort of immediately do. Is this a cat or a dog? Is this... you know, what would you say if somebody asked you, blah blah blah? The kinds of things that they can't do is things like, okay, run code.

in your head. Humans don't manage to run code in their brains, LLMs don't manage to run code in their brains, so to speak. They can call out

To some external thing that does the computation and runs the code.

I'd like to think that the technology we've built over the last 40 years is, by a rather large margin, sort of the best thing you can call out to, because we have this computational language, our Wolfram language, that is kind of an effort to represent computationally real things in the world. And LLMs are talking about real things in the world. That's what they've trained about to think about, so to speak. We have

The computational underpinnings for dealing with real things in the world, whether it's data about cities, or computations about, you know, the pressure under the ocean, or whatever it is. It's, we, we have...

We've spent the last 40 years basically accumulating a giant tower of kind of integrated methodology for dealing with those kinds of things, and it's integrated, which is important, because it means that the AI, which is sort of extrapolating in, well, what's a reasonable way that this might work.

the reasonable way it might work, if I and we have done a good job as language designers, it will work the reasonable way it should work, and so the AI will sort of correctly extrapolate to be able to make use of our tool as its sort of computational neural implant, so to speak.

It's a good methodology. It really helps, as it does for humans. You know, humans don't manage to

do all kinds of complicated math or chemistry or whatever in their heads, they use our technology to do that, and so can AIs. And we've been developing, actually, over the past year or so, still greater capability to really plug our stuff together with modern AI technology. Sort of the sort of a very basic level of that is MCP,

capabilities, but there's a lot more beyond that. Really falls under the heading of what we tend to call

CAG, computational augmented generation, kind of instead of having retrieval augmented generation, which is kind of the idea of you've got this material, this textual material, for example, and you're... first of all, before you send things to the LLM, you're saying, well, does some part of the query that I made kind of thematically match something, some snippet.

in this document that I've already... or collection of documents I've already dealt with, let me kind of hint the LLM with that snippet. Well, what we can do with computational mentor generation is hint the LLM not with an existing snippet, but with something that's actually been custom computed from the actual query that was fed in.

But in any case, I think the, the sort of...

picture here is sort of a critical grounding for LLMs is the tools they can use, and I'd like to think that the computational language that we built over the last 40 years is sort of the best broad-spectrum tooling you could really imagine, I think, for LLMs.

And it's getting used, but I would say it could be used a lot more, and should be kind of a routine part of all LLMs as the thing that LLMs can refer to if they need to compute things, or they need specific

Detailed knowledge about things.

And with the correct kind of harnessing of the LLM,

using CAG technology and so on, you can kind of take an LLM that might have just sort of wandered off in some hallucinatory state, and you can be pretty sure that the LLM will be kind of correctly grounded in actual computational knowledge, which is, I think, a powerful thing.

But in terms of, kind of, the picture of, sort of, the LLM industry, as I say, there's sort of the core LLM capabilities, and then there's, kind of, tools.

And then there's things that the LLM can actuate, can do. You can get the LLM to sort of initiate things as sort of an agent, and then it's got to have some sort of actuation layer that says, yes, go make that plane reservation for me, or yes, go delete my file, or whatever it is. And that's kind of another sort of layer of harness around the LLM. I would say that,

There's a... for a lot of different domains of activity, there's sort of complicated harnesses that have to be built, whether you're doing it for medical diagnosis, whether you're doing it for legal tech.

Whether you're doing it as we have been doing it for teaching and education.

We have a big project that will hopefully soon see the light of day, having to do with building kind of an automated teaching system based where LLMs are a component, but there's an awful lot of harness that goes around the LLM to really make it something that can lead a person through a curriculum, interact with them.

kind of give reports on what they're doing. There's a lot of harnessing that's needed beyond the sort of the raw LLM that's inside there.

But, you know, it has to be said that LLMs, like many other kinds of previous technology, they provide a new spark of capability.

to do lots of kinds of things. They provide this sort of linguistic user interface that we've never had before. You know, if we look at the past graphical user interfaces that came in in the 1980s, basically, sort of over the course of a decade or so.

kind of transformed the way that one uses computers. And then later on, when mobile came online, when smartphones came online in, oh, what was that? 2007, 8, 9 kind of timeframe. Maybe a little later than that, even.

The, that kind of graphical user interface paradigm really sort of came of age.

And now there are many, many things where you say, oh, you should use a command line interface. That would be absurd. You would say, well, of course I'm going to use a graphical user interface. So similarly, there are many things today where LLMs provide the possibility of a linguistic user interface. You can just say what you want. And that works pretty well.

For, just like graphical user interfaces, work for a certain level of complexity of thing.

If your graphical user interface involves click, click, click 5 times or something, it's a win. If, to get what you want.

You need a graphical user interface where you're going through 18 pages, and you're having to read a lot of stuff, you're having to drag various things together, it's gonna fall apart.

One's actually seen that for graphical user interfaces when people do programming with graphical user interfaces. When the programs you're writing are very simple.

it works fine. You just click together some blocks, and all's good for a certain number of blocks.

But when it gets more complicated, and there's more complicated control flow in the programs and so on, that has always fallen apart. I mean, we've done all kinds of experiments over the period of time that often languages existed.

now nearly 40 years. And, every time we do it, it's like, well, yes, for the really simple stuff, you know.

Function applied to function applied to function, with maybe a branch or two, works just fine.

But as soon as you're doing something that has kind of the richness that you really need a true language to deal with, it falls apart.

So, each of these kinds of interface modalities has a place, and now we have a new one, a linguistic user interface, and it can do all kinds of things. It will be able to, you know, can have a

conversation with you, and summarize from that conversation this or that thing. That's a very useful capability.

If you want to take your linguistic user interface and say, make a precise engineering specification of something, that's probably not the best use. Like many kind of machine learning-based things.

it's kind of a, you know, at the 80% level, it's going to do great. Getting that last 20% really nailed down is basically impossible for that technology.

So I think the thing that, for me, has been emerging as a really good kind of workflow is you use this linguistic interface, and you say vaguely what you want. That produces Wolfram language code.

That is a precise, succinct computational specification of something.

Then, the whole point is that unlike, like, a programming language, our computational language is set up so you can read the code.

It's something that is really... operates at a human level, so to speak. At least a short piece of it does. And that means you can say, yes, this is what I wanted, the LLM did the right thing, now make that a sort of solid

block that I can use to build up this big tower of capability in some program or whatever that I'm building. So that seems to be a very good workflow. It's something from last year we implemented our notebook assistant. That's something where you can sort of ask the... make a vague statement in natural language.

get a block of Wolfram language code, and then make use of that to build your, sort of, tower of capability.

So...

I suppose those are... those are things that are working really well. Now the question is, well, okay, put a trillion dollars into, kind of, the world of AI, what's going to come out?

Well...

there's sort of a question, is the underlying wild animal of the LLM going to suddenly become bigger, stronger, faster, smarter, whatever, or is it that the harnessing is going to work better? I think that it is clear that there's a lot to do with the harnessing. Whether there's a trillion dollars worth of stuff to do with the harnessing.

is not clear, but, you know, you can always build out, sort of, bigger and bigger companies and so on to do things. The idea that, sort of, somehow you put all that money in, and somehow out will pop this amazing, you know, AGI miracle.

I think is a different... is a different ask.

my own, you know, my own kind of scientific analysis of it is that's just not going to make any sense. The fact is, there are things, for example, when it comes to, will you be able to get an LLM that doesn't hallucinate?

Well, no, because that's just not the nature of the ask, it's not the nature of what a thing like an LLM does. Can you hint it, you know, support it with CAG and RAG and so on? Absolutely. You can make those things be things that you've kind of nailed down to a large extent.

But I think the question of, of, what, what was,

you know, is there going to be a magic thing that suddenly does the kind of magic that ChatGPT did that nobody expected back in late 2022, of being able to sort of produce fluent language? Is there going to be another moment where suddenly the AI just does everything humans do?

I don't think that will happen as such.

I think what will happen is that the things that one wants... the things that an LLM already does include many things that unaided humans do. The thing that one wants to make something that's really a valuable system in the world is something that's much more connected than, just the kind of brain in a box that's talking to you kind of thing. And I think that, again, needs this kind of notion of harnessing, and it's not something where you just pump more money into the raw machine and get something amazing. Now, having said that.

The history of machine learning, from sort of the breakthroughs in deep learning in 2011 and so on, in image identification, to things about speech-to-text, to things about image generation.

To things about, language generation, each of these has been sort of a step breakthrough that one could never have predicted when they would happen.

Are there other breakthroughs like that that can happen? Yes, absolutely. You know, video generation is sort of on its way. I happen to think robotics, humanoid robotics probably, is another kind of thing that has sort of breakthrough potential over the next not very much time.

I don't think it is the nature of the... of the beast to say, oh, there'll be another breakthrough that allows us to immediately do all of math and science, just not how things work.

And the thing that, sort of, from a scientific point of view, is the main thing that gets in the way, is the whole idea of computational irreducibility that I've been talking about for the past 41 years or something.

That is this phenomenon where even though you might know the rules by which a system operates, to know what the system will do is something that requires actually explicitly following those rules. You don't get to jump ahead and just sort of say, hey, I'm smart enough to just figure out what's going to happen.

I think the, and computational irreducibility is what it... there is... that's what's involved in kind of running code in your brain. You have to be able to go through the steps of the code. You can't just say, oh, I know what the code is going to do, I don't have to run it.

So that's the thing which is sort of a core issue in doing lots of kinds of science and math and so on. Not all of it. There's some things where you can have sort of an intuitive jump.

you can kind of thematically learn something from the existing literature of science, and you can kind of jump ahead. But in the end, lots of kinds of things sort of run into this computational irreducibility, something that we've been able to, in practice, somewhat overcome by doing actual computation, but that's kind of a different thing from what you see in the world of neural nets and LLMs.

Could there be something which weaves together computation and what happens in LLMs? I'm not sure. I think that's an interesting possibility. There's something of a kind of a conflict, because what computes more is harder to train.

If you want to be able to train something, you have to kind of, you know, be able to know, sort of, all the way what it's going to do, and say, well, tweak this so you'll do something different. If there's a big sort of stack of computational irreducibility, that becomes not something that can really be expected to work.

Now...

You know, another thing that's very confusing about the current moment of AI is that now that AI is such a successful brand, lots of things that really aren't very AI or neural net-like suddenly get called AI.

Like, a very typical one that I've been pursuing now for, I don't know, 45 years or so, is just searching the computational universe for programs that do relevant things.

Could be programs, could be kind of chemical compositions of materials, could be geometric structures, those are all the same kind of thing, but the key idea is just go out into the computational universe and search for things. It really doesn't have a lot of kind of neural net AI in it. It's something that is born out of, kind of, the notion of computation. The only thing one can hope for, which hasn't really come to pass yet, is that when you're selecting, you know, you can go out and search for a zillion things, find all sorts of interesting cases, or things that you're heuristics say are interesting, maybe an LLM can tell you whether that theorem you just found, or whatever, will be one where mathematicians will say, hey, that's really cool, or they'll just shrug their shoulders and say, that's a random theorem out in the space, metamathematical space of possible theorems, and I don't care about it.

So...

In any case, where does this leave, kind of, the story of AI? You know, I have to say.

I was, I guess, involved in AI in the 1980s. Even my very first company eventually got named Inference Corporation back in 1982. So, when... when the... the kind of... the way we're going to solve all of AI is expert systems, I never believed that.

My, the technology that launched that company was... was something very different from that. But then I remember,

being a consultant at a company called Thinking Machines Corporation that was making massively parallel computers in the kind of mid-1980s, and sort of one of the investor pitches was, you know, we're going to invent the transistor of artificial intelligence, or at least we're going to have the chance to do that.

Well, of course, it didn't work out at that time, and now we're where we are, how many years later, and there have been big surprises, like ChatGPT.

And the question is, if you pump enough money into the problem, what can you achieve?

I think from a business strategy point of view.

it's quite brilliant on the part of, sort of, the larger AI companies to kind of hold out this sort of hope

of, kind of, the AGI possibility, and say, look, just put more, more, more money into this, because, you know, there's a chance there will be a payoff of just absolutely disproportionate, size.

And that's a very rational strategy if you're running such a company. Now, what will actually happen with that money

well, you know, by the time you've pumped hundreds of billions of dollars into something, you can do practically anything. You know, you could be, I don't know, building new data centers for anything. You could be running fiber optic cables all over the world. You could be putting up new, you know, network systems. You could be doing all sorts of things.

that are probably valuable for the world and are made possible by large amounts of money, but they're not specifically, you know, oh, we're going to do sort of the AGI thing.

If one was cynical, one will be reminded of a story, I think, from the Middle Ages, about a person who wanted to teach, or said they wanted to teach a horse to read, and they...

got some king to say, yes, I'll fund you to teach the horse to read, and I'll give you funding for a few years to do that. And at the end of that time, they're like, well, actually, the horse can't read yet, but if you give us more funding, you know, there's a good chance the horse will learn to read.

And you can kind of keep that going for a long time, and if the amounts are large enough, you get to sort of keep that as savings, and then figure out what to do with it later, so to speak. I mean, it sort of reminds me a lot of what sort of happened in the cryptocurrency world, with the whole ICO craze of, what was it, how long ago was that? Maybe, nearly 10 years ago now.

Where, sort of immense amounts of money were being raised, in a... in a way that doesn't even... it's not even sort of an equity-type thing. It's... it's really just a... this is a war chest that can be used for this kind of cryptocurrency. And a lot of those war chests still exist and are quite large. And sort of the challenge in that domain has been to kind of find

The... to sort of backfill with kind of things that are of obvious value beyond just the sort of network value of... of a store of value that blockchain provides, to really, sort of, have... have something to where one can say, yes, we've taken this X number of billions of dollars, and over this period of time, we've built things that are really solid things in the world that one can readily see have value.

that was a thing that, for cryptocurrencies, that's what's still sort of underway, and those treasuries are still quite large, and still, you know, there's probably, I don't know, a 20-year kind of window, a 20-year, you know, runway, for that to happen in that case. I think in the kind of AI world.

the runway will be considerably shorter than that, just because, although the companies are not public, typically, most of them are not, the, it's still, I think, that the dynamics of, sort of, the investment world will demand returns, or will will take the valuations and make them... it's not the case that you just have the money, and nobody's really asking you how is it valued, so to speak, which is more what happens with treasuries of cryptocurrencies.

So, you know, what will happen?

and, you know, what will cause different dynamics to occur, I think, you know, as there's... there's new data that eventually comes out about, you know, what's really working, what's not. I tend to think that there's enough energy around, sort of, we're going to make AI work in, you know, and we're going to reach AGI and so on.

there's enough energy around that that any negative news is really not going to be that visible. I think that's a, you know, that's a stampede

that is not, that won't really dissipate, easily. Now, of course, things could happen. I mean, things happen with the, deep-seq.

thing at the beginning of this year, where it was like, hey, wait a minute, this doesn't, you know, that was a thing that got enough visibility, that was kind of a shot across the bow of maybe some of this stuff doesn't make sense. But I think it got,

Kind of, got sort of, the stampede trampled over it fairly quickly, and has kept going.

And, so... but I think the important question is, you know, is it the case... you know, my feeling, as I'm saying, is that sort of the core technology isn't... is going to have some aha moments, but in different domains.

And, you know, for example, solving robotics is important. That's a, you know, that's a huge business to be solved. I think the... we're just going to make it smarter, it's going to be more brain-like.

I don't think that has so far to run.

I think that, sort of making the brain able to use tools, like our technology, for example, that's important. Kind of connecting it to sort of actuation channels, that's part of this kind of harness mechanism. I think that the vast majority of the value comes in this harness process, which is, in a sense, very traditional software engineering and

Those kinds of things. It's a place you can spend a lot of money, potentially wisely, to make things that are really useful.

But I think that the real sort of, it's worth putting trillions into this comes from the idea that there's sort of a magic philosopher's stone to be found there.

Which I don't think... I don't think there is. That doesn't mean it's irrational to put that money in, because, you know, it creates a war chest that lets many other things be done, including maybe developments like, let's say, AI for robotics that aren't

specifically on the, on the, kind of, on the path of, quotes, AGI and that sort of philosopher's Stone hope, but are still things that will be very useful.

So anyway, that's a few thoughts about the current state of the AI industry.

I will say that,

as somebody quite involved in all of this, it's, you know, people are telling me these things, oh, we're gonna teach AIs to do math, we're going to teach them to do this, that, and the other, and it's like, some of them, it's like, yeah, there are things to do, but the... the big shining hope

It's kind of, like, we sort of know scientifically that that just can't work.

And... but it doesn't seem to... nobody's enthusiasm seems to be dulled by the fact that it doesn't seem

scientifically plausible. Now, you know, having said that, there are plenty of things which people have said, oh, that's impossible, and it turns out to be possible. I don't think I've been a person who's said a lot, that's impossible, and it's turned out to be possible.

And so I think I... I perhaps, at least for myself, I feel like I have sort of more credibility in being able to say, this is not going to work. There's a thing close to it, maybe, that will work, but the thing you're aiming for is... is not a winner, so to speak.

But, the part that I think is, as I say, very much a winner is sort of putting tooling, harnessing, and so on around this kind of new wild animal that we discovered in LLMs.

Let's see...

Okay, there's a question from LC.

regarding AI and LLM plateaus, beyond capabilities, but rather efficiency-wise, what do you think will unlock order of magnitude efficiency gains? Architectures, algorithms, hardware, concepts, and so on?

I think... Well, I think it'll be a combination of things.

And I don't know, you know, there will be continued engineering progress. This is the story of what happens in engineering, you know, you look at microprocessors, you could have argued Back in the 1970s, we got what we've got. But no, there are more engineering tweaks you can make, more clever ideas in engineering, and the result is we've got things, you know, a million times faster or whatever.

And, I think that will also be true, even given that you're running, essentially, an LLM, It's like, well, let's make a better GPU pipeline, let's make better hardware for the GPUs, let's use reversible computing to reduce the power usage, let's, make kind of little pieces of sort of symbolic

computation interwoven with, what happens with the, with the LLMs. That's kind of more of our kind of area.

there are lots of things which will allow various kinds of optimization.

what will be the most important piece of optimization? I'm not sure. I mean, I think whether it was convolutional neural nets in the 2010s, or transformers in the 2020s, these were kind of ideas about how to organize kind of neural nets that turned out to be very valuable

In terms of going from something which could, in principle, have been done.

but could not realistically be done with the hardware and training times that people had in mind.

So I think, you know, I think there's going to be a mixture of things. I think, I mean, with respect to the hardware.

There's certainly power consumption constraints. There's a certain amount of burn capabilities into hardware that previously required, sort of, software choices. That's maybe a factor of 10, maybe even a factor of 100. There's,

simplify the neural nets, because a lot of what's there isn't necessary, and I think one of the really dramatic things will be just, you know, use our technology to get a lot of the knowledge and sort of computation out, and have the... really make the LLM

be doing what it does best, which is this kind of linguistic interface, common sense wrapper, rather than all of this detailed stuff. I mean, the idea that one is grinding in the training of an LLM

to teach it the third digit of the population of Paris or something is absurd. That's not how it should work. It should instead be getting that kind of information from its sort of neural implant, say, from our technology, and then doing what it does best of providing this kind of smoothing. kind of linguistic interface to the whole thing. I have to say, I think that for some use cases, that will have a dramatic effect.

I mean, I will say there's a project that we're working on that I'm not really talking about a lot yet, we've been working on for a few years, that, will provide kind of a more of a symbolic backbone

for a lot of what LLMs do. We don't know it's going to work, but it's a thing where one's really having a structure on which you can kind of put sort of the wrapping of this linguistic interface. I think that's,

that's a direction I... I'm not... I think that will mostly introduce a lot of new use cases. I... it will also make some things a lot faster, but I don't think its main story is being a lot faster. I think its main story is a lot of new use cases.

I think that,

So I suppose in Elsie's question, that would be more on the conceptual end of what can be done to optimize things. And yes, there will be important thresholds, like when can you really expect to run a good, state-of-the-art LLM on your own computer? When can you expect to run it on your own smartphone?

Those things will absolutely happen, and they will be sped up by using our tech to sort of offload the computational knowledge side of what has to be done. And, yes, there are a bunch of projects going on along those lines right now.

Let's see...

Char says, I see a big problem that there is no out-of-band communication with LLMs. How will we solve the prompt injection

Problem without it.

Yeah, I mean, it's an interesting point.

That, if you're dealing with a computer.

There are all kinds you can do, things you can do. You can have interrupts, you can have, you know, different, sort of things that watch the, in the operating system that watch different kinds of things. In LLM, it's rather black box. It's kind of like the neural net is doing what it does, and you can't be expected to understand it. In fact, I think I have pretty good evidence that a lot of what's going on is sort of lumps of irreducible computation

being put together, and they really are things where it just... the computation does what it does, and it comes out a certain way, and if you say, why does it come out that way? There isn't really a good human narrative for what's happening there.

So, yes, I think it's an interesting point that with LLMs, we don't have this kind of window into what they're doing, which allows us to sort of say, hey, wait a minute, LLM, you're going off in the wrong direction.

It's something where we really have to build harnesses around the LLM and say, yeah, LLM, you did the wrong thing, now go do it again, for example, rather than getting inside the LLM to see what's happening and to put it back on course.

Let's see...

Gary asks, for AI-generated videos and pictures, do you think there should be legal requirements, like a watermark, to denote this is AI-generated content?

That is an incredible can of worms.

You know, for a long time, cameras have had, you know, red-eye removal.

Oh, actually, they don't need it anymore, because they're not using Flash.

But, you know, your average camera today allows even a lousy photographer like me to take a pretty good picture, because it's making a bunch of adjustments, which you could think of as... which are a bunch of AI heuristic adjustments, but yet you say, but it really is kind of just a camera taking the picture, but it's a camera kind of enhanced with lots of AI things.

Now the question is, well, you know, if I really went a bit further, and, you know, people often do that when they have filters, when they're using, oh, I don't know, Zoom or whatever else, they go a bit further, and they're like, well, actually, let me make it so that I look better on whatever it is, or look like I have, you know, horns, or something of this kind.

It's, you know, then there's sort of AI addition.

And what's... it's kind of a sliding scale. There's no, you know, zero AI or sort of 100% AI.

Maybe there is 100% AI, and maybe zero AI would be the raw bits, the raw pixels that come from the camera without any processing, which is not happening at all at this point.

So...

you know, at what point should you label it, sort of, AI-made? It's similar to the question of, you know, when you buy something, and it says, you know, made in France, or something like this, what does that actually mean? You know, how much of it was made in France?

You know, was it assembled there? Was it, you know, did the... did the raw material come from there? Did every part of the raw material come from there? And there are a bunch of rules about when you can sort of advertise made in country X.

Or when you can not advertise Made in Country X.

So, you know, I think it's a bit complicated. I mean, there's a bit to describe in what was done. you know, I think...

in... in things like academic papers, you know, it always used to be the case you write a paper, you just put your name at the top, you put a bunch of references at the bottom, and that is

everything that's said. It became a little bit more conventional to explain a bit more of the story of

person A did this, person B did that. Now it is a little bit more of, well, we used a bit of AI to do this, that, and the other. I always think that's kind of a... it's a crazy thing, because, you know, people who use our tech, which an awful lot of people who write scientific papers do.

you know, our tech has what, in some respects, you would call AI inside. I mean, it's very, you know, kind of organized AI. It's not going to go off and make up some wild thing for you.

But it's nevertheless the case that if you do, I don't know, image processing with Waltham Language, or you do feature space

plots and so on, that's using AI to decide how things are organized, or using something like AI to decide that. And if you, you know, sort of change something, it could change. And so I think it's

a... it's a very complicated issue.

you know, what counts as the raw story versus a thing that has been processed by some level of AI.

I think,

Yeah, I mean, you know, I suppose it's a little bit different if it's, like, you just gave it a textual prompt, and then it generated a picture.

And, you know, one could imagine a case where if the... if the raw material is a different modality, like text, for example, and the image

is sort of generated from that. One could imagine that that is a sort of different case. If the raw modality was an image from a camera that was then processed, that's one thing.

If the thing that went in was just text, that maybe is a different thing, and maybe one could make a distinction there between things where there are, you know, 10 million pixels.

versus two sentences of text. Those are different levels of amounts of, of, kind of raw information that are going in, and I suppose one might be able to make some distinction based on that.

What...

kind of labeling, I mean, it very much reminds one of the kind of health labelings and so on that get put on lots of kinds of products, and, I suppose in some respect, boy, those are... those are... probably achieve... achieve some things, but, you know.

To have,

every picture that's sort of a... you know... Okay, there's another point, which is there are settings in which

This picture is supposed to be a real picture, and there are settings in which it's not expected to be a real picture. If it's the icon, if it's sort of an illumination for a section of a blog that's just sort of a highlight picture.

You know, nobody expects that to be real.

You know, nobody expects that to be a... it's just a thing sort of accenting, it's like a medieval illumination in a manuscript. It's just a sort of a picture, it's sort of eye candy, to help sort of theme the thing. Very different from, you know, here's the result of our scientific experiment where we

Tried to, you know, teach a mouse to paint, and this is the painting it made.

and you show the painting, and oh, actually, that wasn't generated by a mouse, it was generated by an AI imitating a mouse painting, or something like this. That would be a very different kind of story. And I think, you know, I think that's another point, that there are places where you might reasonably expect that it's something real. You know, it's a news article.

there's a big report at the top about how, you know, Person X was in place Y, and talking to person Z, and you know, shaking their hand or something, and that never happened. And, you know, you have an implication that this is a real picture, but it doesn't play out that way.

So, it's a complicated thing, and I... and I think the,

You know, it, sometimes, you know, for images, there's sort of various efforts to.

use, actually, blockchains and things, or at least digital signatures, to be able to sort of assert that this picture was taken in this place with just Geo...

geolocation was taken at least before this time that's relevant to putting things on blockchains and so on. So there are definitely things that can be done there. I have to say, I thought that would take off in a more serious way a few years ago. Hasn't really taken off that strongly, possibly because

there's sort of, in many cases, it's kind of like, well, we kind of know this isn't a real thing, and so on. And, you know, I don't know what the... what the front lines of... of kind of, the... the cat and mouse game between, you know, kids in school.

generating essays with AI, and teachers trying to say, that's an AI-generated essay. I don't know what that, what the current state of the... those battle lines actually are. I tend to think, you know, when I get things, which I do with a depressing frequency, where people send me these, these letters and theories and so on, where it's just like, this was written by an AI, and you can tell pretty much immediately that it was written by an AI. There's certain kinds of things that just... there's almost... it's like the, this carpet was made by a machine, it has no errors in it, type thing.

And, to some extent, that's the same with a lot of AI-generated, kind of very anodyne content.

Let's see...

Well, there's a question here.

From Ludson.

That, about, hardware for neuromorphic computing and AI, would it drift towards analog computing?

And, or will it stay digital?

You know... There have been many attempts to make Analog computation really be useful?

And in the past, it's always failed.

It's... it's kind of like...

You know, when you're doing a simple operation, you're making a power supply, something like this, you're trying to transform from one voltage to another, well, then it's sort of a one-shot thing.

If you're trying to make a giant, sort of, computational tower, then...

you kind of need each step to be very accurate, or this tower is going to fall over. And insofar as the things you want to do involve certain amounts of irreducible computation, you kind of have to have these accurate bricks to make your sort of computational tower out of. And I think the, the thing that tends to go wrong with analog computation is you just don't quite know what you have, so you can't make that kind of tower. You can say, if all I want is to get roughly that voltage out, great, do it with analog. I have to say, I once in my life, I...

I used an analog computer in a serious way at Los Alamos National Lab. This was in the early 1980s.

And I was kind of like, I'm going to study dynamical systems and differential equations and so on with this analog computer.

And after a while, it was like, this is just too unreliable. I, you know, I don't really know what's gone into this, I don't really know what's coming out, I'm just gonna run it on a digital computer. They've been...

You know, in terms of... there are many kinds of things that you can optimize a bit by saying, let's not have that be a perfectly formed, you know, 1 or 0, let's have it be a bit analog, and we'll fix it up in a little while to make our sort of solid brick of computation.

But so far, I'm not really seeing kind of a big story of, analog computation. I mean, there are various efforts to use, kind of, random noise and things like this to help with,

kind of the neural net type algorithms that seem to involve randomness and so on. I don't know how valuable that is. I haven't really thought that through properly.

But I have to say, history does not support the idea that analog computation is going to work. Now, you know, history doesn't support the idea that one's going to be able to make something like an LLM.

So eventually, even though there were sort of precursors to LLMs going back 50 years with language models and so on, nobody knew at what point it would sort of fluently achieve liftoff, so to speak, and maybe there'll be such a point for analog computation, but I haven't seen that coming yet.

As, question here from, C.

Is there any credence to the idea that there might be a kind of emerging agency in LLMs analogous to the notion of emergent agency in bubble sort?

Not sure how much agency there is in bubble sort. Bubble sort is an algorithm for, sort of locally, locally reordering things so that in the end, after enough bubbles have gone through, you will have reordered a long sequence of, of items.

I think the question of what agency is, is a question I was,

Just recently talking to a philosopher I've known for a long time, who's spent much of his life thinking about that question. What is agency? I'm not sure I know a good, really good answer to that.

So I think it's one of these things that's a rather a slippery concept of, you know, does... is agency achieved? Well, we don't really know quite what agency is, so that's something that is going to be very hard to... to discuss.

And if it's, like.

oh, I didn't prompt the LLM to do blah blah blah, but it did it anyway. It has some internal agency that it's going to do it. Well, that's not really a convincing argument, because it's had all that training data, and if everybody says, well, in that situation you do X, then that's what the LLM is going to trot out. It's not something which it had some internal agency to do. Now, you can start arguing about humans.

Do we have the internal agency

to, to do this or that thing, independent of, kind of, what we've learned. You know, if somebody, if you're going to, you know, hold a door open for somebody.

Is that something we would intrinsically do, just from our own sense of agency, or is that something we've learned as a social convention, so to speak, from our actual experience in the world? I think in that case, it's very overwhelmingly the second thing

Not that we have some built-in or deducible-from-nothing belief that that's how we should act.

Zach asks, have you seen any new AI-generated math that actually works?

not really, but let's qualify that a bit.

I did get something to work, in a thing that I wrote quite recently about lambda calculus. I had a sequence of integers.

And I wanted to know, is there a formula for this sequence of integers? And I asked an LLM, actually asked it for 100 different sequences. For about 2 of those sequences, it was a win. For some of those sequences, tools that we have already, computational tools we have, or lookups from the Encyclopedia of Integer Sequences or something, already get the answer. But there were 2 out of 100 of the things I tried were sequences where I couldn't get a sort of a nice formula by any of those methods, and by golly, an LLM thought for a long time, and it gave me a formula. And the formula was... well, actually, I think it was more or less correct. It was easy to fix, insofar as it wasn't quite correct.

What did it do?

you know, it was kind of... kind of remarkable to me. It's like, cool, this is really a thing. But I think what was really going on

was...

In the end, there is a paper out there which probably has a sequence like that and a formula like that, and this is a thing that's coming from kind of a thematic searching of the scientific literature, sort of surfacing this particular kind of result. But I... that's a case for me, where it was a real use case where it did something valuable. And I mean, there...

a bunch of things where there are problems, where if you can sort of guess the answer, you can check that it's right. And in fact, in the version of language that's coming out soon, we have a number of things that use kind of AI methods to sort of heuristically guess answers in more ambitious ways than we've ever done before, and be able to then go back and check, is it in fact the right answer?

If yes, well, here's your answer. If no, well, you've got to do something else.

Now, you know, when it comes to, sort of, discovering New math.

There are a couple of issues.

One is... kind of...

when do you just sort of go out into metamathematical space and try to, you know, find, kind of theorems that people are going to care about? It's very easy to generate theorems

Question is, which ones might people care about? And usually, you have to have a seed of an idea. Then you have to say, well, let's go find a thing that matches this kind of idea, that matches this objective I have.

The one case where I did that very successfully was 25 years ago, in the year 2000, I did a thing where I was searching for the very simplest possible axiom for logic for Boolean algebra, and I had a guess that it was fairly simple, simple enough that you could just enumerate possible axioms and find one that worked.

And by golly, I managed to do that, and found this very tiny axiom, which we know is the simplest axiom for Boolean algebra. And I used automated theorem proving, sophisticated computational technique, to prove that that axiom was correct.

Could LLMs help with that?

earlier this year, I thought, well, maybe they can't help with actually finding the axiom and the proof, but maybe they can help in explaining the proof of that axiom.

And, so the proof, as found by automated theorem proving, is long and complicated, and no human has ever understood what it really says. You go just even a couple of lemons into it, and it's quite incomprehensible. So...

I did a project earlier this year trying to get LLMs to sort of find, kind of anchors, find waypoints that were sort of human-explainable features of what was going on. It was a flop. Didn't work. I mean, I found a lot of other things about the structure of that proof.

and what's going on in it, and it really is a sort of computationally irreducible proof. It really is something where it just happens to work, but there isn't really an easy, compressible narrative for what's going on.

So, you know, there are a bunch of companies, I've been involved with some of them, that are trying to sort of take LLMs and sort of assemble proof assistant code from the LLM. So far, I haven't really seen something spectacular from that. It's sort of the best effort.

Is to say, well, can we take something which is an informally written thing, and can we make something which is sort of proof-assessment style out of it?

The problem with that is that,

people have had to build lots of libraries for proof assistance and things without good design for those libraries, without the kind of thing I've spent the last 40 years doing for Wolfram Language and Mathematica.

you really end up with something which is... which is not easy to understand. And it's like, when you have a proof that's like the automated proof that I generated, and it's that complicated, and it's like, what do you really have here? You know the result is true, but the proof doesn't really add anything to that.

And I think...

that's, you know, I think the... perhaps the underappreciated piece in this whole story is, if you're gonna sort of formalize things from, you know, have an LLM do the formalization, you should have a target that is really a well-designed target. And that's really a language design problem, and yeah, you can get the LLMs to help a little bit with that, but really that's a human problem to define how that should work.

Now, an example that we have that worked really well is Wolfram Alpha.

Wolfram Alpha can take, sort of, natural math of things random people type in at the rate of zillions every day, and, you know, they're kind of shoddy, and they have spaces missing, and wedges in weird places, and parens in weird places, and so on.

And we can do a very good job of the 99 point something percent level of taking that kind of natural math and parsing it into our precise symbolic language from which we can then do computations. So that's an example where once you have the target.

Once you have that good, solid, designed target.

Then you can expect to take, well, sort of.

AI-ish methods and do that parsing. Now, Wolfram Alpha doesn't use LLM methods, it uses other methods that we invented years ago.

for doing that natural math understanding, I don't think it would benefit a lot. We've done a whole bunch of experiments to see whether it would benefit from modern LLM methods, and so far, those experiments have all been a bust.

Really, the things we have are more reliable, stronger, faster, etc, even for doing that step of going from the kind of vague natural math, or natural language, talking about math, down to that symbolic level.

That's not always true. For example, when it comes to things like word problems, where you say, you know, Bob has 17 marbles and gives 4 to

to Jemima, and, who puts them in a bag, you know, how many marbles are there out now? That's something that has a lot of, sort of, fluff around it.

that an LLM can... can help turn into an actual equation, not a very complicated one in that case. And that's something we've been doing lots of experiments on, and I think will actually be rolling out as a part of Malfa in the not-too-distant future. But that's kind of a corner... a corner case.

When it comes to taking a, sort of a mathematical theorem.

And, or a, kind of a proof that somebody wrote in a paper, and turning it into a, something more formalized, that's... that's something where you've got to have a target. And if the target is good, you might get something really useful. If the target is a huge mess, I don't think it's going to be that useful. I mean, there's a big effort right now to take Fermat's Last Theorem, the proof that Andrew Wiles gave for that, and formalize that proof.

And to me, it's a, it's an interesting, you know, it's interesting to try to do something like that. There's a lot of questions about what it really means. For example, the axioms that are being used for the proof assistant are not traditional axioms of mathematics. They're computational axioms that are to do with the structure of the proof assistant from type theory and so on.

And so that's even a question of what does it really mean if we proved, kind of, Fermat's Law's theorem from this type theory thing.

It's interesting, I mean, it means there probably isn't an outright simple mistake, but could it be the case that from the axioms of piano arithmetic, the sort of lowest-level axioms in mathematics, is that enough to prove Fermat's Law's theorem? If you're going to, sort of, foundational questions, that's the type of question you should be able to answer, and I'm not sure you will be able to in this particular case.

So I think, in, You know, there's a... Well...

it is, actually, having been away for a month, I haven't had as much chance to check in with the various companies that,

we've been interacting with who are trying to do, kind of, math with AI, and I'll... maybe I'll have more to, to say about that when I see what they've managed to achieve, but I'm... so far, I've not been optimistic about what's... what's been being done there.

Right, maybe a couple more things, Doug.

Okay.

Let's see, maybe something a little different here.

Let's see... Katie asks, how likely is it?

That the rate of scientific and technological progress of the last two centuries will be sustained in future, or even accelerate.

There have been long periods of stagnation in history.

You know, it's always complicated to know what you mean by progress.

The fact is, if you've got one measure of progress, because look, this is the thing that's been happening. More and more apps have been loaded up into app stores. We've got progress, so to speak. Then you realize, well, actually, they were all pretty boring in the end.

Or something. And then you say, but now the new thing is something completely different that wasn't yet a thing you were measuring.

So, you know, if you say, well, we've got so much progress, there are lots and lots of scientific but academic papers being published, well, except, yeah, but a lot of those papers are really dull, and maybe complete nonsense, and so on. So, you know, that metric doesn't really indicate a lot of progress.

The fact is, I think that it is probably subjectively and objectively true that the rate of new things being invented has increased, I think.

But sometimes, you know, the kinds of things that are invented are kind of organizational structures, or social dynamics, or things like this, that weren't part of the thing you had in your crosshairs, where you were thinking about, you know, are we inventing more science, or whatever?

So I think, I mean, my...

You know, the general pattern is.

New methodologies invented, opens up new areas, lots of new things get figured out.

Will new methodologies be invented? Yes, they'll be invented forever. That's kind of a result of, kind of, ideas about computational irreducibility and the pockets of reducibility that inevitably exist there. It's inevitable that there will always be more inventions to be found. Now the question is more of a societal one.

Does anybody care? It could be the case that you get to a point where people say, hey, I've got enough.

you know, my, the... I don't know, the piece of paper I'm using. It works just fine. You don't need innovation in papermaking.

or the, you know, slightly more, the smartphone I got 5 years ago is still fine.

for me, practical one, the printer I got 10 years ago is still printing just fine. I don't need a new one. Actually, I just got a new one for random reasons, but generally, that's something that where, you know, the technology curve has kind of leveled off.

And, but then there'll be a new thing you need. You know, I didn't think I needed a smartwatch until I did, so to speak. Or I didn't think I needed some other kind of thing until I did.

So, you know, I think this question of whether people will keep thinking they need new things, that's more a societal question than it is a technological question. I think it's true in science as well.

You know, if you left it.

to the sort of giant institutional structures of science, I'm not sure there's a lot of motivation.

to do truly innovative things. The big institutional structures of science are machines that keep grinding along and spending lots of money, and it kind of circulates around certain collections of people, making incremental progress, to be clear.

But if you say, is there going to be some great innovation made? Well, that machine isn't about making great innovations, and in fact, that machine makes it more difficult.

to make great innovations, because you have to say, hey, wait a minute, I want to divert the big machine from what it was trying to do to do something new. It's a challenge always... it's a challenge in companies to do new things, it's a challenge in institutionalized science to do new things. And so, in a sense, you can be the victim of your own success. You start a science and it becomes so big that it has a very definite direction, and then if you say, well.

wait a minute, there's a new thing that could be done here. The, that, that's not a thing.

that the machine of that science can provide. I mean, this sort of happened to me with things I'd done in complex systems research and so on, back in the 80s, initiating that, and then sort of returning to the field and injecting new things in the beginning of the 2000s. It was, like, it was a little difficult to deal with the machine that had sort of built up in my absence.

When it came to injecting new kinds of things like that.

So... I mean...

there's also... so, as I say, I think it is more a societal thing, whether there's sort of a push to have the new. I mean, realistically, when go... when times are tough.

There's more push to innovate.

And, you know, if people are fighting with each other in wars, or if there's some, you know, terrible natural situation of, you know, the Ice Age is coming on, or whatever else it is, there's more motivation for innovation, because the value of innovation relative to survival for example, is higher. And, so...

So that tends to cause more innovation to happen. In times when, sort of, everything's just fine, the... maybe there's less push for... for innovation. I... I tend to think that, there is always a certain segment of the population, or always has been, that sort of want to do innovation, that kind of have this crazy idea that, hey, we can discover stuff, we can invent stuff. They may or may not be embedded in an environment where that's realistically possible. And it may be that, you know, they're just having to spend all their time, you know, foraging for food to get enough to eat, or whatever, or that they're in a situation where, yes, you've got to go with the flow, or you'll lose your job, or be exterminated, or whatever else.

So I think the,

You know, that, that would be my, and, okay, from a societal point of view, we have come to be used to progress, at least in many in many parts of the world, there's the idea that technological progress, in particular, is a... is a... is an everlasting thing, that there will always be technological progress, we can always expect technological progress. A more interesting question is social progress. There are things which people have felt with social progress.

That maybe weren't terribly good ideas, and maybe kind of

Sort of went against the human condition as it is, you know, based on our sort of biological history.

But nevertheless, there has been a notion in the last hundred years, particularly, that there's sort of, just like there's seemingly inevitable technological progress, seemingly inevitable scientific progress, that there should be seemingly inevitable social progress. And...

You know, one can be sort of ideological about the whole thing, and I would say that some things that have gone under the notion of social progress, I personally think, are pretty good ideas. Other ones, I think, are pretty terrible ideas. I think that, but the notion that there will always be social progress is not nearly as clear.

as the notion that there should... that there has been and will be technological progress. I think, you know, we... we have a certain human condition, and, you know, to do things that

Kind of, that... Sort of work with our human condition, great. To do things where

You know, it's like, the new isn't always such a good idea, because we're stuck with a human condition that is what it is. And maybe we can educate ourselves out of some parts of that.

But I tend to feel like that's, it's a different thing from the case of technological progress or scientific progress, but there's a very obviously, sort of endless frontier.

Let's see... Gist asks, how do we deal with progress being so fast that nobody understands it anymore?

Well... I don't know, I certainly try. I don't feel... I don't feel...

like, at least in the domains that I pay attention to, I don't have this feeling, oh my gosh, I'm chasing after the progress and I can't keep up.

I feel like...

once one has a good base level of knowledge about a lot of kinds of things, that the things one sees, yes, something got solved. Well, actually, I've known about that for 30 years, finally it got solved, great. Now I understand that piece of progress.

I think that there are certainly domains where, you know, in...

well, having said that, you know, you probably could catch me out on saying, you know, some new thing in video games, or some new thing in social media, which are not things I'm not exposed to. You know, can I keep up with that? I suppose I have a harder time keeping up with, you know, the social media dynamics of the kids, so to speak, than I do with, kind of, the new technologies in optics, or something like this.

I think... That, you know, is... if the...

This question of... of can one kind of...

is there a leading edge of progress that is sort of incomprehensibly fast? It certainly doesn't seem that way to me.

It, sometimes, to me, it actually seems like things have been painfully slow. One's known this is a possibility for a long time. I suppose what makes it all seem a lot less, sort of, shockingly fast is having really learnt about a lot of possibilities for a long time.

And so when they arrive, it's not like, oh my gosh, you know, we just can't keep up. Something new is arriving that we didn't expect. It's like, well, we did kind of expect that. Now there are surprises.

You know, ChatGPT was a surprise. Nobody expected it. It was, you know, one had... I had tracked language models for, I don't know, 40 years or so. They'd been not that interesting, not that... not that impressive.

You know, just, you know, good for filling in words if you were doing, you know, speech-to-text or something like that. Good for doing, kind of, maybe a little bit of spell checking, grammar checking.

things like this, but not the kind of, I'll write a fluent essay type thing that was achieved with ChatGPT. That was very unexpected, and seemed very disorienting. And once... once a thing like that happens, you start worrying, oh my gosh, everything's gonna run away from us. I don't think that's...

Really, what is... has happened, or is going to happen, but there are surprises from time to time. Let's see...

Elsie asks, what would be the last traits of a top researcher to be subsumed by AI, if at all, say, taking myself as an example, thank you for the top researcher plug.

Is it going to be agency, novel connections, insights, directions for inquiry? I think, fundamentally, the non-automatable thing is what is it you want it that you want to do?

In other words, the question of what to study

And kind of what question is worth asking is a very human question. You know, the AI can say, oh, I've figured out these 10,000 things. You know, when I go out and kind of do ruliology and go out and sort of sample things in the computational universe, I can get 10,000 very wild, interesting, kind of, sets of behavior and so on.

But this question of, do I really care about these? You know, how do they relate to, sort of, objectives that I might have?

I think that's the thing, and really is... For a good researcher.

The question of what to research and the strategy for doing it is the absolute determiner of whether you're really doing something important. The actual mechanics of doing it, yes, that's important, but it pales in comparison with the strategy of what to do and sort of the outline of how to do it.

Now, you know, for myself, I think I've said this before, you know, I consider myself to have lived the AI dream for 40 years, because in research that I've been doing, my role is figuring out what I want to do, and then a big part of the doing of it

I've been able to delegate to technology. I've also had actual humans who help me, which is very nice too, but by far the dominant effect is the stack of technology that I built, and ideas around that technology that's embodied in Wolfram language.

And so, you know, for me, I can go very fluently from, you know, an idea about what I want to do to sort of automating the actual doing of that thing. And that's... that's something where I don't feel like, oh my gosh, I'm... I'm being sort of overwhelmed.

by... but instead, it's like I have this incredibly powerful sort of superpower that I can just move my finger a little bit, and this giant thing will move around, that is the kind of the Wolfram language, computational language.

thing that, that, I am, I am making happen in the direction of the strategy that I've defined.

Alright, I should probably wrap up there.

I like to remind people that...

Well, coming up next week.

Wednesday through Friday is our annual Wolfram Technology Conference, which is online this year, and people are encouraged to

to, to come check it out.

I will be doing a keynote

Apparently on Wednesday at 1pm Eastern Time.

Where I'll be talking about the latest in our technology and the directions that we're going. I'm going to do some work on that. It's always a challenge to pull together kind of a report on what we're up to, because we're doing an awful lot of stuff.

And but anyway, that will be a sort of a coming attraction of something I'll be doing as a livestream, next Wednesday.

I might comment that, before I left on my trip here, I'd been doing a rather long, kind of, storytelling of, sort of, stories from my life, sequence of livestreams, and I'm happy to, hear that people have liked those, and,

I'm kind of up for doing, slightly refactoring the kinds of livestreams that I do, and I'd be interested in input about what people would like to see. I've considered a variety of different kinds of things, and

the main constraint for me is some of those things are much harder work than others. What I am able to do in these Q&A livestreams

is, you guys give me all kinds of interesting questions, which I get an opportunity to think about, but all I do is kind of sit down in front of that camera, look at it, and start saying what I can to answer these things. There's no homework required.

And there's, I would say.

that, well, perhaps because I've been doing it for a while, this is a low-stress activity, even though, at some level, it's kind of like, my gosh, what questions are people going to ask, and how am I going to think my way through and around the different kinds of ways to answer them?

But, for me, this has become a, I have to say, a very pleasant and relaxing part of my time. And so some things that I might consider doing, feel like a lot more work, but maybe if I got into doing them, they wouldn't... wouldn't seem like that.

For example, doing more

that actually involves real-time, live coding and so on. Seeing computations happen in real time is one thing that I might consider, consider doing.

Another thing that is a lot more work is having discussions with other people, actual interview kinds of things. I've done a few of those in the past. Those can be great, but they can be a lot of work to do a good job preparing

as I hope people will prepare when I do interviews with them, and and so on. Anyway, coming up next week.

our technology conference. Hope you'll be able to join that live stream, and it's, I look forward to chatting with you all about all kinds of things in the future.

Thanks very much for joining me, and bye for now.