

Transcript | Artificial Intelligence 101

1

00:00:09.240 --> 00:00:18.290

Lowell Weisbord: Hello! And welcome to artificial Intelligence 101. We're going to give everyone a few minutes to stream in here.

2

00:00:18.520 --> 00:00:35.140

Lowell Weisbord: but over the next hour we'll unpack what AI really is, where it's already showing up in government, and what every public servant should know to navigate it confidently again. Hello, and welcome as folks stream in here. We have a really great

3

00:00:35.170 --> 00:01:02.930

Lowell Weisbord: panel for you all on AI. 101 again. This is a 101 event. So I'd like to underline. We're all very much here to learn. There are no bad questions. This is nascent technology that we're talking about that. Nobody is really sure about how to use or how it will be used 5, 10 years from now. So don't be shy. We're going to have a really brilliant Q&A. With our 3 speakers here, so feel free to keep the questions coming

4

00:01:03.140 --> 00:01:04.940

Lowell Weisbord: as we get into it.

5

00:01:05.080 --> 00:01:30.229

Lowell Weisbord: We'll give people a few more minutes, but my name is Lowell Weisbord. I'm a senior product manager here at Apolitical. I work on some of our AI stuff, but it would be really great to hear from everyone where you all are calling in from today. Appreciate somebody. Just got it going in the chat. Welcome, Sadia, from Toronto. It'd be great to hear where everyone else is coming in from. I'm currently in London on this

6

00:01:30.320 --> 00:01:39.670

Lowell Weisbord: kind of humid afternoon. Hello! From Indonesia, Cardiff, Brussels, Bermuda, England. We have really good start, beautifully international crowd.

7

00:01:40.110 --> 00:02:02.789

Lowell Weisbord: For those of you who may be joining us for the 1st time. Welcome! Apolitical is the world's largest community for government, used by over a quarter of 1 million public

servants and policymakers from over 170 countries, of which I think we have really good representation today already in the Chat, Brazil, Canada. Good morning. Lots of Canadians, Newfoundland beautiful.

8

00:02:03.660 --> 00:02:17.830

Lowell Weisbord: Our mission is to make government smarter. We do this by putting the best knowledge and skills at the fingertips of public servants around the world. This event is part of our Government AI campus which aims to train a million public servants in order to create an AI ready public service.

9

00:02:17.900 --> 00:02:45.409

Lowell Weisbord: The campus is a trusted knowledge and learning hub for governments, and the world's 200 million civil servants. On governing and applying the technology, you can take up your kind of place on the campus which has a lot of resources, courses events like this one. By following the link that will pop up on your screen in a second. You can check out that link in the chat there for the rest of our resources. Just a couple quick technical notes to get out of the way

10

00:02:45.520 --> 00:03:05.960

Lowell Weisbord: this event is recorded, for on demand viewing it will be on the apolitical website tomorrow. Please don't record it yourself. It will be free and accessible to everyone who works in government contact. Details for privacy concerns are in our apolitical homepage footer and closed captions can be available if you toggle them on in your screen kind of zoom settings.

11

00:03:07.910 --> 00:03:29.939

Lowell Weisbord: very excited to get into it. But before we introduce our speakers we're going to run a few quick polls to kind of get a sense of the room. Everyone's current use of AI people's current understanding. And that'll kind of help shape the conversation today. So our 1st question is, how often are you using generative AI tools like Chat Gpt Gemini, copilot in your work currently.

12

00:03:30.540 --> 00:03:44.830

Lowell Weisbord: And we'll give everyone a bit of time to answer this again. We'll talk about what generative AI even means probably in this, in this you know event and and what the kind of differences, and why? This is a question, and

13

00:03:45.020 --> 00:03:48.220

Lowell Weisbord: how all these tools came onto the block.

14

00:03:48.840 --> 00:03:53.009

Lowell Weisbord: but already it looks like we have quite a balanced

15

00:03:53.510 --> 00:03:58.100

Lowell Weisbord: response here, which is, which is really interesting. We have.

16

00:03:58.880 --> 00:04:25.889

Lowell Weisbord: I don't know if people can see the results. Yet I can. Okay, now we can see the results. So we have really, really balanced results here, which is pretty interesting. 28% daily, 26% weekly, 28%. Not that often in 19% never use them. So we have the full spectrum of folks on this call, so plenty of learning from person to person. And I know that this might be for a bunch of reasons whether it's what you're allowed to use at work, and we'll talk about that, I'm sure.

17

00:04:26.040 --> 00:04:35.259

Lowell Weisbord: What tools are available to you, etc. The next quick question we have here is, have you received any specific training on how to use AI in your work?

18

00:04:36.700 --> 00:04:41.780

Lowell Weisbord: This is quite an important one, too. I think a lot of people are aware of concerns,

19

00:04:43.090 --> 00:04:51.080

Lowell Weisbord: information, security data, lots of things. We are getting the results coming in very quickly.

20

00:04:54.310 --> 00:05:16.230

Lowell Weisbord: we have basically 3 quarters of folks saying, No, I haven't received any specific training, despite us having 80% of people saying they've used AI in their work. So that's pretty stark right there. And that's also kind of what we're. We're all about here with the AI campus and the point of kind of getting this going.

21

00:05:17.600 --> 00:05:34.570

Lowell Weisbord: poll 3. It's different from the second one. Have you received official guidance on how you're allowed to use generative AI, so not necessarily training on how best to use it?

But real clear guidance on what you might be allowed to put in, not put in what tools, specifically, you might be allowed to use.

22

00:05:36.820 --> 00:05:40.669

Lowell Weisbord: It'd be really interesting to to kind of hear what people are

23

00:05:40.860 --> 00:05:43.319

Lowell Weisbord: getting on. Got some more hellos.

24

00:05:43.640 --> 00:05:46.650

Lowell Weisbord: Good day, everyone. I'm in Portland, Oregon. Welcome, Leia

25

00:05:48.420 --> 00:05:57.369

Lowell Weisbord: only co-pilot allowed at my org. That is a really common thing across many governments at the minute. Actually, there's a lot of co-pilot co-pilot users.

26

00:05:58.047 --> 00:06:01.389

Lowell Weisbord: We'll see what people are saying here, though.

27

00:06:02.418 --> 00:06:13.400

Lowell Weisbord: so we have 50% again, with yes, on official guidance, 38%, no, 12%. I'm not sure again, worth kind of looking at that difference between

28

00:06:13.530 --> 00:06:28.821

Lowell Weisbord: 38% of people saying there's no official guidance. But obviously, people are using these tools in their work. They're becoming kind of important pretty quickly, and that's interesting, and that'll be some of the framing for why we're having this conversation, and where we might start

29

00:06:29.740 --> 00:06:54.190

Lowell Weisbord: But more importantly than that, I'm going to talk to you about our 3 great speakers who are going to be talking to you all and answering all your questions today. We're fortunate to have 3 guests, well versed in kind of artificial intelligence here, and how it might apply to government in the kind of latest wave. Barbara, Joseph and Cesare, Barbara Bramova is a data and AI project analyst at Undp.

30

00:06:54.190 --> 00:07:17.189

Lowell Weisbord: the development program and helps governments and local partners develop and scale community driven data, governance solutions for responsible artificial intelligence.

Welcome, Barbara. Dr. Joe. Cattle is a technology consultant strategist and engineer has had leadership positions at the White House and the Us. Gsa. The General Services Administration. Welcome, Joe, appreciate having you.

31

00:07:17.240 --> 00:07:40.409

Lowell Weisbord: And finally, Cesare Gesykowski is a senior advisor in AI and Innovation capability development lead at agriculture and agri-food Canada, where he has developed and is implementing an innovation automation and AI acceleration program. Welcome Cesare. You can read their full bios. I'll drop a link in the chat if you're more curious about our speakers.

32

00:07:42.350 --> 00:07:47.079

Lowell Weisbord: But a quick icebreaker for all 3 of our panelists before we get into the juicier questions.

33

00:07:47.210 --> 00:08:03.779

Lowell Weisbord: I'd like to hear everyone else's thoughts in the chat here as well. So feel free to respond to this. But can you tell us about a novel use of AI that you've seen recently that's impressed or amused you, and I'll go in the order I just started in. So, Barbara, do you mind kicking us off.

34

00:08:04.400 --> 00:08:24.700

Barbora Bromová: Absolutely. Thank you, Laura, and it's great to be here today. I've been really getting into voice technologies lately, which is maybe not a groundbreakingly novel use. But I've had to do quite a lot of personal admin stuff recently, and so like, call assistance on my phone, who are on hold, for you have been very helpful in that, so I would recommend.

35

00:08:26.020 --> 00:08:38.419

Lowell Weisbord: Yeah, really interesting call assistance and and kind of hold assistance, I think would make would make everyone's lives better. But also, maybe we will get into some of that of how it's being used on the front line in government as well in terms of customers facing things.

36

00:08:39.371 --> 00:08:42.280

Lowell Weisbord: Joe, do you want to go next.

37

00:08:43.289 --> 00:09:07.709

Joseph Castle: Sure happy to I had. I had one of mine. The one that I use that maybe isn't the newest is is a co-pilot, but built into visual studio code. So I do. Quite a bit of python development. You know, getting into AI modeling. And and what have you predictive analytics? And what have you? And and that? That is like a game

38

00:09:07.829 --> 00:09:26.979

Joseph Castle: change. Your vibe coding is the term what's interesting is that it hallucinates like everything else. And so you have to get really good at the prompting. And and you have to know some code so that you can validate that. It's doing what you want, want it to do and doesn't get stuck in this loop. But yeah, I use it on pretty much every day.

39

00:09:28.180 --> 00:09:52.469

Lowell Weisbord: Yeah, brilliant. Worth noting for, you know the non engineers on the call, which I think will be a lot of us. Is that. That's the number one use of generative AI right now, it's it's used by in, you know, developing, you can look at the kind of AI economic index that anthropics put out, and it's used in terms of software engineering more than anything else by a long shot in in, you know. July second, 2025, maybe not maybe not for forever.

40

00:09:52.470 --> 00:10:15.990

Joseph Castle: Lola, if I can. Real fast, too. I didn't even think about that until you just said that for the non engineering users visual studio code is an application to develop to to basically write code, develop software, and it's it's an open source product it's built by Microsoft and it's pro. It's probably the most you used potentially. I don't have any numbers on that. But most well known as use. Probably these days.

41

00:10:16.710 --> 00:10:23.269

Lowell Weisbord: Yeah, thanks, Joe. And finally, Caesare, an interesting, amusing use case you've seen recently.

42

00:10:23.690 --> 00:10:32.330

Cezary Gesikowski: Yeah. So recently, I've been testing out the very advanced reasoning models. And I'm very privileged to have access to open AI's o 3 pro model.

43

00:10:32.670 --> 00:10:37.309

Cezary Gesikowski: which is a major upgrade to the one you might have heard the o 1 pro model.

44

00:10:37.330 --> 00:11:00.089

Cezary Gesikowski: and this is great for especially deep multi-step reasoning and visuals and web tools. So in my former life I was a professional photographer, and I also was a digital photography instructor. So I tried feeding it. One of my old digital photos in the O. 3 pro model, and in a few rounds of just like very simple prompts. Nothing sophisticated. It gave me a few full curator's critique

45

00:11:00.090 --> 00:11:18.969

Cezary Gesikowski: like it, reimagined it from multiple curatorial perspectives, and even crafted a creative, low budget marketing plan. So it felt like having an expert art team in your pocket. So it's a little bit outside of the use for public servants. But I think these models are really promising.

46

00:11:18.970 --> 00:11:31.239

Cezary Gesikowski: and they may be very expensive, but for those who are really looking to take full advantage of them, it may actually be a very smart investment in the future. We'll see how that works and develops.

47

00:11:31.770 --> 00:11:46.639

Lowell Weisbord: Yeah, that's really interesting. There's a wide spectrum of use. It can analyze photos now. And things like this, we're getting some really interesting use cases in the chat as well to help with my autism, to help with own personal development and interesting things.

48

00:11:46.680 --> 00:12:07.499

Lowell Weisbord: Write lessons for a post project review report, I used AI to develop an AI training manual AI to learn. AI. I think that is kind of the name of the game a little bit. You have to use it to to figure it out troubleshooting code, and then a few government use cases which are really interesting, which we'll get into. But let's take a few steps back. Now

49

00:12:07.790 --> 00:12:24.719

Lowell Weisbord: again, this is AI 101. We have people who've never used any kind of generative. AI. We have people who use it every day. So we're going to try to cover the spectrum. And again, we're going to get into everyone's questions. No questions wrong. So feel free to keep writing them in the chat, and we're going to get to them in the second half of this

50

00:12:24.830 --> 00:12:48.339

Lowell Weisbord: session. But let's let's start at the most basic level. And I'm going to shoot this over to you, Barbara. When people say, AI in a government context or in any context, what what do they mean? What do I mean? What? What's this call about? Break down? Maybe some of the key terms, like an algorithm, a model. Cesare just used the word model machine learning. Give us a little bit of the basics here.

51

00:12:49.330 --> 00:13:11.369

Barbora Bromová: Absolutely. I really pulled the short stick here, as we say, in check with the definitional question that, trust me, even the professionals sometimes really struggle with, partly because, despite it being a very ubiquitous term, AI is, as a concept still quite controversial, particularly among the scientists, there's always an ongoing discussion as to what it is, what it isn't.

52

00:13:11.370 --> 00:13:21.680

Barbora Bromová: and whether it's even a useful term to be using. Actually, there used to be an old joke that I'm sure many of you have heard that AI is whatever the computers cannot do yet.

53

00:13:21.700 --> 00:13:31.810

Barbora Bromová: because the field is very much sort of always chasing tasks that are considered to be very difficult for traditional informatics, let's say.

54

00:13:32.140 --> 00:13:53.109

Barbora Bromová: as itself, it's pretty old, actually, as far as computer terms go. It emerged already in the sixties where it split off from a field of cybernetics, and generally currently it is used as a bit of a catch-all umbrella term for a suit of technologies that use machine learning or deep learning.

55

00:13:53.110 --> 00:14:04.380

Barbora Bromová: which essentially means that they leverage the availability of machines to improve iteratively, learn from data and also independently make decisions based on certain data inputs from the real world.

56

00:14:05.200 --> 00:14:17.469

Barbora Bromová: Now, that is a kind of purposefully broad definition that is sometimes not entirely helpful. But what I like to do is essentially, really focus on this machine learning piece

57

00:14:17.590 --> 00:14:31.460

Barbora Bromová: and think of it as more of a capability of these machines. Similarly to, I mean, it is inspired by neuroscience in some of its technical and sort of conceptual aspects. But thinking about it as a.

58

00:14:31.460 --> 00:14:54.570

Barbora Bromová: it's a capability that can be applied across various use cases and domains or data types can be quite helpful. So it is essentially a methodology that you can use to process text voice data images videos as well as quantitative data. So a lot of them are actually predictive models that try to extrapolate future developments from existing trends and existing

59

00:14:54.590 --> 00:14:55.810

Barbora Bromová: data sets.

60

00:14:55.900 --> 00:15:20.490

Barbora Bromová: then, sort of at the different secondary level. The words model or algorithm generally tend to refer to specific particular implementations of the technology in the context of a particular use case. So with certain adjusted parameters or methodologies to work optimally in the environment where it's actually going to be deployed

61

00:15:20.880 --> 00:15:46.479

Barbora Bromová: in government specifically. And there's a lot of conversations at the moment about generative technologies that we've already touched on. Of course, a lot of excitement around using them, particularly for citizen facing services when it comes to processing. Citizen requests. I think we've already seen that in the chat, but also writing and processing grant application and generally reducing some of the menial work that I'm sure we're all familiar with in government.

62

00:15:46.480 --> 00:16:11.429

Barbora Bromová: However, I think it's also very interesting and important to note that most of the systems deployed that can be classified as AI are still classificatory, particularly in government and in Europe. We've seen quite a lot of conversations about those the Netherlands is a little bit of a pioneer in this sense, where they had some success with deploying these systems, but also some controversies with how they have been working.

63

00:16:11.917 --> 00:16:27.189

Barbora Bromová: When I say a classificatory system you can think of, for example, welfare, allocation or processing applications for citizen services where the system essentially goes

through them for you and assesses, whether they meet certain criteria or where they fall on a certain spectrum.

64

00:16:29.050 --> 00:16:53.409

Lowell Weisbord: Thanks so much. That was a really in-depth answer. I think my only follow up, which I think is an important thing for, and part of the reason we're having this conversation and kind of the current boom of AI is this thing generative? AI, right, which we've already mentioned quite a few times. And I think just as simply as you can, Barbara, do you want to tell us just what's generative? AI, and how does that fit into this AI and machine learning thing you're talking about.

65

00:16:53.410 --> 00:17:16.400

Barbora Bromová: Sure. So a generative AI system is a computer system or a machine that can generate new content that has not necessarily existed before. And it is doing it, and I know that some of my fellow panelists will go deeper into this later, so I will not bore you with the specifics, but it can essentially do it by approximating from former examples.

66

00:17:16.569 --> 00:17:19.220

Barbora Bromová: And they are working with, yeah, okay, great. Then I.

67

00:17:19.220 --> 00:17:19.760

Lowell Weisbord: No, that.

68

00:17:19.760 --> 00:17:20.270

Barbora Bromová: Leave it up!

69

00:17:20.270 --> 00:17:46.976

Lowell Weisbord: Perfect. And and you know, we're not gonna go too deep on anything here. But we're gonna try to, you know, get people's understand up to stuff on on what's going on here. And and that really leads us to our next question. A little bit more about generative AI. So again. Generative. AI. It's creating things from past examples. There's data involved. We have tools like Chat Gpt, and you know, Microsoft's copilot and Claude and and Whatnot,

70

00:17:47.310 --> 00:18:11.890

Lowell Weisbord: And we hear about this term large language models which are related here, and we're going to get more from Cesare on this. But how do large language models like Chatgpt or similar tools, actually work, and again assume our audience is mostly non-technical.

Here, as simply as you can explain, Cesare, how are these generating answers or outputs. What do we need to understand to help us kind of be able to actually work with them and wrangle them.

71

00:18:13.030 --> 00:18:30.189

Cezary Gesikowski: Okay, great. Thank you very much. So yes, large language models or LLMs, as they quite often referred to in the acronym, are a type of generative AI and ChatGPT is the one that kind of emerged in first place. And you've mentioned some of the other ones.

72

00:18:30.260 --> 00:18:54.690

Cezary Gesikowski: and they use neural network technology. There are some that use other forms. But you can really, really look up. What neural networks is is a type of a technology where, you know, information is kind of fed through what looks like almost like a brain kind of a model, where all the information is going back and forth, and kind of a comparing. But basically in a nutshell, they're sophisticated tools that

73

00:18:54.700 --> 00:19:10.530

Cezary Gesikowski: learn patterns, and they analyze enormous amounts of text, such as websites and books and articles, and can be from the Internet, from specific databases, and you can provide it your own sources. So when you ask a question

74

00:19:10.540 --> 00:19:22.730

Cezary Gesikowski: like you do in the chat GPT the model kind of predicts the next word, and then another and another forming sentences that sound naturally human.

75

00:19:22.730 --> 00:19:44.630

Cezary Gesikowski: You can imagine it as a supercharged version of the predictive text feature when you're typing on your smartphone. Sometimes it's really frustrating because it gets you the wrong word and straight to autocorrect. But it's basically that type of technology. But it's crucial to remember that these models they don't truly understand the concepts or reasons like people do.

76

00:19:44.690 --> 00:20:00.380

Cezary Gesikowski: They don't have the model of the world only what was fed into them. They're really extremely skilled at pattern recognitions, you know, and making educated guesses about the probability rather than like genuine comprehension or common sense.

77

00:20:00.380 --> 00:20:23.579

Cezary Gesikowski: So you can think of them as clever prediction machines. And sometimes they're mocked or called so-called stochastic parrots, just kind of repeating what they hear, like the pirate with the parrot. That appears to be very smart. But, as you know, there are some humans like that as well, so who knows what the future holds for these large language models? That's in a nutshell? What it is.

78

00:20:24.100 --> 00:20:32.089

Lowell Weisbord: Yeah, brilliant. So sorry. I think that's that's the key thing is that the you know they're an Llm's limitation is to some extent the data put in right. And that's

79

00:20:32.250 --> 00:20:48.040

Lowell Weisbord: it doesn't have the same kind of data that we have as people. So, Joe, one thing Cesare mentioned, and we've already talked a little bit about is, we often hear about models in AI we hear about. You can select different models when I'm using Chat Gpt, we hear about models more broadly in the context of

80

00:20:48.040 --> 00:21:04.499

Lowell Weisbord: AI. We know everybody kind of knows what a model is, maybe a mathematical model or something like this, but in practical terms, for this conversation. And you know, in the context of generative AI, it may be at least in part. How does a model learn or improve? What's a model?

81

00:21:04.500 --> 00:21:08.989

Lowell Weisbord: Could you maybe share an example? Even where, where a model's being used in in a government context.

82

00:21:08.990 --> 00:21:17.449

Joseph Castle: Sure sure can definitely do that. Yes, as we did a great example of Gen. AI. I want to step back for a second and say.

83

00:21:17.858 --> 00:21:41.380

Joseph Castle: Typically the way I've seen AI explained in some sense, is is doing analysis of prediction right? So predictive capabilities. And then, you know. Generating right? So you're predicting things, making decisions on those predictions. And then the generative aspects of of content. I think most people engage with AI from a generative standpoint.

84

00:21:41.380 --> 00:21:57.390

Joseph Castle: But in reality there's a lot of predictive. The thing that Cesare talked about with the next word could potentially be a predictive analysis as well. It kind of goes one and the other, even though it's generating some models real fast. The way that I break them down

85

00:21:57.390 --> 00:22:22.329

Joseph Castle: usually start more at that base layer, predictive. Think of advanced statistics. Everyone had a statistics class, I would assume sometime through their academic career, and probably most of us did not enjoy it. Unfortunately, the reality is that that's sort of the underlying method or mechanism of what's happening with AI. And this pattern of recognition. So typically think about regression models. And there's different

86

00:22:22.330 --> 00:22:39.720

Joseph Castle: examples. I'm not going to jump into every single model. And because we could spend hours on this, think about regression models for numeric calculations, things like stock prices and growth rates. There's classification models for predicting categorical outcomes like email spam and fraud and sentiment

87

00:22:40.042 --> 00:22:45.840

Joseph Castle: clustering models. And I'm going to go through a couple looking for hidden patterns like customer segmentation

88

00:22:45.840 --> 00:23:10.730

Joseph Castle: music genres. If you think about like apple music and like, here's your jazz, and here's your country. And it's kind of categorizing these music entities together. Time series models, reinforcement learning. And then you get to generative AI, so usually within generative AI, you're looking at data, generation, image, generation and audio. And what have you? So a good example?

89

00:23:10.730 --> 00:23:32.529

Joseph Castle: We actually had a project I used to work for an, and so I was in the government. For a long time. I used to work for an analytics and AI company before my my current gig and one of the things we did. We were working with an environmental agency. I won't name names, but one of our major environmental agencies in the United States. And one of the things we were doing. It was sort of a collaboration between energy

90

00:23:32.530 --> 00:23:51.650

Joseph Castle: and environment and energy was moving chemicals, gases around the country and storing those gases in certain locations in certain States in the United I'm in the United States. So in certain States, in the United States, and we would do a public. They they would do public forums, ask for public input about what

91

00:23:51.650 --> 00:24:11.549

Joseph Castle: individuals thought about the movement and storage of these these gases and chemicals. And and we would collect all of that information. We'd put it into a database if you will run natural language processing. It's basically scanning like reading through all the comments and doing sentimental analysis on these comments of

92

00:24:11.550 --> 00:24:21.840

Joseph Castle: of basically what the public, how they felt about either the messaging, but also the actual like reality of things being stored so that it's just 1 1 example.

93

00:24:22.260 --> 00:24:43.769

Lowell Weisbord: No, that's that's really useful. I think it's at the main piece of, you know, a model underline statistics. You can gather data, use it to predict something. Use it to tell this. You know sentiment of a text. You can use it to do quite a few different things so long as you kind of collect the data, have an idea of what you want to do. That was really really useful explanation.

94

00:24:43.820 --> 00:25:07.160

Lowell Weisbord: So we have some sense of these things insofar as the majority of our audience here is somewhat acquainted themselves with, maybe Chatgpt or copilot in word, or something like this. Cesare. What other types of AI. Powered tools are already in use across public sector organizations. Can you share a few examples where, like, they're delivering real value? And I know already in the chat we have.

95

00:25:07.160 --> 00:25:19.979

Lowell Weisbord: I think we'll get much more into this conversation of how are these things being used? Where do we need to be careful. Where should we have? You know, ethical concerns which I'm already seeing in the chat. But where are we getting really great value at the minute? Cesare.

96

00:25:20.830 --> 00:25:45.119

Cezary Gesikowski: Great. Thank you very much. I just want to say Joe's comment. My statistics, Professor, was really fantastic. He was kind of crazy. He worked at statistics, Canada,

and he was teaching a graduate course at Harvard University, where I attended, and he was so passionate about statistics, and he actually told us in the future what this is going to transpire as AI, and I find it really fascinating, because I don't

97

00:25:45.120 --> 00:25:59.950

Cezary Gesikowski: much about the theory, but the fact that he mentioned that these will be used, and also that they have already been using a lot of kind of various AI models. They didn't call them AI. They were basically calling them statistical models.

98

00:26:00.000 --> 00:26:22.970

Cezary Gesikowski: and they were predicting different things already. So there's lots of names for them. There's algorithms that have been developed for very long time internally, and they never really bothered to call them anything in particular. But in Canada, for example, I can tell you that organizations often use customized. And right now we have open source AI tools designed for specific departmental needs.

99

00:26:22.980 --> 00:26:48.569

Cezary Gesikowski: We have Canchat, for example, one of the models that is being developed by shared services. Canada. I have seen a model that's like a multimodal platform radio from our other departments and agencies, and there are a few other ones they have. They're basically custom models that are looking at the emphasis recently on AI sovereignty. You might have heard this term.

100

00:26:48.570 --> 00:27:07.220

Cezary Gesikowski: So some departments and agencies are partnering with specialized providers here in Canada, like Cohere AI, to ensure that sensitive data remains secure, and within national borders, and even on premises. So it doesn't, you know, leave even your organizational servers, and so on.

101

00:27:07.220 --> 00:27:28.940

Cezary Gesikowski: So many governments are streamlining routine administration tasks to enhance citizen services through virtual assistance, like we have agpal chat here at Agriculture, Canada, that was already developed a couple of years ago by a team in the department before I joined, which is really fantastic, received a lot of awards. It's kind of externally facing.

102

00:27:28.940 --> 00:27:58.100

Cezary Gesikowski: And we are now looking at more internally kind of facing solutions that rapidly analyze massive data or virtual assistance. And beyond co-pilots, we're looking at

different models that we can do this. But I just want to emphasize that it's really important, like at agriculture, Canada. Right now we're looking at. We have a dedicated team and we have very strong executive support to look at how we can actively deploy AI and automation tools

103

00:27:58.190 --> 00:28:24.089

Cezary Gesikowski: across the department to deliver them to everybody. So it's a very exciting time. And I give you one more example from Poland. I'm also a Polish citizen before I came to Canada, and, as you can imagine, it's kind of hard to train models in a particular regional language, you can't outsource it somewhere else. So community members and communities have been helping train a model called

104

00:28:24.090 --> 00:28:47.689

Cezary Gesikowski: which is based on Mistral Open source model. And it's really specific to polish language data, offering culturally nuanced insights that you won't get from other generic models. So there's lots of different approaches. Customization specific needs leveraging open source and leveraging this sovereign. AI kind of approach, as well.

105

00:28:48.140 --> 00:29:01.669

Lowell Weisbord: Yeah, thanks, Caesar, and I guess just to quickly follow up, just to understand, like the a little bit more on the tools. Actually, what? In terms of what value they're kind of giving to governments right now. So you mentioned the Ag call the ag cal chat.

106

00:29:01.670 --> 00:29:02.160

Cezary Gesikowski: Yeah.

107

00:29:02.160 --> 00:29:08.529

Lowell Weisbord: What? What is that? What just in, you know, shortly like, what is that doing right now for government? How is that helping explicitly.

108

00:29:09.150 --> 00:29:33.729

Cezary Gesikowski: Okay, I have to be very careful. So the team that's watching this that built this that holds me sort of. But please look it up agpal chat. Basically, it's a based on Openai model. And it's customized to the data that's available to anyone who's searching for more information. Externally, it's like a very powerful and smart search engine that analyzes the information that we have already available.

109

00:29:33.730 --> 00:29:43.979

Lowell Weisbord: For example, to farmers and growers, and whoever is working in agricultural sector to find out what kind of access they have to different programs that are regional.

110

00:29:43.980 --> 00:30:03.590

Cezary Gesikowski: And lots of information to help them in their daily work, and for scientists and agricultural workers. So that's an example of the externally facing. And there's lots of kind of tools that I'm building, for example, agents with in copilot inside, as we're testing it for specific

111

00:30:03.590 --> 00:30:27.330

Cezary Gesikowski: tools and specific reasons. I've heard about different projects, like, for example, we have access to information and privacy requests that are coming in externally, they're processed internally or consultations, and I have heard that already some organizations have been using customized models to process large volumes of data so that they can synthesize it and prepare responses

112

00:30:27.330 --> 00:30:52.309

Cezary Gesikowski: quickly, and that way they can save a lot of time for what would be like very tedious going hand by hand. And and there's also an important thing. Everything is human verified. So although these systems are there, they're also still we're kind of going parallel. We have human in the loop always, so that everything that we produce that especially that's used for decision making, we have very strict rules about that. So if AI is used

113

00:30:52.310 --> 00:31:17.569

Cezary Gesikowski: for decision making, there's a impact AI impact assessment tool that has to be considered if it's impacting citizens, or if it's impacting anyone outside of government and inside as well, it has to be assessed what the impacts are, how they can be sort of minimized if there, if there's negative ones and humans are always in the loop, verifying and checking everything.

114

00:31:18.150 --> 00:31:27.720

Lowell Weisbord: Yeah, thanks, Caesare. So just touching on the 1st piece, which is that one of the use cases here that I think is really common right now with with government is about how getting

115

00:31:27.720 --> 00:31:48.339

Lowell Weisbord: making it easier to access information. It sounds like that's what the Ag PAL piece is about. It's about making it easier for citizens, or whoever the customer in this context is to access information in a super easy and intuitive way, and then second piece on decision

making. There were some questions around here. So I think we'll get more into that later, because I think that's something that people are really curious.

116

00:31:48.340 --> 00:31:59.850

Cezary Gesikowski: Has been shared in the chat to Akpal Chat. I see already. So check it out. It's a wonderful interesting user case. And if you have more questions. I can direct you to the team that has worked on it. So cheers.

117

00:32:00.200 --> 00:32:01.080

Lowell Weisbord: Nice

118

00:32:03.290 --> 00:32:26.769

Lowell Weisbord: another one. So just kind of getting more into these use cases. There are ways. AI is helping governments to improve services operations, internal decision making. You know that that's kind of what we were just talking about. Joe, do you want to share another kind of simple, specific example more on these kind of services, operations or or decision making assistance. That you've seen across government.

119

00:32:26.930 --> 00:32:46.976

Joseph Castle: Sure. Yeah, and and I'll start by saying to you that I just saw an article yesterday in the States. it's an outlet called Fed Scoop and this administration is carrying forward what the previous one had done. In the States, anyway. So every agency we have 24 major agencies government agencies. And

120

00:32:47.742 --> 00:33:05.309

Joseph Castle: every agency has to create an inventory of their use cases and post it on their website. It's aggregated on Github and you know, you can see across the whole government of what's happening. So actually, I have many. The one that I think is pretty neat, and I would imagine a lot of

121

00:33:05.805 --> 00:33:24.689

Joseph Castle: other you know, private organizations use. But our our United States Postal service with route optimization. So delivering the mail delivering packages. Think Amazon. Think you know, Fedex, anyone that Walmart, others that would also use this. But basically using route optimization with reinforcement learning.

122

00:33:25.022 --> 00:33:47.650

Joseph Castle: Where you're basically maximizing interactions or the interactions are maximized for rewards in the sense that going to one location and another location, you know, versus like an ABC. Like which one do you save the most amount of time? And where should you go? That's typically not done in in real time. But it but it's usually an analysis that could be done after

123

00:33:47.958 --> 00:34:15.139

Joseph Castle: in the defense. I'll just give one more kind of a similar idea in in Us. Defense. And I imagine it happens around the World event. Stream processing is a process of gathering data in real time putting hardware on drones and other. You know, you could use it from planes. And what have you? And you're basically feeding information back in real time. So analysts who are sitting behind their desk in in various places. They could be in the States. They could be around the world.

124

00:34:15.545 --> 00:34:22.853

Joseph Castle: Can can identify vehicles and people. And you know, various buildings. And what have you?

125

00:34:23.280 --> 00:34:32.992

Joseph Castle: And and can basically take action in real time. And so that happens, of basically kind of the same idea, this this

126

00:34:34.205 --> 00:34:39.830

Joseph Castle: identifying points of interest sort of a classification system. And it kind of speeds up

127

00:34:40.184 --> 00:35:02.689

Joseph Castle: what's happening? I've seen it mostly with I'll just say a good. It sounds horrible because it sounds like you're you know the you know spy in the in the in the sky. Kind of thing, but it's really it was for I I had seen it for drug interdiction. So you know, identifying you know, drug labs in jungles and those sorts of things. So there, there is good use for it.

128

00:35:03.720 --> 00:35:31.749

Lowell Weisbord: Yeah, really interesting. I mean, 2, 2 different use cases here, both one more around identification really interesting, but also kind of service optimization through the Post Office example. And I think that's 1 that probably we, you know, we generally have a lot of people who work in some form of service delivery. And that is, you know, the majority of government and and kind of how do you take the data you have and and the live data in some

cases and optimize it, to be able to basically deliver services more effectively, to make our post routes

129

00:35:31.800 --> 00:35:40.770

Lowell Weisbord: more efficient, or or whatever kind of your use case is, I think that's really interesting. And and are some of the really inspiring use cases we're seeing immediately of of

130

00:35:40.850 --> 00:35:43.472

Lowell Weisbord: just doing our work better right?

131

00:35:44.890 --> 00:35:51.900

Lowell Weisbord: this has already come up. But bias, you know, we talked about human human in the loop has been mentioned as a term. I'd love

132

00:35:51.980 --> 00:36:17.239

Lowell Weisbord: Barbara. You did. I'm going to touch on bias here, but I'd love somebody to say what that means for anybody who doesn't know I'd love to. We were talking a little bit about decision making. And you know, having who makes the final decision, does an AI make a decision? Does a human make decision? And then a 3rd thing is is bias in data. You know, we talked about how important data is to train these systems, these models, these Lims.

133

00:36:17.290 --> 00:36:34.380

Lowell Weisbord: the data matters. Right? So people are worrying about bias in these systems. What creates bias in AI, and what role can governments and the people in this room play in kind of spotting and reducing these risks, or where might we start? If that's an ambitious question.

134

00:36:35.460 --> 00:36:57.960

Barbora Bromová: That is a very ambitious question indeed, and bias is also once again a massive term. Oftentimes there are some structural factors that translate into bias that I'll get into. But I also wanted to start by saying that oftentimes bias is a little bit of a feature as well as a bug, because sometimes you can arrive at a biased model, a biased

135

00:36:57.960 --> 00:37:05.800

Barbora Bromová: system by chasing a certain optimization. So sometimes the bias is less of a bias in itself, and more of a

136

00:37:05.800 --> 00:37:16.360

Barbora Bromová: and, let's say, overfitting of the model, we would call it. It's it's too much tailoring to a particular case, or to a particular subgroup of the use cases that you're using them for.

137

00:37:16.410 --> 00:37:23.099

Barbora Bromová: So that would be more of a design flaw. A lot of what is the concern in government, however, is

138

00:37:23.130 --> 00:37:49.219

Barbora Bromová: data scarcity. So in public service contexts, a lot of what sort of results in what we call bias would be the not only flaws in data collection data management systems that are, for whatever reason, not up to par, but also long term inclusion issues that tend to lead to certain blind spots in the training data that the model then has available for itself to make those inferences that we were talking about.

139

00:37:49.340 --> 00:38:07.880

Barbora Bromová: I saw already, Sophie mentioning in the chat that this is very common, actually, across all different technologies and statistics. For example, even in medicine, as Sophie mentioned, a lot of studies are just made on male bodies, which then results in certain inefficiencies, let's say, for women and other people.

140

00:38:08.020 --> 00:38:36.150

Barbora Bromová: However, also globally, we're seeing that as quite a big problem and a big, actually sort of a break on how useful AI can be at the global scale, because a lot of it is developed in Western worlds and in English languages. I'm really glad that Cesare mentioned the Polish example because that is actually quite close to my day job at the UN development program, where I focus particularly on multilingual applications of AI models

141

00:38:36.190 --> 00:38:57.900

Barbora Bromová: trying to resource some of those indigenous and local languages in Africa, but also in Latin America, where speakers of languages that are just not well represented online in those data sets currently do not have access to those systems that speak their language. And in addition, certain languages are, of course, not

142

00:38:57.920 --> 00:39:11.329

Barbora Bromová: equally reliant on text. For example, the way that English, French, or Spanish would be working much more with voice, and that then, can help also populations that are maybe not

143

00:39:11.560 --> 00:39:33.089

Barbora Bromová: entirely literate. And that is a big avenue through which government services can reach new populations also in the developing world. So there's a little bit on what I do in my job on bias. I kind of wanted to emphasize the old economic saying that I like to think about when people bring out AI models as well is that

144

00:39:33.210 --> 00:39:41.850

Barbora Bromová: all models are wrong and some models are useful. I think that does apply in this case as well, where thinking quite critically about

145

00:39:41.870 --> 00:39:53.760

Barbora Bromová: how exactly the model was trained, what exactly was it designed to do? And how is it now getting applied and monitored and evaluated in its outputs is really important. And the same way that you would

146

00:39:53.760 --> 00:40:12.039

Barbora Bromová: do with any other data use or technology use in government on human. In the loop I will push back a little bit. It is a term that you hear quite frequently in these conversations. It generally means that a human should be involved in the decision

147

00:40:12.040 --> 00:40:18.760

Barbora Bromová: at the point that you, at a certain point in the metaphorical loop of that process.

148

00:40:21.150 --> 00:40:29.159

Barbora Bromová: it is not in itself a wrong idea. This is not what I'm trying to say, and human involvement and human.

149

00:40:29.220 --> 00:40:58.560

Barbora Bromová: Yeah, human involvement is incredibly important in these automated decisions. However, sometimes I worry whether it is a little bit of a shorthand for safeguards and for more meaningful auditing of these systems, because something that I think quite a lot about is meaningful transparency, meaningful explainability of model decisions and outputs, which is

something that cannot necessarily just be hand waved away by saying, Oh, yes, we had a person check these outputs. So therefore they are okay.

150

00:40:58.560 --> 00:41:21.110

Barbora Bromová: Sometimes I believe that a lot more effort and thinking should be done in terms of the whole lifecycle, and so as opposed to human in the loop, there should be more human in the design, and thinking a bit more about who is represented in training data, but also who is affected by the outputs

151

00:41:21.110 --> 00:41:37.129

Barbora Bromová: focusing a little bit more on user interviews. The way you would do in product development and seeing the actual social impacts of the system as opposed to evaluate them purely, technically. But that is more of my opinion coming from my background, and I'm very happy to be challenged on that one.

152

00:41:37.990 --> 00:41:52.610

Lowell Weisbord: No, that's brilliant. And yeah, let's keep the challenges coming and the conversation going back and forth. I think it's a really fair point. We're human in the loop. We don't wanna use that as a catch. All or not worry about like fundamental things like the data going into the model, or

153

00:41:52.650 --> 00:42:17.730

Lowell Weisbord: you know, especially perhaps, as we become more trusting of these systems, which kind of goes to Debbie, your last point there, how long do you think it will take to get to the point where we can trust AI not to make up slash fabricate information? I think a lot of people are becoming. I mean, there's there's different ends to this equation. But some people definitely are becoming more trusting as they use it more in their day to day, and that's where we need to like kind of double down on. Where do you safeguard? What does human in the loop mean?

154

00:42:17.730 --> 00:42:24.700

Lowell Weisbord: And and when, especially as as public servants. We, you know most of the people in this call have a very different kind of duty

155

00:42:24.700 --> 00:42:31.299

Lowell Weisbord: to uphold with the kind of risk tolerance we can take for a decision or a system that we design.

156

00:42:31.300 --> 00:42:48.669

Lowell Weisbord: and that takes special care. I had a couple more questions for our panelists, but but I'm going to hold on that and kind of get into the Q. And A, and I think we'll probably touch on these ideas as we get into things so feel free to. There's a Q&A function on

157

00:42:48.800 --> 00:43:01.620

Lowell Weisbord: zoom. So I'd love if you. If you don't mind putting them into there. It makes it easier for me to see, but feel free to. If you don't see that, keep them coming in the chat. We have a bunch of really great questions. So I'm gonna get started.

158

00:43:03.570 --> 00:43:04.530

Lowell Weisbord: With

159

00:43:05.040 --> 00:43:31.619

Lowell Weisbord: one that I think comes up a lot. I think it's worth addressing. But most of the use cases for government are about efficiency and productivity. I think that is generally what we've talked about. How can government also consider sustainability, impact and inclusion as outcomes when considering use of AI, and this kind of touches on the last point around data. But what are the other kind of ends that AI can lead us towards. I think it's a really fair question.

160

00:43:31.917 --> 00:43:38.750

Lowell Weisbord: Cesare, I might throw this over to you, and anyone else feel free to jump in. If you're open to digging into that.

161

00:43:39.180 --> 00:43:39.930

Lowell Weisbord: This is from.

162

00:43:39.930 --> 00:44:04.459

Cezary Gesikowski: Absolutely, absolutely. I think I think it's it's great and kind of ties to what I was saying before, because I mentioned human in the loop, and this is sort of there's not none. This is a very good challenge that Barbara put forward. And you know right now, that's all we have. So that we have to put the humans in there to check our results. But eventually we want to develop systems

163

00:44:04.470 --> 00:44:17.640

Cezary Gesikowski: that will be really kind of reliable. Okay? And there's a great question. How long do you think it will take to the point where we can trust AI not to make up or fabricate information. Or

164

00:44:17.700 --> 00:44:19.479

Cezary Gesikowski: to tell us the truth. Yeah.

165

00:44:19.510 --> 00:44:41.219

Cezary Gesikowski: I think the again facetious answer is like, when we all, as humans agree, what truth is, then maybe we can develop a system that will actually be able to tell us. And we could all agree that that's the truth, right? This is very contentious, philosophical. I love it, we can discuss it, but there's another term that I would like to introduce. You might have not heard, and it's human in the loop.

166

00:44:41.220 --> 00:44:50.829

Cezary Gesikowski: and I shared a link from common good. AI and common good. AI is based on this really interesting technology provided by

167

00:44:50.830 --> 00:45:01.650

Cezary Gesikowski: at the company from the United States which I've kind of worked with for a while, and what they do is actually they bring people together in conversations, asymmetrical kind of

168

00:45:02.000 --> 00:45:22.109

Cezary Gesikowski: asynchronous discussions, where they can talk to each other and rank each other's answers on any kind of topic. So it's almost like a crowdsourcing kind of intelligence. Human intelligence. Okay. And this is leveraged by AI. That kind of keeps track of everything and looks at the answers, suggests different

169

00:45:22.110 --> 00:45:33.620

Cezary Gesikowski: considerations. And so on. And it's humans working together, thinking about problems and solutions and AI kind of being there in the group.

170

00:45:33.620 --> 00:45:59.070

Cezary Gesikowski: helping them synthesize and summarize things. And I think it's a very revolutionary way. We're going to use AI beyond this kind of a human in the loop. Okay, look up common. Good. AI, they're doing some really interesting work. And it explains how that

technology works. So I think that's what I'm really excited about right now. I hope that answered essentially the question that you've put forward. There.

171

00:45:59.630 --> 00:46:05.357

Lowell Weisbord: Yeah, that that's that's definitely helpful, Seth. I appreciate that.

172

00:46:05.930 --> 00:46:33.810

Joseph Castle: Can I just add real real fast? 2 things that come to mind with we're we're talking about truthfulness. But really, 2 things that come to mind are validity and and probability. Right? Remember, these are statistical models. So it's it's, you know. It always goes back to confidence levels confidence intervals. In some sense, right? It's like, that's 99% accurate, right? But there's so. And then that goes to my validity. Comment is, a lot of times when I use

173

00:46:34.116 --> 00:46:55.350

Joseph Castle: more more in the Gpt realm. But you know, when I'm when I'm you know, thinking of things, asking questions, prompting, etc. A lot of times, I'll you know, get the results, and I'll kind of look at it, and I'm like that. Seems kind of right. I'm not sure about that part. Let me go try another gpt for a minute, or let me just go look it up and find like an actual source, like, you know, Google scholar or something similar.

174

00:46:55.611 --> 00:47:07.900

Joseph Castle: But it's the validation piece. And I feel like we're we kind of miss that across like, and I won't go too far with this, obviously. But like social media, and just the amount of information that we receive, and there's not a lot of questioning

175

00:47:07.990 --> 00:47:29.310

Joseph Castle: is what I find and not a lot of discussion. You know, potentially with colleagues or friends or family or things like that. So I would just say, validation is huge, and it applies to AI because it is data driven and so yeah, don't. Don't always believe what you read. Don't always believe what you see right like those standard comments, and remember the confidence intervals. You know.

176

00:47:29.990 --> 00:47:47.459

Lowell Weisbord: And I guess, Joe, if you don't mind me following up on that, we got 2 questions here about one about the hacking of AI, and one about the systems being at risk for bad actors and personal data, I think more, you know, feel free to address these questions to some extent, but also more broadly like, how do we know

177

00:47:47.460 --> 00:48:07.890

Lowell Weisbord: what's safe to use as a government as a public servant who maybe hasn't received any guidance? What are kind of the risks of these systems. And and how should I think about the risks, whether it's, you know, the risk to the data that I'm inputting the data that I'm using to train, or just as someone you know, who's trying to understand the landscape.

178

00:48:08.360 --> 00:48:16.809

Joseph Castle: Yeah, that's it's interesting, because it's it's the amount. It's the quantification of risk, right? And and how much risk are you willing to

179

00:48:16.810 --> 00:48:42.799

Joseph Castle: to take on which kind of goes back to to my validity, and then confidence statements. One of the things I would, I would consider is, I know. In the beginning we had the questions about using AI daily, and then official guidance was like 50 as a public servant. I know, working in government agencies for a long time. And and even this happens in industry. I just talked to someone the other day about developing policy around co-pilots. And what can we use? Because if we put

180

00:48:42.800 --> 00:48:57.340

Joseph Castle: if we put our own code into a co-pilot, now, what are the terms of use and like? Can they take that code. And if that's proprietary and we're selling that software now, it's potentially open to the marketplace. So definitely, you know, figure out what's acceptable in your organization.

181

00:48:57.640 --> 00:49:15.070

Joseph Castle: Maybe if you know some of those folks like the Chief Information Security Officer AI officer like, ask them why, if you're really that interested in some sense there might be some scanning going on or something. I would assume some analysis where you sort of white list like, okay, these products are okay in in this capacity.

182

00:49:15.511 --> 00:49:18.159

Joseph Castle: These are not for various reasons.

183

00:49:18.230 --> 00:49:45.200

Joseph Castle: But I I would definitely start there. Then, then, as far as actually using things, if you're not a a daily user. It looks like 60% are. But you know, the other 40%, or even, you know, I maybe not. Daily. Start with the low risk items potentially like like we had a. We had a chat bot

at Gsa. General services like, I don't. I want to say, like 5 to 7 years ago, or something. So we started very low risk. It was all of our all of our content that we put into it.

184

00:49:45.456 --> 00:49:54.433

Joseph Castle: You know, web based chat, bot and you know, helped us with simple things like, like, you know, how do you fill out a lead form. So I can go on vacation right?

185

00:49:54.690 --> 00:50:12.089

Joseph Castle: you know some very basic things. Who do I talk to in legal about, you know, a potential ethics topic, you know, for public speaking, or something right? And so start very simple. If you you know, based on what you're allowed to do, not allowed to do, then venture into the Gpts, then venture into the education of

186

00:50:12.090 --> 00:50:26.139

Joseph Castle: how do I just use this stuff better like what is it? How do I use it better? I'm actually doing a class coming up for us defense personnel, on prompt engineering. So sort of an AI overview and prompt engineering. And how can they use it every day?

187

00:50:26.428 --> 00:50:33.939

Joseph Castle: You know, if if they want to in their day jobs to basically complement, you know what they're doing, you know, on a regular basis.

188

00:50:35.310 --> 00:50:56.019

Lowell Weisbord: Super helpful. Thank you, Joe, to kind of this prompting of, you know prompting is the act of writing something to an Llm. Right to Chatgpt, or to whatever you know Microsoft Copilot. In some context, one of the questions we got from Bob, why does AI fabricate answers? We also call this. A hallucination is what it's being called.

189

00:50:56.020 --> 00:51:21.430

Lowell Weisbord: If you ask AI to find a piece of data on the Internet, why does it give you a wrong answer sometimes rather than just returning the correct answer. It's a really fair question of why does this happen? Why can it give you all these right answers, and then give you something that's completely made up. Barbara. You talked a little bit about how this works earlier. Do you want to address that piece of why are there hallucinations, or maybe any guidance on how to prompt effectively? Perhaps.

190

00:51:22.390 --> 00:51:36.479

Barbora Bromová: Absolutely so. We did already touch on this. It goes back to something that Cesare said in the very beginning, where oftentimes these generative models are quite probabilistic. So they are doing, they're predicting the next word. And they are doing this even potentially.

191

00:51:36.900 --> 00:51:48.300

Barbora Bromová: This is the technology is changing and improving at this, and a lot of very, very smart people are trying to fix this as we speak, or have been fixing it so a lot of this might be more of the

192

00:51:48.580 --> 00:52:09.010

Barbora Bromová: I'm simplifying a little bit, but essentially they are still predicting what, for example, an academic citation would look like if you are asking for an academic source. They are oftentimes drawing from databases, but not necessarily retrieving information, but simulating what that information might look like, and doing a little bit of a

193

00:52:09.080 --> 00:52:32.540

Barbora Bromová: a melange of what that link, or what that citation would be. Hence. Why, sometimes it would be probably quite accurate if you're asking about something relatively generic. But if you're looking for a very specific information, they might not. They might essentially just simulate what it might look like as opposed to actually retrieve it for you.

194

00:52:32.570 --> 00:52:55.940

Barbora Bromová: Now, in more advanced models, specifically by the usual suspect companies like Openai and Tharpic and the others, they have been actually integrating search models and Internet connections into those outputs trying to solve this a little bit by combining generative systems with search and other more information, retrieval technologies to solve exactly this issue.

195

00:52:55.980 --> 00:53:13.150

Barbora Bromová: That being said, it is because it's a little bit of a combination and a multimodal multi-methodology approach. It's not yet 100% reliable, so evaluation and sort of confirmation of all this information is

196

00:53:13.260 --> 00:53:14.719

Barbora Bromová: always advised.

197

00:53:15.510 --> 00:53:35.850

Lowell Weisbord: Yeah, that's super useful. And I think that final point is a key thing of the way that they're kind of solving. This problem is trying to retrieve a piece of information and then put it into the model and to get more accurate results. That's also kind of what we can do. If you ask it about a paper just generically, you'll get worse results than pasting in the paper and asking it about. So it has

198

00:53:35.910 --> 00:53:55.900

Lowell Weisbord: kind of the context right there. It's a lot less likely to hallucinate and kind of. I think somebody else said this earlier, too, which is just this point that think about what data it has, you know, is the question that I have something that's well answered on the Internet that's been used as training data for this model. I think that's like, always a really interesting question.

199

00:53:57.200 --> 00:54:09.049

Lowell Weisbord: we have time for only one or 2 more questions, and I think I want to get kind of the most bang for a buck out of this, because we've had a really really fruitful discussion here. So appreciate all the

200

00:54:09.170 --> 00:54:12.900

Lowell Weisbord: questions coming in. But I'm going to ask something kind of across

201

00:54:13.020 --> 00:54:15.119

Lowell Weisbord: the panel here, which is just

202

00:54:17.040 --> 00:54:41.189

Lowell Weisbord: looking ahead. What's the most important skill or mindset public servants need to develop, to work effectively with AI in the next few years, because it seems like this thing's coming. It's really here for most people. And if it's not here for you. It will be what's the specific feel free to go specific on a scale or a way to think about this? That's going to be critical. I will start with you, Cesare.

203

00:54:43.230 --> 00:54:57.959

Cezary Gesikowski: Thank you. That's very dear to me as well. I always say it's not about technology. Think human. First.st Okay, AI revolution is not just about tools. It's about how we individually and collectively navigate different opportunities and manage risks.

204

00:54:58.230 --> 00:55:13.360

Cezary Gesikowski: And all these biases and hallucinations and maintain a critical perspective. I find it really important to adopt an experimental mindset. Be curious, flexible, and and willing to continuously learn and adopt and take charge of your own knowledge

205

00:55:13.500 --> 00:55:37.700

Cezary Gesikowski: and kind of validate with others, working effectively with AI means like embracing all these changes and feeling comfortable, exploring, unfamiliar territory. And the last one, I would say, is just really to prioritize relationships and collaboration like technology alone won't solve all our problems. Strong interpersonal relationships and understanding your culture. The human connections are essential.

206

00:55:37.700 --> 00:55:56.040

Cezary Gesikowski: as the old kind of a proverb goes, if you want to go fast, go alone. But if you want to go far, go together. Remember your human fellows. We humans are the ones that are going to make changes and really make an impact. Technology will help us or impede us. But it's our to decide.

207

00:55:56.440 --> 00:56:01.399

Lowell Weisbord: I really appreciate that accessory people, curiosity relationships. Joe.

208

00:56:02.580 --> 00:56:11.920

Joseph Castle: Yes, definitely, definitely echo some. And some of those words, as we said, are in my notes as well. But I but I would say it like the bar

209

00:56:12.303 --> 00:56:28.630

Joseph Castle: is really the assistive. I call it assistive, but augmented, you know, augmented technology. What have you? So it's it's it's don't be left behind because of lack of understanding. Right like like harness, the the the chance to to learn about the technology and use it in your daily

210

00:56:28.979 --> 00:56:50.009

Joseph Castle: life. There's always this sort of in the States, anyway. We have sort of have this dialogue going between innovation and and you know adverse effects, if you will right the the good and the bad. And can we have both at the same time? And you know, depending on what? You know, government administration is in order. It's what are we doing?

211

00:56:50.322 --> 00:57:18.110

Joseph Castle: I will say, as public servants. We're on the front lines for human protection, citizen protection. And it is our responsibility, at least in my opinion, that we do things responsibly. You know, we want innovation. We want help. We want to learn these technologies to be able to harness them. And I fully encourage that. At the same time, we need to be you know, continuously curious and you know, wary of of adverse effects and and adjust for those things. That's our job.

212

00:57:19.220 --> 00:57:24.090

Lowell Weisbord: Curious and wary. I like that as well. And and finally, Barbara.

213

00:57:24.750 --> 00:57:40.710

Barbora Bromová: Yeah, it's incredibly hard to follow up. I would definitely sign under both of those contributions. Maybe the one thing that I would add is beyond a sort of attitude of experimentation. I would also advise to keep

214

00:57:41.150 --> 00:57:49.879

Barbora Bromová: a big emphasis, and it's sort of a combination of both of the previous answers actually, but keep an emphasis on evaluation. And

215

00:57:50.220 --> 00:58:19.020

Barbora Bromová: do think about so. For example, one of the principles that we have at Undp digital with my team where I work is, maybe don't build it, because we see a lot of teams falling into this sort of sunk cost fallacy where a high level leader has really committed to a certain digital project, and they have invested a lot of man hours and a lot of money. Oftentimes a lot of infrastructure into making this happen with all these hopes and dreams of how it's going to make everything better. And upon evaluation.

216

00:58:19.120 --> 00:58:44.640

Barbora Bromová: maybe it's not necessarily benefiting the workplace as much. Maybe it's not really benefiting the citizens. And it's introducing a lot of friction as people get up to speed in terms of digital literacy. So really evaluating continuously throughout the lifecycle of the technology, evaluating the whole system as opposed to just the technology itself. But just also, how is it used? How is it integrated with other data sources, other data processes

217

00:58:44.700 --> 00:58:50.820

Barbora Bromová: and not being afraid to say, Okay, that didn't work. We try something else next time, and it will be okay.

218

00:58:51.610 --> 00:59:03.660

Lowell Weisbord: Yeah, thank you so much, Barbara, and we're getting to the end. So just ask everyone to stay with us for the last couple of minutes. But most importantly, let's give a big thank you. And round of applause to our 3 panelists, Cesare, Joe, Barbara.

219

00:59:03.660 --> 00:59:22.470

Lowell Weisbord: really, really, really fantastic to have you, and also to all the attendees. We've had over 350, maybe 400 people throughout the call. Very, very engaged, appreciate everyone coming and asking really phenomenal questions has been really, really good. There's lots more questions. So feel free to take them to our AI and government community as well. We'll drop the link.

220

00:59:22.470 --> 00:59:25.369

Lowell Weisbord: Keep the conversation going. I know Cesare's in there.

221

00:59:25.370 --> 00:59:43.009

Lowell Weisbord: and we also have a poll which probably would have just popped up on your screen of just was this a good use of your time? 0 to 10? We like to to use public servants time? Well, because there's not much of it, and there's lots of important things to do, so we do not want to waste that.

222

00:59:43.480 --> 00:59:46.340

Lowell Weisbord: My final announcement is

223

00:59:46.720 --> 01:00:16.600

Lowell Weisbord: to introduce. We have an AI readiness check, which is a very short 5 min assessment. It's just a great next step. If you're curious about where you're at with AI and need some direction on how to kind of continue your development as a public servant, so we'll also make sure to drop that in the chat, and I'm sure it will be in a follow up email as well. But this is all the time we have for today. So thank you again, Barbara Joe Cesare, we hope you found it helpful. We want to say a huge thank you again to everyone for taking time out of your busy days

224

01:00:16.600 --> 01:00:23.170

Lowell Weisbord: and look for any kind of follow-ups. But thanks again, everyone, and have a great rest of your day wherever you are.

