



HOUSE OF LORDS

Communications and Digital Committee

Corrected oral evidence: Large language models

Tuesday 24 October 2023

2.30 pm

[Watch the meeting](#)

Members present: Baroness Stowell of Beeston (The Chair); Lord Foster of Bath; Baroness Fraser of Craigmaddie; Lord Hall of Birkenhead; Baroness Healy of Primrose Hill; Lord Kamall; Lord Bishop of Leeds; Lord Lipsey.

Evidence Session No. 6

Heard in Public

Questions 46 - 50

Witnesses

I: Professor Anu Bradford, Professor of Law and International Organisation, Columbia Law School; Dr Mark MacCarthy, Senior Fellow, Institute for Technology Law and Policy, Georgetown Law; Paul Triolo, Senior Associate with the Trustee Chair in Chinese Business and Economics, Center for Strategic and International Studies.

USE OF THE TRANSCRIPT

This is a corrected transcript of evidence taken in public and webcast on www.parliamentlive.tv.

Examination of witnesses

Professor Anu Bradford, Dr Mark MacCarthy and Paul Triolo.

Q46 **The Chair:** This is the Communications and Digital Committee. We are continuing our inquiry into large language models. We are very pleased this afternoon to be joined by three witnesses, who are all joining us virtually because they are somewhere other than the UK. I am very grateful to all of you for giving up your time.

We are going to be looking at international comparisons this afternoon on the approach being taken to the regulation of AI in different territories. It will not surprise those tuning in to hear me say that the focus will be on Europe, the US and China. We have experts to cover each of those areas. I will ask them to introduce themselves briefly and, if they are representing an organisation, to say which organisation that is.

Professor Anu Bradford: Esteemed Members of the House of Lords, thank you for the opportunity to engage in this conversation with you. I am a professor at Columbia Law School in New York and my research focuses on digital regulation and European Union law, among other topics.

Dr Mark MacCarthy: I teach at Georgetown University and am a non-resident senior fellow at the Brookings Institution. I do research and publish articles in technology regulation. I am the author of a book called *Regulating Digital Industries* forthcoming from the Brookings Institution.

Paul Triolo: I work for Albright Stonebridge Group, a political risk consultancy. We work very closely with companies in the AI sector and cover very closely the regulatory developments in all three of the areas you are interested in. I am very glad to be here today.

The Chair: Thank you again to all three of you.

Q47 **Baroness Fraser of Craigmaddie:** Professor Bradford, I would like to start with you as an expert on the EU approach to regulation. We are very interested to hear about the strengths and weaknesses of a privacy human rights model. It has been described to us as well intentioned, somewhat naive and too heavyweight. I am really keen to understand whether it will work, what the controversies are and what you expect the outcome to be.

Professor Anu Bradford: Would it be helpful for me to outline the regulatory approach briefly or are the members of the committee familiar with the details?

Baroness Fraser of Craigmaddie: We have had a fabulous briefing on this so we are familiar with the details. We want to explore the implications of the approach.

Professor Anu Bradford: Given the complexity and fast-evolving nature of technology, especially generative AI and large language models, there has been hesitation over whether the Europeans are

jumping into this task too early. We do not yet have a full understanding of how the technology will evolve. The Europeans are now defaulting to their instinct to engage in what many would say is pre-emptive risk-based regulation, whereby a precautionary approach, in the interest of preserving fundamental rights and democratic structures, takes centre stage. There are a couple of trade-offs here. I see benefits in acting faster, despite the challenges, and the European Union now seems to be guided by that view that favours faster action.

First, there are questions about timing and the competence of the legislators to regulate. In terms of competence, lawmakers in the EU and many other places surely do not understand the intricate details of large language models, but that is no reason for them not to get involved. Lawmakers frequently regulate complex domains of the economy. They regulate airline safety without knowing how to build planes. They do not know how to develop vaccines, but we are very comfortable with regulations on medical safety.

For AI, especially generative AI, you need to think about the associated risks alongside the benefits. The argument would be that this is about not just regulating technology but the implications of this technology for the fundamental rights of individuals and the democratic structures of society. Those are issues where the tech companies are not experts. They are some of the least-qualified entities to be put in charge of our democracy. That alone earns EU regulators not only a seat at the table but to be able to make the key decisions, so that the regulation is based on the rule of law and binding obligations, and subject to democratic oversight.

On timing, this technology is in the early stages and it will continue to evolve. I am rather worried about the delay even with the EU's approach. Even if EU lawmakers finalise their legislation during this calendar year, it is likely to be two years—some have suggested 12 to 16 months, if we really try to expedite it—before these obligations start applying to companies. That is a lifetime when we think about generative AI. We have very important political elections facing us next year, including in the UK, the European Parliament and the United States. It is very dangerous for these models not to be subject to oversight during that time.

There is also a cost to delay. One reason why the Europeans are weighing in early is that the more we wait, the more we entrench the power of the existing frontrunners by allowing them to continue to develop this technology without guardrails. By the time there is eventually a willingness to regulate in many parts of the world, this dominance, which is already very extensive and entrenched, will be even harder to peel back.

Having articulated why it is correct that they have this role, that there is binding regulation in the works and that this regulation is coming sooner rather than later, there are a few concerns on the specifics of the regulation. All regulation is not optimal, but neither is all innovation. That is guiding the European approach. The goal is to influence the pathways

of innovation so that resources are dedicated towards generating socially beneficial innovations that help economic growth in ways that are consistent with the values of democratic societies and, at the same time, limiting the potential harms and downsides of harmful innovation.

One concern that I have is that regulatory obligations are always disproportionately costly for smaller players in this field. Many European AI developers, whether in the EU or the UK, are small and medium-sized companies. I am more worried about their compliance costs than about big tech companies' ability to comply. The EU's response in the AI Act is to offer regulatory sandboxes, which are controlled environments within which the technology can be explored and regulatory compliance can be tested. The question is whether that will work.

I am slightly concerned about some issues to do with the specifics of the regulation of foundation models and generative AI. First, in the version of the regulation on generative AI that came out of the Parliament—we all know the trilogue is meeting today so we might actually be wiser tomorrow—there were obligations such as the need to disclose the copyrighted data used to train these AI models. As a scholar who produces copyrighted work, I am very sympathetic to this. At the same time, the most powerful and capable models will rely on an extensive set of data and the amount of copyrighted material used will be almost limitless. I worry whether it is feasible to expect that all this copyrighted material will be correctly identified and listed.

Secondly, there is an obligation to be transparent; the developers of generative AI ought to indicate, using a watermark or a label, that we are dealing with AI-generated content. In principle I am in agreement, but, if so much content is at least partially generated by generative AI, will it all be labelled? How will we distinguish between risky content, of which we want to alert users, and non-risky content, so that this labelling regime remains effective?

Some of the obligations, for good reason, are still very vague. According to the draft that came out of the Parliament, foundation models need to assess and mitigate risks associated with those models. This is very open-ended. The drafting of this regulation is inevitably in very general terms, which probably means that it will require further standards or other specification down the road. I will leave it there and we can move on to some more specific questions, but I hope that was responsive at the outset.

Baroness Fraser of Craigmaddie: Thank you very much, Professor Bradford. I hear your point about delay being its own risk, but the legislation has already been delayed because of the development of ChatGPT and has had to be rewritten. Is that an illustration that the approach has fatal flaws and that it will not work? I want to push you further on where you think it might land.

Professor Anu Bradford: I share that concern. The institutions of the European Union had been working on the legislation for about two years. It was almost ready to be released when ChatGPT was launched. There

was then a concern that the legislation might already be outdated, but the European Union institutions responded rather rapidly.

This is still one of the points of contention. How do we incorporate generative AI into the regulation when it does not neatly fit into the categories for the level of risk that determine the level of obligation to which a developer is subject? There was little time, proportionately, but it has been dedicated to a deep discussion of generative AI.

If we think about the alternatives, one is that we just do not include it within this regulatory framework. That is also a very dangerous path to travel. We at least need the kind of transparency that EU lawmakers are calling for. We need accountability, so that developments at this stage are directed towards not just maximising profits and the potential for high-performance capabilities but a thoughtful consideration of risk mitigation and governance aspects.

That way, every developer of generative AI knows that they do not have a free pass and that this is very much on legislators' radars. Ultimately, it invites them to engage in a productive dialogue about how we govern this. The governance is forthcoming; the regulatory obligations will be there. How do we make sure that they work, that they are feasible and that they ultimately succeed?

There is a feeling of unease among consumers. There has been a rather high-profile conversation about potential existential threats and the idea that these models could spin out of control and have a life of their own. It is a problem for innovation if consumers and business users are hesitant to take on technologies because they are too concerned about the risks involved. Mindful and thoughtful regulation could set guardrails that will create more trust in these technologies. It would also then yield benefits in the form of better development and greater uptake.

Baroness Fraser of Craigmaddie: I would like to explore the point about helping the smaller companies without stifling innovation. The other thing I want to come on to is enforcement. If consumers are worrying about it and the EU might enforce this at different national levels, will that be an even greater challenge for innovation, companies or whoever?

Professor Anu Bradford: You are asking exactly the questions that concern me as well. The biggest failure or disappointment with GDPR was exactly the asymmetrical cost for smaller companies. The EU seems to have learned its lesson. In subsequent digital regulation, including the Digital Markets Act and Digital Services Act, it has imposed asymmetrical obligations, targeting the bigger companies with a greater potential to harm but also a greater ability to comply.

I would like to apply that philosophy in the context of AI as well, but it is somewhat more complicated. It is not immediately clear that larger companies and the kinds of applications they produce are necessarily always proportionately more harmful. Exempting the smaller companies altogether from the obligations might not be wise or feasible. The

question is then about how much resource is dedicated to enhancing the ability of these companies to meet the compliance burdens.

Another related issue about which there is lively discussion—to my knowledge, it has not been fully resolved—is where the liabilities rest in this chain of development. If the bigger companies with a greater ability to comply still continue to produce most foundation models, will we retain the liability mainly at that level and not at the level of the smaller companies that apply those models and develop on the foundations laid by the bigger tech companies?

One option is to make sure that we do not just release the initial developers of generative AI from this liability and pass it on to the next stage. The large tech companies have made the argument that, if they develop a foundation model without developing it for a specific purpose, they should not be held liable if it is used in a risky setting. There is a counterargument. Whoever is deploying this model, including the smaller companies, may not have an intricate understanding of how the model was developed. They may not be in the best position to bear the compliance burden. Therefore, you have the opportunity to adjust the level of obligation based on who is in the best position to comply.

Secondly, there is the question of enforcement. I have been less worried about the burdensome European digital regulation or the cost on innovation. There is another question about why the Europeans are not innovating at the level of the Americans. I would attribute very little of that to digital regulation. I have been much more worried about under-regulation and the inconsistent track record of enforcement by the European Union.

That has been the big issue with GDPR, which has underdelivered on the promise that was made to citizens and society at large. The big test case is whether the AI Act will be enforced well. The concern you raise about fragmentation was the problem with GDPR. When implementation was delegated to member state level, individual agencies, including the Irish Data Protection Commission, were overwhelmed by the size of their remit.

The EU is now considering a model whereby we still empower national-level enforcement agencies but with greater co-ordination by an EU-level AI office. There is also a potential model whereby the biggest companies would be regulated at the EU level. That is the model followed by the Digital Services Act. There is a great danger in promulgating regulations that are not consistently and effectively enforced because it erodes the legitimacy and credibility of the Government and the deterrent effect of the legislation.

I am very grateful that you have posed those two questions because they also concern me.

The Chair: Lord Foster signalled that he had a supplementary, but I am very conscious that we are already behind.

Lord Foster of Bath: I can wait; it may be picked up. I will not ask the question, but will just say that I am very interested in Professor

Bradford's remarks about the difficulty of carrying through the transparency requirements in terms of intellectual property. Other colleagues might find a way of referring to that and getting an answer to how we might deal with it in their remarks.

The Chair: That is very helpful. We will move to Baroness Healy, who will concentrate on the US. I am conscious that our other panellists have not yet had an opportunity to comment on anything relating to the EU. I am sure that, as you answer the questions put to you, you will be able to refer back to other territories.

Q48 **Baroness Healy of Primrose Hill:** Welcome, Mr MacCarthy. How is the US addressing generative AI? What are the strengths and weaknesses of that? The question I am really interested in whether there will be a role for legislation finally in the United States. At the moment, it looks like it will be a voluntary but very interesting agreement. I would like to know what you think the way forward is.

Dr Mark MacCarthy: The way the issue was set up by Professor Bradford is an excellent way to think about it. I am struck that, in the United States, Europe and the United Kingdom, there is a sense that this new technology cannot move forward without some sort of regulatory framework. The United States has taken a particular approach in that area. It may seem like a purely voluntary approach, but in fact for the last three Administrations the United States has focused on regulating artificial intelligence in context using existing authority.

If artificial intelligence is used to evaluate creditworthiness, for example, that issue goes to the Consumer Financial Protection Bureau to protect consumers against unfair and biased decision-making. If artificial intelligence is used in a way that amounts to consumer manipulation, it goes to the Federal Trade Commission to make sure that the practices are not unfair or deceptive. If it is used in the workplace to evaluate someone for employment or for promotion, any issues of fairness or bias go to the Equal Employment Opportunity Commission in the United States.

The head of the Federal Trade Commission summarised this approach very recently by saying that "there's no AI exemption" from current law. The idea of regulating artificial intelligence as used is entrenched in the United States system. That is, in part at least, the best place for this. Most of the risks and benefits of artificial intelligence arise as the technology is used in a particular context and a particular business operation. The general rules that might apply to that have meaning only in the context of that particular use.

For example, the AI Bill of Rights that was proposed in the United States in December 2022 talks in general about safe and effective systems, algorithmic discrimination protections and data privacy. It requires some sort of notice and explanation. All of those are then moved to the particular agencies for application. The same is true of the NIST risk management framework that was developed in January 2023. It looks at whether these systems are valid and reliable; whether they are safe;

whether they are secure and resilient; whether they are accountable and transparent; whether they are explainable and interpretable; whether they enhance privacy or intrude into private information; and whether they are fair and manage harmful bias properly. All those questions then turn to the particular agencies.

That approach is how the United States will almost certainly continue to regulate artificial intelligence. It appears as though the United States is going beyond voluntary commitments, but those are all voluntary commitments rather than regulatory requirements.

In Senator Chuck Schumer's artificial intelligence forums, there have been some calls for an artificial intelligence agency like the Federal Communications Commission or the Federal Aviation Administration. Some legislators, such as Senator Blumenthal and Senator Hawley, have proposed a licensing framework administered by an independent oversight body. By the way, there is a new session today at Senator Schumer's AI forum with Marc Andreessen and some other researchers. We may have a report on that very soon.

Senator Schumer's own legislation is all in the area of reports. He asks the financial regulators in the United States to issue a report on artificial intelligence as used in their domain and he asks various parts of the Department of Defense to issue reports, to develop a bug bounty programme and so on. There is also a coming executive order. We do not know in detail what it is going to be about, but it is almost certainly going to focus in general on assessing risks, transparency and avoiding bias.

I do not see in the United States a move towards a single artificial intelligence agency with regulatory power. That might threaten to happen under the European approach. I am a little worried that at that level a general artificial intelligence agency might not be able to be an effective regulator.

I want to say two more things and then I would be very interested in answering your questions. First, there are serious issues to do with copyright that might need to be addressed. We do not know, for example, what the application of copyright law would be for the input of material for the development of artificial intelligence. Should AI developers pay copyright owners for the use of their work? There is no question that it is copying, but is it fair use and no compensation is needed? If an output is similar to copyrighted material, should the owner be compensated? Another question is whether the output is copyrightable at all. Maybe it is just material that goes into the public domain and no one has the right or the opportunity to place it under copyright.

In the United States, the US Copyright Office is conducting a study right now on these questions—I imagine the equivalent in the United Kingdom should be doing the same thing, if it is not already—and asking whether new laws are needed to deal with these issues or whether we should just let this play out in the courts and then adjust to whatever results from court decisions.

This is connected to the liability issues that Professor Bradford mentioned. It is pretty clear in the United States that what we call Section 230 liability, which makes social media companies and others immune from the content of their users, does not apply to the deployers of artificial intelligence systems. Yet if I ask a question to a software program such as ChatGPT and an illegal comment comes forth, it may be very difficult to assign responsibility for that illegal content.

We will hear from Paul Triolo in a moment, but China says that in that circumstance the operator is responsible. As Professor Bradford asked, what about the developers? They know an awful lot about the system that is being used by the deployer. Why should they not share some of the responsibility? Indeed, what if the user circumvents protections that are built into the system? Should that not create some obligations on his part for harmful or illegal use?

Those questions need to be addressed but, in general, on the question of regulating artificial intelligence, the best thing to do is provide the existing agencies with sufficient authority so that they can do their job properly. If that is not sufficient and if, as seems to be the case in some agencies in the United States, they cannot get at the models or the data used by companies under their jurisdiction, that should be fixed. They should be given full authority to inspect the data and the models to make sure that they are operating in a fashion that is consistent with the rules that they are required to implement.

The supplemental approach of giving existing regulators more authority is the better way to go, rather than creating a brand-new agency that would attempt to regulate artificial intelligence across the board in all of its manifestations and contexts.

Baroness Healy of Primrose Hill: That is very clear and interesting. How is the White House approaching competition issues? We know that it is very keen that innovation should bloom.

Dr Mark MacCarthy: Yes, it is very interested in competition issues. It has not reached a result on this yet. There is a danger that threatens in the competitive area. It was brought to my attention most clearly by the Competition and Markets Authority in the United Kingdom. It recently did a report that said, "If you're not going to allow open systems"—systems that are not proprietary, which other entities can build on and use for their own purposes—"you are in danger of preventing innovation and the development of competition". The United States is looking very seriously at whether the calls for licensing and restricting open systems might have an anti-competitive effect.

Baroness Healy of Primrose Hill: How concerned is the US about China? Is there a focus on outcompeting China rather than on regulatory intervention?

Dr Mark MacCarthy: These days, many policymakers throughout the world are concerned about competing with China. As you will hear from Paul Triolo, China is doing a pretty good job of regulating artificial

intelligence, which suggests that regulation need not be antithetical to innovation and the development of competitive systems.

The wrong kind of regulation—the kind that imposes a very large-scale set of bureaucratic overlays, focuses on artificial intelligence as such and attempts to have a single agency providing rules for all uses of artificial intelligence—could very well have an effect on innovation and threaten the ability of US companies to compete effectively with China. Regulation is not a problem in principle but, if done in the wrong way, it can create difficulties and problems that would limit the ability of the US to compete effectively with China.

The Chair: I am going to move on to Lord Hall to ask questions of Mr Triolo, which will focus on China. Then we will come to some broader questions about the implications for the UK of what is happening in other parts of the world. When we get there, we may want to press on the very comprehensive analysis that you have provided and get more of your own opinions on this, which would be really helpful.

Q49 Lord Hall of Birkenhead: Paul Triolo, your evidence has been led into very well by the previous speaker. I was much struck by one expert I was looking at, who said that we need to have a much better understanding of the choices that China is making in terms of regulating generative AI, which is really the cue for you. How is China addressing generative AI? In what way does that differ from the two rather different approaches that we have heard so far this afternoon?

Paul Triolo: That is a great question. Let me quickly make a couple of comments on the great comments made by Mark.

In contrasting China's approach, it is important to note that the US approach to generative AI in particular has been driven by national security concerns. There is a lot of concern around the use of generative AI models, for example, for overcoming cybersecurity or for collating information around bioweapons and nuclear weapons, even what would be considered classified. A lot of the recent US efforts are also being driven very much from a national security point of view.

The other issue that the US and China are both grappling with is the open sourcing of some of these models. That goes to a concern about who gets their hands on the models and what they do with them. That is a really knotty issue right now.

There is also a sense that any new agency in the US could focus just on frontier AI. As Mark rightly notes, there has been this approach and preference for sectoral regulators. One of the discussions in DC is about a new agency that would focus just on frontier AI and generative AI models, because they are arguably different from some of the other applications of AI at a sectoral level.

In China, the situation is very different, with a lot of different drivers and motivations. First, it is important to note that China has also been dealing with the regulatory issues around AI for some time. Arguably, it has been thinking about it as long as the EU. Back in 2017, the Chinese

Government put out a national AI development plan. That plan sought to both encourage innovation in AI and give a major nod to developing a regulatory framework domestically and engaging at the international level. The Chinese regulators have been thinking about this issue for a long time.

Secondly, it is important to note that the Chinese system in general is very pro-technology. There is a sense that technology is a positive force that brings benefits to people. That has meant a couple of things. There has not been this huge debate in China about the downsides of AI, including, at the tactical level, things such as facial recognition. People generally think that some of these applications are useful. You can walk into a store in China, scan your face and buy things without having to engage with a human. People like that, in general. It has not got the same sort of attention as it has in other jurisdictions.

There has not been this robust debate that we have seen over the last six months about the potential existential risks around AI. In China, very few people signed the one-sentence risk letter from the Center for AI Safety in May. Most of the debate in China is around the positive elements of AI. There is not really a huge focus on what Yann LeCun has called the doomsday camp. Those are important distinctions.

The other difference relates to the legal issues that were mentioned by both Professor Bradford and Mark about platforms being held liable and competition issues. That sort of legal and litigatory approach to regulation is not really part of the Chinese regulatory domain. While companies in China, and platform companies in particular, have undergone a fairly serious round of regulation of their business models over the last two years, they are not facing a steady stream of litigation around their competitive practices, as companies in the US are.

Let me just touch briefly on the approach China has taken here, which is really important because it is an evolutionary approach that has developed pretty quickly. As I noted, even before the advent of ChatGPT, Chinese regulators, for example, were looking at some sort of licensing system for AI algorithms. They were looking at regulation around things such as deepfakes.

The advent of generative AI really threw the regulatory system into a bit of overdrive in China. In the last six months, for example, we have seen China quickly developing interim measures for generative AI. There was an initial draft that came out in April. There was a considerable amount of industry input into those draft measures. They quickly issued the draft measures in July, and they went into effect on 15 August. There has been a pretty quick regulatory approach to this.

These are called interim measures for a reason. Like regulators everywhere, the Chinese regulators are still trying to get their heads around the different issues around AI, particularly generative AI, and how to regulate it. The other important factor here, though, is the very clear sense you get from reading the generative AI measures, for example, that the Chinese Government are trying to balance the innovation and encouragement part with regulation.

Some of the measures in the generative AI provisions, for example, were toned down pretty considerably in the final version that came out this summer. That reflected a very high level of input from the leading Chinese AI companies and the academic community in China around AI. There has been a very vigorous back and forth between the legislators, the regulators, the companies and the AI community in China.

In addition, the Chinese system is arguably set up to do more regulation. The focus of the primary AI regulator, the Cyberspace Administration of China, is very much on content because of the priorities of the Chinese Government around controlling content and ensuring that new technologies such as generative AI are going to comport with existing rules around censorship, for example. The Cyberspace Administration of China is the lead regulator, at least at this point, because of this high focus on content. But even the CAC, which is concerned about content, has still included in the interim measures a nod to encouraging companies to collaborate and encouraging the overall innovation of the system. It is not just about content regulation.

The other important issue to note is that the Chinese are taking an integrated approach. It is a mix of legislative approaches, administrative regulation and things such as technical standards. China has a very well-developed division of labour in this regard, so it is also important to focus on things such as standards. In the last couple of months, the Chinese standards body for information technology and security has issued several standards around things such as watermarking, how the input data for models should be tested and how to test the outputs of models for accuracy.

These build on the interim measures that I mentioned. There is an iterative process whereby the standards bodies are fleshing out some of the vague guidance that was provided in the interim measures and saying, "Here's what you really have to do to assess the validity of your data and the types of data you're using as inputs. Here are some very detailed metrics about how you assess the output of your models".

All of this builds up to what will probably happen later in the year, which is a broader draft AI law that will draw in part on this iterative approach that the Chinese Government have taken to regulation of specific areas such as content. Again, this is typical of the Chinese system: beginning with some administrative measures, giving companies a sense of where things are headed and then eventually codifying the various pieces of regulation and, in this case, standards into some kind of broader law that will provide a very well-developed framework. In developing that law, there will be lots of on-the-ground experience, through dealing with the companies, talking to them and trying to understand where the sector is going in China.

It is important to note that in China there is no real discussion of frontier AI. It is a tricky issue. That term is not generally used in China, so a distinction is being made in the regulatory space. In the upcoming UK AI summit, for example, the focus is very much on frontier AI. In China, that is not really part of the discussion. That term has not really been

translated into Chinese and used as part of the discussion. The focus in China is on generative AI. They are not necessarily thinking about frontier AI almost as a separate category.

This well-developed system is leading up to what will eventually be this broader law. The important takeaway here is that it is really drawing on all elements of the system. Over the last few months, I have talked with most of the leading AI companies in China about how they are interacting. They are very pleased with the way they have been able to influence the Government's approach to regulation in this space. They are being listened to. The standards process, for example, also involves almost all of the leading AI companies and the other key standard-setting bodies within China. There is a lot of internal debate on the nitty-gritty of how you do regulation and how you tackle that.

The bigger issue is how China fits into the broader global debate around the regulation of frontier AI or generative AI. Last week the system spat out a little clue as to China's thinking on this with the release of the global AI governance initiative, which was put out by the Cyberspace Administration of China but was probably generated by other parts of the Chinese bureaucracy, such as the Ministry of Science and Technology and the Ministry of Foreign Affairs.

This is an attempt by China to lay down a marker on how it believes AI should be regulated. It is also a message to the global South that China is committed to carrying its interests and priorities. It is also an anticipation of the UK AI summit next week, which will tackle some of these issues and which China will likely participate in, at least on day one.

The issue of how China plays internationally is important. Again, it is coming at this with a fairly well-developed set of tools. It has a process in train domestically and is now prepared to determine how it will engage in the broader global debate about regulating frontier AI. The Chinese approach is very different and well developed, but there are definitely some areas of overlap with the US and the EU, which we can discuss.

Finally, taking the voluntary commitments from the White House as a good example, when you talk to Chinese academics, they say that these are things that China could accept because they are pretty well known and not that surprising, but they would prefer it not to be branded as being from the US. They would prefer it to be branded as global best practice as part of that process. The Chinese Government are definitely willing to engage globally, on the basis of a fairly well-developed domestic approach to regulating AI.

Lord Hall of Birkenhead: I was going to ask you about international collaboration, but you have answered that question. How does the business landscape in China for large language models compare to that of the US, for example?

Paul Triolo: That is a great question. There are similarities and a few differences. The leading developers of foundation models in China, as in the US, are largely the big cloud companies, the hyper-scale companies

that have large datasets and the compute resources to develop large language models. There is certainly a similarity there.

In the case of generative AI in China, from very early on there was not this sense of rushing to release models for public use. From their long experience, Chinese companies realised that, because of content and censorship-related issues, they would likely get in trouble if they released these models, like OpenAI did, suddenly to the world and generated a lot of concern.

Chinese companies have been focused over the last year on enterprise deployments of AI models. Even before many of them were approved by the Cyberspace Administration of China in August and September to release their generative AI models for public chatbot-like interactions, they had begun to license and allow enterprises to access those models and to use them to begin thinking about how they would be applied in enterprise-type applications.

The other point I would make is not really a difference; it is a similarity. Some of the integrated Chinese companies such as Baidu, Huawei and Alibaba are also focusing on using their generative AI models as part of their existing product line. For example, Baidu is integrating its latest and greatest model, which it released last week, ERNIE 4.0, into all its products: its search engines, its mapping programs and those kinds of things. Chinese companies are first looking at how they apply generative AI to their existing product line and then focusing on the enterprise side of things and licensing their models for enterprises. That is also being done by US companies, such as OpenAI, Anthropic and Inflection AI, as well as Cohere in Canada. These companies are also internally very focused on the B2B space.

The other similarity is this open-sourcing of models, which I mentioned. Over the last few months, most of the major players in China, with some interesting exceptions, have open-sourced their models. By the way, they have also hosted other open-source models, such as Llama 2 from Meta, on their cloud services, which has caused concern in some circles in Washington because it is allowing Chinese companies and developers to have access to fairly advanced models. Once Meta had open-sourced Llama 2, a lot of Chinese companies decided that they would follow suit. They are very closely watching and learning from developments in the US.

The other big issue which Chinese companies are concerned about, as US companies are but for different reasons, is access to advanced compute. In the US, when you go to Silicon Valley, you find that every venture capital company out there is very concerned about access to advanced compute. All of these companies, both small and large, are attempting to train their models. There is a real shortage of GPUs, particularly from Nvidia. Everybody is scrambling to figure that out.

Of course, just last week the US Government made it more difficult for all the major players in China to access advanced GPUs, particularly from Nvidia, with the release of the new export control package. Chinese companies are very much plugged into the global debate, Silicon Valley

and many aspects of generative AI model development. They are now very concerned about their long-term ability to access a reliable source of GPUs to do advanced compute and train their models. They are certainly more concerned with that right now than with government regulation around AI, for example.

The Chair: Just before I move on to Lord Kamall to kick us off on the final category of questions, on that last point, Mr Triolo, how ready is China to create its own GPUs?

Paul Triolo: Over the last two years, a sizeable number of GPU start-ups in China have begun to attempt to develop an ecosystem of hardware and software to compete head-on with companies such as Nvidia and AMD in the US.

On 17 October, two of those leading companies, Biren Technology and Moore Threads, were put on the US entity list. That came as a bit of a shock to the Chinese industry. Those companies were doing fairly advanced GPU designs. They were having those manufactured in Taiwan at TSMC. All of the advanced GPUs in the world are manufactured and then packaged within a small radius of Tainan in Taiwan at TSMC's fabs.

As of last week, a number of Chinese companies were moving pretty quickly to develop capabilities that were certainly capable of competing with Nvidia at the cutting edge. One of the challenges in this space, of course, is the software development environment. Nvidia has a very well-advanced software development environment that is used widely by AI developers around the world. Chinese companies face a big challenge there.

Some of them were developing systems that would be compatible with some of the Nvidia software. Some of those companies are peopled by engineers and managers from leading US companies. The US Government's actions as of last week in particular have made it very difficult and have put some big obstacles in front of Chinese companies that are competing in the GPU space.

Huawei is going to be a big player in all this because it is developing its own semiconductors that can be used for training AI models. Huawei is working very closely with Chinese semiconductor manufacturers to try to overcome the US controls on semiconductor manufacturing tools in particular. The situation is very much in flux in China right now, but certainly there are many good companies in China that have a lot of expertise in this area. They will be pursuing this development despite some of these recent controls.

Q50 Lord Kamall: I am going to ask the obvious question. It is a sweep-up question, if you like. What are the key lessons for UK policymakers and regulators to take away from the three international counterparts you have talked about? I also want to ask a slight variation of that. Are there other countries that we should be looking at—India, Japan or Russia? Are they irrelevant? Are they going to fall in line with one of the three?

Then, on the role of international organisations, I noticed that the OECD,

G7 and G20 have general non-binding principles. The UN still has the multistakeholder High-Level Advisory Body on Artificial Intelligence. We also see the EU-US conversations. That might well cause some nervousness in the global South and in China. The global South might hide behind China for that.

What should the UK learn from these three countries? Are there other countries? What about the role of some of the international governmental organisations?

Professor Anu Bradford: It is absolutely the right question to ask. In terms of the obvious models to look at, I would start with these questions. What are the values? What is the kind of digital society and AI revolution that serves UK citizens, the country's economic and political goals and its societal interests?

The UK and the EU are not that far apart. If you look at how individuals describe their approach to a digital society in various public opinion surveys, there seems to be a general demand for the protection of fundamental rights and the preservation of democratic structures. The EU model has a lot to give the UK there.

There is another reason. It is just pragmatic or, if I may use the word, inevitable for the UK to look closely at what the EU is doing. My earlier book, *The Brussels Effect: How the European Union Rules the World*, explains how EU regulations often have an influence outside of the EU because it is, especially for the UK, a large and proximate market that is essential for UK companies' exports.

Simply by regulating the single market, the EU has often managed to export its regulations around the world because the companies that depend on access to the EU market will need to comply with the EU AI Act regardless. They often choose to apply EU regulations across their global production or global conduct because they want to avoid the cost of producing multiple variations for different regulatory regimes.

There is a logic in large language models: the more data you deploy for your model, the more powerful your generative AI tools will be. If UK companies want to offer their products within the EU, which is a very big and proximate market to carve out if you would rather trade elsewhere, they need to comply with the EU's AI Act. That AI Act is forthcoming. It will be constraining and will impose obligations that affect any company doing business.

If UK companies need to comply with the EU's AI Act anyway, the question is whether they will want to comply with a different set of obligations for the UK market. If they want to avoid the EU's regulations, they will need to retrain those models for use in other markets based on data that excludes European data. That would mean throwing away a lot of valuable data. There is a gravity model or economic logic to having close dynamic regulatory alignment with the EU for that reason alone.

Let me use this opportunity to say more about ideological differences. We have heard today about the Chinese approach, the American approach and the European approach. There is a common value

foundation. Even though there are differences between the US approach and the EU approach, as Mark MacCarthy explained, there is still a deep commitment to ensure that liberal democracy is preserved and that civil liberties and fundamental rights will be protected in the AI age.

There is a case for close collaboration among the world's techno-democracies, in particular to have a united front and to offer an alternative. Paul Triolo was talking about China's approach to global AI. There are some fundamental differences, such as the use of facial recognition for mass surveillance.

I commend the UK for hosting the global AI summit. There needs to be closer engagement among like-minded countries, including the techno-democracies, as well as with those that are key players, including China. The UK is in a very good position to be central to those conversations. The main takeaway is that it is very hard for the UK to avoid the economic logic that comes from its proximity and economic opportunities tied to the EU.

Dr Mark MacCarthy: That is the key question. I strongly recommend that the UK continues to engage with China on these questions about AI regulation. I just came out with a paper for Brookings arguing that it would be in the interests of the United States, the UK and the EU to continue to engage in that area despite our differences in connection with digital authoritarianism versus digital democracy.

If you look at the mechanisms that China has developed, it may have gone down a direction that we could learn from or one that we think is mistaken. When it put out regulations about responsibility being put directly on the operator who deals with the public, that was a little too narrow. When we approach that question, we should open it up and think about why it should not be the developer or the user. It might be useful to talk with the Chinese regulators and ask, "Why did you go that way rather than some other way?" We might be able to learn from them.

On the question of access to data and models, China is pretty clear: if a company uses a model in its business activity, that model has to be registered with the regulator. To the extent that the UK or the United States is looking at regulatory authority over companies that use AI models, they clearly have to be given the authority to gain access to the underlying data and models.

The UK is running into that issue as it begins to develop its AI safety approach. An institute has been set up to talk about frontier AI models. It is an advisory group; it is not a regulatory group. Clearly, that group wants to get access to models and to evaluate the risks. The related idea from Eric Schmidt is that we should have an equivalent of the Intergovernmental Panel on Climate Change to say what is going on, what the capabilities are and what the implications are in order to give good information to the UK, the United States and other agencies.

Those two different kinds of groups, one of which is a capabilities assessment group and one of which is a risk assessment group, will be

crucial. These will be voluntary agencies that can advise Governments. To do their job properly, they need access to models and data.

The proper regulatory approach is not to go in the direction of regulating general AI but to assess the capabilities and risks of general AI. As these systems are used in practice, we need to make sure that the competent sectoral regulators have adequate authority to do their job.

Paul Triolo: It is important to focus on where we are at. The current debate around frontier AI reflects the growing concern about the national security risks, but significant questions remain about what specific policies will be effective in mitigating these risks, how they could work in practice and the potential costs that a more restrictive environment could impose on beneficial AI innovation. All Governments, not just the UK, are trying to balance these factors.

The situation is complicated by the fact that, while frontier AI risks are easy to imagine, they are not very well understood. Leading AI developers themselves, for example, do not fully understand everything their models are capable of or all the ways in which they could be exploited by bad actors. That gets to some of Mark's concerns.

Over time, both AI systems and policies for managing their risks will evolve as companies and regulators gain more experience with technology. One key step here is gaining more experience with the technology. There is no substitute for that. As the process unfolds, Governments will have to answer difficult questions about what degree of AI risk they are willing to tolerate and how to design effective policies for addressing risks without stifling the beneficial uses of AI.

Bringing China back into this, any regulatory framework is going to have to consider how to control access to both proprietary and open-source models and datasets, as well as to the significant compute capabilities required to train and deploy these advanced models. Any regulatory framework that arises out of this is going to have to deal with interoperability. Again, that is where China comes in.

I agree with Mark's comments. Finding some level of interoperability and figuring out how to understand things such as licensing and the provenance of data will be really critical to that process of determining how systems are going to be interoperable, to avoid a world where we have two different regulatory approaches to frontier AI.

In terms of specific recommendations for what the UK Government should do, this is already happening, but any regulatory approach should be developed in close collaboration with industry at many levels, as well as with academics and other institutions and stakeholders in the regulatory process. When devising approaches, there is a need to address the long-term national security-related issues that were mentioned earlier. That should be a major focus of Governments. That is where Governments have an edge, as it were. This should not be the only focus. There are a host of other shorter-term issues, which are well known, such as bias, discrimination and transparency over how AI algorithms work.

It is really important to note that the regulatory approach should be balanced, as I have noted, with efforts to encourage innovation. The needs of the large players should be balanced with those of new entrants in the space, as was addressed very well by Professor Bradford. The important areas where Governments can play a positive role include developing a national approach to advanced compute in order to lower the costs of entry, given the high cost of training large language models.

Again, China is very much in the lead on this. I did not get a chance to mention that, but one of the approaches that the Government are taking is, in addition to regulation, to develop a very robust advanced compute strategy, which includes national initiatives to develop data centres and provide public access to advanced compute. Again, that is another area where Governments can play a particularly strong role.

Finally, there is this really knotty issue of how to integrate some of these diverse approaches to regulating. The UK has a bit of an advantage here in not being part of the EU. The AI summit coming up next week is a really important opportunity for the UK to be a convener—I hope it continues to be a convener after the event—and bring these tough issues, particularly around the integration of China and Chinese companies, into the broader discussion.

I know China will be attending on day one. The establishment of an experts group is likely and the willingness of China to participate in that will be a really important marker. The UK can uniquely provide that convening venue and continue to engage at a high level on these issues with countries that are not among the like-minded democratic countries that Professor Bradford alluded to.

It is important to keep China in the conversation going forward and not to exclude it right off the bat. That would be very counterproductive over the long run in getting control of frontier AI systems and developing some kind of common framework that all countries can agree to, recognising that there is going to be a fair amount of difference in the way these regulations are implemented domestically.

Lord Kamall: Just very quickly, to Dr MacCarthy and Professor Bradford, Mr Triolo mentioned the UK AI summit. What are your hopes and criticisms for it? What are you hoping will come out of this AI summit? What would be a good thing? What are your concerns?

Dr Mark MacCarthy: One of the best things that could come out would be some understanding of an international or global AI safety institute that would focus on evaluating both risks and capabilities. That would be one good deliverable. The second would be establishing that China is a significant player in this area, which can participate successfully and make a substantial contribution.

Professor Anu Bradford: The dialogue itself is already valuable. I am not as hopeful that this will be the beginning of an international AI treaty. The differences and interests are just too far apart. It would be naive and set the UK up for failure if we expected this to be the beginning of something truly constraining or ambitious.

At the same time, even if we do not pursue agreement on issues for which we know agreement is not feasible, this is an opportunity to launch a global gathering of better information and knowledge about AI. There have been proposals to have something like the Intergovernmental Panel on Climate Change. That way, we would all be working with the same baseline of information and objective scientific research, which would be removed from the political disagreements surrounding AI. That would be something concrete, and it could be initiated. That would be a real contribution for the UK to have that happen in a meeting convened by the UK.

The Chair: Thank you very much to all three of you. I am very grateful for the time that you have given us this afternoon from whichever time zone you are dialling in from. We have covered a lot of ground. Clearly, there are lots of differences between you, but none the less you have given us very rich material to reflect on as we think about the role of the UK, heading into next week's AI summit. On that note, I will draw this to a close. Thank you again. I hope to meet you in person at a future opportunity.