



HOUSES OF PARLIAMENT

Joint Committee on Human Rights

Uncorrected oral evidence: Human Rights and the regulation of artificial intelligence, HC 1262

Wednesday 25 February 2026

3.50 pm

Watch the meeting

Members present: Lord Alton of Liverpool (The Chair); Baroness Chakrabarti; Baroness Hamwee; Lord Murray of Blidworth; Lord Rook; Lord Sewell of Sanderstead; Peter Swallow; Sir Desmond Swayne.

Also present: Dame Chi Onwurah.

Questions 99 - 111

Witness

I: Kanishka Narayan MP, Parliamentary Under-Secretary of State, Department for Science, Innovation and Technology, and Minister for AI Opportunities.

USE OF THE TRANSCRIPT

1. This is an uncorrected transcript of evidence taken in public and webcast on www.parliamentlive.tv.
2. Any public use of, or reference to, the contents should make clear that neither Members nor witnesses have had the opportunity to correct the record. If in doubt as to the propriety of using the transcript, please contact the Clerk of the Committee.
3. Members and witnesses are asked to send corrections to the Clerk of the Committee within 14 days of receipt.

Examination of witness

Kanishka Narayan MP.

Q99 The Chair: Welcome back to the Joint Committee on Human Rights and the session that we are holding today on artificial intelligence and human rights. We have heard already this afternoon from two of the big players—the corporations involved in this. We were listening to Microsoft and to Meta; we previously heard from Google. Now we come to the ultimate part of our oral sessions, and that is to hear the Minister, Kanishka Narayan, who is Minister for AI in the Department for Science, Innovation and Technology. He was appointed as Parliamentary Under-Secretary of State in the Department on 7 September 2025, and was elected as Member of Parliament for Vale of Glamorgan in July 2024.

Minister, we know that you have also recently been in India, in New Delhi, at the conference that was being held there—it was referred to in our earlier session—so we know that you are probably a bit jet-lagged. For that reason alone, we are very grateful to you for spending time and coming here today to answer our questions. You know many of my colleagues from around the table; we have six members from the Commons and six members from the Lords.

As the name implies, we are a committee whose mandate is to look at human rights, where it is written into the DNA of each and every department, and especially when sometimes departments themselves feel that their only preoccupation should be with the subject matter that they are dealing with. We occasionally have to say to them that they have other responsibilities, too. You probably know what I am driving at in that remark, and I know that one or two of my colleagues will explore it further in terms of how we ensure that there is cross-departmental interest in human rights questions.

I will ask a curtain-raising question before we turn first to my colleague Baroness Chakrabarti, who will ask you about upstreaming risk and pre-deployment testing. I would like to ask a general question about the Government's approach to AI regulation. We heard a lot about the AI Opportunities Action Plan in earlier evidence to the committee. It is being criticised for prioritising innovation at the cost of effective regulation. I wonder whether you agree that the desire for innovation should not come at the cost of human rights protections, whether you agree that strong guardrails in AI systems can prevent harm and promote public trust, and that without that public trust, the benefits of artificial intelligence for the whole of the community would not be properly realised. Minister, the floor is yours.

Kanishka Narayan: Great. Thank you, Chair, and thank you also for both your generous welcome and empathy with my jet lag. The enthusiasm to be here to talk to you has carried me through the last 48 hours or so. It is because the question of assessment before the committee and the questions that you ask are very much the priority for me and the Government, which is to make sure that, yes, of course, the

United Kingdom is at the frontier of capturing the opportunity of AI, but while recognising that obviously we do so only if people adopt it with trust and we are able to make sure that we carry the entire country with us in that path.

I have just been in India, as you said, and the central question there was how we treat the question of adoption and innovation in sync with the question of trust, security and safety. Too often, we think of these as involving trade-offs. I am much more interested in finding the very significant overlaps in those two pursuits as well. In that spirit, what we have been focused on in government has been understanding the building blocks of trust; first, capability in the public sector, not least through the AI Security Institute; secondly, having built that capability, being very quick to make sure that where things are going wrong, we are able to move quickly in regulating—I am happy to talk about some of the aspects of online safety in that context—and thirdly, alongside that, to recognise that static forms of regulation, where we have been historically, may not be best adapted to where we are likely to go. That is why we have moved very promptly in some initiatives, not least through the Regulatory Innovation Office as well as the initiative we have with the AI Growth Lab, to make sure that our regulatory context builds trust in a dynamic way over time.

The Chair: Thank you very much indeed. Lady Chakrabarti, and then we have some further questions.

Q100 **Baroness Chakrabarti:** Minister, the committee has already heard from several witnesses who suggest that managing risk upstream—so right at the design and development stage of the AI life cycle—is the best approach to preventing human rights harms. Could you tell us how that upstream risk is managed and monitored in the UK?

Kanishka Narayan: One of the things I would flag is that we should talk about risks that are specific. One of the views we have taken is that when we are able to talk about the particular risks and harms that might be of concern, we are able to regulate it better. AI developers will also be subject to the series of regulations that we have, whether that is from a privacy point of view through our existing legislation, data protection, the Human Rights Act, the public sector equality duty when it is being deployed in the public sector, and indeed the Equalities Act more generally.

The starting point is that there is no difference whether a technology is AI technology, linear regression technology, or other statistical methods as regards liability and responsibilities on the existing pieces of legislation.

The second point is that we have to make sure that we still build public sector capability to apply those pieces of legislation and those regimes of regulation appropriately to AI as well. As I mentioned in my prior response, that has been a big part of our focus. The AI Security Institute complements the regulatory regimes we have with an even greater depth of understanding of these models prior to deployment.

You will have heard from some of the regulators, I hope, but I am aware that the AI Security Institute, for example, co-ordinates and often speaks to representatives in Ofcom—as do I—to try to ensure that Ofcom’s regulatory functions are being appropriately deployed upstream, as well as subsequently, in light of the information that it understands from other regulators and government bodies.

Baroness Chakrabarti: In your department’s written evidence to the committee, it stated that “organisations deploying AI remain accountable for discriminatory outputs”. I am assuming that includes government departments when they are deploying AI systems. Is that your position?

Kanishka Narayan: Yes, both in legislation and, of course, in terms of the public sector equality duty, which very much applies as well.

Baroness Chakrabarti: The public sector equality duty is one thing. The Human Rights Act that you mentioned applies only to public authorities, so there is a potential gap there in relation to private deployers of AI.

Kanishka Narayan: As you will be way more aware than me, private sector bodies—particularly in an employment context—are very much subject to the equalities provisions.

Baroness Chakrabarti: In the employment context, okay. At present, are there any mandatory conditions that must be met before AI systems can be deployed either in the public or the private sector: for example, disclosure requirements, impact assessments, declarations of compliance and so on?

Kanishka Narayan: Model developers are subject to a series of regulatory regimes that we have across the harm categories that I mentioned. In most of those cases they have a liability to ensure that they do not do anything that falls foul of the law. I do not have a single answer across the whole range of regimes as to whether all or some of them require mandatory disclosures ahead of time, but they are regulated in the usual ways across the range of harm categories.

There is one area, the ability to create child sexual abuse material, where we have deemed it particularly important to clarify that there is a distinct criminal liability position that model developers will be subject to, in addition to the wide range of regulatory burdens on them already.

Baroness Chakrabarti: Which UK regulators can currently block AI systems from being deployed if significant upstream risks have been identified?

Kanishka Narayan: The primary aspect that comes to mind is that, under the provisions of the Data (Use and Access) Act 2025, we have specific conditions under which automated decision-making can take place. That feels like an obvious place where there is a pre-emptive requirement to satisfy very particular conditions of deployment. In the absence of those conditions being met, automated decision-making deployment should not be taking place.

If I think about other areas, a lot of the other areas of regulation rely on a reliability regime where the regulator then assesses ex post compliance rather than ex ante, as is the case with other technologies, too.

Baroness Chakrabarti: I think that this is a concern, and it is something that we discussed with the regulators when they were sitting in your place a little while ago. We have some very potent technology here, which I agree is potentially very beneficial, but it is potentially very harmful as well. For example, unlike medicines or other potent technology that can be used for good or ill, there is not really a prior approval regime in this country, is there?

Kanishka Narayan: What we have broadly done is, as with other statistical methods—and artificial intelligence is effectively a general-purpose way of analysing large volumes of data—that we have a series of regulatory regimes, as I have mentioned. In most of those cases, the focus is on then assessing compliance with the regulation post release.

What we have done very distinctly here is to say, while it is not medicine—it is much more similar to, say, the internet—we will none the less build public sector capability, very uniquely, globally, in the AI Security Institute, to be able to work very closely with large model developers internationally, but still assess them pre-release and be able to support their work in mitigating risks pre-release. That combination of robust regulatory enforcement ex post and this understanding and co-ordination ex ante feels to me an appropriate starting position but, of course, I take into account the points you are making.

The Chair: Thank you. I think, Minister, you are taking us into the next area of questioning. Sir Desmond will ask you a question in a moment on the Council of Europe framework documents. I should say that a number of members of the committee—Lord Murray and others—have been exploring this during a visit to Strasbourg and, since then, with members of the Council of Europe. Then I know that Lord Rook wants to take us on to further questions about the global arrangements that we can make and perhaps touch on your recent experiences in India. First of all, Sir Desmond.

Q101 **Sir Desmond Swayne:** We signed up to the Council of Europe Framework Convention 18 months ago. We have yet to ratify it—to be fair, none of the other signatories have ratified it yet. Is there anything that you are aware of that stands in our way in terms of being able to deliver on the commitments that currently prevent us from proceeding with that ratification?

Kanishka Narayan: I appreciate that question. I have two points to make. The introductory comment I would make is that across government we are focused on the UK playing a significant role in international co-ordination on the question of risk mitigation in AI. Therefore, the fact that it was the first legally binding agreement globally on AI grounded on human rights, democracy and the rule of law was an

important moment. I would be delighted to talk about how we have continued those conversations in India.

On the particular question, as you will know very well, the convention largely focuses on existing human rights regimes within countries and to ensure that countries comply with them when it comes to the question of AI. To that extent, I feel no significant obstacles in us being able to put that into effect. It is just that the programme of work required to do so—given the sweeping scope of it—is pretty extensive, so we are making sure that we are working hard on that.

Sir Desmond Swayne: The extension of scope, as I understand it, is where it places an obligation on governments not just to regulate public authorities and the private sector, acting on behalf of public authorities, but on all private sector obligations. How is that to be achieved?

Kanishka Narayan: We are still working through how we make sure that we adapt what we have committed to—which is that, as a country, we want to make sure that AI is grounded in human rights, democracy and law—to our existing law but also the regulatory regime we have, which is primarily sector-led rather than sweeping across the economy horizontally. Whether it ends up being the scope expansion that you are describing is still an open question.

Baroness Chakrabarti: Do you see a gap in our current regulation of AI, AI products, and deployment of those products?

Kanishka Narayan: Yes. Taking a very harms-focused approach, when it came to the question of online safety—not least through conversations across Parliament—we felt that there was a potential gap related to risk for AI chatbots that were not in scope of the Online Safety Act. Currently, user-to-user engagement, search services and regulated porn services are in scope of the Online Safety Act, but AI chatbots that do not involve user-to-user engagement or do not involve live search are not.

Now, largely speaking, a lot of the risk association with chatbots was already in scope of the Act, but we felt that there was a theoretical possibility that there would be some gaps in that. Therefore, we have sought permissive powers in the Crime and Policing Bill to be able to make sure that we have the powers to close any of those gaps as well.

The Chair: Those gaps include the tragedies that we have seen: the people consulting AI as to how to commit suicide, and then one young person doing precisely that.

Kanishka Narayan: Yes. The first thing I would say is that it is absolutely horrific. One of the motivations to making sure that we are able to move fast on reviewing any gaps and then seeking the power to move very fast on it was through the testimony of both parents in particular. The Secretary of State and I met a series of bereaved families who had those experiences. I have to clarify that that was not in the context of AI chatbots but more generally online experiences.

Secondly, one of the first things the Secretary of State did when the two of us were appointed was to deem self-harm and suicide content in the Online Safety Act a priority offence. That means that platforms do not just have a retrospective duty to react to that content but a proactive requirement to look for it and to root it out. That also applies to AI chatbots in scope of the Online Safety Act. Where AI chatbots are not in scope, as I mentioned, we have sought the permissive power to now be able to bring them in scope.

Q102 Lord Rook: Minister, you and I have had the conversation before about the pride we can take in this country leading on AI safety, and the previous Government's work to convene the AI Safety Summit at Bletchley Park. I had the privilege of working particularly with faith communities and civil society groups to look at how we could harness AI for human flourishing, but also to protect against potential harms. Clearly, as a country, we were at the forefront of trying to convene the world in that conversation.

I am interested to hear from you—straight back from India, despite the jet lag—where you think that conversation is right now. Are we still using international collaboration to deal with some of these challenges and opportunities? Also, are we as a country still taking a leading role in trying to push that conversation forward?

Kanishka Narayan: Yes. I appreciate this question because it goes to the heart of how we are able to move together, especially in a technological context, which is much more global than specific sectors. I will make one overarching observation on how the conversation has shifted and will then answer the specific question on our commitment.

In terms of the overarching development, the thing I experienced in India was that where we were at Bletchley Park, which was a remarkable moment of UK leadership—I have to commend, cross-party and cross-Governments, the work of the last Government in doing that—and where we have moved to now is that the conversation is much more focused on the day-to-day experience of people rather than much more abstract, very long-term questions of how AI might either fundamentally transform the economy long term or fundamentally might pose risks. It is much more rooted in the lived experience of people now. That is exactly what I felt in India.

On the question of adoption and opportunity, the argument was not how you create massive large language models with billions and billions of investment, but how you create small models deployed on people's devices, giving them privacy as well as the power in their palms for billions of people in conditions of economic deprivation.

Similarly, on the risk question, the focus was very much on not just saying, "Hey, AI is a cross-sector, massive, global, long-term risk", but much more about what the risks are that we are experiencing today, looking at the lens of harm and then regulating or making progress on mitigating those harms.

We had some very fruitful conversations with the Indian Government and a series of other Governments to look at questions of, for example, child safety and violence against women and girls—the very acute risks that people are experiencing here and now—and all of us are trying to move promptly on those questions. I found that a very helpful shift from where I think the conversation was to where the conversation is now, which is much more in sync with the lived experience of people.

The second thing on the question of British commitment is that we are absolutely committed to making sure that the UK continues to play a leading role in the international conversation here. Part of that is through the ongoing capability of the AI Security Institute and an alignment fund that we have now committed to with partner countries, which remains at the forefront of a lot of the state of the science in understanding risk. However, part of it is also engaging very deeply with allies in thinking about how we are moving in an appropriate way when it comes to wider questions of multilateral governance.

On that question, I simply say that the question of how we engage all society—particularly spiritual groups and the wider communities that we have—is increasingly at the heart of it. I felt that the prior conversation at Bletchley was a very insular, elite conversation. I think that there were 250 people or so there talking about super long-term risk and the opportunities of AI. At the Indian summit there were 250,000 people, including faith leaders. I think that was reflected in the fact that we are now talking much more about specific opportunities, specific harms, and specific mitigations.

Lord Rook: That is very helpful. Thank you. Some of us are hearing rumours—which we hope are just rumours—that our eagerness in the UK to progress a tech prosperity deal with the USA might be leading us to be less connected with other countries around the world and reduce our international engagement and just back the US as the main horse in this race. I am interested to know your perspective on that. Are we attaching ourselves more strongly to the US or are we still working with other global actors, particularly on issues of AI safety? Also, are we actively involved in looking at how AI safety protects human rights in different constituencies around the world?

Kanishka Narayan: I appreciate that question. The United States is, of course, a very key ally for us when it comes to the question of technology. That said, I have personally spent time in both South Korea and now India. When I was in India, I did not just meet the Indian Government but met the Australian Minister, the French Minister, the Norwegian Minister, and beyond as well. That was precisely because we want to make sure that the UK continues to play a role, not just in our bilateral engagement with the United States but much more broadly as well.

On the broader question you raised, I will flag that the work that we are doing on AI safety and security more generally is very much conscious of the impacts that other countries and people right across the world are

experiencing, too. That is part of the reason why, not just through the alignment fund but more generally through a series of engagements, one of the things I did in India was that the UK led the effort, alongside Professor Yoshua Bengio, to make sure that we launched the International AI Safety Report. It was a report when the secretariat was in the UK, and both Professor Bengio and I were present at the event to show that the UK remains an important contributor to that conversation globally.

Q103 Dame Chi Onwurah: Kanishka, it is great to have you here. I know you have a lot of experience in AI from before you entered Parliament. I want to ask you about the current regulatory framework, but did I hear you say that AI was just another statistical analysis model? I think that you might cause an AI market crash if you are saying it is analysis as opposed to content generation.

Kanishka Narayan: What I was saying is that it is general purpose in the way that, say, statistical methods are, by which I mean effectively that it is a way of having data inputs that would be used to generate material and a range of outputs. Now, of course, linear regression is much more limited in its scope as a statistical method, but we should not think of this as some radically new magical thing. It is simply a pretty significant advance in capability but, at the heart of it, an algorithmic method that is not novel.

Dame Chi Onwurah: It may not be radically new as a technology, and I would agree with that, but its impacts are new.

Kanishka Narayan: Sure.

Dame Chi Onwurah: Throughout the inquiry, we have heard evidence that the current regulatory framework for AI in the UK is fragmented. Back in 2023, the predecessor Science, Innovation and Technology Committee asked for a regulatory gap analysis for AI, and when we had the Secretary of State, Liz Kendall, before the current committee, she promised to bring in appropriate legislation to address any gaps in AI regulation and in a framework. Has that regulatory gap analysis been completed, to your knowledge and, if so, what were the findings and recommendations?

Kanishka Narayan: The main thing I would say on this, Chi, is also that looking from a harms-based point of view, we have said, "Where are there gaps?" In this context, for example, on online safety, as I mentioned earlier, we found that there was a potential gap and we put in place legislative measures to be able to act on it very promptly, as you will know very well.

Dame Chi Onwurah: You have done a harms-based regulatory gap analysis.

Kanishka Narayan: Well, I have not seen a single look that says, "Where are the gaps on AI?" What we have done is said on online safety, which is—

Dame Chi Onwurah: You found one gap and identified it.

Kanishka Narayan: On online safety, we have said that we have a significant regulatory regime in place. Are there any gaps that are correlated to real-world risk? We found one and we acted on it. Therefore, that has largely been the approach, which is to take the example of where there are experiences of people that we think are not currently regulated appropriately and to act promptly to stamp them out.

Dame Chi Onwurah: Are you identifying where there are experiences of people that are not currently regulated?

Kanishka Narayan: Yes, in a combination of ways. One is in the context of online safety, for example, which is—

Dame Chi Onwurah: You have talked about that, but I am thinking about other examples. That is one example.

Kanishka Narayan: I talked about child safety and online safety, but if you want me to park that, I can park the rest of online safety.

Dame Chi Onwurah: Yes.

Kanishka Narayan: Broadly speaking, in parts of the portfolio that I am responsible for—as I said, we take a generally sector-specific view—what we have said is that through engagement with the public, and I engage very often with civil society groups as well, we have tried to say, “Where are there gaps and how can we plug them?”

Secondly, in most of our regulatory regimes we have a systematic process of review, monitoring and evaluation, which also yields a clear understanding of where there are risks that are currently not being effectively countered. Therefore, thanks to the combination of those things: systematic reviews alongside a pretty clear sense of being attuned to ongoing feedback, I think that we are in a position to regulate where it is needed.

Dame Chi Onwurah: So every regulator has done a review of gaps in current regulation with regard to AI harms.

Kanishka Narayan: I certainly would not be in a position to tell you if the regulators completely outside the scope of my department would have done so, but I am happy to engage with them and try to write to you on it.

Dame Chi Onwurah: I suppose that is part of the challenge, is it not? While you are the Minister for AI, AI impacts every sector of every department.

Kanishka Narayan: Yes.

Dame Chi Onwurah: I understand that you are saying to us that you are perhaps confident about potential regulatory gaps for current harms within DSIT, but not within other sectors?

Kanishka Narayan: I think that is right. As I hope you will have heard, the privacy regulator says, “I regulate privacy as a harm vector, but I do not regulate specific technologies”. I would hope that the privacy regulator looks at both AI and non-AI harms. If you think about equalities, I hope that the equalities regulator looks at the possible vector of exposure from AI, but from non-AI as well.

Dame Chi Onwurah: I understand that, Minister, but I do feel that AI—we are looking at AI—can be complex and its implications should be considered across government. I suppose that brings me to your level of responsibility. Do you have a responsibility for AI across government or purely within the Department for Science, Innovation and Technology?

Kanishka Narayan: I am just going to clarify the question. What I am definitely not responsible for is public service deployment of AI. My colleague, Minister Murray, is in charge of public service deployment of AI. Of course, I would be happy to try to answer any questions on that.

Dame Chi Onwurah: No, that is fine.

Kanishka Narayan: If I think of the question more broadly, there are some aspects—labour market impacts and the future of work unit—which do sit in my portfolio, but that is not just across the public sector. It is a cross-economy question.

Dame Chi Onwurah: For example, on a question on the ethics of using AI in benefits analysis, is there anyone in government who is responsible for high standards of ethics or any standards of ethics in AI deployment across government or across sectors?

Kanishka Narayan: Very importantly so, yes. On the question of benefit analysis, the DWP has a series of requirements to make sure that when it is deploying technology, including AI, it is able to cater to not just regulation but the public sector equality duty. If the question is more broadly: is someone responsible not for sector-specific aspects but for cross-public sector aspects—

Dame Chi Onwurah: No, it is for cross-government aspects. Is anyone responsible for ensuring certain standards of ethics or other best practice across government when it comes to AI?

Kanishka Narayan: Hoping not to repeat this too much, it is very much the case that each department is responsible for ethics and regulations. On top of that, the Government Digital Service works on a series of standards across government—including an AI playbook that it is working on—to make sure that departments are able to share best practice as well.

We have intensive sector-specific deployment responsibilities for departments, but alongside that the Government Digital Service in its usual role, with oversight from Minister Murray, my colleague, has a sense of looking at standards across government as well.

Dame Chi Onwurah: What powers does it have to enforce those standards across government?

Kanishka Narayan: I am happy to write to you on that. It is not in my portfolio, so I am afraid I do not have a full view of it.

The Chair: Thank you, Minister. That is very helpful. Do you have any further points you want to make, Dame Chi?

Q104 **Dame Chi Onwurah:** I did just want to ask a quick question on the AISI and the December 2025 report, when every single system was overcome by jailbreak testing. Do you feel confident that the UK can address the challenges posed by AI frontier systems, and is there a plan to place it on a statutory footing? In responding to that question, could you just say why they changed the name from the AI Safety Institute to the Security Institute?

Kanishka Narayan: On the three questions, I will take the last first. Of course, it preceded me so I will largely pass on second-hand information. The change in name was reflective of the fact that AI, as a general-purpose technology, increasingly had national security implications and we wanted to make sure that those implications were very much front of mind for the institute as well, as it increasingly concentrated its work not just on questions of imminent safety but on questions of imminent national security.

As a reflection of that, the name was changed. I am sure it also had some implications on our ability to attract particular types of talent. You will have no doubt seen that the new leadership at the AI Security Institute has come from contexts that have experienced national security assessments in depth as well. We felt that was a very important function to have inside government.

On the first question of confidence in light of the report that you mentioned, the main thing is that the AI Security Institute has some of the world's most exceptional talent in being able to assess models. To that extent, the fact that the UK Government have built, retained and grown that capability is a very unique source of strength for us in understanding risk, engaging early in upstream and being able to work with developers and labs to mitigate those risks and to build cross-government understanding—not least, as I mentioned, in the child safety context—for us to take action where we think more action is merited and appropriate.

Therefore, given that, I would never be in a position of complacency to tell you that I feel deeply confident. Of course, I worry about this as I worry about a number of risks across the portfolio, but I think that the combination of those functions gives me some assurance.

Then on your second question about statutory, I feel at the moment that the AI Security Institute is unparalleled globally in its access to models pre-deployment globally as well. I have seen no reason to suggest that its institutional design requires reform to enhance that access and its

capabilities in any way, but of course I will be very keen on keeping it under review if there are alternative views on that question.

The Chair: That is very helpful. I am sorry, Dame Chi, that we need to move on from there, but if you are keeping that under review, I think we would want to hear more about it. Dr Swallow, and then after that we will turn to Lord Murray.

Q105 **Peter Swallow:** Thank you, Minister. I am going to ask you a very simple and single question: what is sovereign AI?

Kanishka Narayan: It is always a pleasure to answer a single, simple question to which there is no simple, single response, but I shall try none the less. My personal dictionary definition is that it is the ability for a state to have strategic leverage when it comes to this technology, such that it can ensure ongoing access to critical inputs, and ongoing assurance that its wider economic and national security objectives can be met more broadly.

Peter Swallow: We have set up the UK's Sovereign AI Unit. How many UK-based companies have been contracted to deliver services by that unit?

Kanishka Narayan: I wish I had a more assuring answer for you, but the Sovereign AI Unit is a very recent phenomenon. We have just secured money for it and we are hoping to officially launch it in a few weeks, therefore so far the work of the unit has been very limited. It is just very nascent. At the moment, we are very much focused on hiring, on governance, and on setting up the unit.

That said, in order to get those design questions right, we have already started to engage with some companies. For example, just a few months ago I announced an engagement with a couple of British AI for science companies, specifically in terms of a compute facility that the unit had, but also through our wider data initiatives on protein folding and open bind we have engaged with UK researchers and start-ups as well. I would be very happy to write to the committee with a list of companies that have so far been engaged.

Peter Swallow: That is helpful, and you will appreciate, Minister, why I am asking these questions—you will know where I am going with this. Lord Rook referenced it earlier. There is a view that the UK cannot possibly hope to compete with the US and China when it comes to AI. There is a view that the UK does not have the kinds of companies that we need to have in order to have a genuinely sovereign AI capability.

I know from speaking to you before this—and from my own experience as a constituency MP with some fantastic AI companies in my patch—that that is not necessarily the case, but obviously there is a role for government here to foster and grow the AI capability that we have in the UK. What is your view on this? How do we get the balance right between fostering genuinely sovereign AI capability while also recognising that of course we need to work with international companies, particularly those

based in allied states? Where do you see the balance between those two?

Kanishka Narayan: This is a absolutely critical question, so the first thing to say is that this is the most important question that I think about when it comes to the trajectory of this country on AI. Increasingly, not just economic power but hard material power and a huge series of areas of influence are determined by the quality of our position on this technology. That is the first thing to say.

The second thing is that when I think about sovereignty as I define it—which is strategic leverage more broadly—you get strategic leverage through three steps on the ladder. One is just to have enough of the critical inputs. I will come to all the other aspects of where it is from but just have enough of them—have enough of the chips so you can even do anything with AI in the first instance. Therefore, in that context, we were very keen that we were able to secure the level of capital investment that meant Britain was at least on the table.

Once you are on the table, the second part of it is to make sure that you have some diversification in who you are procuring critical inputs from so that you can bargain effectively. We are the party of labour; we understand that who has power matters as much as what the powers are. Therefore, for that context, one of the first things I did was to engage with a series of companies in every part of the stack so that we were able to build some diversity into the landscape.

The third rung of the ladder is ultimately to build British—to make sure that we have the full-fat version of sovereign capability here in parts of the stack. You are right to say that we do not have that across the entire stack. I would say gently that I think we do have it in parts of the stack. Arm is the leading chip design company globally and is still headquartered in Cambridge. We have fantastic companies growing up in the inference AI chip part of it, Fractile and Olix being two of them.

When it comes to models, we have huge strengths, not just in the fact that a lot of the Gemini teams sit in DeepMind in King's Cross, but in AI for science, and in AI for autonomous vehicles. Wayve just raised the largest funding round in Europe this year to date, of \$1.5 billion. It is a fantastic company looking at embodied AI and vision models. Therefore, there are lots of specialist areas of science in particular where we have very remarkable capability that gives us leverage.

The third aspect is applied AI. We have some very significant strengths, not least because we have been building areas of trustworthy AI and assured AI as well. The second thing to say is that I think there is some mix of procuring from outside to have enough of this stuff, procuring from some diversified sources, and then, yes, ultimately building British.

Peter Swallow: Would you accept that historically we have not been good enough at procuring from UK companies?

Kanishka Narayan: Without a doubt. If you look at it over the commercial cloud wave, frankly, I am very glad that there are lots of companies that we benefit from which are international companies and we will continue to do so. However, the fact that we have been so susceptible to a very limited presence in that part of the stack from the UK has been a challenge. Does that mean that we should attack that part of the stack? Not necessarily; it is an economic decision, and ultimately, we have finite resources and we have to pick our areas of focus. However, it is pretty clear to me that, broadly speaking, Britain lost out in the last 15 years of software as a service and in cloud development and I am very determined that we do not do so now.

The third part of your question is: what can government do? The Sovereign AI Unit is specifically focused on using both capital and compute to anchor British companies here, but that is only a part of the solution. We have not just a role in terms of compute and capital through sovereign AI and British business bank sources, but we have a role also as a customer.

In that context, one of the early things that we did, just three or four months ago, was to announce an advanced market commitment, where the Government will say to British start-ups and novel chips, "As you unlock technical risk by developing new technologies, you will unlock incremental government revenue. When you do that, you can raise a lot more money from private markets as a result of doing it". Therefore, the Government have a huge role to play in being much more innovative as a customer as well.

Peter Swallow: Briefly on AI growth labs, where the Government are looking to allow AI systems to be tested within safe sandboxes, where some of the regulation is taken away in order to test novel uses of AI, are you confident that human rights concerns are being tested as part of that process, as part of that sandboxing?

Kanishka Narayan: That is a really important question. One of the things that we consulted on in the consultation was to make sure that we solicited views on regulatory red lines. Human rights, workers' rights and intellectual property rights were three that were top of mind for me when I thought about that. However, we want to make sure that we are creating, yes, permission in cases where we think innovation in an assured, responsible way can progress the economy, but to do so with no compromise on fundamental values and rights as well.

The Chair: Thank you very much, Minister. I am being advised that there may be a Division in the Lords in due course—an occupational hazard. We may have to suspend at that point. I want to turn now to Lord Murray—and after that we will hear from Lord Sewell—building on the answer you have just given to Dr Swallow. When things go wrong, what happens about remedies? I think that Lord Murray wants to pursue that.

Q106 **Lord Murray of Blidworth:** In your view, Minister, are the United Kingdom's current arrangements to provide redress for harms caused by

AI, as required by the framework convention, presently sufficient?

Kanishka Narayan: I am always wary of assuring the committee that I have looked across the entire economy but, broadly speaking, there is a series of civil liability regimes under existing regulatory frameworks as well. Of course, if you suffer loss or injury in connection with the use of AI, compensation through the courts—through both contract and tort—can be solved there.

Alongside that, we have a series of frameworks—the data protection frameworks, for example—which apply to AI when it involves the processing of personal data. In that context, as you will know very well, data controllers have to make sure that any processing is lawful, fair and transparent. The ICO has a relatively clear redressal process, as well as a process for subject access requests, which allows individuals to seek both transparency on their data and redressal in the event of non-compliance with law. I do not propose to go through it harm by harm but, broadly speaking, we have redressal mechanisms across the regulatory regimes that apply to AI. Some of them work much better from a capacity and efficiency point of view than others, and it remains an important task for us to make sure that we are levelling all of them up as well.

Lord Murray of Blidworth: What you are suggesting is there are really two routes: there is a civil litigation route and an Information Commissioner data-processing route.

Kanishka Narayan: That is only because I stopped going through each category of harm, but if you wanted to take another instance: for example, if you look at online safety, platforms have a liability to ensure that they are acting on illegal content as well as child safety duties. In the context of redressal, individuals and trusted flaggers can flag harms in those contexts as well, and where they relate to systems and processes, Ofcom can then take them up. We have created this trusted flagger mechanism for redressal there as well. If I then think about employment, of course, in the employment context, if there is non-compliance there are redressable mechanisms through the law there as well.

Lord Murray of Blidworth: Obviously, it is another route of litigation.

Kanishka Narayan: Right, yes.

Lord Murray of Blidworth: So, of those three principal routes, are you satisfied that there is sufficient capacity within those mechanisms to address what might be massive scalability of harm from AI systems, given the increase in the use of AI?

Kanishka Narayan: That is an important question. One of the things I worry about is that automation on the offender side moves faster than automation and capacity as a result on the defender side. Therefore, one of the things that I am really focused on is how we build capability in the regulatory enforcement mechanisms that we have to keep up with that.

We recently announced an AI capability fund through the Regulatory Innovation Office to make sure that regulators are able to do so. Personally, in the context that I know best—which is Ofcom and the online safety part of my portfolio—I have seen that the hiring of machine learning engineers has very significantly increased in that regulator. Of course, it is a question that I reflect on quite a bit and remains an ongoing area of concern.

Lord Murray of Blidworth: The department's assessment of the adequacy of private law and consumer protection legislation is on the ministerial dashboard, but it is blinking yellow rather than green.

Kanishka Narayan: As the Minister who is not in charge of how the process works, I would not wish to comment on my Justice colleagues' behalf, but certainly from a cross-sector point of view, my instinctive view is that we need to keep a very close eye on this question.

The Chair: Thank you. Do you want to continue the questions?

Q107 **Lord Murray of Blidworth:** I was just going to move briefly on to procurement and the use by the Government of AI in procurement processes. Are you satisfied that government procurement processes are sufficiently using AI as a tool to improve procurement processes, and how do you go about monitoring that use, as DSIT, to ensure that it is being used well and uniformly across government departments?

Kanishka Narayan: I will take each part of that question: are we using it enough and are we using it well? On the question of whether we are using it well, again I am reluctant to opine on behalf of my colleague who has oversight of the GDS—the Government Digital Service. Minister Murray would probably be in a better position to answer it. Broadly speaking, the Government Digital Service ensures that there are common standards and a playbook for the deployment of AI in different parts of the public sector, as a way of ensuring that good standards and best practice are shared across the board. Indeed, as part of that, compliance with a range of responsibilities that government bodies have, whether that is data protection and privacy under UK GDPR, information security and auditability under provisions and standards of the NCSC, or specific rights in the context of either equality or fairness, broadly it ensures that those are complied with. That is the framework. I do not have a specific position on whether an assessment has been made of the quality of that.

On the question of whether it has been used enough, my broad view is that the opportunity is vast here and we are absolutely keen on grabbing it. It would be very remiss of me to say that I am satisfied with progress as of today; I want us to go much further, much faster.

Lord Murray of Blidworth: In terms of the perspective of this inquiry in relation to this committee, are you satisfied that the mechanisms within the procurement processes suitably consider the question of due diligence measures to protect human rights? I am thinking particularly of procurement of tech in the most obvious sense.

Kanishka Narayan: Yes. I am very happy to share again—I do not mean to avoid responsibility but simply to avoid credit—that having not had oversight of public sector procurement and AI in that context, I will flag that across HMG we use standardised legal contracts for procurement of software services across the board. In doing so, we have particular regard to strict mandatory clauses in those contracts, standards that are shared across government in those contracts, and a robust risk management framework that looks specifically at data protection privacy, information security and auditability, and at the questions of privacy and fairness in a standardised way.

Therefore, my expectation is that that is a relatively robust, mandatory, standardised approach across government. Of course, in all these contexts, it is always as good as the quality of the people and the enforcement of it. I am very happy to write to the committee on the question of where there is any assessment of that.

Lord Murray of Blidworth: Obviously, you are familiar with the DWP 2024 example, when an algorithm used to identify potential instances of benefit fraud was found to be producing discriminatory outcomes. Clearly, the reason that came to the Government's attention was because of the efforts of NGOs challenging that. What lessons have been fed into government procurement processes to prevent that kind of situation recurring?

Kanishka Narayan: At a very high level, I have seen the DWP instance. I would say that in this particular context, the public sector equality duty places a particular duty on government bodies. Again, not being the responsible Minister, I cannot remember whether that was an AI model or whether it was actually just a linear regression model. I am very happy to write back to the committee having consulted the DWP on it. But the big picture lesson is, of course, very clear: we want to make sure that the standards that we have, and the standardised contracts that we have, are properly enforced across government departments. Those include a very specific focus on ensuring compliance with the equality duty and avoiding any discriminatory outcomes of the sort that has been described.

The Chair: Minister, I am going to turn to Sir Desmond Swayne in a moment because I know he wants to ask about guardrails. However, before we leave this area entirely, Dr Swallow asked you earlier on about sandboxes. I am advised that human rights regulations will not be turned off in sandboxes, but will AI systems be tested for possible risks they pose to human rights? Again, if you do not feel able to answer that now, I am very happy for that to be added to the note to the committee subsequently.

Kanishka Narayan: The specific sandbox that I described, the AI Growth Lab, has not been designed in certainty yet. We have just consulted on the design of it. One of the key principles of the consultation was to solicit views on how we ensure that regulatory red lines, as they relate to rights in particular, were maintained in the design of it. Therefore, when we complete the design of the sandbox in that context, I

will be very happy to write to the committee at that point to assure the committee of where the red lines have been.

The Chair: Thank you very much indeed. Let us go to Sir Desmond and then Lord Sewell.

- Q108 **Sir Desmond Swayne:** What guardrails are or would be in place to prevent the disruption to a public service were the processing discovered to be non-compliant in respect of the AI that was being used? How would we ensure that the public service continued to be provided?

Kanishka Narayan: I do not have the fullest legal view on this, but if I just play the hypothetical scenario forward, I think that the situation being described is of AI being used in some context, it being found through its live deployment that the processing of personal data was non-compliant, a report having been made, I guess, to the ICO in that particular instance of harm, and the ICO finding that the technology was non-compliant. In that instance, I would assume that the ICO would then go through its enforcement process, which would include the possibility of enforcing fines and corrective remedies.

If it were the case that none of the corrective remedies for the processing of personal data was possible alongside the maintenance of the service provision, I suspect that there would have to be a period of transition from that system to another system, in parallel, to avoid breaks in continuity while maintaining the enforcement process. All that is to say that I am very much veering outside of my scope of understanding the legal context here, but I would assume that in that case we would have to make sure that we were running both the enforcement process and the service transition process in parallel to avoid a material break in continuity where it would threaten continuity of an essential service.

Sir Desmond Swayne: I was thinking more in terms of having thought through the need for a back-up system to replace that processing were it found to be non-compliant. Let us say it is a military recruitment system or something of that sort. There would need to be in place a back-up—a recovery system.

Kanishka Narayan: That is an important point about resilience of procurement context more generally as well, and I entirely agree with it.

- Q109 **Lord Sewell of Sanderstead:** The Government have clarified their intention to build a national data library, and this is a way in which they can have some leverage in health data held by the NHS as an asset to promote collaboration between the UK and investors and developers. How will the national data library ensure two things: that individuals' privacy rights are not violated, and that bias within national datasets does not translate into discriminatory outputs?

Kanishka Narayan: That is a central question. The starting point is that the national data library is a very ambitious government programme. We have backed it up with over £100 million, and the intent there is to create a very efficient, effective way of allowing public datasets to be used, for

public research in particular, and the Health Data Research Service is a very important part of that in the first instance.

On the question of protections of privacy, we have set out new guidelines for AI-ready government datasets more broadly, and those guidelines are precisely focused on how responsible development takes place using those datasets. They require human oversight at key stages to protect privacy rights, a duty to prevent bias, and that people, rather than just machines, are accountable for decisions that matter as well. That is all alongside the range of existing regulatory responsibilities that apply, of course, to the public sector use as well as private sector use of any data.

Therefore, we are very conscious that, yes, we have existing duties but, in the case of data usage for AI, we have decided that it is worthwhile to set out even further guidance to ensure that those regulatory responsibilities are carried out with granular guidance as well.

Lord Sewell of Sanderstead: Genuinely, are you really happy with that framework?

Kanishka Narayan: It is untested at the moment because the national data library is very new, but we have worked incredibly hard and intensively across different parts of government, as well as with experts externally, and have engaged very widely. Of course, if there are questions that others want to raise about the quality of that guidance, I will always be very keen to hear them.

Q110 **Baroness Chakrabarti:** The Data (Use and Access) Act removes most of the previous GDPR restrictions on automated decision-making, so these restrictions now apply only to the most sensitive types of personal data. Minister, why was this decision made?

Kanishka Narayan: Just to clarify, the Data (Use and Access) Act changes effectively say that we will not place specific scope constraints on where automated decision-making is allowed, but of all the protections that existed prior, many of them will continue to apply in a very material way: first, the protections around people being informed of when solely automated decision-making is taking place; secondly, their ability to contest it when they feel a decision has not been appropriate; and thirdly, the ability to secure human intervention in the use of that decision-making.

Therefore, the robust set of entitlements and protections that people have had prior with regard to the use of solely automated decision-making continues to apply, alongside, of course, the existing set of responsibilities in terms of equalities, privacy, and beyond. The intent there was effectively to say that we did not feel—especially for this new technology—that prescribing on the face of legislation where it should or should not be applied was the right way to go about it. The right way to go about it was to make sure that where it is applied, the appropriate protections were available in terms of rights to individuals to ensure their safe and responsible usage.

Baroness Chakrabarti: Therefore, you are satisfied of the robustness of human rights protection despite that deregulation?

Kanishka Narayan: On the face of it, what we have done is to maintain the most important parts of the regulation, which were focused on the protection of human rights in any of these cases, and we are absolutely focused on ensuring that they are carried out. As to the question of whether I am satisfied on the whole, a huge amount of this remains a question of enforcement quality as well, and I would not want to be in a complacent position and say that it is a one-and-done question—we continue to be very focused on it. But I hope that I have set out a rationale that says we have put the protection of rights in the provision of the law front and centre.

Baroness Chakrabarti: Are there any areas of life where you think decisions about people should not be automated?

Kanishka Narayan: I think that the existing data protection framework that we have effectively says that we should only consider processing, and that automated systems will require large-scale processing of data only where it is proportionate and fair. There are a number of areas—I would not propose to go through all of them—where those tests are not likely to be met and, even when they are met, they should be done so only in a way that still allows individuals to be informed, to contest, and to secure human intervention as well.

Baroness Chakrabarti: Are the Government considering removing any other rights protections or considering any further deregulation as part of their AI ambitions?

Kanishka Narayan: Just to be very clear, Baroness, with regard to the first framing of your question, in my view this was not a removal of any rights. As I mentioned, the protections for individual rights remain very much in place in the law. It is deeply important that we continue that.

On the broader question of deregulation, of course, we are looking at a series of instances where we think regulation can be smarter still, and can allow regulators to move dynamically and to be better off in being able to enforce regulatory outcomes in a way that is different. The cross-economy sandbox, which I mentioned, is very much part of that. I hope that some of the individual sector regulators have talked to you about their sandbox environments as well, the intent of which is effectively to achieve robust regulatory outcomes without necessarily static prescription of activity.

Baroness Chakrabarti: What the regulators did say—three of them sat there—was that not one of them has the power to say about a new product, “This product should not be used”. Not one of them is able to give approval or disapproval for a new AI product.

Kanishka Narayan: Yes. As I mentioned earlier in our conversation, a huge amount of focus in this area and others has been on liability

regimes—regulatory regimes that operate robustly in being able to tackle harms. A number of those are ex post. What we have done very uniquely here is to create a capability for ex ante assessment and co-ordination with the companies to make sure we are mitigating some of the most acute harms.

Q111 Lord Murray of Blidworth: On automated decision-making, there are four safeguards that are meant to be in place, including the right to meaningful information and the right to contest the output. In practice, how are these four safeguards monitored?

Kanishka Narayan: Lord Murray, you are very much asking me to veer outside of my realm of understanding, so I am happy to write to you on this question as to how they are enforced in practice and any assessments that have taken place of compliance with it.

Lord Murray of Blidworth: When you write, perhaps you could also cover this question, which is: if the safeguards are not adequately provided, how can individuals complain and seek redress?

Kanishka Narayan: Yes, I will.

Lord Murray of Blidworth: I am also conscious that a Division has just been called in the Lords.

The Chair: Before I suspend the sitting, though, out of politeness to the Minister, who not only has been jet-lagged but has been giving very freely of his time this afternoon—I am told there may be two Divisions, so we do not want to have him hanging around to answer just two more questions. Minister, you have very kindly agreed to write to us about a number of things. I hope colleagues agree that, rather than keeping you here, if we were to give you those two other questions—one is Dame Chi's and one is mine—and if they could be added to your letter, can I prevail upon you to ensure that we get a reply to those questions fairly quickly? We need to move on to our draft report and we need your replies to be able to do that. If your officials could help in ensuring that we get a reply within a reasonable timeframe, in the next week or so, that would be really appreciated.

If you are happy on that basis, I will be very happy to now suspend the sitting by saying thank you very much indeed, Minister Narayan, for giving us your time today. I have been very personally impressed by the way you have answered our questions, the competence that you have shown in your portfolio, and the deep knowledge you clearly have. I think that people should feel comforted that someone who is on top of his brief has been appearing before this Select Committee today.

Kanishka Narayan: I very much appreciate that, Chair. Thank you for your generosity in considering my time. I hope to reciprocate the generosity by my timely response as well.

The Chair: Thank you very much indeed. In that case, I can formally end today's meeting because we dealt with all the private business earlier on.

If any members wish—we do not need to be quorate—to come back after the Divisions are over, Professor Yeung is available to come and talk to us, as our special adviser on this committee on this inquiry; Karen is sitting modestly over there. She will be happy to brief us about some of the replies that we have heard today. With those words, let me end this session.