



Communications and Digital Select Committee

Corrected oral evidence: Large language models

Tuesday 28 November 2023

3.20 pm

[Watch the meeting](#)

Members present: Baroness Stowell of Beeston (The Chair); Baroness Featherstone; Lord Foster of Bath; Baroness Fraser of Craigmaddie; Lord Griffiths of Burry Port; Lord Hall of Birkenhead; Baroness Harding of Winscombe; Baroness Healy of Primrose Hill; Lord Kamall; The Lord Bishop of Leeds; Lord Lipsey; Lord Young of Norwood Green.

Evidence Session No. 13

Heard in Public

Questions 130 – 145

Witnesses

I: Viscount Camrose, Minister for AI and Intellectual Property, Department for Science, Innovation and Technology; Lizzie Greenhalgh, Deputy Director of AI Regulation, AI Policy Directorate, Department for Science, Innovation and Technology; Sam Cannicott, Deputy Director of AI Enablers and Institutions, AI policy directorate, Department for Science, Innovation and Technology.

USE OF THE TRANSCRIPT

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

Examination of witnesses

Viscount Camrose, Lizzie Greenhalgh and Sam Cannicott.

Q130 **The Chair:** This is our second witness session today, and the final one, of our inquiry on large language models. I am pleased to welcome the Minister for AI, Viscount Camrose, and two officials from the Department for Science, Innovation and Technology. May I ask the officials to state the positions they hold in the department?

Sam Cannicott: I am the deputy director for AI capability.

Lizzie Greenhalgh: I am the deputy director for AI regulation.

Q131 **The Chair:** As you are probably aware, Minister, there has been some comment about the way in which the Government have seemingly shifted their strategic approach to large language models this year, starting with the White Paper in March talking about large language models or AI being an opportunity and wanting to be pro-innovation and to take an agile approach, with lots of emphasis on the opportunities from the technology. But once the taskforce was established, there was a shift more to safety and risk. What do you see as the Government's strategic approach to LLMs? Is it to make the UK a hub for safety or to capitalise on the opportunities?

Viscount Camrose: Thank you, Chair, and thank you to the committee. There may be a false dichotomy between regulation, safety and innovation, and if I speculate, that goes back to the earlier days of the internet. Our view is first and foremost that large language models and indeed AI in general are an opportunity for innovation that can generate prosperity, health and wealth across society that can solve a huge range of societal problems and bring enormous benefits.

But, in order for AI to be adopted, AI must be safe and not only trusted but worthy of people's trust. So I do not think there is an either/or between safety and innovation. Safety is a necessary precondition to innovation. I regret that the language often veers either too much to the side of innovation or too much to the side of safety, and I wish we could all collectively find a form of language that allowed us to speak of the importance of both sides, because both are extremely important.

The Chair: I take your point about there not being a false dichotomy, but can you at least acknowledge that there was, none the less, a shift in emphasis by the Government very much to safety first after they had published their White Paper, which was not safety first?

Viscount Camrose: I do acknowledge that. Perhaps the tone, if not the substance, of the White Paper as originally published stressed innovation perhaps at the expense of safety, although I stress that its content did not. So much of our effort since then has been devoted to what was then called the Frontier AI Taskforce—now the AISI—building the safety institute and conducting the Bletchley Park summit that, necessarily, a lot of the dialogue coming from us has been about safety. If there is

something that I wish we could consistently do better, it is to talk with equal emphasis about safety and innovation.

Q132 **The Chair:** There have been continual name changes of that task force. You did not have to call it whatever it was before it became the AI Safety Institute, and it is curious that you made that decision.

To develop a bit more on the opportunities, what is the Government's thinking on developing a sovereign large language model? We have heard arguments about the huge cost of doing one from scratch, and there are big question marks as to whether that would be a realistic prospect, but others have highlighted opportunities in the Government perhaps commissioning one of the developers to build something specific for the UK Government for deployment using our own public sector datasets, which therefore would not require the transfer of that to an international model. Is that something that the Government are currently exploring?

Viscount Camrose: A theme that we will no doubt come to often in this session is action based on evidence. The Government certainly see in principle the advantages of having a sovereign large language model capability. As to whether it makes sense to do that or to consume other large language models in our own way as a Government, we have to wait for the evidence to tell us and point us in that direction.

However, given the extensive investment that we are making into compute—on top of the £900 million already committed, there was a further half a billion in the Autumn Statement last week that will be invested into exascale and AI capability in our Bristol, Cambridge and Edinburgh centres—the sovereign compute puts us in a good position to elect, at some future point, to acquire the capability of a large language model for sovereign purposes. That is one of many examples of retaining optionality and being agile in our response to the opportunities of AI.

The Chair: When do you think you might be in a position to make those sorts of decisions?

Viscount Camrose: It is down to the capabilities. The real questions are: what are the capabilities of emerging models for AI, which sectors can they help, and what use cases might we be able to support with them that are not delivered on favourable commercial terms elsewhere? AI models supported by the newer generation of chips will start to emerge over the next year. At that point, I expect us to be presented with an interesting range of models and possibilities for a sovereign LLM model. That is not to say that we would necessarily go that route, but there will be compelling opportunities.

The Chair: Regardless of the model or how the Government might want to deploy this technology in the public service, what consideration is there at the moment of when and how some public services might use large language models more? Where are discussions about that taking place?

Viscount Camrose: There are a range of possibilities in different departments for how those might come about. The ones that are frequently discussed include, obviously, the Department for Education in the reduction of the administrative burden on schoolteachers. Then there is Health, of course; because of the extensive data available to it, there are opportunities there. We are now seeing them, department by department, identifying approaches that they might choose to pursue for their own AI models.

I do not think we will see cross-cutting AI models that might affect multiple departments until sometime next year as the more frontier-capable models start to emerge.

The Chair: Before I move on to my colleagues, who want to talk more specifically about the AI Safety Institute—we will come on to regulation and other matters—I want to go back to the Government's pivot away and shift in emphasis. At one point, when the Government came out with their White Paper, the emphasis was very much on opportunity, and that coincided with a change in the structures, such as the AI Council, that were supporting and advising the Government on technology. There was a suggestion that those bodies were not sufficiently seized of the opportunity in or the rapid development of that technology, and one reason why the Government decided to set up the task force was because there had been no real input or notification from the AI Council. Is that your understanding?

Viscount Camrose: No. That is not a fair characterisation of the decision to disband the AI Council. The reason was that the members of the council had come to the end of their agreed terms. Over their terms, they provided extremely valuable and valued expertise and opinion. The problem was, and is, that this kind of structure does not lend itself at all to the agility that we need. I have come to this again and again: what we need for the technologies of AI in particular is adaptiveness—that is, the ability to adapt and to make decisions at speed.

The problem with the structure of the council was by no means the people on it. It was simply the fact that you structured the whole work of the department around quarterly meetings, so you were always preparing for the next meeting or responding to the previous one. In a way, although the material from any one meeting may be useful and valuable, that really slows you down. So we have moved to a pool of experts—for want of a better term—structure comprising former members of the AI Council and a range of others whom we can consult both individually and collectively on an as-needed basis. That is far more suited to the pace at which we are attempting to work here in order to accelerate from a position of large language models being something of a novelty in policy circles to extremely fast-moving and adaptive regulation and strategic response to AI.

The Chair: That is probably a neat segue for me to move on to Baroness Harding, who wants to talk to you about the AI Safety Institute.

Q133 **Baroness Harding of Winscombe:** Can I be cheeky and ask a quick

supplementary to your question first, Chair?

Thank you so much for coming today, Viscount Camrose. You have talked a bit about investment in compute, and the Chair has pressed you on whether we should be building sovereign large language models, but we have not talked about chip manufacturing. I am struck that we never talk about Graphcore, a UK processor company. Should we have a national strategy in that part of the value chain?

Viscount Camrose: We have a national semiconductor strategy. What is very clear—it is slightly beyond the pure AI discussion, because it is applicable in a lot of areas—is that, in the Government’s opinion, it would be extremely risky to attempt to create a semiconductor foundry. It would cost multiple tens of billions of pounds and would have very uncertain chances of success.

You mentioned Graphcore. We have a range of very good design and manufacturing capabilities in compound semiconductors in this country. It is not immediately clear that that will ever be a very large business, on the scale of GPU manufacturing or design, but it is an important specialty area in the semiconductor industry.

Q134 **Baroness Harding of Winscombe:** That is helpful, thank you. Sorry about the technical bit.

Let me go back to the question I was meant to be asking. Can you give us an update on the work and forward plans of the AI Safety Institute? It has a budget of £100 million per annum, I think. How much of that has been spent so far this year? What has it been spent on, and what do we expect the rest to be spent on going forward?

Viscount Camrose: Let me start with what it is for and what we have got from it so far. The AISI has three distinct purposes, all of which overlap to some extent. The first is to develop and conduct evaluations on frontier AI systems, specifically to characterise their risk and level of safety.

The second is to drive foundational AI safety research and, indeed, conduct that research itself or engage with other organisations that do so.

The third is a deeply important purpose: facilitating information exchange both nationally and internationally. The international point is particularly important. We engage with a huge number of international multilateral fora and organisations. Any global AI regulatory system in any nation has to be highly interoperable with the regulatory schemes of other nations and internationally. The AISI will be our principal interface internationally to, if I may be permitted the metaphor, construct that orchestra and get them singing in harmony.

As for what the AISI has delivered thus far, bear in mind that, in its present form, it has not been around for very long—about a month now.

Baroness Harding of Winscombe: That is because its name has been changed twice.

Viscount Camrose: Well, no, its form has also been changed. It is working on the next AI Safety Summit, which will be in Korea, and is continuing its work on evaluation. It has been setting up its team and was in the process of appointing a CEO. It has taken care of international engagement thus far—principally with the US and Singapore, given that those were the two countries that indicated an immediate willingness to proceed on this basis at the AI Safety Summit. It is hosting the State of AI science report, led by Dr Yoshua Bengio, which was also initiated as part of the summit.

Lastly, you asked about its budget. I think it is about £65 million a year for the remainder of the decade. That comes to £400 million committed in total, but obviously it is subject to the institute being able to demonstrate ongoing value added for the work it does. That goes principally on people, of course—both our own people and the external experts it deals with. It also goes on compute and on giving them a highly secure working environment in which to perform their function. If you need a more detailed breakdown, it might be better if we shared that in writing, if that would be all right.

Q135 Baroness Harding of Winscombe: It would be extremely helpful if you could give us, in particular, the split between the people costs and the capital or compute costs. That would be fine.

This may be my misunderstanding, but I think I am right in saying that the AI Safety Institute is different from the central risk function that the AI White Paper set out. Am I right that those are two different organisations?

Viscount Camrose: They are two different things. I share any view that says our acronyms are not ideal, but the central AI risk function sits within DSIT. Its role is to scan the horizon for emerging risks, focused on where we are with today's technology, and to liaise with existing regulators that form the backbone of our current regulatory model for AI. The set-up is that existing regulators currently regulate AI because they understand the context in which they operate and can regulate the AI most effectively of anyone within the sector in which they operate.

Supplementary to that, at the central level is a range of further organisations that we have now set up. One is the central AI risk function, which performs scanning and liaises between the different regulators and advises them on how to enhance their capability. Others include the Centre for Data Ethics and Innovation, the monitoring and evaluation function and the AI Standards Hub, which develops and owns the expertise in AI standards for the evaluation for existing models. The AISI provides more of a leadership role, in addition to that.

Baroness Harding of Winscombe: So who in all that is the group of people the regulators would go to for guidance, support or technical expertise to implement the regulation of AI in their sector?

Lizzie Greenhalgh: Further to what the Minister said, through the White Paper we committed to establishing a central risk function. That was one of the proposals in the White Paper that were very popular. People

welcomed our context-based approach to regulating AI but said that co-ordination, supporting knowledge exchange and making sure that we had a holistic overview of risks across the landscape would be critical. So we set up the central function, prioritising the risk function in the first instance.

It is probably helpful to note that the risk function considers all AI risks. It is not just focused on frontier AI; it is also considering other factors that drive risk, such as adoption and vulnerability across society. We are in the process of establishing the other elements of that central function, which will include supporting regulators and knowledge exchange co-ordination. We plan to say more on that in the White Paper response.

Baroness Harding of Winscombe: There is a £400 million budget for the AI Safety Institute. What is the budget for the central risk function, and how many people are working in it?

Lizzie Greenhalgh: I do not know the specific budget, but on the risk function we have a brilliant team at the moment, with people from regulators and from legal backgrounds, economists and AI forecasters. They are an impressive bunch of people. We can write to the committee with more details on the specifics, but they have been instrumental in increasing the Government's understanding.

Baroness Harding of Winscombe: What will its output be over the next six months?

Lizzie Greenhalgh: We committed in the White Paper to publishing an AI risk register around 12 months from the White Paper's publication. That is something we are looking closely at at the moment.

Baroness Harding of Winscombe: It would be fantastic if you could share with the committee the budget and scale possibly not just of the central AI risk function but of the other pieces that you mentioned, so that we can understand the full jigsaw puzzle.

Sam Cannicott: I will just add something on the role of the institute and how it has an important input into that risk work. One of the things the institute will do is evaluate frontier models. It will build and understand the capability of those models and therefore the potential risks and harms that they pose. That is a piece of work that at the moment really only sits in the hands of industry, but we are building the capability to have an independent view, which will be a crucial input into the wider risk work.

Lord Kamall: Sam talked there about inputs. How do you make sure that you get those inputs? Do you just sit there and wait, or will you be proactively going out to the people you want to hear from? How do you capture the people who you may not have thought about but who might have some valuable input?

Viscount Camrose: We have talked throughout the process—before, during and after the Safety Summit—about continuous engagement, particularly around the Safety Summit. Before then, collectively the ministerial team had 100 individual engagements building up to that, and

in addition we have committed to ongoing engagements with innovators, civil society, academics and others. That is an ongoing process. The central AI risk function, as part of its horizon-scanning function, will have to continue to perform in that way as well. There is always a risk that there will be something that does not occur to them, but their role is to make sure that they think of everything.

Q136 The Lord Bishop of Leeds: You referred earlier to the Centre for Data Ethics and Innovation. That board was disbanded along with the AI Council. How have ethics been transferred over into the safety institute—or have they not?

Viscount Camrose: The AI Safety Institute would certainly have a view on ethics. It was the board of the CDEI that was disbanded, for very much the same reasons as the AI Council was: it had come to the end of its allotted time, and we had to move it away from was working around a series of quarterly meetings and into a more agile approach.

The Lord Bishop of Leeds: I understand that. My question is really about how ethical questions are being represented in the new institute. Do you have ethicists engaged with it?

Viscount Camrose: I do not know about the breakdown of the institute regarding ethical specialists.

Sam Cannicott: One of the areas of evaluation that the institute will look at will be potential societal harms posed by models, which will cover things like manipulation and trust. There will be important ethical questions there so the institute will need to involve ethicists in that, particularly when it comes to the alignment question of whether AI aligns with our values. A couple of the people who were on the CDEI board are now closely involved with the institute, including Rumman Chowdhury, who looks at some of these issues and is now part of the safety institute board. The institute has a commitment to work widely with people, and I think that will include ethicists, given their focus on societal harms.

The Lord Bishop of Leeds: That is a utilitarian understanding of ethics. I will leave that there, but I think there is more to be said.

Lord Griffiths of Burry Port: Carry on, Bishop, carry on.

The Lord Bishop of Leeds: Another time.

The Chair: The mention of the AI Council prompts me to remember something I was going to ask you before, Minister. To go back to what I was saying before when we were talking about the AI Council, there was a suggestion when it was disbanded that it had not alerted the Government to developments in AI. Dame Wendy Hall, in one of our previous sessions, said that the council had written to the relevant Minister in 2022—I suppose it would have been a different department at that point—and rather challenged what had been said publicly about what the AI Council had done. I wondered if you wanted to acknowledge that what she told us was right.

Viscount Camrose: The AI Council was instrumental in the principal key deliverable before I took on the role, which was the AI White Paper. Without the wisdom and expertise of the AI Council, the White Paper would not have been written in that form. We were absolutely reliant on the council. The White Paper was created in that form with all relevant knowledge, and it would not be fair to say that the AI Council in any way had failed to provide us with up-to-date, relevant insights and information.

The Chair: So any suggestion to the contrary would be wrong.

Lizzie Greenhalgh: We would also point to the insights from the AI Council and others in driving the establishment of the Foundation Model Taskforce. The UK Government would point out that they have been quite front-footed in looking at some of the risks of the most advanced AI systems. That is thanks to the AI Council and other experts we are working with, who have been instrumental in driving the Government's response including most recently, convening the summit on those issues.

The Chair: So it is a bit strange that that impression was allowed to emerge from government when it did.

Viscount Camrose: It is certainly unfortunate that it emerged. We very much tried to handle the disbanding of the AI Council in the right way, and I regret that anybody would conclude from that that its members had not done their job effectively, because we owe them more thanks than anything else.

The Chair: That is good to hear. Baroness Harding talked earlier about the AI Safety Institute, and I should point out that, when Mr Hogarth was before us at the beginning of this inquiry, we were interested in conflicts of interest and the management of those. I am grateful for the ongoing dialogue that I am having and the engagement of the Perm Sec about this matter.

Q137 **The Lord Bishop of Leeds:** Several weeks ago, we had a number of regulators here giving evidence to the committee. One of the interesting things that emerged was the gap between the regulators saying that they have the capacity to deal with whatever is coming down the road and saying that they have no idea of what is coming down the road. That begged certain questions. Would you support, and are there any plans for, further action to improve regulatory oversight, for example through mandatory safety tests for high-risk models and better support for regulators—we know that everyone wants more money and more people—and standardised practices for auditing models?

Viscount Camrose: Let me start my answer by briefly describing how the adaptive regulatory model is structured today. I do not want to repeat earlier answers, but essentially, as you say, the existing regulators are charged with being the teeth of the regulation and intervening, in regulatory terms, where needed in their specific sectors. We also have the central bodies, which we discussed earlier, and thought-leadership bodies above that.

I am sorry for the long-winded answer, but it is worth setting this out. That set-up gives us a number of advantages. First, it allows us, and has allowed us, to move at speed without setting up too many brand-new bodies that need to stand on their own. It makes us adaptive to context and to change in the technology environment. It also diminishes the risk that new bodies start to overlap with existing bodies in a confusing way.

The key element to this approach is adaptiveness based on the emergence of evidence. The Government have no philosophical objection of any kind to legislation in this of any kind, provided that that legislation is supported by knowledge and evidence. Particularly in the case of frontier AI, that knowledge and evidence do not currently exist—not just in the United Kingdom, but anywhere in the world. We have invested very extensively, actually more than any nation, in acquiring this knowledge and evidence, so that we are able to move guided by it. Whether that is new legislation, amendments to existing legislation or non-statutory measures, we will be placed to move forward in a targeted way. There is no set timeline for when we will do it, but I suspect there will inevitably come a point when binding requirements will be placed on both the regulators and the regulated.

Q138 The Lord Bishop of Leeds: I noticed the language of adaptivity based on evidence. Are the voluntary commitments for safety testing adequate and, if so, how are they based on evidence?

Viscount Camrose: We have undertaken the voluntary commitments with the largest AI labs, which have all agreed to hand over their safety plans and models to us for evaluation. As with everything else that we are doing, that is a scheme to test and iterate, as opposed to one that goes straight to legislation. That carries two principal advantages. The first is that, by testing and iteration, we get to a point where we can create smarter interventions as we need them. Secondly, it brings the industry with us. This is not some top-down new set of rules that everybody has to obey; it is a collaborative environment for the creation of wise interventions into the safe conduct of AI.

The Lord Bishop of Leeds: To a layman like me, that sounds as if we can get a degree of control over the technology, but not over the commercial drive behind some of the development of technologies.

Viscount Camrose: I would not characterise that as the principal driver. This is a brand-new, very fast-emerging technology, fully understood by precisely nobody. We therefore need to be wise about how we approach any plan to regulate it. That means collaborating together in a spirit of openness. I believe that all these companies are working with us in such a spirit, because, voluntarily and uniquely to us, they agreed to hand over their models for our safety evaluation.

Q139 The Lord Bishop of Leeds: That illustrates the need for ethicists to be involved in this. Finally and briefly, what powers do the regulators have to require businesses—as opposed to asking them voluntarily—to engage with pre-release safety tests or to determine the standards for those tests?

Viscount Camrose: I am speculating slightly in my answer, but that depends very much on the context. What is permissible in, say, the police force might not be permissible in a marketing organisation, and so on. This illustrates the real importance of giving the existing regulators as big a role as possible: they understand the rich context in their sector and can act accordingly.

The Lord Bishop of Leeds: What sanction is there if businesses decline to implement the Government's safety suggestions?

Viscount Camrose: They may well be breaking an existing regulation or rule, in which case there are existing sanctions available. If not, it would be a piece of evidence to say that we need to regulate in this case and move quickly to do so.

Lizzie Greenhalgh: One advantage of our non-statutory approach is that we can do things right now, today. As the Minister indicated, we are not afraid to legislate and it has not been ruled out. But some of the steps we are taking—for example, the Secretary of State requesting that companies hand over their AI safety policies and voluntary agreements—mean that the AI Safety Institute can crack on with things right now that help make the world safer today, rather than wait for legislation that we know can take a while. It does not mean that those things are off the table, but this is a fast-moving technology and there is a lot for us to be doing right now.

Q140 **Baroness Fraser of Cragmaddie:** Minister, you will remember that, in the Online Safety Bill, there was a lot of debate about safety by design. In this space, it is all very well innovators handing over their models at that stage, and it is all very well regulators looking at the harm downstream, but is there a danger that we will lose where the liabilities and accountabilities need to be because everybody thinks it is someone else's responsibility?

Viscount Camrose: First, I will make an illustrative point about safety by design, and then I will come to this important issue of liabilities. On safety by design, I had the pleasure yesterday of launching a document—that does not sound very exciting—with the guidelines for the cybersecure development of AI across the life cycle. Of course, cybersecurity is one of the key risks that we need to take care of. I bring this document up, because although the UK took the lead on it, it has been signed and agreed to by 23 agencies in 18 countries, so it is a good example of the international interoperability of a cross-life cycle standard.

There is no doubt that LLMs and other forms of AI will cloud the issue of liability. This is a deeply complex issue, to the extent that I understand it—I am the first to admit that is a fairly limited extent—and it is absolutely one of the areas that the AISI is looking into to give us the evidence and opinion to guide our approach to this. But all jurisdictions that worry about AI worry about this liability issue, and it will be something for which both national and international solutions will be key.

Q141 The Chair: I have a couple of almost schizophrenic questions before we move on. Have Chinese companies also signed up to the safety testing agreements, and maybe the other ones that you mentioned?

Viscount Camrose: No—not with us, anyway. I am not privy to their arrangements with their own Government. The ones on the list—I am happy to cite them—are, for the most part, American, except for Mistral AI in France.

The Chair: The other thing I was keen to understand is what arrangements you have in the department and what you are doing. How do you feel, as Minister, about avoiding regulatory capture by the industry?

Viscount Camrose: I would like to say on record that we will not have done a good job if the companies that end up succeeding in the world of AI are simply the same companies that succeeded previously. It is important that new innovators are allowed to arrive on the scene. I know there was some criticism of the guest list at Bletchley Park, given that a preponderance of the AI companies that came were at the larger end—that was purely a function of the fact that they were the closest to producing frontier AI models—but we are very much looking at ways to remove barriers to innovation for smaller companies. A few of those are: making sure that we have a good skills base here in the UK, providing access to compute for UK-based start-ups and, perhaps above all, providing a sense in the UK that AI is trustworthy and, therefore, that it is okay to adopt it. There will then be demand for new companies.

One point that is very much debated in relation to allowing innovation to occur is about open source. One of the benefits of allowing open source more is that smaller companies have greater opportunities to take advantage and can therefore rise to prominence. We can associate a number of disadvantages or risks with open source, which the IASA is looking at and has characterised as an extremely complex problem.

The Chair: But are you, as a Government, open to open source?

Viscount Camrose: We are open to evidence on open source. If we could somehow cryogenically freeze AI in its current state of technology, we would be very much in favour of it. With open source, the worry is that, as models become an order of magnitude more powerful, the risk of open-source models getting into the wrong hands becomes very much more serious. So there are safety arguments against it, and indeed for it, but there are innovation arguments both ways as well.

The Chair: This is something that you are—

Viscount Camrose: —wrestling with.

Q142 Baroness Healy of Primrose Hill: We have heard lots of evidence so far on copyright, and it is a matter of great concern because there is so much controversy around it. Developers believe that there is no infringement of copyright, but the publishers are saying that it is on a massive scale. So what is the Government's current legal position on the application of copyright law to the use of copyrighted materials in LLM

and training data?

Viscount Camrose: Again, this is a difficult area. The first thing to say is that it is highly contested on both sides, between the rights holders, who feel that they are being infringed, and the innovators—the AI labs—which feel that any attempt to stop them creates far too significant a drag on their ability to innovate.

It is difficult to give an overarching answer on the Government's legal position, because this operates in so many different contexts. We need to solve two particular problems here. The first is finding the appropriate balance—the landing zone—between two hotly competing sides in this debate. The second is the need to make whatever solution we come up with internationally operable. We cannot have a set of strict rules over here that then allow people to go and train their models elsewhere. I do not think that would help anyone.

There are a number of these cases in the courts globally—I point particularly to *Getty v Stability*—and we need to see how that will end up. Overall, the best outcome—we are pushing hard for this—is a voluntary agreement between both sides that recognises the needs of both sides. To that end, we have established a working group, which we continue to operate with. It is led by the IPO, which continues to engage with the innovators and the rights holders to try to find a successful outcome. Our current focus is on developing the set of principles around which we may or may not be able to operate, and then turning that into a code of conduct. Ideally, this would operate on a voluntary basis because legislation on this basis runs the risk of sending people to operate overseas in jurisdictions over which we have no control.

Baroness Healy of Primrose Hill: Is there any timescale for when this code might be announced?

Viscount Camrose: We had hoped that we would have that code by the end of this year. The participants on both sides have made strong representations to us, saying, "Please do not go ahead until we've properly argued this out. Speed is a secondary consideration to getting this right". That said, we will not get into an endless talking shop about this, but we need to talk it out. Should it, sadly, emerge that there is no landing zone that all parties will agree to, we will have to look at other means, which may include legislation but I very much hope we will not have to go there.

Baroness Healy of Primrose Hill: There is a lot of concern in the publishing world about what will happen. Would you support new transparency requirements for model developers that would allow rights holders to check if their copyrighted data has been used in a dataset? That could be on a voluntary basis, perhaps, via a third-party auditor. There has to be a solution to this.

Viscount Camrose: Indeed. If, in this country or anywhere else, you create a work of the mind, you should continue to have an expectation of reasonable reward for doing so. It is really important that that continues

to be the case. There are a number of avenues that we might go down, and I would certainly consider looking at that one.

I would also point out that some technology opportunities may well end up being part of the solution, such as automated watermarking technologies that may be invisible to the human eye but visible to an AI. There is a range of other possibilities. Google, Adobe and, no doubt, a great many others are conducting work on this. I do not feel particularly comfortable with the idea that we would totally rely on a technology solution to solve all our problems, but the endgame is a mixture of voluntary agreement about a landing zone, technology and any legislation that we absolutely have to put in—as well as, most of all, good will. Deep down, I do not believe that infringing the rights of copyright holders is a necessary precondition for developing successful AI globally.

Baroness Healy of Primrose Hill: The White House agreement talks about watermarking. Have you had many dealings with what it proposed? Were you informed?

Viscount Camrose: I know my officials had dealings, so I will hand over to them.

Sam Cannicott: The voluntary agreements there were looking more through the safety lens and at transparency in a slightly different way. They looked at how we can have more transparency around inputs into training data, so there is a link there. We need to be careful about thinking about it in two separate ways. A couple of weeks ago, the US executive order, which covered a lot of AI issues, included this issue as well and it tasked its copyright office to look at this. So we are not the only country dealing with this: all jurisdictions are having to look at how we navigate this.

Lord Young of Norwood Green: I want to know how closely you are able to work with Europe on the general issue of safety and copyright.

The Chair: Let us discharge with copyright before we come back to that.

Baroness Harding of Winscombe: Specifically on copyright, I absolutely hear your point that this is a complex issue to resolve, but I want to put a couple of things to you and get your reaction. I am not a copyright lawyer, but I would hazard that you could have made the same argument—that it is international and that all the innovation will move overseas—with any other form of copyright infringement. I wonder what it is about large language models that means we think they are best dealt with in voluntary agreements, when in all other areas we have discovered that you need law that is black and white.

Viscount Camrose: In principle, that is an absolutely fair point. Obviously, the distinction with AI is that it can copy an awful lot of information quickly, inexpensively and in new ways that have not been available to copyright infringers before. So it is the same risk of copyright infringement, but it is happening many millions of times faster, which is why it is more complex. It is quite straightforward for someone who intends to infringe copyright to train their model in a different

jurisdiction, in ways that, for existing breaches of copyright law, are more complex and expensive.

Q143 Lord Foster of Bath: I have a simple question. Notwithstanding the court cases going on, like *Getty v Stability AI*, what is your understanding of the law currently in the UK? Do you believe that someone developing a large language model should pay for the data they use to train it?

The Chair: In addition to whatever answer you give Lord Foster, would the Government make a legal statement on their position on copyright in the context of large language models?

Viscount Camrose: In general, if you are copying a copyrighted piece of material, you are infringing that unless you have the permission of the owner, through a licence or other means, or there is an exception. Clearly that is a very general position but, depending on the specifics of the large language model involved, that can vary a great deal, because the creative content is copyrighted but the data or the information within the created content is not. Depending on the specific model you are using, that might give rise to a different answer from the overall question. That is why we are waiting for the courts' interpretation of these necessarily complex matters. That said, with regard to the Chair's follow-up question, overall there is no philosophical problem with legislation, but it would be difficult for us to design a piece of legislation that—

The Chair: I am not necessarily looking for legislation, just a formal statement from the Government of their interpretation of the law.

Viscount Camrose: Sorry, I misunderstood. I absolutely see the purpose of the question, but I worry about committing to that, because the uses and the context in which these potential infringements are occurring are so wide that a statement that affected one might not be applicable to another. That is my concern.

Q144 Lord Foster of Bath: If we cannot get clarity on that, can we get clarity on another issue? On the principle that "garbage in, garbage out" applies to the development of large language models, clearly the safety and usefulness of an LLM depends on the quality of the data that is input during the training. If we accept that, do you therefore accept that there is significant benefit to ensuring security and confidence by insisting on the transparency of the data that is input, whether or not you have to pay for it?

Viscount Camrose: The premise of your question is well made. The CDEI and Central Digital Data Office has recently published the algorithmic transparency recording standard. A range of products are produced by the CDEI to guide not only the safe design of inputs into a data model but how to make those transparent, which is not a trivial question. As for enshrining that in law, our view is to test, iterate and come to a decision about legislation based on observation of how that is

implemented, how that goes and what evidence we can draw from the experience of doing that.

The Chair: It seems from the evidence that we have taken that agreement on this is further away than you might hope, so further work needs to be done.

Baroness Featherstone: I have a quick point. Minister, you referred to cases before the courts.

The Chair: Do not forget that we are all sub judice here on specific cases.

Baroness Featherstone: I was not going to mention any cases. I will simply say that other witnesses have said that that will take a long time. I want to know what the Government are going to in advance of that, because I do not think people who own IP can wait for these legal cases.

Viscount Camrose: I agree, and I did not want to give the impression that we were waiting on the outcome of any of those cases. We are pushing hard to get to a point of voluntary agreement with the players here. If we are unfortunately unable to do so, we will have to carefully consider what we do about that, but in my view that really would be unfortunate.

Q145 **Lord Young of Norwood Green:** I want to know how closely you are able to work with Europe on this issue.

Viscount Camrose: Across AI, we work closely with the EU. Both sides recognise that international interoperability of regulation is key. It is clear that we have slightly different philosophies on how best to regulate AI, but we continue to engage with them, not just on copyright but across the whole question of how to make AI safe.

The Chair: I am grateful to all three of you for being here. There were a couple of questions that you agreed to follow up in writing. We look forward to that further information as soon as you are able, because, as I said at the beginning, this is our final public hearing of this inquiry. I thank everyone who has given evidence to us, both those who gave us written submissions and those who have appeared before us in the various different hearings that we have had. We will now go away, put a wet towel around our heads and try to arrive at some clear conclusions and recommendations heavily directed at the Government. We aim to publish our report early in the new year.