



Communications and Digital Committee

Corrected oral evidence: AI and copyright

Tuesday 9 December 2025

3.35 pm

[Watch the meeting](#)

Members present: Baroness Keeley (The Chair); Viscount Colville of Culross; Baroness Elliott of Whitburn Bay; Baroness Fleet; Baroness Healy of Primrose Hill; Lord Holmes of Richmond; Lord Knight of Weymouth; Lord McNally; Lord Storey; Baroness Wheatcroft.

Evidence Session No. 5

Heard in Public

Questions 79 – 105

Witnesses

I: Professor John Collomosse, Professor of Computer Vision and AI, University of Surrey; Ed Conolly, Vice President Engineering, Cloudflare; Eugene Huang, Senior Strategy Advisor and Co-founder, ProRata.ai.

USE OF THE TRANSCRIPT

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

Examination of witnesses

Professor John Collomosse, Ed Conolly and Eugene Huang.

Q79 The Chair: I think you caught a few minutes or a little of our first session. We now start the second half of our session today and I welcome our witnesses. In this part, we are interested to explore further some of the challenges on enforcing copyright in relation to AI systems and the technical solutions that help to address those challenges. The first question is: which stages of the generative AI life cycle, from data access to outputs, raise the greatest practical challenges for right holders seeking control and remuneration, and for developers seeking access to content? I am sorry; I should have asked you to introduce yourselves, if you could do that first.

Professor John Collomosse: I am a professor of artificial intelligence at the University of Surrey, and one of the co-investigators of the UK AI CoSTAR National Lab, which was responsible for the recent *Time to ACCCT* report, which sets out some technical and policy recommendations for the copyright and AI challenge. I am also director of DECaDE, the UK Centre for the Decentralised Digital Economy, which brings together researchers across the University of Surrey, Edinburgh and the Digital Catapult. Our work at DECaDE fed into the ACCCT report, looked at digital supply chains and how media provenance technologies can support consent and compensation in the creative sector in the age of generative AI.

Eugene Huang: Thank you for the invitation to provide evidence to the committee. I am the senior strategy adviser and co-founder of a start-up named ProRata.ai.

Ed Conolly: Thank you for the invite here today. It is an honour. I am the vice-president of engineering at Cloudflare. We are a large internet infrastructure provider, and I look after the teams that provide security and performance products to millions of websites all around the world, including—relevant to today—all our bot protection and identification teams and systems.

Q80 The Chair: Thank you. Let me repeat the question. We are interested in which stages of the generative AI life cycle, from data access through to outputs, raise the greatest practical challenges for rights holders seeking control and remuneration, and obviously for developers seeking access to content.

Professor John Collomosse: Data is accessed at several stages in the AI life cycle: during the large foundation model training, during model fine tuning, and sometimes at runtime when a model might retrieve additional data to help to answer a query. At all these stages, the challenge for rights holders is the same: how can they meaningfully control whether their work is used and under what terms, and how can they gain an opportunity to share in the value created by the model outputs? Today we have standards and tools to deliver that control via

blanket opt-in and opt-out indicators. That is a great start, but in the workshops we ran with creatives through the ACCCT project it was clear that creators wanted much more granular control, particularly the need to specify who may use which works, on what terms and for what purpose—in other words, essentially an asset-level licensing system. In ACCCT, which is an acronym, we summarise this as the three Cs: control, consent and compensation.

On the developer side, the challenge is of course getting access to sufficient scale of content at sufficient quality and knowing where it has come from, so it is about provenance. Provenance matters because synthetic or low-quality data can harm training and because, without a clear chain of custody, it is difficult to handle content in a rights-respecting way. That has now driven developers to make deals with large collections of content. That is an opportunity being mirrored by some of the policy initiatives around the Creative Content Exchange to derive value from large cultural-heritage collections. The challenge there is that the creative sector is inherently decentralised. More than half the global creative GVA is from small businesses and freelancers that do not really have practical, low-friction ways to license content at scale. My evidence today is that attaching provenance signals to assets can help to address the challenges of both sides. They can give creators better agency and give developers scalable access to the content that they need but in a rights-preserving way.

The Chair: That is helpful.

Eugene Huang: I agree with what John has said, especially about the three Cs—consent, control and compensation—especially in terms of a framework that this committee would find helpful for how we understand the tensions that exist between rights holders on one side and technology companies on the other.

Let us look at those three Cs for rights holders. Fundamentally, consent means whether or not your content should be used in training. Control means, even if you have given that consent, whether your content should be used in the output of a generative AI model. On compensation, I know the previous panel talked about it, but compensation really is at the crux of the issues being debated around AI and content. Ultimately, most of the value is being generated on the output side, not the input side. I know that Ed will most likely talk about some of the work that they are doing on the blocking side but, when it comes to setting up a market mechanism, the challenge that exists today is about the value on the input side from the training perspective as well as on the output side. If you are a content creator, how do you take advantage of all the incredible value being created when it comes to AI and AI generation? Ultimately, if you are just being paid on the input side, you are not taking advantage of any of the great value that is being created by AI today.

Q81 **Ed Conolly:** Building on the previous points, I agree very much on the consent point. It is critical that we find the right ways to let content

creators indicate which ways their content is intended to be used, whether it be for AI training, for AI searches or for traditional search. We also believe that transparency of the bots and crawlers that are accessing sites, and for which purpose they are doing so, is critical. We believe that if we have those things in place then a fair marketplace can exist to compensate content creators correctly. We certainly feel like we have a solution that works on the data access and input side, and that that is one of the simpler places to implement such a system.

However, we also believe that, in order to get that right, we need a level and even playing field. A lot of the AI companies that exist are in a reasonable place to do that, but some of the existing large companies have some unfair advantages here. I will call out Google specifically, which currently combines its traditional search crawler with the crawler that accesses AI and performs AI summaries of contents and live AI searches too. We strongly believe that, unless we address that problem in some way, we cannot create a fair marketplace for everyone.

Q82 The Chair: The follow-on point is that we have to weigh up whether problems could be addressed through technical measures—we have looked at that in a number of sessions and want to carry on doing so—and where we might need legal, regulatory or market interventions. I might start with you, Mr Conolly, because the problem you have just described with Google might be one of those.

Ed Conolly: Precisely. From a technical framework standpoint, many companies today have the ability; the transparency that we need is in place. We have a marketplace framework, and we have working end-to-end payment solutions. So we pretty much have an end-to-end solution that works for most people.

We hear from content creators that they are interested in participating in the platform because of course they want to be remunerated fairly, and we hear from lots of AI companies that are also interested in being part of the platform. The big hold-up seems to be the elephant in the room: is Google here? Effectively, it has this unfair advantage, and if Google is not paying then no one else wants to pay either. That is a problem that we need to find a way to address, and I am not sure that just leaving it to market forces, as we currently are, is going to resolve that important problem which we need to solve.

Q83 The Chair: So we need a legal approach to a regulatory problem.

Ed Conolly: I believe so.

Eugene Huang: To give the committee some framework for my comments, ProRata has developed a technology that allows compensation based on the output generated from AI large language models. From our perspective, there is technology that works today, which we have demonstrated, that begins to show how output-based compensation has the ability to work. We are at the early stages of this, but one of the things we are proud of is that we have partnered with 150 rights holders covering over 1,000 media properties. Here in the UK, that

includes Sky News, the *Guardian*, the *Independent*, the *Telegraph*, the Daily Mail Group, which is an investor, and a long tail of smaller organisations. Our business model is that it is a 50-50 rev share, so 50% of all our revenues go out pro rata—hence the name—to rights holders when their content is used.

The question is: what is the challenge? As Ed partially mentioned, there is a cold-start problem in some of this when it comes to companies like Google. For full transparency, I am a former Google employee, but I will say that Google is a bad actor in this, as Ed pointed out, because if Google does not have to pay for content then no one else is going to. I know that the CMA has recently designated Google as a firm with strategic market status in search, and I applaud that; its actions when it comes to marrying up its web crawler with AI is an example of how it is unfairly leveraging their position.

However, on the output side, when it comes to our technology, one of the challenges in getting this market started—as I think your previous panel was alluding to—is that if you are going to rely on fair-use arguments, especially fair-use arguments that try to protect you from liability in terms of how you have used the content, then of course you are going to say that these technologies cannot exist or do not work because it undermines the argument surrounding how you have trained your content and whether you can compensate fairly. From ProRata's perspective, yes, there should be market mechanisms, but there should also be a recognition that some of the market mechanisms may actually be market failures today because of how the market has evolved vis-à-vis what we see today from an AI-licensing perspective.

Professor John Collomosse: We see AI and copyright, fundamentally, as a content supply chain problem. The technical measure that you are asking about is really to address this problem and agree on a machine-readable, interoperable way to communicate where or who contents come from and what rights are associated with it as it moves around the internet from creation to consumption, ultimately to potential use in AI models.

As soon as we start talking about agreeing on an interoperable way for machines to communicate, we are really talking about a technical standard; if it is an open technical standard, that is all the better because that will reduce friction and cost in adoption. I am a strong advocate for open standards as a big part of the solution to the copyright and AI challenge. I am sure that we will talk about examples of this later—there are examples, such as C2PA, that can be built on for rights and licensing—but technology pieces exist today to signal that rights can be built on for licensing.

The question is definitely a sociotechnical one: how can markets and policy drive the adoption of these technologies? How can we set an expectation on content platforms not to strip away these signals as they move around the internet in the supply chain? Crucially, how do we ensure that developers respect these granular signals? There is little

point in going to all this trouble of robustly attaching licensing and rights information to content if it is ignored.

This could be approached through incentive-based compliance—a UK AI kitemark was suggested in our workshops on the ACCCT report, backed by an audit mechanism in the spirit of ISO 42001—focusing on processes and data governance rather than on the disclosure of training details, which I know was raised in the previous panel as a concern for some AI developers in general.

Q84 Lord Knight of Weymouth: I have a set of questions about technical mechanisms at the training end of the process, so I will probably focus on some of you more than on others. I will probably ask stupid questions that require short, stupid answers—actually, clever answers—and we will go from there.

In terms of crawling, scraping, TDM and so on, it feels as though there is a moving target when it comes to whether what is coming and looking at your website is a bot or a human. If we wanted part of the regulatory regime to differentiate between the two, is it possible, with the moving targets and the sophistication of bots, to differentiate between where they are coming from, whether or not they are bots, what their game is and what their intent is?

Ed Conolly: That is an excellent question; it is one that we talk about a lot internally. The way I think about the bot versus human distinction is that, if you try to make an abstract classification, there are some things that fall clearly into one camp and others that fall into another, and there is also a grey area. The way we like to think about it is more in terms of the intent of the visit. If the intent is to download or scrape content for some purpose, the question of whether the actor doing it is a bot or a human is perhaps less important than the intent of the visit.

Q85 Lord Knight of Weymouth: There has been some discussion of robots.txt, machine-readable attribution and so on. Technically, there is no reason why we could not move to something more sophisticated than robots.txt in order to differentiate between crawling and scraping and different intent. Is that right?

Ed Conolly: That is correct, although we should take a second to talk about some of the limitations of robots.txt. Let me focus on Google for a moment. I can tell it not to use data for training purposes or for traditional web search purposes, but, with robots.txt, I am delegating the authority to make that decision to Google; the actual bot that visits my website is the same bot, irrespective of what I use. That makes a difference when we talk about transparency in auditability. As a content creator, when I see Googlebot appear in my logs, I have no idea whether the intent of that visit was to scrape my content for the purposes of AI training or to scrape that content for the purposes of the traditional web. We believe that this is a big issue for transparency and for rights holders.

Q86 Lord Knight of Weymouth: Do you think that solving that particular

problem is something with which the CMA, as the competition regulator, should be grappling? Or is it something that another part of government should be looking at, in terms of ensuring that companies do not use their market power to have the same bot doing the same thing? Should we separate off searching and AI training?

Ed Conolly: I am not sure that I know the answer to that question.

Lord Knight of Weymouth: But do you think someone should do it?

Ed Conolly: Yes—possibly the CMA and possibly others. I would suggest that those who are interested in the incentives, in terms of ensuring that copyright holders are fairly compensated for their works, are the people who should drive that forward.

Q87 Lord Knight of Weymouth: With some of these things, we may want to impose some technical standards that we want people to abide by here in the UK, as we might do for transparency. We have heard from previous witnesses assertions that, in essence, this is not realistic because everyone would just go somewhere else; that we would be wasting our time; that this is a global market with global sets of technology; and that we should just give in to everything US tech tells us.

Ed Conolly: I fundamentally disagree with that. The world is looking for a solution to this problem right now, not just here. This is an opportunity for us to lead the way, set an example of what can be done and, potentially, work with some other countries in defining it. People are looking for someone to come up with an answer to this, so it is an opportunity for us to lead the way and define a pattern that could be used. We at Cloudflare are already trying to set and define some of those standards.

Q88 Lord Knight of Weymouth: My final question for now, before I get John's reflections on some of these things, is about our discussion with the previous witnesses of how realistic it is to have granular transparency data recorded automatically in some kind of register, as the crawling and scraping takes place. The assertion is that that is ridiculous and unrealistic. Do you have a view?

Ed Conolly: Some of this comes back to the difference between charging for inputs and charging for outputs; Eugene is better placed to answer questions about charging for outputs. What I can tell you about charging for inputs is that it is not significantly complicated to implement a system that does this. As a bot visiting a website, I can make a check against a purpose—I am here to crawl content or whatever—and you get a permissions check at that point. You are not saving the data in your database, trying to attribute it and storing billions of lines on where I am allowed to go and data attribution. You make a single permissions check as you enter the gate into the site, as it were; if you are allowed access, you go in, but, if you are not, you do not. I do not believe that that is technically complicated.

Q89 Lord Knight of Weymouth: That brings me to where I was going to end up anyway—I was going to bring John in but I might as well finish this off—which is the asset-level, or unit-level, approach. If I am uploading my song, am I better off trying to get some protection for it or going to a site that has a regime of protection at the site level?

Ed Conolly: We have a workable solution in our current state; it is, in effect, a site-level solution. An asset-level solution is probably better—that can come in time—but I do not see it as a requirement for us to make progress under the current set of rules.

Professor John Collomosse: I would love to come in on that because there has been a lot of discussion about site-based methods, which are an important part of the puzzle. In essence, they are a voluntary access control mechanism. If a creator uploads their work to a platform but it then gets reposted or reshared elsewhere on the internet, the preferences attached to that content are not going to move with it; that is why we need asset-level controls to complement the site-level signals.

An asset-level signal requires open standards so that it is interoperable when it moves around and can be read across different platforms. C2PA content credentials are an example of an open standard in this space. It is a provenance standard, so it is a way of describing where content has come from, who made it and what has been done to it. Through that mechanism, you can assign ownership and start talking about layering on for AI; opt-in or opt-out preferences; and, ultimately, granular licensing control, which we were proposing.

This mechanism is complementary to the access control mechanisms because, in an asset-level system, you already have access to the asset. It is now a case of asking: what rights are associated with it? What are the opportunities for compensating the rights holder for that asset if it gets used? Then there is the possibility of layering open standards on top of things such as C2PA or rights-description languages, such as ODRL. People are already starting to explore this licensing scheme, but that is not as adopted as a site-level scheme.

Q90 Lord Knight of Weymouth: We are relatively familiar with Creative Commons, where you set different licensing arrangements according to the content. All we are talking about here is essentially something similar for the AI world.

Professor John Collomosse: I would say that it is more granular, in the sense that you can describe who owns a piece of content and then they can set up their own preferences and licensing schemes on top of that using these other standards, which would enable, for example, remuneration models to be layered on top rather than blanket licensing agreements that can be assigned on to content, which would be something like Creative Commons. Again, it is another piece of the puzzle.

Q91 Lord Knight of Weymouth: Finally, is that best put into regulation or into standards, where we are creating some kind of market

encouragement in terms of what good looks like?

Professor John Collomosse: Standards, as I have said, are a major part of the solution. Industry, the implementers, are the right people to lead on the development of standards. We have seen that through, for example, Content Credentials industry partners coming together to create a standard voluntarily and implement that at scale. That did not require any form of regulation, although the adoption of that standard into content platforms largely around AI labelling was likely encouraged by regulation. As I say, there are stick and carrot approaches. The incentive-led approach would be something along the lines of a kitemark that could be issued around the training of an AI model where, through some audit process, it could be demonstrated that these markers that move around with assets had been respected. That would be another way to do it that would not necessarily require regulation.

Q92 **Lord Knight of Weymouth:** Do you have a preference for regulation or for the kitemark-type approach?

Professor John Collomosse: I am not here to set out policy options, but my feeling is that I have seen a lot of success in industry-led adoption of open standards through provenance standards. I would hope that that route could be similarly followed.

Q93 **Baroness Wheatcroft:** How well do you all think the provenance marking that you are talking about could work, given that so many small contributors will find it difficult to check whether their work has been used?

Professor John Collomosse: It is precisely the decentralised nature of open standards, the fact that anyone could write provenance markers into content, that makes them attractive for the creative industries, where the majority of practitioners are freelancers or SMEs.

Baroness Wheatcroft: They are busy.

Professor John Collomosse: Yes. There are free tools that enable this. There are open-source toolkits that would enable smaller developers that do not have access to large development teams to implement these standards. I do not see there being a technical barrier. The sort of barriers that would exist are more around larger platforms honouring and retaining these signals as they moved around through the content supply chain. But, again, these provenance signals that are expressed through data that moves around inside the asset—what we call the metadata or supplementary data—can be reinforced or made sticky to assets through things like watermarking, where you inject an invisible identifier inside an asset like an image which lets you reidentify the content to look up and recover that metadata if it is stripped away. There is a similar technology in fingerprinting. All these things come together. These three technologies—we call them the three pillars—make this a practical technology for adoption.

Q94 **Baroness Wheatcroft:** Essentially, you think the obligation or decision

to watermark, or whatever, should be with the creative rather than the onus being with the AI company using it.

Professor John Collomosse: We have heard in our workshops that creators want granular control and to be able to state, “You can use my content for this purpose and on these terms and, actually, I’d quite like to be paid for that content use as well”. So taking the active step to issue those rights and licences on your content and make them durable so that, when they pass into the content supply chain, they can then hopefully be respected by, for example, model trainers and AI developers, is the approach that we would advocate for.

Q95 **Baroness Wheatcroft:** Mr Huang, do you agree with that?

Eugene Huang: I do. I would like to use an analogy with the music industry. I say this coming from a place where Universal Music Group is one of our partners. If you think about how the music industry has evolved, especially with new genres like rap and electronic dance music, a lot of that relies on sampling other pieces of music and content. As I understand the music licensing regime, if you are a rights holder, you can upload your content to a platform like a collective licensing agency or some other organisation that will then take the responsibility of helping you to enforce those rights. One part of it, as John indicated, is a responsibility that the content owner has, at least in the music space, to assert their rights associated with a unique composition that they have made. It is about how that assertion of rights then gets prosecuted. You can think about whether existing mechanisms can be used to ensure compensation and make sure that rights holders’ rights are enforced, from the big folks all the way to the small folks. It exists in other places, like it does in music and video. There is no reason why it should not exist in the AI space.

Q96 **Baroness Wheatcroft:** Might I just push you a little on the music issue? Until the recent licensing agreements with AI companies, the sort of cases that made it to court and eventually got some compensation for those whose music had been pinched, effectively, were led by pretty big names.

Eugene Huang: Yes.

Baroness Wheatcroft: Why do you think that was?

Eugene Huang: They had the resources to prosecute it. The reason why I mentioned the collective licensing agreement agencies, though, is the recognition that the smaller rights holders—we have heard this directly from smaller rights holders who signed up with us—want to create whatever they are creating. It is hard for them to think, “How do I manage my rights? How do I go about this, that and the other?” When some of the large AI companies come to august bodies like this and say, “We want a blanket exemption”, that is where I say, “Wait, you want a blanket exemption, not for the technical reasons that you are stating but for a business model reason”. Ultimately, if you think that you can get one of the critical inputs to your business for “free”, all of us here will tell

you that it is not free. There is a cost associated with that and, from our perspectives, we are trying to figure out how we bring some version of technology and market mechanisms to bear to give rights holders the ability to say, “Hey, our content is not free. We’ve put a lot of effort into making that”.

Q97 **Baroness Wheatcroft:** That is very helpful, thank you. Mr Connolly, we finally come to you.

Ed Conolly: I do not have much to add. I agree with all the previous points. What we hear from content creators is a desire to have more granular controls as enforceable mechanisms and make decisions about fair compensation for their content. I agree with all the previous points.

Q98 **Baroness Wheatcroft:** Before we move on, can I come back to you, Professor Collomosse? Can you elaborate a little on your commercial interests?

Professor John Collomosse: As with many academics in the AI sphere, I have a commercial hat. As I declared, I work with Adobe, but I am here to speak on my academic behalf today.

Baroness Wheatcroft: Of course, but it would be helpful to the committee to understand a little more about your background and where your knowledge comes from.

Professor John Collomosse: Sure. I am happy to elaborate. I have been working in the field of artificial intelligence for 25 years.

Baroness Wheatcroft: Your links with Adobe.

Professor John Collomosse: Yes. I work at Adobe Research in a part-time role.

The Chair: We can leave it there. Thank you. We do not need to go into that now. Lord Holmes has the next question.

Q99 **Lord Holmes of Richmond:** Thank you, Chair, and good afternoon. What kinds of licensing, remuneration and market infrastructure can facilitate access to creative works for AI development in the UK while ensuring value flows to those rights holders?

Professor John Collomosse: The best way to view the AI copyright problem is as a supply-chain problem. Provenance standards like C2PA are a way of increasing transparency in the content supply chain. Just as you might want to know how the sausage in your supermarket is made—that is, your food provenance—you also like to know which model made an AI image and which datasets or collections an AI model has been trained on. Provenance data in model inputs opens up the possibility of potential revenue-sharing models at the collection or dataset level. For cases where a model also answers a question by pulling in specific documents or articles, say, from a newspaper publisher, provenance can show which sources were accessed to answer that query, and revenue-sharing agreements could be made on that data access as well. So I see provenance as essentially an infrastructure that can be built on. As I was

previously explaining, it is essentially an open standard, something like C2PA, that is supported by various technologies to make it sticky and durable to assets.

Once you have established the provenance of an asset, you can start talking about its ownership as well as its authenticity, which is another great use case for that technology, and then start talking about defining—using other open standards layered on top—rights and licensing mechanisms. I see that as being a technical stack that is highly applicable to the challenge of copyright in AI.

Eugene Huang: From ProRata’s perspective, we only use licensed content and we are very transparent in terms of our licence, which is that we do not pay any up-fronts but we give 50% of our revenue back to content creators based on how their content is used, or at least how we calculate how their content is used, coming out of an AI LLM output. That is different than what most other individuals in this space have been trying to do, which I recognise is slightly challenging and potentially impossible, which is to trace a piece of content all the way through. What we do instead is statistically reverse engineer what goes on in the black box of the LLM to give a report that says, “We believe that 25% of an article from the *Daily Mail* contributed to this output”.

From a licensing and market perspective—piggybacking off of what John said—we respect all the signals that a company says that it wants to give us in terms of how its content is used or not. We have a blanket compensation scheme, but if we were to engage in a compensation scheme that ended up being unique to a particular publisher then we would adopt that as well. That is a long way of saying that everything John has been talking about we have attempted to implement with our technology.

Q100 Lord Holmes of Richmond: On the percentages, what level of precision could you go to on that—say, to what decimal point—and could you do it irrespective of scale?

Eugene Huang: Right now, in our technology tests, it is north of 98.99%. When it comes to scale, right now we have upwards of 50 million documents that we do our attribution against. We are a start-up, and having come from large companies I recognise that that is a different scale than very large companies, but I think those who have signed up and are backing us go to show that there is faith that we can help to move the market in this way.

Ed Conolly: I will talk about the data input side of this equation and a bit about not just what could exist but what exists today and tools that my company, Cloudflare, has built. Other marketplaces already exist, so it is definitely not just us.

The foundational piece on the input side is that you need a robust bot identification mechanism, and you need a robust set of tools to give to content creators to let them declare the purposes of what their content should be used for, as well as the ability to set a price to access that

content. That tooling already exists. When AI company bots come to visit sites, before they access the content, they make a check to understand that a payment is required in order to access the content, and they can make a decision about whether or not they want to pay the price to access it. Should they accept that, they can crawl the content and, through Cloudflare's marketplace, value can be distributed correctly. Again, many other marketplaces exist. We are supportive of different models existing because it is hard to know exactly where this is going to go. We as a company feel strongly that it is important that we protect the future of the internet and content creators, not just the large media publishers but the long tail of creators too. We believe that having a system that allows everyone—individual contributors alike—to participate in something is really important. That system already exists today and, like I said, we have many content creators and LLM companies lined up ready to work with it, but, at the risk of repeating myself, there is one blocker in the room that we feel we need to address: Google.

Q101 Lord Holmes of Richmond: Thank you, that is really clear. What role is there, if any, for Government and regulators in spinning up this marketplace for licensing?

Ed Conolly: An awful lot is needed. Like I said, from the tools and technology standpoint, I think we are pretty much there. We have all participants ready and willing to go but—I will say it maybe one more time—the only blocker that we have here is one company. That is the one thing we could do with assistance with. Other than that, everyone seems keen to participate, to be honest.

Q102 Lord Holmes of Richmond: Eugene, is there a role for Government?

Eugene Huang: I support what Ed has said. Market mechanisms are developing. As Ed alluded to, there is a sort of market failure, and I would probably encourage the committee to look at those places where there is market failure that is hindering the development of those marketplaces, whether that is Google's behaviour, as Ed pointed out, or the behaviour of other tech giants that are trying to get away with saying, "Okay, we don't have to pay for this because we're relying upon a fair use argument that is outside the jurisdiction of UK copyright law". Perhaps it ends up being a case of, "If you want to participate in the UK economy, you have to abide by whatever rules and regulations that this body decides end up being the right balance between supporting innovation and supporting content creators".

Q103 Lord Holmes of Richmond: On this specific copyright point, are there any other market failures that you would bring to bear for us today?

Eugene Huang: Not off the top of my head.

Lord Holmes of Richmond: Thank you. John?

Professor John Collomosse: I see four areas where government and regulators could help. The first is endorsing open standards. As I say, they are a big part of the solution: signals that are open, using

technology that is free and interoperable, to declare the provenance of assets and track their rights through the content ecosystem. I do not think the Government should name any particular standard in legislation—that is dangerous and it would go out of date very quickly—but there could be a role for regulators to give guidance on best practice in industry at any given time.

The Government could help by encouraging durability and respect for the signals that I have mentioned with content platforms to preserve them and maybe by encouraging, through some sort of kitemark scheme, respect for those signals and recognition when developers have followed them. Some really interesting work is being done around the Creative Content Exchange idea in creating a platform for licensing content. If that were done through open standards, and again this would be leading by example, we would see a blueprint that could shift from large collection owners to the decentralised creative economy—which, as I say, is the majority of the creative economy—and that could be used as a vehicle for small and medium-sized enterprises and freelancers to license their content, if done through these sorts of frameworks. Those are the main areas.

Lord Holmes of Richmond: Thank you, you have all been very clear

Q104 **The Chair:** One last point. You heard the earlier panel, and you probably understand that we are moving through a situation where we are talking to creatives, rights holders, tech companies and AI providers and hearing different perspectives. The last panel said things like 20% to 40% of the future market will be dependent on their version of how a copyright market should be. Just how optimistic or otherwise are you about being able to sort this out? A couple of your organisations might be involved in the technical working groups that the Government are running, and it would be helpful if you could tell us how you feel about this. Is it something that, if we grapple with it, we can get to a place where we can help with remuneration and support the move forward with AI?

Professor John Collomosse: I am very optimistic, and the way to realise that optimism is by moving the argument on from the polarised opt-in/opt-out debate to a place where we recognise that creators, as we have heard, want granular control over who can use content for what purpose and under what terms, and ultimately remuneration models can be established around that. This is not a question of opt-in or opt-out; both technologies—robots.txt for access control and granular licensing technologies—have a place. If we recognise that there is a landscape of solutions and work with industry and open standards bodies to develop those solutions and incentivise developers to adopt them, then I would be very optimistic that we can find ourselves in a better place.

Q105 **The Chair:** And that is technically feasible?

Professor John Collomosse: Yes, I believe so.

Eugene Huang: I am going to piggyback off of what Ed said at the very beginning. I talk to stakeholders around the world, and there is a desire

for leadership in this space. When I say leadership, it is not the “anything goes” US style of leadership when it comes to AI regulation, because that does not take into account any of the concerns that rights holders have—and this is coming from the perspective of a technology company and a technologist. I am talking to rights holders all the time, and I fundamentally believe that in creating AI we have to acknowledge the contributions, both historically as well as on an ongoing basis, that content owners can make when it comes to helping AI to get better. Your previous panel talked about how technology can enable use cases. If all of a sudden AI just takes over the creative industries, what ends up happening then?

I am optimistic that there can be a solution if, as Professor Collomosse said, we can move beyond the polarising yes or no and get to a better place, especially from the vantage point where all of you sit, asking: how do we set up the right market environment, especially when it comes to smart incentives and smart regulation to let this marketplace thrive? That is the type of leadership that, when individuals like Ed and I go and talk to stakeholders right now, we find is lacking, not just here in the UK but on a global level.

Ed Conolly: I am enormously optimistic when I look at the situation today. We have AI companies training in the US with no contribution to anyone creating anything but, if we put in place the new operating model that we have spoken about today with granular controls, not only can we bring AI training to the UK but we can compensate people who are creating works, so we can transition two things at the same time. That is a reason to be enormously optimistic, and technically it is very feasible.

The Chair: That is excellent. It is always good to end on an optimistic note.