

Treasury Committee

Oral evidence: AI in financial services, HC 684

Wednesday 15 October 2025

Ordered by the House of Commons to be published on 15 October 2025.

[Watch the meeting](#)

Members present: Dame Meg Hillier (Chair); Dame Harriett Baldwin; Chris Coghlan; Bobby Dean; John Glen; John Grady; Dame Siobhain McDonagh.

Questions 131-206

Witnesses

I: Jessica Rusu, Chief Data, Information and Intelligence Officer, Financial Conduct Authority; Tom Mutton, Director of Central Bank Digital Currency, Bank of England; Jonathan Hall, External Member, Financial Policy Committee, Bank of England; and David Geale, Executive Director of Payments and Digital Finance, Financial Conduct Authority.

Examination of witnesses

Witnesses: Jessica Rusu, Tom Mutton, Jonathan Hall and David Geale.

Q131 Chair: Welcome to the Treasury Committee on Wednesday 15 October 2025. We are here looking again at artificial intelligence in financial services. This is our third session on the issue, and we are delighted to have witnesses from the Bank of England and the Financial Conduct Authority with us to discuss, among other things, how it will be regulated. I am pleased to welcome Jonathan Hall, who is an external member of the Financial Policy Committee at the Bank of England, and Tom Mutton, who is the director of central bank digital currency at the Bank—he will be speaking more for the Bank, as Jonathan Hall is an independent external member. We also have Jessica Rusu, who is the chief data and information and intelligence officer at the Financial Conduct Authority, and David Geale, who is the executive director of payments and digital finance at the FCA. A very warm welcome to you all.

We have had some very interesting sessions with financial services, academics and others. It would be helpful to gauge what each organisation—so maybe not all four of you—thinks the greatest challenges are for UK financial services in using AI more.

Tom Mutton: I can say a couple of things. First, I do not think anyone will be surprised to hear there is a very significant opportunity around artificial intelligence in financial services. The joint survey we do together with the Financial Conduct Authority suggests that, by and large, artificial intelligence is being used in relatively low-materiality workloads. Firms identified that roughly 70% of workloads are low-impact, low-materiality—things like “know your customer” checks, anti-money laundering and so on. Provided that there is an appropriate regulatory basis and firms understand the artificial intelligence they use, it would be appropriate for it to start to be used in more material workloads, where there may be greater opportunities for productivity, innovation and so on. But there needs to be the right regulation and the right understanding around that.

So there is a significant opportunity there, but there are clearly some very material risks, which we need to be across. The first, evidently, is cyber-risks and broader operational resilience. Our survey suggests that some of the models that firms are using are on the more complex end and that management understanding is imperfect at the moment. I think management have identified that themselves: only 34% of management feel they have a full grasp of the complexities of the models they are using.

Chair: That was your survey?

Tom Mutton: Done jointly with the FCA.

Jon is much more able to talk about this than I am, because he published a very good speech on it, but there is clearly also herding, pro-cyclicality,



HOUSE OF COMMONS

and potential volatility and disruption in critical markets. Then there are the impacts on systemic institutions—banks, financial market infrastructures—if they use artificial intelligence in ways that could threaten their safety and soundness, or threaten outages in their services. We are very focused on all those things.

A couple of points draw those risks together. The first, clearly, is the complexity and opacity of some of the models and the fact that they are very dynamic. The second is management understanding of the models, including—and perhaps in particular—where they are provided by third parties, where management may have imperfect sight into the models. Thirdly, the models depend on data, and many of the issues around them originate in the data rather than in the model. Finally, there are the third-party risks, which I am sure we will talk about.

Q132 Chair: When we did our work on this, we wrote to the top 10 banks about outages, and third-party risk, even in that area, was clearly a bit of a trend. Mr Hall, did you want to add anything to that?

Jonathan Hall: No, I thought that was pretty comprehensive.

Q133 Chair: Okay, we will get into more detail. Jessica Rusu?

Jessica Rusu: We welcome the opportunity to talk about AI again. The FCA has been proactive in thinking about this new technology. We speak quite openly with firms and our stakeholders about having an outcomes and principles-based regime. As you know, the UK has taken a vertical approach to the regulation of AI, and we have therefore taken time to think about this through the lens of consumers, the firms that we regulate and the overall systemic impact. Tom has already highlighted some of the risks.

In terms of the opportunities, we see firms broadly in two camps. There are firms that are progressive in adopting technology and thinking about the opportunities that AI presents. I can talk more about the types of use cases that we have seen. We also have firms that are more reluctant to lean into the benefits of technology, whether AI or other technologies, and have been more backfooted about adopting progressive analytics techniques.

Broadly speaking, we see a lot of opportunities particularly for consumers, such as financial inclusion. We can talk more about use cases, but there are opportunities for the use of machine learning and AI in particular to help consumers who have previously been excluded from credit-type opportunities to be welcomed into the market.

We also see great opportunities for AML, KYC and other types of fincrime prevention. Network analytics in particular is a great area of exploration. There are plenty of use cases and opportunities for firms to think about as they adopt AI.

Tom mentioned a few of the risks: critical third parties, concentration risks and understanding the interworkings of a model—we can explore bias in

more detail later. Certainly, there are data science areas of research and applications and techniques for thinking about bias and how it can be mitigated as a risk through the life cycle of the production of an AI product. All in all, there are systemic risks, as well as idiosyncratic risks for firms.

Q134 Chair: Both your organisations say you are technology agnostic and you are not prohibiting specific technology, but there are risks in that, aren't there? If you are technology agnostic, how can you demonstrate to us that you have the capability to really understand what extra risks there may be in AI compared with other technologies, or even between different AI models? I will start with the FCA again.

Jessica Rusu: In terms of technology agnosticism, it is really important as a principle to not prescriptively state that a certain provider or technique must be utilised to get to a particular outcome. However, you will have seen in our new strategy, published earlier this year, that the FCA intends to be increasingly tech positive. That is very important for good consumer outcomes and good market outcomes. It helps firms adopting KYC, AML and fraud and crime prevention.

The benefits of technology are one of the reasons we are being more assertive about our positioning on technology, both in terms of the ways that the FCA is using technology internally to perform our statutory duties and firms' use of technology.

Q135 Chair: Do you have the skills you need to keep ahead in what is a very fast-emerging set of technologies?

David Geale: Jessica can talk about the skills and training that we have. I would add that, in many cases, the risk already exists and AI perhaps exacerbates it. Take third-party risks, for example: we already see firms reliant on third parties for significant matters. The use of AI can increase that risk, but we already have tools to understand and mitigate the risk.

We also are very keen on being outcomes-focused. We have a series of tools, such as the senior managers regime, that play to accountability. We expect senior managers to focus on the outcomes and test the outcomes. That sort of expectation is already there; I think it just heightens the expectation around the use of AI.

Q136 Chair: How can you test that? One of the points that Mr Mutton made was about data—the imperfections that can get into the system. How can you identify that? This is partly where the skills question comes in. Will you be able to go back and interrogate an AI model or algorithms? You are doing some of this already, presumably, with basic algorithms.

Jessica Rusu: One of the ways we have been proactive in speaking to firms about how they use AI is through our innovation sandbox work. We have two quite unique offerings: one is our supercharged sandbox, and the other is AI live testing. The former leverages our existing digital sandbox capabilities and provides synthetic datasets such as fraud typologies, and other sorts of dataset that a firm might want to utilise to

develop an AI solution. We have gone further by offering them access to Nvidia's AI software suite, so that firms have a level playing field and access to the right tools and technology to develop their solutions. Alongside that, live testing is for firms that are further progressed in their own AI solutions and products that they have developed in their firms.

We opened both of those activities earlier this summer and are ready to begin testing in the next couple of weeks, with cohorts of firms going through both experiences. We think this is the right approach for us to understand the challenges that firms are taking on, as well as how they have addressed the AI life cycle, from the curation of the data to how it will be used, and on to the implementation. That gives us line of sight to where the risks are.

Q137 Chair: That is all great, but can you put it in simple English to explain to the consumer who will end up at the end of the chain of the AI work? When a firm comes in and says it wants to test this with you, can you walk us through what that looks like? What is your intervention at each stage, and what real sight do you get of what is often quite a commercially sensitive model for those firms?

Jessica Rusu: Let me give you a basic example. There is a firm that already has a chatbot—the traditional chatbot that you might type into and ask questions. The firm is looking to move that into a conversational experience, so that you can have a conversation and get information that you might need about products and services. Through the testing life cycle, it will be important for that firm to implement guardrails around the kinds of question that the chatbot might be able to answer. So if you ask it about a personal situation, it would say, "I'm sorry, I am unable to give you advice about that matter."

Q138 Chair: Okay. But you would have full sight of that at each step and they would not have any problems sharing that information with you?

Jessica Rusu: No, they shouldn't.

Jonathan Hall: As you said, AI and this type of technology is moving incredibly fast. I think that is a benefit of being technology agnostic at this stage. That is because being technology agnostic is forward-looking, flexible, outcomes-based. Think about the new technologies that Jessica was talking about: the firms and consumers and other market participants know what is expected of them in an outcomes base, and as they implement a new technology they know that those expectations have not changed. As we move through this phase when there is a lot of uncertainty around the technology, that is really valuable.

Obviously, technology agnostic is not technology blind. We need to make sure that from an FPC perspective financial institutions are able to withstand and are resilient to the shocks that might face them. Obviously, the question for us then is, are the shocks that they might face changing as a function of AI? We can look at that through the framework of stress-testing, for example. Our stress-testing framework may start to incorporate AI-based scenarios in a way that it might not have in the past.

So yes, I think at this stage, while there is still uncertainty and while these are hypothetical possible futures, it is really actually valuable. But it is not hard-coded that we need to remain technology agnostic, so if there is something specific that means we could change that, we will consider that.

Q139 Chair: So if something dangerous emerges, you will consider how you would deal with that.

Jonathan Hall: Yes. The big question is: is there something that is really new, or is there a development of current risks?

Chair: When we had the financial institutions in, they said, "Nothing to see here. It's all fine. We will manage this." Perhaps I am mischaracterising it, so tell me if I am, but what we got from both of your organisations, the Bank and the FCA, is that you are monitoring it. You have the sandbox in your case. You are watching it at the FPC level—at all levels of the organisation you are watching it—but are you alert to what might go wrong? Do you know what you are looking for? Or are you waiting for something to go wrong before you act? That is particularly for the regulator, I suppose.

David Geale: If I can take an example, I would not expect my supervisors to be looking at the individual model. They are not coders and they will not go into that, but under interventions such as the consumer duty, firms are expected to design products with a target market in mind, with outcomes in mind, and they should know what to expect. We then expect them to test the outcomes to make sure the correct outcomes are delivered, and to demonstrate and be able to evidence that. That we can look at and do look at, regardless of whether the solutions are being delivered through technology or through individual advisers or whatever it may be. We can at any point go and look at what it was designed for, what is coming out of the other end and what controls are in place around this—aside from the outcomes, we can ask firms what controls they have in place. Examples of good practice could be firms who have dedicated AI risk committees, for example, looking and making sure that they have the appropriate controls around the technology they are deploying. We have the ability to look at the beginning. We have the ability to look at the end and at the controls that are in place, and we expect firms to act on it if the outcomes are not right.

Q140 Chair: But as the models get more sophisticated and learn from themselves, there could be a glitch that occurs between what was intended and the outcome. How would you have any visibility of that? Again, it perhaps goes to the skills point about how you would know where to interrogate it. As you say, you cannot possibly have enough staff to be expert in a particular AI model or how it develops, because it is just not feasible, is it?

David Geale: No, and I think there is a difference between what happens in sandbox testing, where firms do get a greater degree of scrutiny, versus day-to-day supervision. But still we expect the firms to be able to demonstrate the outcome they were expecting as it is being delivered. We



can look at indicators such as customer complaints. We can look at the ombudsman experience. We can look at trends in terms of products being sold and particular features and so on. There is a lot we can ask firms to look at to demonstrate whether what they expected to happen is happening. I think there is a point there. You are right: the models get more and more sophisticated and self-learn. That is why we are seeing firms step into this cautiously at the moment, because at the end of the day somebody is on the hook through the senior managers regime and somebody is accountable.

Q141 Chair: The chief executive and the chief risk officer, primarily.

David Geale: Yes, or it could be the head of the business unit. It depends where it is, but somebody has to be accountable for the activities, including the deployment of the technology.

Q142 John Glen: I do not want to put ridiculous hypothetical situations, but clearly the amount of data that is being interrogated and the sophisticated building up of new propositions based on almost compounded use of AI is very complicated and can lead to new products being offered to consumers. The Chair talked about the problem of your keeping up with the skills level, but isn't there a real danger as these business models evolve? The very clear notion that you gave us of what you are trying to do—what is the outcome and how do you measure it?—is quite difficult to sustain in a world where some of these firms will be using very sophisticated models that started in the sandbox, but moved on. How do you account for the frontier actors who are obviously pushing some of the understanding and actually stay ahead of the game? Don't you need the more specialist skills of people who can interrogate what they are doing proactively, rather than use that—dare I say it?—older model of looking at how regulation works?

David Geale: For me, principles-based regulation is absolutely right in this space, because it can move. The world is moving so fast that if we try and get too prescriptive, that is where we miss things in terms of the framework that applies. I think this plays right into principles-based regulation across the consumer duty, across the senior managers regime, and across the systems and controls that we expect. We need to make sure that our own people are able to interrogate what they see. I do agree with that. That is different from diving into the model itself and getting into coding. It is about how the model plugs together. Jessica might want to talk about some of the training we are doing to upskill.

Jessica Rusu: I would add a couple of things. From a data scientist perspective, we apply the same approach to testing an AI model. No two models are built with exactly the same parameters so, when testing a model, we have to look at the data that goes in and the data that comes out. If, for example, I want to control for bias, I am looking at the performance of 10 CEOs, where nine of them are male and one female, and I want to predict the performance of a future set of CEOs in a future situation where I have a different gender ratio, then as a data scientist I can use certain techniques to put synthetic data into the training pot, so



that the future application of the model would be non-biased. The way we are able to look at the sophistication and appropriateness of the model is by looking at inputs and outputs, and by measuring whether the impact of what is happening measures up against what the firm expects it to use.

Earlier, Tom highlighted that 70% of applications are internal, back office. Let us take complaints handling as a particular use case that a firm might want to automate. That information will be collated in the same way that a traditional MI or Tableau report might be surfaced up to an adviser, who would review the information and make decisions, with a human in the loop and the ability to interrogate whether the information they are presented with is accurate. In that same way, we are approaching the implementation of AI across the FCA.

Q143 Chair: What if a rogue bit of data gets in? Do you think you would pick that up, Mr Geale, at that early point? That is where it could go quite badly wrong, if we put our sceptical hats on for a moment. There is a lot of talk about how AI will solve every problem in the world, but all it takes is one error, does it not, and then we have a model—even with the best will in the world—where we will not be able to find it. That will be like looking for a needle in a haystack.

David Geale: Genuinely, I do not think that we would be able to spot a rogue piece of data ourselves. We would be looking at the outcome that has been achieved and, if those outcomes are off base or trends are coming in that we or the firm were not expecting, we would expect them to act on that and, first, to rectify it and, secondly, to go back to the cause and deal with it. As these things develop, however, that is absolutely the sort of thing that we have to be alive to. It is a risk.

Q144 Chair: Anything to add from the Bank?

Tom Mutton: I will add two things to that. Already our approach to supervision is to focus on frameworks and management accountability. We cannot check every position in a trading book; we can check the framework that is around it. The question is how to scale that sort of approach and to apply it to artificial intelligence. That is a very active area of work in the Bank and, with the industry, in our Artificial Intelligence Consortium. It is very difficult to say that we are going to place all the emphasis on prevention and detection; I think it has to be about frameworks and accountability. Also, it has to be about dealing with the impact of situations when they occur.

As David said, it is going to be impossible to get to a zero-error approach here. Very much in the space of operational resilience, the way that we approach this is to say that there will be events that cause disruption. How do we respond to them? How do we recover from them? How do we minimise that disruption? How do we stay within our impact tolerances? It is about striking the right balance between prevention, detection, and remediation and dealing with the mitigants. That is about having the right frameworks.



The second thing is about skills. In the Bank, we know that the profile of supervisors and regulators has changed in recent years. When I started in the Bank, I started in supervision, and it was dominated by credit risk and market risk, or people with conventional trading experience. Operational risk and knowing about the technology are now increasingly big parts of the skillset, and we have been making that change over a number of years. If you look at the profile of who works in our supervision and regulatory areas, it has changed quite a lot.

One of the things we are very focused on, though, is the referee-on-the-pitch type of idea. We are not going to train every person to be able to interrogate an AI model—that is just not realistic, and nor would it be appropriate—but we need to make sure that we have enough skills within the organisation to be able properly to scrutinise and to have a proper conversation about the framework and about the critical modelling choices, the critical data choices. That is where we will need to continue to look at the skills we have available to us. Realistically, every organisation in the world is going to say—even the AI developers themselves are going to say it—“We need to stay on top of skills.”

Q145 Chair: That is certainly an issue across the board. Obviously, there has been a lot of discussion in Government, by everyone from the Prime Minister down. The Prime Minister has talked about the UK moving fast “to win the global race” on AI, and the Government have set a secondary objective for the FCA of growth. Is there pressure on you as regulators to have a light-touch approach, so that firms can embrace this, given the very big drive by Government? Do you feel any pressures, Ms Rusu?

Jessica Rusu: At the start of the year, we set out that we would take a balanced approach to regulating AI, and that we—

Chair: In your letter to the Chancellor.

Jessica Rusu: Yes. We agreed that we would take the approach that I have outlined of collaboration with firms, understanding the risks and adjusting. If we reflect on other jurisdictions that have taken other approaches, we could have written a rule that we would now be reflecting on changing—by the time the models are implemented and adjusted, things look different. For example, high-risk use cases were identified, potentially, as credit, where we might say that no way would we want an AI algorithm to be involved in AI credit decisions. We have, however, seen many compelling examples where customers have better outcomes through the use of AI in such situations. That is why we are taking a proactive approach in understanding the opportunities, as well as the risks, in what we have done. I would just add that, in terms of the skillsets across the FCA, we have invested tremendously. Every single colleague has been on training and has had AI education.

Q146 Chair: So you do not see a tension between your primary and secondary objectives, regulation and growth.

Jessica Rusu: I would say that we are always quite balanced between financial crime, consumer outcomes and market stability. Each of those



aspects reinforces each other. On our secondary competitiveness objective and innovation focus, so much has been happening this year. One of the other hats that I wear is my responsibility for innovation and innovation services, and I have never seen such incredible demand for our services as I have seen this year. More than 132 firms have applied to the supercharged sandbox and 50-odd more for AI live testing. We are doing mortgage sprints and looking at synthetic data. The demand that I am seeing from firms and industry for our innovation services is higher than it has ever been.

Q147 Chair: And there is no tension about your ability to regulate as a result. Are you leaning one way, because of that and the Government's drive for growth?

David Geale: No. We are very clear that opportunities and risks come with AI. We are also very clear that people are going to come to us if they see those risks starting to emerge, and we are alive to that. I would say again, though, that for every risk there is an opportunity that almost comes with it, and it is about balancing the two. For example, the ability to give more people advice can be seen as a positive. It has to be done safely, so we have to monitor and be alive to the outcomes, but if we can give more people advice, that is a positive for growth and we want to encourage it, because it matches our objectives and improves people's lives.

Q148 Chair: What about the Bank's perspective? Do you feel pressured to lean in the innovative direction?

Jonathan Hall: From my perspective, my behaviour is governed by the mandate, and obviously that has not changed. Within that mandate, there is a balance between the primary and secondary mandates, or obligations, and between efficiency and resilience. I think here, the easy way to say it is, "Safe adoption and safe innovation," and those two capture both. For example, we were talking about the risks, and we can adopt technology pretty fast if the risks are low—back-office things, such as some kind of assistant—but if there is high materiality, it should be the case and is the case that even the firms themselves would adopt the technology more carefully, more slowly and with more testing. Balancing the gains and the potential risk is not in tension with our mandate.

Q149 Chair: Regulators are established but ultimately Government set the rules and Parliament passes the legislation. Ms Rusu, you were quite emphatic in a speech about the FCA's strategic priorities in July that you have seen our policymaking processes, and how could we possibly write rules that keep up with the speed with which AI is changing? That is a fair point. Mr Hall, the FPC can make formal recommendations to the Treasury for legislative change. Ms Rusu probably does not think so based on what she has said, but do you think that the FPC is likely to recommend new powers for the Bank to regulate AI?

Jonathan Hall: The FPC is looking for systemic risk. That could come through one of three main areas. One is the systemic institutions. Are they, for example, building up significant risk in terms of their capital



HOUSE OF COMMONS

allocation that will put them, and therefore the services that they provide to households and businesses, at risk? That is one potential systemic risk. Another, about which I wrote more in my speech, is systemic markets—if we thought that the gilt market was particularly at risk due to herding or increased interlinkages as a function of AI, for example. The third is systemic services like payment services. There, the risk is often linked to either third-party providers and operational risks or cyber-attacks.

If any of those risks build to the extent that we think they are systemic, obviously the first line of defence is the firms themselves and the micro-regulators, but then we would potentially act on recommendations. At the moment, all those things are hypothetical future worlds. We are watching them. Our job at the moment is to monitor whether risk is growing as we think it might do, as well as whether we are wrong and there is a new risk that we have not thought about or something is turning out differently from how we expected. Then, we will respond. Obviously, it is too early at the moment; we do not think that there is a systemic risk from the AI technology that has already been adopted, so of course we have not made any recommendations yet.

Q150 Chair: Ms Rusu, you make a fair point: our legislative processes move slowly, and once law is in primary legislation, it takes a very long time to change. However, there may well be a need, so do you want to add anything to what you said in July?

Jessica Rusu: I spent many years of my career in risk and risk management roles, and I think it is always very important to take a proactive approach. That is the one that we have taken—it is assertive—in speaking to firms about the opportunities and the risks, understanding and going on a journey together. Our collaboration domestically and internationally with other stakeholders is our best opportunity to mitigate the risks that we face.

Q151 Chair: You cannot foresee any need for Government and Parliament to set some limits around this?

Jessica Rusu: I would not forecast that we will never reach a place where we require regulation. In the approach that we have taken thus far in the four years that I have been at the FCA, AI has gone from being a topic that no one wanted to hear me talk about when I wanted to talk about it—

Chair: Everyone is talking about it now!

Jessica Rusu: Now it is the topic that I get asked about every day. In trying to look ahead and think about all the things that we might need to be on top of—whether it is quantum or whatever the next technology might be—we need to be proactive.

Chair: Thank you. We have a lot of other ground to cover. I will ask John Grady MP to come in.

Q152 John Grady: This is all very exciting but, starting at first principles, the way Sandra Wachter—a professor of technology and regulation at the



HOUSE OF COMMONS

University of Oxford—has described the problem is that not even those who build the system fully know what goes on inside the black box. With a quick yes or no, perhaps working along the table from Mr Hall to Mr Geale, would you agree with that basic premise?

Jonathan Hall: Can I not give a one-word answer? There is a real question here: what do we mean by “understand”? On the one hand, no one understands fully what is going on inside those models. We have the answer that 46% of people say they have a partial understanding, so how can they be implementing that technology? Well, their understanding is great enough for them to implement the technology.

How do you bring those two together? By the risks. If it is a material risk, your level of understanding has to be almost full; if it is a relatively small risk, it doesn't. That is how we can merge that, but of course no one will be able to fully understand, and so the question is how you bound and test the model, how you check for surprises, how you look at the model versus a simplified version of itself, how you interpret it with features, and whether it has a kill switch. All those things can mitigate risk.

Q153 **John Grady:** On the risks piece, is it that if it is a low-risk deployment, you don't need to understand as much, whereas if you are betting a bank's balance sheet on something, you jolly well do need to understand? Is that what you are getting at?

Jonathan Hall: Absolutely. AI as a research assistant is much less risky than AI as a trading agent.

Q154 **John Grady:** Mr Mutton, do you have anything to add?

Tom Mutton: It depends on what the model is and how it is used. The key thing is understanding when it is changed, and that is one of the most difficult pieces in here.

Q155 **John Grady:** How would you recommend that we go about controlling for when the model is changed?

Tom Mutton: There is clearly quite a significant volume of work at the moment, particularly in the AI industry, around documenting and tracing the life cycles of AI models, and understanding the parameters under which they work, how they have been trained and the circumstances in which they were intended to be deployed. The concept of life cycles, traceability, and having transparency in communication through things like systems cards about how those models operate and when they will change, is really important. There are growing bodies of techniques that can explain the model's performance. Those techniques on their own may not be adequate to fully understand it, but they take you a good way towards understanding how the model works. Taken together, those things can give a degree of comfort, but in certain critical applications there will need to be further safeguards like kill switches or other things.

Q156 **John Grady:** What I take from what you said there, Mr Mutton, is that there is work ongoing to understand these issues. It necessarily follows that that work is not complete. Are you worried that things are being

deployed with insufficient understanding in the industry?

Tom Mutton: It is the nature of science and innovation that things will change all the time, and that is often desirable. Nobody will have perfect understanding all the time; that would probably not be a good outcome. What we see from the financial services industry is that there is quite a high level of caution and prudence around how they are using it. The key thing for us to understand is whether it starts to move into more critical markets—more material deployments—without the commensurate increase in understanding and safeguards. We are engaging very actively with financial services institutions, and speaking to academics and technology vendors, to understand that. That will be a real focus of the work we do in the coming years.

Q157 **John Grady:** How are you going about understanding that? Just talk me through that in simple terms.

Tom Mutton: There are three things in particular. The first is that we supervise a set of systemic financial institutions. That means that we can have frequent contact with them and talk to them about what they are doing. We can also use things like our survey, which gives us quite detailed insights, and we talk to technology vendors, so we have that market intelligence.

Secondly, we have bodies such as our Artificial Intelligence Consortium, jointly with the Financial Conduct Authority, which is us working with leading figures to really understand what is going on and what we need to know about it. The third thing is that increasingly using artificial intelligence ourselves gives us that applied knowledge.

Q158 **John Grady:** Obviously, you have put a bit of work into understanding it, but are you content that all the institutions that you regulate have people who have sufficient understanding of this, given what you said about people working on understanding as they implement it?

Tom Mutton: I don't think we can ever say with certainty that every institution understands what is going on. In fact, the institutions themselves—I think 34% of them—said that they don't have a full understanding of what is going on in these models. To some extent, that is always going to be the case when there are third-party models. No, we are never going to be able to say that everybody has full understanding. The key things are to make sure there are very clear expectations, clear accountability, and alignment of incentives. At the moment, I think the incentives are aligned. We want people to innovate, but do so responsibly. Businesses are very aware of the risks, for their reputations and ultimately for their P&L, that come from flawed deployments of artificial intelligence.

Q159 **John Grady:** They can be aware of those risks, but help me out a little bit: do you think businesses actually understand this to a sufficient degree? When you have discussions with the financial services industry, are you confident that the people you are talking to have the requisite level of understanding of AI—the pros and cons, risks and rewards? It is quite subjective, but that is what I want to know.



HOUSE OF COMMONS

Tom Mutton: Right now, what is quite clear is that financial institutions have access to some highly accomplished people who are working on this, because there is a big opportunity around artificial intelligence. We are not going to be able to say with confidence that they will always have the people they need, but right now, when you take the fact that the workloads are on the less material end and the fact that there are very authoritative people in institutions, I think there are reasons to feel that this is something they can stay on top of; but we are never going to be able to say with confidence that at all times they understand everything that is going on in this.

Q160 **John Grady:** This may be a question for the FCA—Ms Rusu or Mr Geale—about the management structures around this to make sure that senior management properly understand the risks. Are you minded to change what you require and have a new senior management function dedicated to AI or something like that?

David Geale: First, “I did not understand it” is not a defence, because you should understand what you are deploying, or you should understand what you are seeking to achieve. I go back to the points about designing products and services to meet a need, having clear outcomes in mind and being able to measure and demonstrate those. I am not sure that needs a new senior manager function—I think it would be captured under the framework that is already there.

Going back to your original question about what happens in a black box if it evolves and people don’t know what to expect coming out of it, we have taken action on that before, not necessarily around AI, but around algorithms being used. In the high-cost, short-term credit market, for example, we have taken enforcement action where the firm did not adequately understand what was happening in the black box and the outcomes that were coming out were not sufficient. So we can and have taken action in those sort of circumstances before, with a combination of the senior managers regime, consumer duty, systems and controls rules and indeed our threshold conditions for firms. I think we have got the tools.

Q161 **John Grady:** That is algorithms, but moving on to AI, are any of you aware of action being taken where you have intervened with a firm and said, “Look, you do not have the right management structures in place,” or, “We have had these discussions with you—you are clearly not across the detail of how this works, and this carries too much risk”?

David Geale: Specific to AI, no, I am not aware of any specific conversations on that front, beyond the fact that we do have conversations with the larger firms. We would seek to interrogate before they go ahead. If they are talking to us about the AI applications they are considering, we seek to understand that they have the controls in place. I would refer back to Tom’s earlier point. What we are seeing is that firms adopt this in a cautious manner, I think because of the responsibility and accountability that goes with the senior managers regime. The buck stops with an individual.

Q162 John Grady: Are you confident that the people you are coming across in the industry fully understand the risks of this?

David Geale: I think I would say the same as Tom, but probably on a wider scale, which is we regulate around 50,000 firms across all sectors and therefore I can never take confidence that everybody fully understands it. What I would say is that the approaches we are seeing at the moment are cautious.

Q163 John Grady: This question is for Mr Geale and Ms Rusu. One of the big things people ask about is transparency and explainability in decisions, which is an aspect of the consumer duty. How would you go about assessing whether a firm's AI decision making had sufficient transparency and explainability?

Jessica Rusu: I think it would depend on what the use case was. For example, I have worked in previous jobs in credit decisioning where you can employ an AI model as you would any other statistical model to make a decision. You would need to understand and record what the variables were that were most predictive of that decision. Then you would need to look at your pricing framework and your decision-based outcome to make sure that you are confident that whatever the model is—whether it is a traditional statistical model or any other AI model—you understand that the decisions you are making are in line with your expectations around consumer behaviour.

I would say that in general we have a broad suite of supervisory tools available to us, not just looking inside black boxes. We monitor complaints. We monitor social media. We monitor the sup hub. We have multiple sources of intelligence that let us know if something that is not okay is happening in a particular firm.

Q164 John Grady: I have two follow-up questions. First, you mentioned a vast amount of demand for your sandboxes. This is pouring through your door, and that is in a way consistent with the chat we see more generally about the huge sums being invested in AI. Can the FCA keep up with this?

Jessica Rusu: I think we are and we have. We have doubled our throughput, in particular, for these cohorts because of the size of the demand. In the past four years, I have doubled the size of my team, and we are looking across the entirety of the FCA with a lot of enthusiasm in AI. So everyone is working on it, not just my division.

Q165 John Grady: Would you have an emergency switch you could pull if you could not keep up with it just because of the volume of demand? You have your day-to-day regulation regulating existing applications.

Jessica Rusu: Maybe I should say it differently. We do not expect or require every firm to go through the sandbox. It is not a requirement; it's an opportunity.

Q166 Bobby Dean: On that point, I remember speaking with fintech people before and they were saying that their experience a few years ago was that getting access to the sandbox was easier. I guess that is probably



HOUSE OF COMMONS

partly related to the increase in demand. I know it is an opportunity and not a requirement, but obviously firms would feel much safer entering the marketplace having gone through that process, because they would have got the reassurance required. Given that the Government want you also to be growth-focused and this could act as a possible constraint on growth—because firms might not want to enter the market until they get a chance to do this—do you think that you need to be resourced better by Government to be able to open up this opportunity?

Jessica Rusu: What we do on the outside of any sprint that we have done is publish the results, so we are intending to have showcase events where we play back to the industry what was tested, what the outcome was and how it worked. That gives an opportunity for others to look at those use cases. Let's say, for example, there was a really successful pilot of a consumer conversational chatbot. That would be published, and everyone would have the opportunity to look at that and think about it and whether it applied to their firm. As I said, it is an opportunity to go through the sandbox. In fact, we were cautious about accepting only the firms that we think have genuine positive impacts on the market and genuine use cases that are worth testing.

Q167 **Bobby Dean:** That is all great, and the collective learning bit in particular. My question, though, is this. If you had more resource, would more people go through it? Would we be seeing more firms entering the market with confidence?

Jessica Rusu: Not necessarily. I have kept the acceptance rates at 30%—pretty flat—over the past four years, independent of this latest growth in AI, and that is purely because the application needs to be of high quality, have an important impact for the market and do something on financial crime, consumers and market outcomes. Those are the reasons why we keep our cohort small. It is also so that the impact of them can be more successful.

Q168 **Chris Coghlan:** For my sins, I was a Deloitte auditor in New York in 2007, auditing collateralised debt obligation models for a major investment bank—one year before it all collapsed. I can assure you that I didn't have a clue what was in those models. I am pretty sure that nobody in the bank did, except possibly the maths PhDs who built them. With respect, the attitude back then—I am sure it is different now—to regulators was one of contempt: "Regulators aren't paid enough to understand what we are doing; anyway, the market is always right." So hearing what you are saying now, I have a slight sense of déjà vu. Of course I understand we need to have growth; that is vital. How do you balance those risks so that the same thing does not occur again, because it sounds quite familiar? Obviously, the issue in 2008 was that no one knew where the debt was or who owned it or the counterparty to it. That, I guess, may be part of the answer, but Mr Hall, would you like to comment?

Jonathan Hall: I was in the financial markets at that time. To me, one of the biggest issues at that time was that when you built up these securitised or collateralised debt obligations, the packaging basically was

trying to get the benefit of diversification in order to get a high rating. How do you get that? You get it by having a low correlation between the different mortgages, for example, in the securitised product. Maybe there are different parts of the US; there are different ages. So you have this correlation that gives you the diversification, which gives it a good rating. What was interesting is that that correlation is obviously something that is subject to change, so basically the system as a whole had huge risk to this correlation assumption. That correlation assumption might have been wrong—that is one point—but the more important point is that the correlation could change significantly. What happened is that the correlation went from very low to very, very high as basically everyone in the US got into the same problem at the same time. Those that lost a lot of money at that time were the ones who not only had mispriced the correlation but did not realise their sensitivity to that parameter. That is why it is a really good question.

If you think about interpretability or SHAP¹ values—trying to work out how your model works—you are getting these features out, but at that time you had this one feature that some institutions really were not focused on, and that was a huge risk. They did not understand their sensitivity to that. When we think about these models being dynamic, evolving and having parameters that are many times of relatively low importance, but in a different market environment might be of high importance, how do we ensure that people understand that they have that sort of sensitivity to this risk if the paradigm shifts? I think it is a really good question—it is really hard, obviously.

One thing that you can do in AI models is think about new types of stress-testing. You can have adversarial tests, where you actually try and use AI techniques to find out where the model might not work, and target that. That is a new thing that the market has not had to do in the past. We need to find sensitivity to parameters—some firms failed to do that then. I think it is a big risk. At the moment, the exposure to those parameters is tiny. As the risks build up, the size of exposure to those parameters will build and it will become more and more important for us to be able to locate that.

Chair: Thank you; that feeds into some further questions we will have later.

Q169 **John Glen:** Can I build on something Mr Grady was talking about around the consumer duty and the FCA's expectations of what that requires firms to do? To clarify, in terms of the transparency and explicability of their decisions on consumer outcomes, there have been no enforcement actions against firms on the basis of not meeting expectations—is that correct?

David Geale: I am not aware of any enforcement actions that have concluded. We have certainly worked with firms. It is very easy to look at enforcement as the only thing we do, but we do a lot through the supervisory lens. If you look at, for example, the work that we did around

¹ Clarification: Mr Hall was referring here to Shapley values.



HOUSE OF COMMONS

savings, we went to firms and said they were not being clear enough with people about the options that were available to them, we have seen some quite significant behavioural change there—I cannot recall the figures, but they are significant—in terms of people moving from instant access into notice accounts where they did not need to be in instant access, and closing the gap in terms of movement between base rates changing and the rates on deposit accounts changing, as well as the disparities between different types of instant access accounts. That is in the transparency space, where we have seen firms not doing enough to explain to consumers what the options are. We did not need to enforce, because we achieved the outcomes that we wanted by using supervisory tools.

Q170 John Glen: That is very important for people to understand, because this is often asked as a measure of whether you have done anything, but of course you are having a dialogue all the time to nudge people into the right places. Is that the right way of looking at it?

David Geale: I think that is exactly right—nudge may even be a bit weak in some cases. Certainly, the first thing we would do is seek to identify the problem with firms, for them to do something about it. If we saw that there was significant harm in place, we would expect them to remediate. In other cases, we would expect them to fix forward.

Q171 John Glen: I will turn to Ms Rusu and ask about the issue that you were beginning to talk about around your prior experience with credit—perhaps we could broaden this to insurance, as well. If you are a firm, you obviously want to make a profit; if you want to make a profit, you want to identify clients that you can make a profit from. But similarly, you probably want to exclude a long tail of consumers that may not be profitable to you. That is rational economic behaviour. The concern is, and some of the evidence we received shows that, that AI can build up and infer attributes at a level that goes beyond your experience, even if it was recent—you said it has all moved very quickly. How is the FCA properly evaluating the risk of consumers essentially being targeted by the assumptions that can lie within AI, and therefore being excluded from access to products?

Jessica Rusu: That is a really good question. As datasets have expanded more and more over the years, firms are also looking at open finance data and it is possible that some might web-scrape data and use all sorts of information to have a more compelling offer. There are two things that happen. First, consumers get better rates, or prices that are more accurate. Generally speaking, the more information a firm has about an individual, the better and more surgical it can be about the offering that it makes.

However, there is a positive opportunity as well. Some consumers might historically have characteristics that associate them with not having access to credit, such as living in a particular postcode that has a high crime rate. Now, they can actually demonstrate, through other information that is made available, that they have a well-paying job and they have other

good behaviours that make them a valuable opportunity for a firm to lend to.

Yes, it is a risk that some individuals could be excluded from credit or other such insurance opportunities by having more and more information, but logically it tends to the conclusion that the more information a firm has, the better it will be at offering credit or insurance.

Q172 John Glen: I understand exactly what you are saying: you can have more data, so you can be more informed, and therefore avoid some of the biases that come from glib, easy assumptions about postcodes. On the flipside of that, how do you avoid a situation where firms do not go down that route, and do not use the best of what is available, to maintain minimal access to generally problematic postcodes?

Jessica Rusu: That could be the status quo today. It could be that some firms have not leaned into data science or technology in general, notwithstanding AI. There are firms that are more advanced than others in how they use data science. We can identify that through the information that we gather, which goes to your previous questions on consumer duty. We look at outcomes-based metrics and we can see the outcomes for vulnerable consumers and other consumer groups. We do a lot of other listening to consumers, and we have supervisory intelligence that helps us identify where those things are happening.

Q173 John Glen: Can I just dig a bit further into that? I think your CEO, Nikhil Rathi, spoke last year about the issues around hyper-personalisation and tailoring products. How can the consumer, and we as a Committee and others, transparently evaluate your scrutiny of what is happening in this space? It is like proving a counterfactual, in terms of what is moving in postcodes, or ethnic minority groups or a range of vulnerable groups. Do you actually publish any data that is meaningful in this regard that we can use to say, "Ah, things have moved forward in the right direction"—or the wrong direction? Is there anything you can point us to that would be helpful to see how this works?

David Geale: We recently published a research note on insurance—I think it was last year—in which we did not find any evidence of systemic bias. For example, we would expect that groups such as younger drivers are higher-risk than some older drivers. Even then, I am not sure that it is a one for one; you can use telematics to help and interrogate that through AI models. We will be doing further work in the insurance area, looking at the interaction of AI with firms' pricing models.

Q174 John Glen: Can you be as specific as possible on when and what that will be, Mr Geale?

David Geale: It is likely to be next year. As part of our programme of work on insurance for 2026, we will be setting out our priorities—probably around February—and for the insurance sector I expect this to be among them. We will be picking up with a series of firms to go in and ask them exactly how they go about this.

Jessica Rusu: I would just add that the analytics team have done several pieces on bias research, and the mortgage bias research article was published a couple of weeks ago.

Q175 **Bobby Dean:** This might be out of bounds for you as regulators—it is possibly more a question for us—but do you ever think about whether there are ethical boundaries to what is going on or push things back to the Government to say, “You might need to think about this”? There is the ability to track more than we have ever tracked before in the insurance industry. I have used the example on the Committee before of not keeping up with your step count and your life insurance premium going up the next month. That is possible, and probably legal, at the moment. Do we want that to happen? I am not sure. As you wade your way through this new technology, are you finding ethical questions that need to be bounced back to us? How does that process work?

David Geale: There certainly will be. What we would expect is that when customers take out any form of insurance and it is linked to things like a step count, firms need to explain that to people on the way in, not halfway through. People need to be clear what they are purchasing and the terms on which they are purchasing it. You can, for example, buy insurance that gives you a benefit for having arrangements like keeping up a step count. Equally, you can buy insurance that does not take that into account. There are different things that firms should do on the way in to make clear to people what they are getting.

It is fair comment that there may be points where the data gets so strong that people from certain communities may be at risk of being excluded. That is a social policy question. We cannot force insurers to insure. We cannot force lenders to lend in particular areas. There may be some social policy questions that come up. The key thing is for us to keep on top of the research and the monitoring.

Q176 **Bobby Dean:** I guess my question was more about the process for that. You are at the front of this, and you are more likely to spot ethical dilemmas that are more philosophical. I know we have a philosopher on the panel somewhere—I think it is you, Jonathan—but they are probably more relevant for us to answer. When you spot those, do you pass them back to Government and go, “We have been confronted with this. At the moment, it is all legal and down to the market, but we think you should think about it.” Is there any process for you to feed back those kinds of things?

David Geale: We have regular dialogue with different parts of Government, including the Treasury as our sponsoring Department. We also have a perimeter report. If we find risks emerging that are outside our remit and that we feel are worthy of being raised, we will raise them.

Chair: We had the insurance companies in, so there is some further work for us to do in that area for sure.

Q177 **John Grady:** This is a question for Mr Hall and Mr Mutton. Is AI-driven trading, in particular autonomous AI market trading, being used? Is it a



HOUSE OF COMMONS

concern? At what point would you intervene in regulating it?

Jonathan Hall: Quantitative and high-frequency market making will probably be one of the areas where we see cutting-edge technology implemented fastest. It is happening right now, so it is very early days. The way I think about it is that there are two risks.

The first risk is that it goes really well and is incredibly successful. We talked before about the financial crisis, but this would be more like an LTCM 1998 scenario where you have super-intelligent traders who have a really good model. They know where fair value is and they are incredibly successful. If that happens, you will find market share going more and more to these models. You will see a very fast transition so that most people are using these models and those who are not using them will stop trading and stop making money. The effect will be that the market becomes more and more efficient, but there is more and more dominance of one kind of model.

The issue with that is that the efficiency eventually comes with a brittleness. Then you can have some kind of shock or some shift in paradigm. In LTCM, it was the Asian crisis and then Russia defaulting. Their model did not become wrong; it is just that their risk was far too great, and a divergence away from fair value meant that they all blew up spectacularly.

Another example of that would be flash crashes, just as you see in the market on a short-term basis, where everything is going well, then something comes in from left field and they all get out at the same time. It is very sudden. That is not systemic, but imagine that at a much larger scale. That is one risk. How could we monitor that? That is really about concentration in the market and diversity of the financial ecosystem. That is quite hard for us to monitor, but we really have to monitor it.

The other risk, which is much more like we discussed before, is that these firms start to implement something that looks like it is good, but it has some flaw. Then you have a shift in paradigm, and that flaw emerges. That is where you need to interrogate the model, look at the underlying parameters of that model and have off switches and so on that can sort of protect you in that. That is more of a model and a responsibility on the managers of those models, whereas the other one is more on the kind of financial ecosystem as a whole, which is a little bit harder, but we absolutely have to monitor.

Q178 **John Grady:** Have you done any contingency planning for an AI-driven market crash?

Jonathan Hall: I do not know if we want to talk about a stress test, but we have done a system-wide stress test. It was not on AI, but that kind of stress test is something that, in the future, we could definitely apply an AI-driven scenario to. It was the first one we did. It was exploratory and used the market as a whole. That is definitely an area where, as these risks evolve and we are able to think of a scenario, we could apply that. That could be pretty valuable.

Q179 Chris Coghlan: I guess, from what I am hearing, that we are doomed anyway, right? As in, there is not a lot you can do about this. You can monitor these risks, but efficient markets mean it is going to happen anyway. As long as credit markets are not being frozen, like what happened in 2008, a crash is a good thing over the long run—that is capitalism, right? Or am I wrong?

Jonathan Hall: At the FPC, we aim to build resilience. For example, if you look at our recent reports, at a global markets level, there are a lot of risks out there. Then you can look at the UK economy—UK households, corporates and banks are actually pretty resilient. In that instance, you can have a shock, but it would not be disastrous like it was in 2007 and 2008. When I am talking about these scenarios, to some extent, it is a thought experiment—it is what might happen—but that informs what we are looking at, what we are worried about and the way we are thinking about future regulation. Hopefully, by coming up with these scenarios, we are able to plan and be a little bit more frontfooted than if we were just saying, “This is complicated.”

Q180 Bobby Dean: I want to move on to the third-party dependency element of this. There is a statistic in front of me here: 90% of UK banks rely on AI from Google, Microsoft and Amazon. I find that slightly terrifying. They are all US tech companies. Just imagine if we had a US President who was willing to throw his economic power around—it would be problematic. The Bank has already highlighted this as a key risk to UK financial services. How big of a worry should this be? We heard from other firms that it is not so much a worry—we are used to dealing with third-party dependencies and we have a way of managing it—but given the political activity of the US President over the last year, are there additional risks to the UK being so heavily reliant on US tech companies?

Tom Mutton: The attention on critical third parties is quite right. It does not really matter where they are geographically located; it is their importance to the UK economy and UK financial stability. The key thing is how important they are, not where they are geographically located. This is not exclusively about artificial intelligence. This conversation started with cloud computing. One thing we have tried to emphasise in a lot of our work is that it is actually quite hard to unbundle this. You cannot really unbundle the AI model from the cloud provision—from the data provision—so you have to look at it as a technology stack. Frankly, they tend to come as a package. That emphasises and underscores the importance of these third-party dependencies.

We have paid really close attention to this since 2017. As I said, that started with the cloud. In 2021, the FPC made a conclusion that this was a material issue for UK financial stability. We asked for a critical third parties regime; we got it, and it is something that, together with the FCA, we can use. That has been in place since 2024.

Q181 Bobby Dean: Can I just clarify whether that is on just the cloud element or all the elements?

Tom Mutton: It will be on any critical third party. Ultimately, designating those is an HM Treasury decision. We can provide information, but it is Treasury that makes the decision.

Q182 **Chair:** When you say the Treasury makes the decision, is that just a normal Treasury decision, or does it have to pass legislation or guidance?

Tom Mutton: It is a Treasury decision. It has published its approach, and that is a matter of public record. As I said, we asked for that regime, we got it, it is important and I think we intend to make use of it, but we need to work through what we may want to apply it to and how we will use it in practice. That is some very important work. That is a very important way of speaking to this issue of third-party dependencies, which are critical. Coming back to your original question, it is less about the geography and more about how critical it is to UK financial stability and the UK economy.

Q183 **Bobby Dean:** I guess you are saying that it is very critical, and I respect why you are making that distinction. I guess the reason why I highlight it is not particularly because I have concerns about the US, but because of the way politics can get involved in economic power. A less predictable political environment may increase vulnerability.

That then makes me think about competition regulation. Like you said, this is an issue that started off with cloud computing, and there are other services that these big tech companies provide. If there are any UK-developed industries, as soon as they get to the point where lots of financial services firms might want to adopt them, they tend to be bought by one of the big US tech companies. Is competition regulation a factor in the UK being able to build some of its own resilience?

Tom Mutton: Jessica or David may want to come in, given that the Financial Conduct Authority has competition responsibilities. From the Bank's perspective, a diversity of providers—providing they can offer effective and resilient services—is a good thing. These markets are also prone to network effects and winner-takes-all dynamics, so they are always going to be concentrated markets. That is why we have the critical third parties regime, and that is why it was so important to us.

Clearly there is a very important role for collaboration with other types of regulators, including the competition authorities. These services—cloud, data and artificial intelligence—are not developed exclusively for financial services. That is why things like the Digital Regulation Co-operation Forum, which the FCA chairs, are so important. It is also about those collaborations, because the third parties are not providing services exclusively to financial institutions.

Jessica Rusu: Tom put it very eloquently, but in addition to those points, there has been a lot of focus on UK sovereign AI. I have participated in discussions at No. 10, and with Nvidia in particular, about how we go about creating a sovereign UK-based AI sandbox, which would help not just financial services but all sectors of the economy, and give us the ability to control more of our own destiny in terms of the entire tech stack. As Tom was saying, it is not just about the software—AI is essentially a

software package—but about cloud, network, infrastructure, telecom and energy. There are many considerations in owning our destiny for technology.

Q184 **John Glen:** Just to clarify, I think that Amazon Web Services, Microsoft Azure and Google Cloud combined have about 60% to 65% of the cloud market. Notwithstanding what you have just said, given those market dynamics how realistic is it that we will ever be in a position of having a meaningful, home-grown alternative, given the embedded relationships you have described in the stack provision that these firms provide?

Jessica Rusu: I do not know that I have an answer to that.

Q185 **John Glen:** It is a good line for all politicians to take, but in reality, if you are a large corporate, a stack of services is provided by one of those firms, and you are embedded in and have a long-term contractual relationship, the notion that somehow you are going to easily move to something else that has resilience is surely a big ask.

Jessica Rusu: Putting my CIO hat on for a minute, when you are making a decision about your tech infrastructure front to back, you are thinking about resiliency. You are thinking, “If this one goes down, which one do I have in the back that can pick up services?” When you are making decisions about cloud providers or critical infrastructure pieces, you are spreading things around. I am sure that David can talk more eloquently about those things in terms of the operational resiliency requirements that we put on firms. Then, when you are making a procurement decision, the procurement process itself will allow you to consider aspects of concentration risk, the relative size of that provider, and whether or not they have resiliency considerations that should come into effect.

Being a CIO requires you to think about all of those things at the same time, so that the services you provide are reliable. And that means that you have to work with more than just one. It might be cost-effective to put all your eggs in one basket, but I can tell you from experience that it is better to have a diversified portfolio.

Jonathan Hall: I have some statistics and then some points about operational resilience. We ran these surveys about AI in general, as you know, in 2022 and 2024. Regarding the top three providers, for cloud that number is slightly reduced, but it is still almost 75%, so it is extremely high.

Q186 **Chair:** For the top three?

Jonathan Hall: For the top three combined. Model use has increased—almost doubled—but is now 45%. And for data, it has increased marginally; let’s say that it is 30%. These three things are super-important for us in terms of the concentration.

We have this test called cyber and operational risk stress test. What it does is that it assumes that one of these things goes down and, if so, how are you able to recover? There are two aspects to that. One is your own internal processes. As Jessica was saying, “What do you roll over to? What

is your back-up? Is your back-up operationally ready? Is it able to handle the volumes?" If it is pen and paper, you are not going to be able to do it. There, maybe you have your second provider ready for you.

The other thing that is very important from our perspective is that firms often take decisions based on what is privately optimal rather than what is optimal for the system as a whole. By doing these tests, we bring them together, in terms of how fast you can recover, how fast you can reconnect and all of those things.

I think we are thinking about it in the right way, but obviously we are watching these numbers pretty closely and I imagine that we will see them go up.

Q187 Bobby Dean: I will move on from my fears of global cyber-wars, which I am sending down our industry. Let's assume that we are in a comfortable place with the critical third parties regime and we want to engage with these companies on a particular regulatory need for the UK. I know, David Geale, that we spoke before about Meta's responses and their responsiveness to dealing with financial crime, for instance. How much engagement have you had with some of the larger AI firms on this developing regulation? Are you expecting them to be responsive to the direction that British regulation is going in?

David Geale: I have not been particularly close to that myself. I do think that the critical third parties regime is hugely important, coupled with the expectations that we have of firms to manage their own exposures, thinking about materiality and concentration, and not just size.

That said, regarding the consultations that we have done—Tom or Jessica may want to say more—what we have found is that both UK firms and the firms who may or may not be within the critical third parties regime are actually very responsive to this and want to lean in. If we think about the number of outages that we have seen, around 38% of those recently have been due to third parties. That does not do them any good any more than it does to the UK banks who are perhaps on the front end—the sharp end—with their customers—

Q188 Bobby Dean: Obviously, you have good relationships with the firms. Is the expectation then on the firms, if there are any issues with what they are using, to run that up the chain, or would you have direct engagement with some of these large AI companies saying, "This is a big problem for British regulation? Can you fix this issue?"

David Geale: In terms of the critical third parties regime versus AI, we do not know who's in that yet; let's just be clear about that one. That gives us the ability to go in with those firms that may not be directly regulated UK firms, in terms of AI.

Jessica Rusu: From a DRCF perspective, if I put that hat on—

Q189 Bobby Dean: "DRCF"?



Jessica Rusu: Digital Regulation Co-operation Forum—so thinking about the Digital Markets Act, for example, and about some of my other responsibilities around looking at online harms and how they manifest in consumer scams.

One AI angle, for example, might be that a consumer has been defrauded because they saw what they thought was Martin Lewis—deepfaked—recommending a crypto scheme. For example, we are monitoring and looking at networks of bad actors and individuals and what they are doing online. Not I, but my enforcement colleagues have conversations with the big tech firms, notifying them about bad actors and taking proactive interventions. Some of the tech firms are more responsive than others in responding to our actions.

Q190 Dame Harriett Baldwin: We have obviously had a really fascinating afternoon and a lot of subjects have been covered, but I am going to try to cover an eclectic bunch of things that have not been so far. The first question is, to what extent are your organisations using AI to reduce the cost of regulation, so that you are reducing your own costs—your own take-out? Do you see potential for that there?

Jessica Rusu: I would be delighted to talk about this for hours. We are looking at several internal use cases. One of the things that we have done horizontally across the organisation is we have rolled out Copilot to all colleagues, and everyone is getting proactive training. All of our SLT is going away for intensive off-site training. We are making sure that absolutely every employee within the FCA is educated on technology, and AI in particular.

In terms of our use cases, we have had some very promising results. I will list a few of them. On sanctions screening, our social media, which I mentioned a few moments ago, looking at the relationships between individuals and entities, is identifying bad actors faster through network analytics and curation of intelligence. Synthetic data generation for anti-money laundering is a very prominent use case. In the supervision hub we have call transcription that helps us to curate that information as it comes in and turn it into intelligence. We have implemented in the supervision hub, in the inbound calls, automatic voice recognition, so that the person calling has the opportunity to be routed to the correct individual at the first opportunity without having to have possibly multiple handoffs.

We are also using it in document review. As you can imagine, we have lots of unstructured information—documents, information that is sent for supervisions or authorisation casework—and we have seen incredibly promising results from the use cases that help us to curate that information much faster. In enforcement outcomes, we are talking about going from its maybe taking months to research something, to being able to identify it in hours. We have been really proactive, and I think you will have seen the FCA being more proactive recently in many of its cases.

Q191 Dame Harriett Baldwin: You have made a start and it sounds like there are more use cases to come, and more training. And at the Bank, Mr

Mutton?

Tom Mutton: Likewise, it is a huge area of opportunity for us, both in terms of efficiency and effectiveness. We deal with huge amounts of information, and the ability to use AI to start to look at unstructured information is really helpful, because that is a very good source for supervisors, those doing market surveillance. We have an AI strategy—we have published that. We are very encouraging of our colleagues thinking about innovative ways to use artificial intelligence. We put safeguards around that in terms of ethical responsibilities, risk management, and we have an oversight committee that does that.

One example that I think is interesting is the consultation paper on the digital pound. We had 50,000 responses. It would have been impossible for my team to read every one of those. So we thought about it really carefully and had a discussion with HM Treasury, and decided that we would use some artificial intelligence tools to help us with that process, and we used natural language processing to help review those. We then had a secondary human review of a sample to ensure that we were satisfied with what had been concluded.

We not only used that to make that review more feasible, because 50,000 is a huge amount to review, but we saw some positive cases around bringing in a level of objectivity in how we reviewed responses. Naturally human readers, particularly where they may have written the consultation themselves, will obviously see to some extent what they want to see in that. We saw a really positive opportunity not just for efficiency in that process, but also effectiveness in terms of improving objectivity. Analytically, it gave us some really interesting insights into how people felt about the digital pound.

That was one of the first times that a regulator or a public body had used artificial intelligence in that way. We got quite considerable external assurance, including getting an external provider to review our model and a KC legal opinion. That is something we are very proud of: using artificial intelligence in a quite innovative, but very responsible way.

Q192 **Dame Harriett Baldwin:** Thank you. Changing subject, the Governor and the FPC recently made a statement about the valuation of AI stocks and the risk of a severe correction. Does that outweigh any of the other risks we have talked about today?

Jonathan Hall: It is interesting that we spend a lot of time talking about long-term hypothetical risks from AI use, but the imminent potential risk is from valuations. Obviously, not only are equity markets pretty highly valued, but credit spreads are also extremely tight at the moment. There seems to be somewhat of a disconnect between valuations and market exuberance and some of the challenges in the world—whether geopolitical or economic.

That said, there is something behind the rise in AI stocks, which is a very significant increase in those companies' revenues. To some extent the forward P/E basis explains the valuations, rather than it being an increase



HOUSE OF COMMONS

in their multiples. We brought that together by saying that we are not trying to time the market, but that valuations are very high and are based on that continued increase in revenue, which is impressive, but expected to continue.

Therefore, what is the risk? What we care about is the downside risk, which is that those high expectations are not met. If they are not, then not only do you have the valuations based on those high expectations, but you also have incredibly high concentration of exposure to that within the market, both within the US and globally. That applies whether it is actual AI stocks, energy producers or chip providers. We are not trying to time the market; we are just trying to say that there is a risk out there. If, for example, optimism faded on the imminent benefits, then the markets are pricing in a pretty optimistic scenario.

Q193 Dame Harriett Baldwin: Thank you. Obviously, one of the risks we touched on is the cyber risk and the stress tests for industry participants. Clearly, one of the biggest worries is that bad actors get hold of some of the data, particularly now we are using so much more facial recognition and there are those databases and things like that. On the risk of the bad guys getting into the system, how much does AI increase that danger? It must be a huge amount.

Jessica Rusu: AI has opportunities for good and bad actors. It increases the velocity, volume and the sophistication of cyber-attacks. We have already seen examples of deepfakes. For example, phishing attacks and sophisticated social engineering are still among the most prominent ways that even unsophisticated cyber-attackers work. AI gives them an additional tool to work with.

That is why it is so important that firms continue to do their desktop exercises, pen testing and invest in their cyber-security infrastructure. Cyber-security technology has been using machine learning as a technique for a long time. The increases that I see in cyber-security software and the patches that are released are incremental in nature, so I would not say it is a massive step change in terms of the technology I am seeing right now in cyber-security.

Q194 Dame Harriett Baldwin: But the risk of the Bank of England, for example, being hacked and undermining our financial confidence as a country is so significant, Mr Mutton, that it must keep you up at night.

Tom Mutton: Fortunately, I am not responsible for our cyber-security, but it definitely will keep our chief technology officer and our head of cyber-security up at night. Yes, it is a huge issue, and we have to make sure that we collaborate with all the parts of Government who are able to help us in this to make sure that—

Q195 Dame Harriett Baldwin: Is your back-up plan pen and paper, as the National Cyber Security Centre recommended this week?

Tom Mutton: I am not the best person to talk about our back-up plan, but I can tell you that we take it incredibly seriously. We know that people



HOUSE OF COMMONS

test us. It is a lived experience for us to make sure that we are secure. We work with all the agencies in government who can help us with that. I am sure we can provide you with some further information about what we do there. It is a huge issue. We work with financial institutions around this because they are also subject to similar issues. As I said at the start of this session, I think non-financial risk, particularly operational risk and cyber-risk, is becoming every bit as important as conventional financial risks.

Q196 Dame Harriett Baldwin: That does not sound very reassuring. May I ask one final question on a slightly different topic? With protectionism and nationalism having been very prominent recently, to what extent does data fall into the area where we should worry about barriers being put in place for the flow of data?

Jessica Rusu: Our national datasets are very important and are quite fragmented currently. Our tax data, NHS data and financial service data are sitting in separate silos. There is an opportunity to bring together this information using API gateways, but that is not currently done.

Q197 Dame Harriett Baldwin: Is that something we should be worried about in the context of artificial intelligence additionally being a structural problem for the UK?

Jessica Rusu: I think it is important that the UK continue to invest in technology across all sectors. That is really important. The opportunities in all sectors would outweigh the risks of not looking at it, so I think it is really important to think about.

Dame Harriett Baldwin: I am not sure whether I feel more or less worried than I did at the beginning of the session, but thank you.

Q198 Chair: There is lots of food for thought in this inquiry generally. Before we finish and say thank you to our witnesses, Mr Geale, I think this is the first time you have been in front of us since May, when the Payment Systems Regulator was rather suddenly subsumed into the FCA, or abolished, as the Government announced it. How is the team settling into the new regime? Is there any change in day-to-day activity for you as a result of the move in May?

David Geale: The Government are currently consulting on the legislation to put that into effect, so at the moment the PSR still exists and we still continue to run it as an independent subsidiary of the FCA.

Chair: Because you are twin-hatting, of course.

David Geale: I am twin-hatting, yes.

Q199 Chair: So there is some change.

David Geale: Yes. I am twin-hatting, so I sit on the FCA ExCo and am also managing director of the PSR. We have an independent board for the PSR, which Jessica is also a member of. That will continue until the legislation changes, because that is required under the legislation. The



HOUSE OF COMMONS

Government's consultation closes next week. I do not know when the legislation will come through, but I anticipate no earlier than—

Q200 **Chair:** Legislation, as Ms Rusu has highlighted, is quite slow.

David Geale: No earlier than the end of next year is what I am expecting. So we will continue to run the PSR as an entity until such time as that legislation changes. We are not sitting still, though. You are right: my role is now double-hatting, which brings together the two payments areas—payment systems and payment firms—in one place. We are seeking to use single governance where we can, but respecting the fact that we have independent boards.

Q201 **Chair:** So the boards still exist.

David Geale: The boards are still separate. In something like open banking, for example, what we have done is bring the team together so that it sits within the FCA and reports to me, and it has both PSR and FCA staff in it. We have also already taken steps, where we can, to look at functions carried out within the PSR that could be carried out within the FCA. A number of colleagues have already moved into the FCA, and the PSR is then effectively paying for services. A good example of that would be some moves we are making very shortly to bring our two comms teams together. What we will actually have is a dedicated payments comms function within the FCA.

Q202 **Chair:** The same number of people?

David Geale: The same number of people in that instance, or thereabouts. This was not necessarily about a reduction in the number of people, but streamlining the regulation.

Q203 **Chair:** I think we managed to get out of you last time that there is not going to be reduction in the number of people or a change in payroll—the same people are going to do more or less the same job in the FCA. Is the headcount the same at the moment?

David Geale: All PSR staff are FCA staff anyway, in terms of contract. What we are seeking to do is, wherever possible, move people across to do something very similar. That is not always possible, and it is more complex in some areas, as the FCA runs different structures. In some areas in the FCA, a function that is centralised in the PSR might be devolved on a hub-and-spoke model, and that is the sort of thing that we have to work through. Some people may end up doing something slightly differently or, by virtue of the size and scale of the organisation, one person in the PSR may be doing the equivalent of three jobs in the FCA. That is the sort of thing that we work through with individuals. There are some roles that are harder to—

Q204 **Chair:** Have you started that work with individuals yet, or are you waiting for the legislation?

David Geale: We have absolutely started that work, and that is why I say that we have made a number of changes already, mainly in the operational functions. We are looking at how close we can get to an end



HOUSE OF COMMONS

state as quickly as we can, both from an efficiency perspective and to provide clarity for staff.

Q205 Chair: One of the reasons, ostensibly, for doing this was efficiency and streamlining regulation, both of which are important things. We had a conversation about how much money would be saved by doing this. At some point, that will appear in the FCA accounts—presumably at the end of this financial year—or will it be after the legislation is introduced that we will be able to winkle out what the savings have been?

David Geale: I think it will probably be when the legislation comes. At the moment, we still do not have final certainty on what that will look like. It is intended that the responsibilities and work of the PSR transfer to the FCA, so we would continue to finish what we started. What we will do is bring things like our horizon scanning and prioritisation together, so that we should have a single version of the truth on payments from the FCA and PSR, as well as a single set of data that will enable us to prioritise accordingly. Firms will already see an element of rationalisation for themselves in that we can do more joint meetings, for example, so they would have one meeting rather than two.

Q206 Chair: So there is some saving in time there for firms.

David Geale: There will be some rationalisation for firms. We will obviously look very closely at the PSR budget, and we will look to spend it wisely. We have an eye on doing things as efficiently as possible.

Chair: I would hope that a good civil servant would say that to this Committee. Thank you very much, and we will obviously keep an eye on that. I thank our witnesses, Jonathan Hall and Tom Mutton from the Bank of England, and Jessica Rusu and David Geale from the Financial Conduct Authority, for their time. We have another session with the Economic Secretary to the Treasury on 4 November, for anyone excited about this. That will conclude our evidence sessions, unless she says something extraordinary that we need to pursue, and we will then produce a report after that in due course. Thank you very much.