

Digital Regulation Co-operation Forum
Riverside House
2a Southwark Bridge Road
London
SE1 9HA

Which? response to the DRCF's call for input on *Consumer interest and AI: Regulatory, policymakers, industry & consumers tools*.

Submission date: 04/09/2026

Response Summary

Introduction

We are grateful for the opportunity to input to the DRCF's work on AI at this early stage, and look forward to continued engagement on the critical consumer protection elements entailed.

Policy work on AI at Which? has initially focussed on identifying the implications, opportunities and risks to consumers from a range of existing and emerging AI use cases, and their underlying technologies, in specific consumer markets, namely **retail financial services** and **general retail/e-commerce**. We believe that developing a deep understanding of the challenges posed by AI to specific consumer protection frameworks will help inform and design mitigations for those gaps in the future.

We are sharing our reflections on the questions posed in the Call for Input (CfI), alongside 'supplementary evidence' containing our overall findings in these two markets. These findings illustrate how those markets' unique characteristics, regulations and governance frameworks raise different challenges for regulating AI-based services from a consumer perspective, which make one-size-fits-all responses to questions around AI regulation extremely challenging.

The findings in this response and the supplementary evidence are based on analysis of only these two specific sectors. While there could potentially be parallels drawn to other markets, there could also be significant differences. Further investigation would therefore be needed before any position was taken forward outside these sectors.

The supplementary evidence presented alongside this response explores the CfI questions within three grouped themes, from which the following positions emerge:

- **Risks and Harms** - The use of AI-based services can bring risks of harm to consumers, and they are ill-equipped to understand, manage or 'tolerate' those risks. All markets must be underpinned by strong regulation to address this.
- **Accountability and Liability** - It is essential that liability and accountability in AI-based services are clearly attributed by regulators, particularly for agentic AI. Consumers should not be expected to bear disproportionate levels of responsibility.
- **Tools and Frameworks** - In these two sectors, existing consumer protections could apply effectively to current consumer-facing AI services, with regulatory oversight, although some adaptation may be required to maximise effect. As AI services become more autonomous, this will place greater strain on tools and frameworks which may require regulatory change.

Policy Recommendations

Based on our assessment of the regulatory frameworks in the two sectors that we have examined so far, and their limitations, we have developed some key policy proposals in each of these sectors. These represent initial practical steps that could be taken towards improving consumer protection in the context of AI.

General Retail / E-commerce:

1. Regulators and the government should **acknowledge that gaps will be exposed in consumer protection frameworks** from AI-based services (notably AI agents) deployed in consumer-facing retail services, and commit to **explore ways to address these gaps**.
2. A **working group should be established by CMA/DBIST** to analyse consumer-facing AI use in retail markets with two primary goals: a). Assess AI service capabilities against existing consumer protections to ensure that they remain robust, clarified and applied sufficiently. b) assess how future agentic services may expose gaps in consumer protection. This group should be empowered to recommend a range of interventions, such as defining codes of practice or charters, clarifying key regulatory guidance, and advocating for wider regulatory reform.

Retail Financial Services:

1. The FCA should **issue updates or guidance to the Consumer Duty and SM&CR to articulate their application** in relation to deployment of AI by financial service firms more clearly.
2. The FCA should **consider how the AI Supercharged Sandbox**, and its other Innovation Hub services, **can be expanded to increase participation and/or made more effective** through sharing of outcomes and 'lessons learned'.

General observations

Before getting to our responses and evidence, we would like to share some general observations from our work in this area and on the questions posed in the Cfl.

Terminology

We appreciate that there is no general consensus in relation to describing AI and its implementations, and an array of commonly used definitions and terms exist. However, there

is considerable variation in the definition, description and usage of even key terms (agent, agentic etc) and frameworks (e.g. 'levels of autonomy') across government and regulatory outputs, and this lack of consistency creates confusion. It makes it challenging for organisations to contribute to discussions on the same topics, and it makes cross-regulatory coherence more complicated to achieve. It will be increasingly challenging for industry and consumers to interpret new guidance, for example, if this issue persists.

Harmonisation and consistent use of key terms across government and regulators going forwards is essential for coordination of regulatory activity and improving market consistency. It will also likely support better compliance, and reduce risks of consumer confusion.

Implementations

The existing regulatory guidance (along with this Cfl's examples) seems to have focused primarily on companies directly deploying AI within their own AI applications or environments - for example, Amazon's Rufus¹ or NatWest's Cora² assistants. It effectively treats all consumer-facing AI as variations of the same product class - chatbot tools. It remains an open question whether this guidance would appropriately cover other implementations we are starting to see in consumer-facing services. For example:

1. **'Intermediary' or 'cross-market' AI and agentic applications** that can operate across markets and handle different stages of the consumer journey. In this case, they both represent the consumer and interact with multiple third-party services, for example Google Shopping³. We believe such agents will put specific strain on consumer protection frameworks that needs to be considered [see supplementary evidence].
2. **AI chatbots set up as tech platforms on which other apps integrate services** to be used by the consumer, such as ChatGPT's 'plug-ins'⁴. These 'apps' work similarly to plug-ins on a web browser, integrating services like [Booking.com](https://www.booking.com) into the chatbot platform experience. They are constantly evolving in functionality and capability.
3. **Consumer level or 'personal AI'** - this would involve consumers directly deploying their own AI agent to work on their behalf, either using 'open-source' technology or via an existing technology platform. We are not aware of any work on considerations specific to these implementations by regulators to date. Given the pace of developments and plans of tech companies⁵ we believe this form of AI needs specific consideration going forwards as it raises additional consumer protection questions⁶.

This is not simply an extension of the terminology observation above; it is a recognition that consumer protection regulations designed for conversational chatbots may well be ill-equipped for divergent business models, varied use cases and particularly AI agents with increasing levels of autonomy. We discuss such examples in our supplementary evidence.

¹ www.aboutamazon.co.uk/news/retail/amazon-rufus-launch-uk-generative-ai-shopping-assistant

² www.natwest.com/support-centre/cora.html

³ www.aboutwayfair.com/category/tech-innovation/wayfair-partners-with-google-to-make-it-easier-to-shop-for-your-home

⁴ <https://chatgpt.com/features/plugins/>

⁵ www.meta.com/thefutureisforeveryone/

⁶ <https://www.bbc.co.uk/news/articles/cn0nww2qlp7o>

Regulatory Coverage

The DRCF Foresight paper⁷ helpfully bears the use case of a retailer's autonomous customer assistant powered by agentic AI, to illustrate how AI technologies can straddle multiple sectors of the economy and therefore different regulatory remits. Such a scenario could be covered by an array of existing consumer protection tools and frameworks, including in relation to potential product safety problems.

Product Safety regulations sit with the Office for Product Safety & Standards (OPSS) who, alongside DBIST, would need to consider whether cross-market AI chatbots should be designated as online marketplaces under new legislation. However, the OPSS is not part of the DRCF or any other similar forum, and so there are risks of decisions being made without full collaboration and insight from other regulators.

It is unclear how alignment with other regulators will be ensured going forwards, or whether the government's existing approach to sector-specific regulation of AI can continue to deliver for consumers. As consumer-facing use cases, particularly regarding agentic AI, continue to fall under a patchwork of different regulations, it is essential that the approach to regulating them is as joined up as possible.

Summary Responses to questions 17–28

We appreciate the request to be as specific as possible in the Cfl, but have found the task to consider all 'AI' far too broad for one engagement. This is in part because of the aforementioned challenges in relation to terminology and implementations. Therefore, we have kept our responses to reflections on the specific use cases we have considered in the two sectors we have analysed, but these do not cover all 'AI'.

Many of the risks associated with AI, AI agents and agentic AI are emergent, and therefore we have found limited empirical evidence in these contexts. Much more research is required on these topics to ensure an informed debate about the consumer implications. We look forward to collaborating and developing a clearer understanding of them going forward.

17. What tools can industry offer to consumers or otherwise implement (risk governance frameworks, impact assessments or others) to manage AI risks?

In relation to tools industry may implement, a variety of AI industry-led standards⁸, frameworks, impact assessments^{9, 10} and tools¹¹ exist across different markets. The extent to which these, typically voluntary, frameworks provide actual consumer protection is unclear. Further, many of these tools are 'pre-deployment' and used in scenarios where AI applications are already operating within certain guardrails.

⁷ <https://www.drcf.org.uk/siteassets/drcf/pdf-files/final-drcf-the-future-of-agentic-ai-foresight-paper.pdf>

⁸ <https://www.iso.org/standard/42001>

⁹ <https://www.microsoft.com/en-gb/ai/tools-practices>

¹⁰ <https://saif.google/risk-self-assessment>

¹¹ <https://deepmind.google/models/model-cards/>

Given that standards are largely voluntary and often developed closely with industry, in our view they are no substitute for clearly defined regulatory rules and expectations.

18. Should consumers be able to opt-out of AI, and if not, what is the legal basis for the use of their data?

It is unclear whether this question refers to opting out of the processing of personal data by AI applications, and/or refers to consumers opting out of AI being used entirely in the products and services they intend to use. Clarification on that point would be helpful in providing a response.

Consumers' understanding of AI and related data processing is currently low, so that could undermine any efforts to create any AI 'opt out'. Even if written in plain English and with clear warnings (which is uncommon in many existing industry opt-out implementations), consumers may struggle to comprehend the consequences of their choices. Having said that, consumers should always have the option to 'opt out' of using some forms of AI in favour of access to traditional channels, in particular when it comes to customer service chatbots versus human customer service operatives. In some circumstances the needs of consumers, particularly those with vulnerable characteristics, are better met by human assistance than an AI chatbot¹².

In addition, any opt-out mechanisms must not become liability shields. Firms must not use disclosures to abdicate statutory duties of care or suggest that consumers can "consent away" their rights to redress for negligence.

As Agentic AI evolves it is unclear how these individual choices would be implemented - such as whether an agent could opt-in/out on a consumer's behalf when interacting with third-parties. New mechanisms of safety governance that ensure accountability, regardless of the consumer's interaction model, may be needed.

19. What tools, frameworks or standards can regulators leverage to safeguard consumer protection? Can these be easily adapted or transferred on a cross-sectoral basis?

There are a number of tools and frameworks that can be employed to protect consumers in the AI context (see specific examples assessed in our supplementary evidence). We have seen limited evidence of established, existing *cross-sectoral* tools, standards or frameworks for consumer protection, although some of the frameworks used in financial services (e.g. the Consumer Duty and SM&CR) have clear structures which have the potential to be adapted for use in other sectors.

Standard setting is often slow and can be industry-dominated. Regulators should mandate meaningful consumer/civil society participation in bodies writing technical safety standards to ensure those perspectives are represented in support of these technical discussions, or they risk compliance thresholds being set by industry rather than reflecting consumer reality.

¹² FCA, Vulnerability review, 2024. [\[Link\]](#).

Overall, in our view, standards are no substitute for clearly defined formal rules and expectations.

See supplementary evidence for our full response.

20. What types / modes of guidance or other regulatory tools are most helpful to industry for achieving compliance – e.g. detailed guidance, sandboxes and others. If any effective international examples come to mind please refer to these.

While detailed guidance, sandboxes and other regulatory tools could certainly be useful, their impact will depend on the specific structure of the sector, approach taken by relevant regulator(s), and the balance between voluntary and mandated elements. Benchmarking is another tool that could also be useful but is one we have not assessed in detail.

See supplementary evidence for our full response.

21. What is the current application of AI standards, if any, and is there international best practice?

As above, we note that the development of standards for AI, particularly for agentic AI and AI agents, are still emergent. We have seen limited evidence of their effectiveness to date.

See supplementary evidence for our full response.

22. Is informed and meaningful ‘consent’ the right lever to ensure a consumer is protected and aware of risks?

It is unclear what exactly is meant by ‘consent’ here. Consent processes can be applied in a host of ways and for varying purposes all throughout the lifecycle of AI use by consumers.

While our work on consent is evolving, a general thought is that current consent approaches are likely to be poorly suited to the dynamic risks posed by AI, especially agentic AI, where models update continuously and impacts evolve. It is difficult to maintain confidence that consumers are continuously informed in such a context. Similar to opt-outs, relying primarily on consent risks creating a "tick-box" compliance culture where firms use disclosures as a liability shield to abdicate their statutory duty of care.

Given current levels of AI literacy, the technical opacity of AI and widespread "consent fatigue", expecting consumers to meaningfully evaluate the risks of complex, probabilistic AI systems is unrealistic and should not be a primary safeguard. A robust consumer protection framework must place safety obligations directly on providers, ensuring accountability is not simply waived away through opaque disclosures. Companies choose to utilise and implement AI knowing those systems are inherently probabilistic. They should not then be able to pass the associated risks on to the consumer.

To illustrate the point, the following is from a hotel chain which had recently launched an AI-powered trip planning chatbot. It emailed its customers requesting them to accept new ‘terms of service’. Customers would have had to read the full 16 page document to find the

following section which, in our view, was not described in an upfront way in the email: *“While we strive to provide accurate and useful information, AI-generated content is based on algorithms and may not always reflect the most current information or consider all variables. All information and content generated by such AI tools are provided on an “as-is” and “as-available” basis. We make no representations of any kind as to the relevancy, accuracy, or completeness and are not responsible for damages or losses arising from your use of or reliance on such content.”*

23. Are there any other levers (e.g. cooling-off periods, right to opt-out of using AI, ‘cookies’-like requests)?

We have analysed a range of consumer protection levers, such as right to withdrawal/right to cancel (which we assume ‘cooling-off periods’ refers to), protections against unfair trading practices, and routes of redress. Some such levers are potentially suitable for addressing the risks of harm to consumers posed by using AI-based services in retail and personal finance. However, that will require effective regulatory oversight, prompt enforcement and potentially adaptation of rules and guidance when and where required.

Our analysis has shown that when AI systems adopt more autonomy, current levers of consumer protection could be placed under strain. Protections could therefore require adaptation or deeper reform in order to address the novel risks introduced to consumers.

Please see supplementary evidence for our full response.

Specifically on ‘opt outs’ for AI, please see response to question 18 above. Regarding ‘cookies-like requests’, we are unclear what exactly this term refers to so would welcome clarification before responding.

24. Are there existing examples in other sectors or jurisdictions who utilise tools efficiently?

Our early analysis has focussed on two specific sectors and therefore we have not yet assessed in detail any regulatory examples in further sectors. In financial services, there is evidence of effective application of AI-specific frameworks in other jurisdictions with Singapore in particular demonstrating a forward-thinking approach in building specific rules on to their general principles-based regulatory system.

Our findings to date suggest that while lessons can be learned inter-sectorally, there are also substantial differences in how markets operate that need to be considered before any cross-sector assumptions can be drawn. However, UK regulators certainly should seek to coordinate their work to align general principles and approaches and to share lessons in order to secure more effective overall regulation across sectors.

See supplementary evidence for our full response.

25. Can outcomes-based approaches, analogous to ‘Consumer Duty’ type frameworks, where firms are expected to actively consider and deliver fair outcomes,

be a vehicle for consumer protection?

The consumer duty offers some protection against AI risks by setting relevant and appropriate standards which should increase the likelihood of AI services being of reasonable quality for consumers. However, the duty suffers from some limitations regarding AI risks, such as its non-specific approach and the lack of attached consumer redress options. These elements mean it needs to be supported by strong regulator enforcement in order to ensure effective protection, which creates potential weakness.

See supplementary evidence for our full response.

26. Who should be held accountable, and to what extent, for AI risks materialising? Consumer, firms (including integrators, developers), or regulators?

Companies should not be able to pass accountability on to consumers. AI can make outcomes unpredictable, but a firm's deployment of standard AI or agentic systems is a choice, because it creates economic or other benefits for them. Where companies have chosen to expose consumers to risk, they should be accountable for doing so.

Accountability becomes even less clear in agentic solutions where multiple actions can happen autonomously between multiple parties/services without the consumer's awareness, and sometimes across different regulator domains. This requires particular consideration.

See supplementary evidence for our full response.

27. Do 'tolerable risks' exist for consumers in these contexts, and if so, how should policymakers/regulators define and operationalise? What about industry?

Before proposing what risks consumers should tolerate, the specific risks need to be properly identified and understood. Tolerance for risks may differ depending on the sectoral and other context of the AI use case.

See supplementary evidence for our full response.

28. What metrics, thresholds or proxy indicators are feasible?

No response.

Supplementary Evidence

1. AI and General retail/e-commerce

1.1 Risks and Harms

AI chatbots

Currently AI in retail/e-commerce focuses largely on enhanced and accessible shopping discovery and research through AI chatbots, such as Google Gemini, Amazon Rufus and ChatGPT's apps/plugin platform. These chatbots will also soon expand into direct purchasing, such as Google's 'agentic checkout'¹³. From here we refer to 'chatbots' as tools that are typically in the Assistant/Operator levels of autonomy¹⁴.

Compared to standard e-commerce tools and traditional search engines, chatbots offer a more 'conversational' interface that enables personalised search and product discovery. Large Language Model (LLM)-based technology can also do more to infer user intent, presenting a curated selection of results. Users can ask follow-up questions, with 55% doing so to refine their needs and criteria, according to Which? data¹⁵.

There are potential risks of harm resulting from chatbot-based tools. These include:

Type of harm	Overview
Financial harms	AI chatbots could misrepresent or 'hallucinate' information about products, leading consumers to make purchases that don't meet their needs. Chatbots carrying advertising or sponsored links could self-preference certain services over alternative options, leading consumers to make purchases not in their best interests.
Physical harms	An AI chatbot could be misled by fake reviews or other deceptive design practices to recommend an unsafe product that later causes the consumer physical harm.
Time harms	For unwanted or incorrect purchases made using AI, consumers are faced with a potentially more complex burden of proof requirement to get redress. For example, retailers could challenge proposed returns if the reason was due to poor information and recommendations given via AI applications.
Psychological harms	With consumers increasingly relying on chatbots to handle search, discovery and even purchasing functions, it lessens their capacity to detect unfair trading practices, or 'manually' locate important information (eg, safety data).

¹³ <https://blog.google/products-and-platforms/products/shopping/google-shopping-cart/>

¹⁴ www.drcf.org.uk/siteassets/drcf/pdf-files/final-drcf-the-future-of-agentic-ai-foresight-paper.pdf

¹⁵ www.which.co.uk/policy-and-insight/article/consumer-use-and-attitudes-towards-ai-search-tools-aTnr81n3FOQI

Data harms	With chatbots increasingly offering options to turn on ‘memory’ functions to ‘improve’ performance, consumers may consent to more personal data sharing with AI companies and third-parties than they had understood or expected.
-------------------	---

Our early assessment is that the existing consumer protections (outlined in the ‘Accountability and Liability’ section) may mitigate some of the risks of harm, **although only as long as rules and expectations are defined, applied and enforced correctly by the market and regulators.**

Regulators have sometimes struggled to tackle harms created by new business models, such as the level of unsafe products in online marketplaces. AI chatbots could create new problems, as well as exacerbate existing ones (such as if a chatbot was rolled out on a big marketplace, potentially deepening existing harms). Therefore, how existing rules are clarified, applied and enforced with AI in retail services will be crucial to maintaining effective consumer protections.

Autonomous AI Agents

There have been some developments in the past year towards the introduction of **AI agents** in retail. Although definitions vary, an AI agent goes beyond simply presenting the consumer with options to buy, and can instead plan, research and execute actions, complete contracts and make transactions on the user’s behalf. For clarity, here we refer to these AI agents as tools that have greater degrees of autonomy, as Collaborator/Actors¹⁶, to differentiate them and their associated risks from chatbots above.

As of writing this, we are not aware of any consumer-facing applications or services on the UK retail market. The market is hard to predict due to high volatility, but once these agents are introduced we anticipate there is likely to be a long period in which agents are gradually introduced alongside other retail, e-commerce and AI channels.

Although evidence is nascent, AI agent risks of harm could include:

Type of harm	Overview
Financial harms	Autonomous agents may proactively buy goods that do not meet the consumer’s needs, aren’t wanted or are not in line with budget requirements. AI agents could become targeted by new forms of online behavioural manipulations that exploit the lack of human oversight in the purchasing process. New forms of scams and cybercrime could also target consumers via the agents they use. Although issues such as self-preferencing of commercial relationships and collusion between agents are in theory covered under consumer law, the fast moving, machine-led nature of autonomous agents could make detecting and identifying such activity significantly harder - both for consumers and

¹⁶ <http://www.drcf.org.uk/siteassets/drcf/pdf-files/final-drcf-the-future-of-agentic-ai-foresight-paper.pdf>

	regulators. With consumers delegating purchasing to their agent(s), they are financially harmed when those agents make purchases on their behalf that they do not actually want, don't meet their needs, are scams, or run against their best interests.
Time harms	Should redress systems not be updated to reflect the use of AI agents, this could result in a deepening complexity of how consumers return goods, make complaints or seek other forms of redress. When things go wrong, it could be unclear whether the consumer (having not made the purchase themselves) should pursue the trader who sold the item, or the agent that bought it for them. This could lead to a heavy and time-consuming burden of proof. Should AI agents lead to an increase in attempts to seek redress, the volume of unwanted purchases and chargeback requests will likely increase. Companies could seek to place friction on routes such as returns, such as levying or increasing charges or obfuscating customer service processes. Consumers then have to spend significant time, energy and stress attempting to seek redress for unwanted or incorrect purchases made on their behalf by agents. This could also result in financial harm if consumers are charged fees for returns, or are unable to return goods they don't want or need.
Physical harms	We know from extensive research the widespread problems that exist with unsafe products being sold on online marketplaces. If AI agents are not effectively trained on how to proactively monitor for product safety information and requirements, that could result in them struggling to secure oversight of safety data, or they could be manipulated by wrong information, such as fake UK CA marks. They may even 'hallucinate' information that doesn't actually exist. With the consumer out of the loop, human oversight is significantly reduced. When AI agents purchase unsafe or dangerous goods on behalf of the consumer, this could lead to physical risks to health, wellbeing or even life.
Psychological harms	As consumers become more removed from research and purchasing decisions in retail this could lead to 'cognitive loss', as consumers can no longer proactively understand the risks of shopping online, and in turn put too much trust in agents that represent them. Consumers may therefore become more susceptible to manipulative practices in their own purchase decisions, and/or fail to realise that their agent was working against their best interests until it was significantly too late. This issue could be heightened among vulnerable groups.
Data Protection	In order to provide more customised, personalised, and potentially useful, experiences for consumers, AI agent solutions could require deeper access to consumer data over traditional retail and e-commerce services. Agents could have access to sensitive financial information (such as bank accounts), health data (eg, food allergies), and other information (such as home floor plans for planning furniture choices), and autonomy over who to use that information towards various objectives. This data

	relationship may be held over a long period of time, and change according to usage. Consumers could struggle to give informed and valid consent that reflects how AI companies and the third-parties they interact with are actually using their data. Alternatively, consumers may be so concerned by that perceived risk that they disengage with AI-based services that could actually work to their benefit.
--	--

Our assessment is that a higher level of risk is posed by autonomous AI agents used in retail. However, consumers are ill-equipped to decide what risk level is ‘tolerable’ before using them. Therefore, policymakers/regulators must ensure that risk is allocated proportionately and fairly among parties at regulatory level, with consumers not overly burdened (see more below under ‘Tools and Frameworks’ section).

1.2 Accountability and liability

AI chatbots

For the position ahead, we have taken as a basis the Competition and Markets Authority (CMA) report on [agentic AI and consumers](#), backed by guidance for [Complying with consumer law when using AI agents](#). Although the reports refer to ‘agentic’ and ‘AI agents’, the CMA is broad with its definitions, and most scenarios given are, in our view, not actually autonomous agents. In our view these guidance do not clearly extend to the Collaborator/Autonomous Actor levels of autonomy, and we have assessed them in the context of chatbots (Assistant/Operator levels of autonomy).

The CMA guidance notes that companies deploying AI take full responsibility for the outcomes, even if someone else designed or supplied the AI technology¹⁷. As there is no specific AI legislation in the UK, AI use is therefore governed by existing technology-agnostic legislation, specific regulatory focus (including the provision of guidance, such as via the CMA), and the application of long-standing legal principles.

If a company deploys an AI chatbot to assist with customer-facing retail activities, then it must adhere to consumer law and protections, including (but not limited to):

- [Consumer Rights Act 2015](#) (CRA)
- [Consumer Contract Regulations 2013](#) (CCRs)
- [Digital Markets, Competition and Consumer Act 2024](#) (DMCCA)
- [Product Regulation and Metrology Act 2025](#) (PRaM)
- [Online Safety Act 2023](#) (OSA)
- [General Data Protection Regulations \(GDPR\) Data Protection Act \(DPA\) 2018](#), [Privacy in Electronic Communications Regulations 2003](#) (PECR)

The legislation above applies where a business uses an AI-based application (such as a chatbot) for a variety of retail activities, including handling customer queries, processing

¹⁷ This dovetails with the 2024 Air Canada case on chatbot liability, stemming from a Canadian tribunal ruling that the airline’s chatbot had given misleading information over bereavement fares to a grieving passenger in late 2022. The tribunal dismissed Air Canada’s claim that it was not liable because the chatbot was an ‘independent legal entity’ responsible alone for its actions. The case is viewed as pivotal in showing companies deploying AI applications are strictly liable for those applications and their actions.

refunds, recommending products, and managing marketing campaigns. This means that various harms resulting from AI chatbot use in retail should in theory be covered under existing consumer protection frameworks: For example:

- **Financial harm via hallucinations** - If an AI chatbot invents or 'hallucinates' a product purchase option or even an entire retailer that doesn't exist, and the consumer is given an option to purchase that product, then the AI supplier and/or retailer might be liable under the DMCCA or via misrepresentation as a legal 'cause of action'. The consumer will also have CCR rights to cancel the transaction depending on the type of sale.
- **Financial harm via self-preferencing** - If an AI chatbot presents products that have paid advertising or commercial relationships behind them, but without disclosing this, that could represent a breach of the DMCCA.
- **Physical harms via misinformation** - If an AI chatbot misrepresents important product information - including its availability, quantity and specification, benefits or risks, origin of the product etc. - it should be liable under the unfair commercial practices provisions of the DMCCA and/or the CRA.
- **Physical harms via fake reviews** - Should an AI chatbot present a retail site or product without flagging high levels of customer complaints or poor reviews regarding a safety issue, that could count as a misleading action (DMCCA). The AI chatbot could also be classed as a marketplace and have liability under the PRaM Act (see more below).

In addition, English contractual law may also apply. For example, under principles of agency law, contracts formed through AI applications (such as a contract to purchase an item) would be attributed to the natural or legal person who deploys them (in most instances, the consumer). So, if an AI system malfunctions or produces erroneous outputs, then legal 'causes of action' (such as mistake, negligent misstatement and/or misrepresentation) might be available to the consumer, depending on the circumstances and facts. However, this remains untested and needs to be clarified in specific legal contexts.

'Cross-market' chatbots and the Product Regulation and Metrology Act 2025 (PRaM)

The CMA's guidance covered above has, in our view, focused primarily on companies directly deploying AI applications. However, it remains an open question whether this guidance would appropriately cover AI applications that can operate on a cross-market basis (meaning they can interact with multiple third-party services).

Liability must be clearly defined when the AI acts as not just a directly deployed tool of the company or organisation, but also an intermediary facilitating shopping activity between third parties. We have seen with online marketplaces how emerging retail platforms with ill-defined liability can create new risks to consumers, such as with the proliferation of the sale of unsafe products online. Cross-market chatbots could become a new frontier of this problem.

If we examine [Google's 'Agentic Checkout'](#) - a Wayfair service in Gemini utilises Google's integration of the Universal Commerce Protocol (UCP) standard. The shopping research phase (eg, 'I need a new lamp for a spare bedroom') is conducted in the Gemini interface,

and the payment execution, too. However, Google is very clear that Wayfair remains the 'merchant of record'. This means that the payment phase, while conducted via Google Pay and Google Wallet (backed by the Agent Payments, AP2, protocol) uses the Wayfair backend and so all responsibility and liability is retained by Wayfair. Wayfair has also [confirmed](#) that it remains the 'merchant of record'.

We believe therefore that Wayfair retains responsibility to meet its existing obligations under consumer regulation requirements outlined above, including fair trading, redress, liability, etc. This applies even though the consumer is purchasing through Google's AI chatbot. However, questions remain over what liability Google should assume. The 'merchant of record' seems analogous to the set-up related to online marketplaces, where the legal contract is with the individual third-party seller who provides the item, not the marketplace itself.

The Product Regulation and Metrology Act 2025 (PRaM) would give the UK Secretary of State for Business, Innovation, Science and Trade the ability to regulate where there are gaps in product safety regulation. A definition of online marketplaces is due in secondary regulations, and this is expected to cover '*any other platform by means of which information is made available over the internet*'. In addition, there is a catch-all provision whereby product safety requirements may be imposed on '*any other person carrying out activities in relation to a product*' (section 2(3)(i)).

We have previously flagged to OPSS and DBIST that the duties and supporting requirements need to take account of how AI business models are evolving, such as chatbots. PRaM regulator OPSS is understood to view a tailored approach to different business models as the best way forward. In our view, it is vital that AI use cases as identified in this response are taken into account when designating online marketplaces, particularly when it comes to binding specific requirements for liability and accountability (eg, monitoring/requirements to remove unsafe products).

As covered in summary, there are other AI-based business and technology models, such as ChatGPT's plug-ins platform, that also potentially require specific regulatory considerations.

Accountability and liability: conclusion

Overall, based on our analysis, the near-term outlook for AI in retail appears to be focusing on chatbots that are limited in scope, heavily guardrailed and largely preserve existing e-commerce infrastructure, payment rails and liability lines. Existing consumer legislation and guidance¹⁸ should therefore be sufficient as a baseline of consumer protections for the harms posed by AI chatbots in retail.

However, this **comes with a significant caveat** - that will only work as long as the market is effectively monitored, further guidance issued where required (such as with the cross-market chatbot gap identified above) and enforcement action taken promptly when transgressions are identified.

18

<https://www.gov.uk/government/publications/complying-with-consumer-law-when-using-ai-agents/complying-with-consumer-law-when-using-ai-agents>

Although autonomous AI agents haven't yet reached the consumer-facing retail market, Which? and others (such as the [European Law Institute](#)) expect the deployment of agents to expose gaps in consumer protection, including in relation to accountability and liability. This is because the fundamental basis of consumer law assumes the human is primarily in the loop. Instead, autonomous agents will take the lead in purchasing chains, and reduce the active role of the consumer. These gaps need to be fully understood, and mitigations considered, before agents become widely used by consumers. See 'Tools and Frameworks' section for more.

1.3 Tools and Frameworks

AI agents will pose challenges to existing regulatory frameworks

We note existing regulatory statements (such as in this [DRCF report](#)) that, "although agentic AI systems operate with a degree of autonomy, this does not remove organisational responsibility for data processing and compliance with the law".

While 'human-in-the-loop' governance is often suggested, there is no clear indication as to how this will practically work when agents adopt autonomy for outcomes. It is also unclear how consumers will navigate this new layer of complex and largely unpredictable AI systems in retail without strengthened regulatory guardrails being implemented first.

A greater risk is that the fundamental basis of existing consumer protection frameworks will be challenged by the use of AI agents. Based on our assessment (in retail) to date, below is a non-exhaustive list of key challenges to existing regulatory frameworks for consumer protection we feel agents will expose.

Human-in-the-loop assumptions under consumer law will need to be reconsidered for AI agent transactions, especially when agents adopt a greater level of autonomy.

The majority of consumer law is set up under the basis that a human is 'in the loop' when a contract is formed. AI agents could change this, with the AI instead executing contracts on the consumer's behalf. In some scenarios, the consumer could only be minimally involved, if at all. This is an existential risk to the basis of various consumer protection regulations and legislation. It's a root-and-branch problem with substantial, far-reaching implications. This gap could result in financial and physical harms caused by AI agents buying products that do not meet consumer needs. It could also result in psychological harms with consumers moved out of the loop of purchasing decisions, and de-skilled in the process.

The 'average person' must be recategorised in the context of the use of AI agents.

UK law has a legal benchmark defined as an average person who is reasonably well-informed, observant, and circumspect insofar as the prominence of contract terms (section 64(4) CRA) and unfair business practices (under the DMCCA) is concerned. The DMCCA also includes a definition for vulnerable persons. AI agents could increase the capacity of the average person due to the use of a higher-spec technology for search, discovery, execution, etc. However, agents could also lead to 'cognitive loss' in consumers, by moving them away from the deliberation and consideration phases, and de-skilling them. Relatedly, for some of the unfair commercial practice provisions, the transactional decision test (a standard applied by regulators to determine if a commercial practice is legally

misleading or unfair) is used. Strain could also be seen to this test in the context of AI agents making decisions on behalf of the consumer, as the context for determining misleading or unfair practices, and their outcomes, would likely shift.

Risk allocation must be defined fairly across AI agent transaction chains, and not fall overly on the consumer.

Consumers should not be expected to tolerate any risk that goes beyond levels to which they are currently exposed. With new contractual relationships forming from the use of AI agents, it is vital that risk is spread fairly across the chain. There is a danger that consumers will walk unknowingly into a 'use at your own risk' basis, with disclaimers placed in T&Cs that the consumer may not fully read. Suppliers of cross-market, intermediary AI agents, will most likely want to shift risk onto the retailer and consumer. However, that could create an 'uneven playing field' from the outset of the market, with risk unfairly allocated across all parties in the transaction chain. We have already seen similar issues occur in online marketplaces. If risk is not fairly allocated, it could result in financial, physical and time harms as consumers are exposed to greater levels of risk than they should be. There could also be psychological harm from consumers having to adopt a higher risk position than they are comfortable with.

Professional Diligence safeguards should be considered to apply in AI agent chains, including agent supplier and end-point merchant (if different).

Under Section 229 of DMCCA, professional diligence applies when traders have "knowingly or recklessly" breached commonly-held standards in the market. Contravening the requirements of professional diligence is prohibited if the average consumer is likely to take a different decision as a result. With the use of autonomous AI agents arguably bearing an automatically high level of risk (due to the persistent danger of financial loss on behalf of the consumer), it could be considered appropriate to apply a de-facto professional diligence requirement on the retailer and AI supplier (if different), imposing a higher expectation of due diligence. However, professional diligence is a very high bar to prove legally, and successful cases based on it are rare. There are financial, physical and time based harms resulting from this, should agent services not be set up with the highest possible standards of due diligence by the suppliers, yet professional diligence not being an effective safeguard that consumers can adopt to seek redress.

AI agents might fail to meet various transparency requirements, including disclosure of working on behalf of the consumer, transparency over decision making, and clarity over commercial relationships.

AI agents might not transparently disclose that they are working on behalf of consumers when entering contracts with retailers. This might be to get around data protection concerns, or it might be if they are 'bulk buying' items for multiple users at the same time (such as with concert tickets or seasonal items). However, non-disclosure might complicate issues around liability. It would make it harder to understand if the agent had acted outside of its authority, and so complicate redress. Retailers may also struggle to prove that they had complied with regulations in non-disclosed agent transactions. Further, agents may not be sufficiently clear over how decisions were reached, or any commercial partnerships or sponsorship agreements that were involved in options considered or presented to the consumer. This reduces transparency and accountability. Consumers could suffer financial harm by agents making purchases for them without full transparency, making it harder to claim redress. That

could result in time harms, and psychological stress from trying to sort problems out that had been caused by agents. Similar issues arise with a lack of transparency regarding decision making, with consumers left unaware of what agents had done on their behalf, and why.

As consumers delegate product selection and purchasing to agents, the ‘as described’ component of the Consumer Rights Act (CRA) will need to be rethought, based on the assumption that the agent, not the consumer, will be assessing the goods/listings pre-purchase.

A central tenet of the CRA is that goods will match the description so that a consumer knows what they are buying. However, with AI agents interjecting an intermediary that assesses such descriptions for the consumer, this places a lot of responsibility on the agent to assess, relay or potentially execute a decision based on the information it has seen. The consumer could theoretically receive an item having never seen a description of it at all, and so they are wholly reliant on the agent doing what is required by law, but also safeguarding their best interests. This could result in financial harm of purchases not meeting consumer requirements, and physical harms of agents mistakenly buying unsafe goods. This could also complicate redress processes, leading to time harms. It might be that retailers honour commitments to put things right but this could depend on whether the use of the AI has been disclosed or not, and the level of autonomy in use. Finally, consumers could be psychologically harmed by losing skills to interrogate product descriptions for themselves (or not being able to do that at all should online retail tools shift primarily to cater for machine-to-machine systems).

Consumer redress mechanisms must be rethought to accommodate AI agent transactions, including robust clarifications on how these will be managed, and by whom.

Even if we clarify issues around liability and risk with autonomous AI agents, there is a danger that redress mechanisms become overly burdensome on the consumer. Retailers might not accept a redress basis of ‘my agent bought this for me’, and the tools needed to unpick decision-making processes may not keep pace with agents being intermediaries in the process. While higher levels of documentation in agent transactions could help, that might not matter if the burden of proof is placed at a higher level by retailers and traders, or even the courts. The primary harm is time, with consumers facing lengthy, challenging and burdensome processes potentially to seek redress for unwanted, wrong or even dangerous purchases made for them by their agents.

The right to refund/right to cancel put under strain with AI agents purchases.

Consumers have the ‘right to a refund/right to cancel’ under the CCRs, depending on the type of sale. However, there are a number of risks when using AI agents, such as if the agent purchases something that the consumer does not want, but they only realise after the standard returns period has expired (for example, they may be late to spot an email confirmation or transaction on a bank statement). AI agents could be useful in helping consumers to cancel purchases, but they could prove to be a complicating factor, too. The primary harm here is time, but there is also likely to be financial harm should consumers find it hard to get their money back for unwanted purchases.

Tools and frameworks: Conclusion

The first distance selling regulations were introduced in the UK in October 2000, around the time Amazon and eBay first joined the e-commerce market. These early regulations set the blueprint for some of the protections we retain today - cooling off periods, legally required pre-contract information and options for cancellations and refunds.

Creating robust consumer protection at an early juncture of a market shift can actually benefit growth and innovation by giving consumers confidence to engage with services knowing there are protections to fall back on. While the market for AI agents in retail and e-commerce is nascent, this does not lessen, in our view, the importance of defining just as early how they should be regulated.

2. AI and Retail financial services

2.1 Risks and Harms

AI adoption is relatively high, and growing, across financial services firms in all markets and for varied purposes¹⁹. Potential benefits from the deployment of AI in financial services are identifiable, related in particular to increased accessibility of financial services and potential cost savings passed on by financial services providers from the creation of efficiencies. However, benefits are not certain; competition and quality of AI must be sufficient to see them flow through to consumers.²⁰

Meanwhile, there are also clear risks of harm that could flow to consumers from the use of AI. Our analysis of core existing and expected future consumer-directed use cases of AI in financial services has identified the following potential consumer harms. Given the current state of the technology and signs of risk, we believe these harms are reasonably likely to occur. Many of these harms could be severe, as they involve financial decisions that significantly impact consumers' lives.

AI Use Case	Potential Harm	Harm Type(s)
Customer support (e.g. customer service chatbots)	Poor quality AI service/outputs leading to worse customer choices and/or customer frustration (e.g. difficulty navigating processes and admin tasks).	Financial Psychological Time
Service access & pricing (e.g. predictive AI models used by firms for credit scoring, loan application decisions, or insurance pricing)	Unfair denial of access to products/ services due to assumptions made by AI models.	Financial Psychological
	Discriminatory or unequal pricing decisions, i.e. offering unfairly high prices to specific consumers due to assumptions made by AI models.	Financial

¹⁹ BoE/FCA, 'Artificial intelligence in UK financial services - 2024' [\[link\]](#).

²⁰ FCA, 'The Mills Review', July 2026.

	Sub-optimal decisions made by consumers due to manipulative use of AI for marketing.	Financial Psychological
Advisory services (e.g. Robo-advisory apps which provide budgeting, tax or investment advice)	Sub-optimal decisions made by consumers due to poor advice given by services.	Financial Psychological
Agentic AI (e.g. financial assistants carrying out tasks such as money transfers or product switching on behalf of the consumer)	Agents implement decisions which do not match the intent of the human user.	Financial Psychological Time
Cross use case	Difficulty obtaining redress or issue resolution - due to complex supply chain, agentic interactions, and/or low explainability / transparency of AI.	Psychological Financial
	Reduction in financial literacy as delegation to AI increases	Psychological Financial
	Data privacy - data used for purposes without specific & informed consent; exclusion of high-privacy consumers who opt out of data sharing.	Psychological
	AI-driven financial services-based fraud.	Financial Psychological

In general, the lack of understanding of AI technologies among consumers, and therefore of the risks entailed, means that they cannot and should not be expected to be able to protect themselves, or bear the burden of doing so where businesses choose to deploy AI because it brings some organisational benefit. The inherently high stakes when individuals' personal finances are involved increases the severity of the risks of harm for most of the key use cases in this sector. As a result, there is a particularly low level of permissible risk tolerance related to the use of AI applications in retail financial services.

Whilst consumers should never be exposed to high levels of risk or be expected to take on inappropriate burdens of self-protection, risk tolerance may vary sector by sector, and service by service, given the range of contexts AI is likely to be utilised in.

2.2 Accountability and liability

The FCA has stated that financial service firms remain responsible for outcomes produced by AI that they deploy in the same way as for any other outcome, under existing FCA regulations²¹. The CMA has also emphasised that businesses are still liable under consumer law when customers interact with an AI-based application²². This means that liability to the consumer for harms caused to them clearly sits with the financial services firms themselves,

²¹ <https://www.fca.org.uk/firms/innovation/ai-approach>

²² CMA, 'Complying with consumer protection law', March 2026 [\[Link\]](#).

not with model developers, AI interface providers or any other elements of the AI supply chain. We support this position as the most accessible route for consumers to seek assistance when things go wrong. Financial services firms are tightly regulated because of the importance of the services they provide and the role they play in society, and it is appropriate that they are held to high standards. The Mills Review highlighted that allocation of liability will become more complicated with increases in the autonomy of AI²³ and we note this as an area of potential future concern where accountability may need to be re-clarified.

There is potential for business-to-business redress routes from financial services providers upstream to model providers or other parties, via commercial contracts or legal action for example. However, any such arrangement should not alter the consumer-financial services provider liability relationship. It may be reasonable to limit liability to some extent based on 'average' or 'reasonable' consumer behaviour expectations, in line with existing consumer protection laws, but the general low level of AI understanding amongst consumers²⁴ makes it inappropriate to expect consumers to take on the burden of protection when using AI-driven services.

Clarity regarding which party is liable for any harms caused by AI is important. This has been cited by a range of financial services stakeholders, including industry firms and trade bodies, as key to enabling their adoption of AI (and therefore the realisation of any benefits for consumers). In addition, consumers should never be let down simply by a lack of clarity about where they can go for redress when they suffer harm. Whoever is decided to be accountable in different AI-related scenarios, this should always be made extremely clear in sectoral and wider regulation; this is beneficial for all parties involved.

Given the range of consumer-facing use cases that will be deployed as AI adoption grows, it will be essential for firms to have formal responsibilities aimed at ensuring safety and to be incentivised to meet these through strong accountability frameworks. It is not practical to expect the regulator to safeguard consumers alone without service providers themselves taking proactively and appropriate action, although oversight and enforcement are important to ensure protections have a strong basis and are respected by firms.

2.3 Tools and frameworks

FCA Regulation

Retail finance provides several interesting examples of specific regulatory tools and frameworks that are different from those seen in other sectors and that are intended to be applicable to AI deployment by authorised firms. We have assessed the effectiveness of these in the context of AI. A summary of our analysis of some of the key tools employed by the FCA is set out below.

a. Consumer duty / FCA Handbook

The consumer duty offers some protection against potential risks of AI services by setting relevant and appropriate standards which should increase the likelihood of AI services being of reasonable quality in terms of their outputs. The duty is also appropriately flexible in that

²³ FCA, [‘The Mills Review’](#), July 2026.

²⁴ Eg, FIS, [‘Banks must educate as they innovate’](#), 2025.

the same outcomes and principles it enshrines are appropriate in the context of different AI tools and can continue to apply as AI use evolves.

However, the duty has some limitations. First, it offers no private right of action for consumers who suffer harm if it is not met. As a result, its impact relies heavily on effective supervision and enforcement by the FCA. Additionally, some stakeholders have raised concerns over the lack of specific articulation and clarity regarding how the duty applies in the context of AI-enabled services²⁵; this lack of clarity and stakeholder confidence could reduce the duty's ability to ensure standards. Finally, as the autonomy of Agentic AI use cases increases, it will likely become difficult for firms to evidence consumer duty outcomes such as consumer understanding and it may need to be adapted to continue to apply²⁶.

This overall assessment is also true for other aspects of the FCA Handbook (of which the Consumer Duty is part) when applied to AI use - for example, several of the other Principles for Business which cover aspects like due skill and care, regard for customers' interests and fair treatment of customers, and appropriate communication with clients.

b. Senior Managers & Certification Regime (SMCR)

It is important to have strong governance and assignment of responsibilities around AI given its complexity and the severity of its associated potential harms. The SM&CR fulfils this need to some extent by ensuring that liability for poor practice is assigned personally to managers within regulated firms, which acts as a strong motivator for ensuring high standards and safe deployment of AI, and a deterrent against over-adoption of AI without sufficient safeguards.

However, the SM&CR would be a better enabler and protector of responsible AI adoption if it was updated with specific articulation of its application to the use of AI. Submissions to the Mills Review from many financial services firms and trade bodies highlighted that an AI-specific articulation would be helpful or give them greater confidence safely governing AI use in their organisation, noting they currently see gaps in how it operates in this context²⁷. In addition, it will be harder to exercise oversight and accountability requirements under the regime in the case of much more autonomous AI tools so the framework may need to be adapted in future to cover such AI use effectively²⁸.

In addition, trusting managers to uphold standards for AI functions and truly be accountable is complicated by the lack of understanding within firms of how the AI technology they are deploying works - the Bank of England's 2024 survey of financial services firms using AI showed that only 34% consider themselves to have 'complete understanding' of the AI technologies they use, with 46% admitting having only 'partial understanding'²⁹.

There is a question around whether it is appropriate or effective to expect individuals who do not understand the technology they are deploying to provide their services well to ensure the right standards are maintained and embedded. This points again to the need for the SM&CR

²⁵ E.g. CFA Institute, 'Response to FCA Mills Review', Feb 2026 [\[link\]](#); HMT Financial Services AI Adoption Plan, July 2026 [\[link\]](#).

²⁶ FCA, [The Mills Review](#), July 2026.

²⁷ FCA, [The Mills Review](#), 'Annex II - Engagement paper responses' at p121, July 2026.

²⁸ FCA, [The Mills Review](#) Figure 16, July 2026.

²⁹ BoE/FCA, 'Artificial intelligence in UK financial services - 2024' [\[link\]](#).

to set out specific requirements and unique allocation of responsibility for AI within firms to an appropriate staff-member, who may differ from the appropriate person assigned responsibility for other elements under the SM&CR.

c. Supercharged Sandbox

The FCA's AI sandbox has seen excellent uptake since its introduction in 2025 and is an effective way of improving safety of tech for consumers before it reaches the market. Participants can use the sandbox to develop, test and refine AI-enabled service propositions in a secure and controlled environment³⁰ - this enables risks to come to light, and the effectiveness, safety and feasibility of services to be improved prior to deployment towards consumers. However, the sandbox has volume limitations; 132 firms applied to join the first cohort, but only 22 were able to take part, while 21 firms are participating in the second cohort from 199 applicants. The high applicant numbers, which grew from cohort one to two by 51%, demonstrate the appetite for this type of tool and therefore its potential for success but this is limited at this time by understandable capacity constraints at this early stage.

The effectiveness of the sandbox could be maximised, especially while its capacity is limited, by better and more open sharing of lessons learned by those participating in each cohort. The popular sandbox showcases already being run by the FCA are a good example of such an approach³¹. Additional structures for recording and sharing with other financial services firms information on how testing progressed and particularly what risks materialised in relation to the propositions in the sandbox (whilst avoiding sharing of proprietary model/product details) could increase the reach of the sandbox's influence to develop safer AI for the sector at a faster rate. The FCA has stated that they will publish an 'AI Live Testing evaluation report' in late 2026 or early 2027 to share the experience of their innovation tests³² and will be running a further sandbox showcase this year. These are welcome next steps towards building a more comprehensive framework around tools like the sandbox.

Other sectors and regulators could consider how aspects of the three tools discussed above could be applied to AI regulation in their industries as useful examples to learn from:

- The Handbook, Consumer Duty and SM&CR could work as template schemes, with elements of their overarching approaches and structures transferred to other sectors. Specific principles and wording would need to be adapted to the sectoral contexts, but the structure and the particular outcomes and tools employed within the frameworks could be mirrored. However, there is a question as to how effectively such frameworks could be implemented and upheld in the absence of a specific regulator to oversee them and the necessary infrastructure to allow adherence to be monitored - so their usefulness may be limited to more organised and closely-governed sectors with a dedicated regulator.
- Sandboxes and other innovation-focussed approaches could be successfully rolled out for other sectors, with some central sectoral co-ordination for their setup and administration. The usefulness of sandboxes in allowing safe development of AI is applicable to any sector. In addition, lessons learned and use cases tested in specific sectors where such initiatives are implemented (e.g. by the FCA for financial services

³⁰ <https://www.fca.org.uk/firms/innovation/supercharged-sandbox>

³¹ E.g. https://events.fcainnovation.co.uk/supercharged_sandbox_showcase_c1

³² TSC, '[AI in financial services: Responses to the Committee's Fifteenth Report](#)', April 2026.

firms) could flow through to other sectors if sufficient infrastructure for cross-sector knowledge-sharing was set up. For example, customer service AI needs may be similar enough that models could be adapted to various sectors following development in one sector's sandbox.

Other approaches to financial services regulation

Industry standards are part of the discussion around managing AI in financial services (e.g. they are touched on by Mills Review recommendations³³). However, we are not aware of any evidence of effective industry standards developed for AI use in financial services to date.

Specifically regarding *voluntary* industry standards, we do not believe they would be a very effective tool in governing AI use in financial services. There is a risk of their development being captured by the biggest firms, which can result in minimisation of protections and/or reductions in competition, weak adherence by participants, or little or slow progress in their development. This is true in financial services where there are large incumbents within each market, and could well be the case in other established sectors.

In terms of *international regulatory examples* for AI, Singapore is often held as a good example of regulatory approach in relation to AI in financial services. They have a principles-based approach (similar to the UK) but with more specific and practical tools and frameworks built on top to cover AI use. These include the FEAT principles³⁴, Veritas protocol³⁵, clear accountability and escalation routes, and AI-specific data protections. The UK could learn from some elements of Singapore's regulation, where they have flexed their approach in order to boost clarity and safety around AI use. This includes, for example: requiring regulated firms to maintain inventories of AI systems; publishing guidance explaining how existing rules apply to AI; developing open-source fairness and explainability testing methods; clarifying the boundary between customer support tools and regulated advice; and ensuring consumers who decline non-essential AI profiling continue to receive access to core financial services.

However, ultimately there are cultural and structural differences between other jurisdictions and the UK when it comes to financial services regulation, and other areas of consumer protection, which influence what regulation is effective. There is no single best practice approach yet established internationally. It is important that the UK develops tools which work well for our unique markets, whilst collaborating with global counterparts in terms of general AI standards.

³³ FCA, '[The Mills Review](#)', July 2026.

³⁴ <https://www.mas.gov.sg/publications/monographs-or-information-paper/2018/feat>

³⁵ <https://www.mas.gov.sg/schemes-and-initiatives/veritas>

About Which?

Which? is the UK's consumer champion, here to make life simpler, fairer and safer for everyone. Our research gets to the heart of consumer issues, our advice is impartial, and our rigorous product tests lead to expert recommendations. We're the independent consumer voice that works with politicians and lawmakers, investigates, holds businesses to account and makes change happen. As an organisation we're not for profit and all for making consumers more powerful.

For more information contact:

Cressida O'Donoghue and **Andrew Laughlin**

Senior Policy Advisors, Digital Policy

Cressida.odonoghue@which.co.uk; Andrew.laughlin@which.co.uk.

[September 2026]