

Response to Call for Input

Consumer interest and AI

CCIA is an international, not-for-profit trade association representing a broad cross section of communications and technology firms. For more than 50 years, CCIA has promoted open markets, open systems, and open networks. CCIA members employ more than 1.6 million workers, invest more than £75bn in research and development, and contribute trillions of dollars in productivity to the global economy. Many CCIA members are developing generative and agentic AI services and/or integrating them to improve existing devices and services.

This submission responds to the Digital Regulation Cooperation Forum call for input [Consumer interest and AI](#), particularly Part 1 on Consumer Attitudes. CCIA would be glad to act as a resource to the DRCF as this work develops, including by convening members and sharing evidence on how consumers actually use and respond to AI services. The aim should be a UK framework in which consumers adopt trustworthy AI with confidence, and in which specific harms are addressed without discouraging adoption that delivers many benefits.

1. To what extent should we expect consumers to understand the full risks associated with AI and manage those risks?

CCIA shares the objective of UK consumers being able to weigh the benefits and risks of AI with confidence. Users are familiar with AI in outline and have often used it to some extent, but might have limited awareness of the details. This will include upsides (e.g. the different ways in which consumers might use AI tools) and downsides (i.e. any risks). This knowledge should grow over time as consumers use AI tools more and become better aware of the full range of associated risks and opportunities.

In many settings consumers fully understanding any risks is not required for consumer choice to drive risk mitigation, however. The most popular AI services are often large brands either in themselves (e.g. ChatGPT) or as part of wider product offerings (e.g. Google). Those services have an incentive to maintain those brands, promoting consumer trust and acceptance, by not taking excessive risks, even leaving aside any regulatory intervention. Other commercial actors such as app stores and operating system providers can also play an appropriate role. Generally speaking, users' confidence will be greater if trusted parties play a role in safeguarding them from potential risks.

Regulators should consider the risk that interventions which undermine mainstream AI services may push consumers toward less-regulated alternatives with fewer safety guardrails, where there may be less capacity and incentive to manage risk appropriately, or undermine the potential for other players in digital ecosystems to create a trusted environment in which smaller developers can benefit from the resulting consumer trust.

In many settings AI will not be encountered by consumers as a freestanding tool, but integrated in existing goods and services, e.g. as discrete features in applications or mobile ecosystems. This may mean, for example, that not only the services themselves but also device manufacturers, operating system or app store operators and other ecosystem participants compete to attract consumers by providing a more trustworthy environment for using AI.

It is also important to note that consumers, just like policymakers, rationally balance their attention across different risks and opportunities, which might be more or less directly related to AI. [Freshwater polling for CCIA](#) suggests, for example, that the impact of AI on society is prioritised by around 5% of voters, versus 55% that prioritise inflation and the cost of living. If attention to AI risks becomes disproportionate to the point that it undermines adoption of a technology that could support real-terms economic growth (easing pressures on household cost of living) then it might not reflect consumer priorities overall. Targeted protection against real harms remains important, but a regulatory posture that treats AI as a leading consumer threat in and of itself would not reflect how people themselves rank what matters to them.

2. Is there a particular category of consumers who may be at higher risks of harm (e.g. specific categories of vulnerable consumers)?

AI might change the risk profile of the devices and services in which it is embedded less than immediately assumed.

Users tend to view AI more positively as they use it more. This is reflected in several findings in the Freshwater polling described above, for example that younger voters, those with degrees, full-time employees and those on higher incomes are more likely to say that it would be concerning if the UK regulated AI more than other countries. While this does not preclude harms, it fits a pattern where for most users the experience is a positive one.

To the extent that AI use diffuses through the population as other digital technologies have done before, voters are likely to understand it better. If people are not making significant use of AI, that will mean they are likely to understand it less, but also that they are less likely to be exposed to any risks.

Other than that, vulnerabilities will often reflect existing vulnerabilities. Fraud, for example, is a greater risk for those with diminished cognitive capacity. AI might create means both for bad actors to exacerbate those risks and responsible parties to mitigate those risks. A diverse range of organisations might play a role in mitigating the risks that users face in mature digital ecosystems including agentic and generative AI systems, for example mobile ecosystem operators implementing operating system-level safeguards against developers accessing data inappropriately.

Finally it is important not to conflate differences in attitudes with risks of harm. Attitudes to automated decision making, for example, vary across the population but being less

comfortable with it in principle does not mean a given consumer group is more likely to be harmed by its use.

3. Are there particular values that consumers feel more strongly should be upheld when interacting with AI?

While certain values do appear to loom larger when voters are asked about their priorities for AI (e.g. privacy - see the Freshwater polling cited above), regulators should be wary of drawing broad conclusions. This might simply reflect a still-forming understanding of the risks and benefits. Values are also likely to vary between markets and aligning an otherwise safe product with consumers' values is beyond any regulator's realistic scope. The focus for regulators should remain on minimising barriers to realising the benefits of AI while addressing specific potential harms, reflecting wider societal values rather than specific values that currently seem relevant to consumers experiencing a still nascent technology.

4. To what extent do you see consumers being willing to tolerate risks from generative or agentic AI in the face of perceived benefits, and how do they articulate this trade-off?

The Freshwater research cited above considered this question. 51% reported that it would be fairly or very concerning if the UK regulates AI more than other countries, versus 35% who felt that it would be not very concerning or not at all concerning.

In the same polling, the risks where there was the strongest net agreement were: that AI systems might be unreliable or make dangerous errors; that they might be a risk to privacy or data security; that humans could lose control; and the risks that AI might produce biased or unfair decisions.

Many of these risks fundamentally speak to the quality of AI models. Regulatory decision-makers should therefore proceed carefully with interventions that might, for example, diminish access to training data or the commercial incentive to improve models over time. One downside of the UK's copyright regime providing less protection for commercial text and data mining than other major economies, for example, is that it will diminish training in the UK and mean models are less well-adapted to British users' needs. Many risks will be diminished with better models.

Generally speaking, regulators taking an approach which frustrates AI development may therefore inadvertently exacerbate risks that concern consumers and this should be considered in regulatory impact assessments and other policy development. By contrast, a more open and collaborative approach can yield much better results without holding back development unnecessarily, e.g. the voluntary testing of models with the UK AI Security Institute.

Others identified risks should be addressed through existing regulatory frameworks such as general data regulation. Regulators should generally err on the side of letting companies

innovate and apply those existing regulatory frameworks appropriately for new settings. AI is being developed with far more proactive regulatory engagement than earlier technologies. In the UK, this includes DRCF's own Innovation Hub, ICO/FCA sandboxes, and voluntary commitments (e.g. with the AI Security Institute). Where there are opportunities for new approaches, better-suited to new technological contexts, regulators should generally seek to work with companies able to pioneer novel approaches, but then treat that as an example to adapt the regulatory framework across the board, not restrict innovation to pilot projects or regulatory sandboxes.

Some survey respondents indicated concerns that relate to AI misuse and misalignment, although as noted above most reported that it would be fairly or overly concerning if the UK regulated more than other countries. This underscores the importance of seeking international consensus, informed by specialists (e.g. the AI Security Institute)

The benefits cited include advancing scientific research and progress; improving public services; improving efficiency and productivity in the workplace. Regulators should work to ensure that any interventions do not impede these benefits, particularly as in many cases these address long-standing problems that have proved difficult to address through other means, e.g. sustained productivity growth.

Scientific research and progress, and the embodiment of that research in practical innovation, will often produce hard-to-predict long-term benefits that are easy to neglect in assessments of the proportionality of regulatory interventions that might complicate innovation. Given that analytical challenge, these long-term benefits deserve real weight when the proportionality of an intervention is assessed, particularly where it bears on the pace of innovation. In particular, proposals to extend outcomes-based obligations (such as the FCA Consumer Duty) cross-sectorally should be approached with caution, given the difficulty of defining and measuring "good outcomes" for diverse AI applications serving heterogeneous consumer populations.

Regulatory barriers to research and innovation will often compound each other across different regulations, regulators and policy priorities within those regulators. Organisations with a broader scope such as DRCF itself, the Regulatory Innovation Office and Parliamentary committees might have a role to the extent that there is a need to ensure that regulators do not impose an impractical cumulative burden on AI innovation across multiple domains.

Regulations can also conflict with one another, exacerbating the challenges to introducing innovative services. This risk is illustrated by challenges in the EU, where competition interventions (particularly under the Digital Markets Act) have been introduced that [conflict](#) with privacy and cybersecurity regulations. There is a clear risk that this could be repeated in the UK with changing rules to mitigate AI risks conflicting with conduct requirements under the Digital Markets, Competition and Consumers (DMCC) Act.

5. How much human intervention or oversight do consumers expect when using AI (if any)?

Consumers' expectations will vary enormously with context. In many cases, where AI is embedded in consumer tools, e.g. chatbots, they are unlikely to expect or want human

involvement. The same is likely to be the case for decisions that seem sufficiently routine, e.g. personalisation of digital services, where human involvement is impractical and likely to seem intrusive if anything.

In the context of more substantial automated decision-making, there is still likely to be real nuance. Involving humans may in many cases give consumers a sense of accountability for the decisions that affect them and, for now at least, be associated in many consumers' minds with more reliable decisions. However, as AI systems become normalised, and particularly in contexts where services are required to adjudicate between two parties, human intervention could also be seen as an unfair intervention in the system, overruling a neutral decision in one party's favour. Regulators should be wary of over-interpreting high level views expressed by consumers, who might prefer human intervention in general and in theory, when their actual reaction might vary enormously depending on context and the associated trade-offs (e.g. does requiring human intervention mean services are less responsive or more expensive).

Generally speaking, regulators should be willing to engage with organisations that seek to pursue good faith innovation over how AI is integrated in their processes, rather than having a blanket presumption in favour or against human intervention and oversight. Effective oversight will evolve as AI matures, for example, from human-in-the-loop to human-on-the-loop approaches involving verification, sampling, and exception handling, and should be meaningful and proportionate to context, not merely procedural.

6. What do you believe is the current level of AI consumer literacy?

Without clear metrics for literacy, any consideration of whether consumers are “AI literate” is inherently subjective. Many consumers have a good practical high-level understanding of AI, in that 68% report having used the technology and 29% report using it often in the Freshwater polling cited above.

While it is always possible to quibble, respondents gave reasonable descriptions unprompted such as “a technology enabling computers to perform tasks needing human intelligence like learning deduction making and solving problems” or an “automated system that can think like a human but faster” or a “machine designed to mimic human intelligence through learning and problem solving”. The most common description was a system that simulates or mimics human intelligence.

This does not mean, of course, that people will understand the strengths, weaknesses and appropriate use of every device and service using AI, any more than is the case with devices and services prior to the development of contemporary AI. There should be a balanced approach where providers are responsible for doing what they can to promote safe and appropriate use, without overly limiting the function of AI-powered services. Not only would this limiting their function undermine the potential benefits of these services, and progress towards the Government's target for the UK to lead the major economies in AI deployment, it would also undermine learning about how to use the services appropriately.

7. Is there a perception of data protection risks across consumers, and is that outweighed by perceived benefits?

Potential privacy breaches are one area where consumers do have concerns around AI, according to the Freshwater research cited earlier. Whether this is outweighed by perceived benefits will vary depending on the application and the consumer in question. However the relationship may also be more complex in terms of practical policy questions.

In some cases, regulatory requirements might increase privacy and security risks. For example, policymakers have proposed “transparency” requirements for AI developers as part of the debate over copyright reform which would make the functionality of AI models public. This would create opportunities for people to exploit those models through poisoning the inputs and therefore potential security breaches.

Competition interventions, for example, which can undermine the oversight role for providers of digital ecosystems, have the potential to exacerbate AI-related privacy and security risks. The Digital Markets Act’s prioritisation of interoperability has been criticised for exacerbating privacy and security risks. Regulators should engage closely around potential interventions to ensure that they do not impede AI security. In other cases, regulatory impositions that undermine the attractiveness or viability of reputable AI services could lead to circumvention that again increases privacy risks.

Data protection risks are an example of where coherent regulatory oversight is important. A single AI deployment may engage ICO, CMA, Ofcom and FCA simultaneously. Lead-regulator designation should be encouraged where possible, regulators should avoid duplicative reporting and ensure there is a coherent hierarchy where multiple duties apply to the same system. Regulators should align with international approaches where possible. Significant divergence between UK, EU, and other frameworks risks imposing disproportionate compliance costs on firms operating globally without delivering additional consumer benefit.

8. Would standardised disclosures about AI data use and behaviour (e.g. “does not retain data”) improve trust? Which disclosures matter most?

Many AI-powered services already include visible disclosures, particularly relating to the potential for inaccuracies. Regulators should be very wary of promoting standardisation in disclosures while many AI services are still nascent and diverse. Disclosures that attempt to comprehensively explain the full potential function of a service could change regularly and be so detailed and technical that it could easily overwhelm the average user.

Experience from other disclosure regimes notably cookie consent banners under ePrivacy demonstrates that proliferating consent requirements can degrade rather than enhance consumer understanding, leading to reflexive click-through rather than informed choice. Industry experience with AI content labelling (e.g. using C2PA content provenance

metadata) suggests that contextual, in-product transparency mechanisms are more effective than static consent-based disclosures.

The impact of any disclosure is likely to be highly sensitive to context that will change over time: how popular and well-understood is the service; what risks are consumers most concerned about; and what technical commitments would actually contribute most to diminishing risk. Standardisation risks disclosures that are not well-aligned to the most meaningful risks and are therefore ignored, undermining progress on educating users about more salient risks.

More than that, because many disclosures carry with them technical commitments, a commitment not to retain data in the example given for this consultation question, which will constrain the development of these models. Over the longer term, it would be unwelcome if innovation were delayed because of its implications for standardised disclosures, more than any meaningful change in its actual risk profile.

The best context for standardised disclosures might be areas in which a regulator can provide meaningful “safe harbour” to enable innovation, but appropriate disclosures are needed as a result. In this case, the benefits of the regulatory certainty provided by a safe harbour might outweigh the inflexibility associated with standardised disclosures.

9. To what extent do consumers meaningfully engage with AI-related disclosures, and what formats (e.g. reading terms and conditions, summaries, labels, prompts) improve engagement?

CCIA does not have any specific evidence on this point. We would be cautious of over-interpreting findings about, for example, the percentage of people that read a given disclosure as a signal of serious risks. That overall percentage will include many people using tools in contexts where the risks are low and it is not irrational for them to simply try these services out and learn their strengths and weaknesses. Low risk experimentation is as legitimate a way for users to learn the strengths and weaknesses of a service as reading disclosures. Such findings are instead a reason to be careful about overly-burdening users with too many disclosures, or disclosures that bear little relevance to meaningful risks.

10. To what extent do consumers perceive the benefits (e.g. more tailored advice / service) versus the risks (e.g. data breaches, discrimination, fraud etc.) of AI becoming increasingly personalised? Is a higher level of personalisation a desirable outcome?

Generally speaking, consumers benefit from personalisation as an essential means to find content, goods and services among the almost infinite options theoretically available online which are relevant to them, and not much broader categories (e.g. viewers of a certain portal). While personalisation should take place with appropriate protection to mitigate privacy, safety and security risks, consumers therefore tend to prefer personalised experiences. Deloitte [research](#) in 2024, for example, found that 80% of consumers prefer brands that offer personalised experiences. Ofcom's research has [found](#) that AI chatbot adoption has grown from 31% to 54% of UK adults in a year, suggesting a strong consumer appetite for personalised AI services. While AI services will vary enormously, and consumer attitudes may vary in turn, there is little reason to think that AI will be any different in that it will enable AI services to surface more meaningful responses to people's questions and other requirements.

Personalisation is also important as a practical means to achieve many regulatory objectives. Age-appropriate experiences for children, for example, are a welcome form of personalisation. Any regulatory obstacles to personalisation will often be counterproductive as a result.

11. What would give consumers confidence that it is appropriate to trust an AI system? E.g. government-backed trust marks or access via trusted public platforms, verifiable accreditations (registration numbers, public registries, certification schemes).

Usage of AI systems is growing rapidly and those using AI services generally have the most positive attitudes towards their use. It is not clear that there is a gap that a trust mark or accreditation scheme could solve. To the extent that such standards do develop, it is best that they do so globally and are industry-led, such that the accreditation requirements for developers are proportionate and the scheme can adapt as the technology itself matures.

Existing industry-led standards that are already building trust, such as C2PA (Coalition for Content Provenance and Authenticity) for content provenance, AI labelling systems deployed by major platforms, and model cards. These are examples of how the sector is already investing in trust-building mechanisms, reflecting a range of incentives including the commercial incentive to promote adoption.

11a. To what extent could simple, user facing controls over agent behaviour (e.g. “approve before action”; “never make purchases without confirmation”) help to build trust?

User controls are a part of the development of AI services and important both in terms of building trust and as a basic means for consumers to have the models do what they want. Effective control will be a legitimate part of how agentic models compete for consumer adoption and regulators should therefore be cautious in assuming a market failure here that would justify intervention. In sophisticated agentic AI systems, user facing controls over every step in a process are not likely to be viable and a range of organisations can have a role in ensuring consumers can still protect themselves against potential risks, including the services themselves and other ecosystem participants (e.g. operating system and app store operators).

12. Do consumers’ perceptions of risk and trust vary depending on the type of AI they are using - for example, whether the AI is embedded within a specific firm’s service (such as a built-in chatbot) or provided by an external system, such as ChatGPT?

CCIA does not have any specific evidence on this point. It would be important to bear in mind that this is often related to whether or not a service is truly general-purpose, or tied to a specific set of use cases (e.g. around e-commerce, or productivity software). These are generally lower-risk environments in which the levels of risk are genuinely lower and therefore consumer trust is easier.

13. What defines high-quality AI service (e.g. accuracy, reliability, fairness, transparency, user control), and how should it be measured?

There will not be a meaningful general answer to this question as it will vary enormously by use case and therefore the associated risks and rewards associated with using the AI system. It is also important to note that regulators should not generally see their objective as promoting high-quality AI services in itself. While regulatory barriers to raising the quality of AI services, or risks relating to low-quality AI services, might be a legitimate concern, there will be genuine trade-offs between quality and cost, for example, where regulators should not seek to put their thumb on the scale. Regulation should not require, and the market does not need, convergence on a single, uniform set of privacy and security standards for platform-agent interactions. Allowing platforms to compete on the design,

rigour, and architecture of their protective measures promotes innovation, user choice, and long-term resilience.

14. How do existing terms and conditions include or reflect risks of AI?

No response.

15. What redress mechanisms exist today, how accessible are they, and do consumers trust them to resolve AI-related harms?

No response.

15a. Does the availability of clear complaint routes and escalation to an ombudsman or regulator increase consumer confidence?

Generally speaking, AI should not change the question of whether or not regulated complaint and escalation routes are appropriate. They are likely to remain appropriate in institutional contexts in which they are expected now, and not in others where they would be a disproportionate institutional burden given the risks.

15b. Should regulators mandate strict liability in order to facilitate redress measures?

There is unlikely to be a general answer to this question which will depend heavily on the circumstances, but there is no reason to believe AI should generally imply strict liability. Its risk profile and the ability of potentially liable partners to avoid those risks will vary.

The European Commission [withdrew](#) a strict liability proposal for AI in February 2025 as part of an attempt to simplify regulation in the digital economy.

16. Who do consumers expect to be accountable when AI causes harm (provider, developer, user)? Does this differ for generative vs agentic AI?

While a diverse range of organisations have a role to play in mitigating risk, in cases where there is the potential for liability to be created by misuse of AI tools, liability should sit with the user misusing the service. Intermediary liability, by contrast, will tend to result in high

costs as intermediaries then have to shut down legitimate activity in order to avoid hard-to-anticipate liability risks (e.g. curbing free expression where intermediary liability regimes are established for content).

Where a model is released openly or via API and a third party deploys it in a specific context, accountability should follow the party with control over the deployment decision and the end-user relationship, not the upstream model developer who has limited visibility into downstream use cases. Similarly, where platforms host or distribute AI-generated content or services created by third parties, their accountability should reflect their role as intermediaries, not as the decision-makers.