

Response to Call for Input

Regulatory, policymakers, industry & consumer tools

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 2 on Regulatory, policymakers, industry and consumers tools.

17. What tools can industry offer to consumers or otherwise implement to manage AI risks?

The most fundamental but important contribution industry can make is simply better models, deployed responsibly and more broadly accessible. In part because many AI risks relate to the potential for errors, which should become less likely as models improve. And in part because there are many applications of AI in mitigating consumer risks. AI will on balance tend to favour the defence in ‘needle in a haystack’ tasks defending consumers against online risks such as fraud, where a platform needs to identify fraudulent content hiding amongst large volumes of legitimate content. This will be strengthened both by improvements in AI services over time and by broad access to those tools beyond a narrow set of institutions and early adopters. Regulators should therefore be cautious about any intervention which might undermine this innovative process and AI deployment overall.

Beyond that, AI companies can innovate in a responsible manner. In *Understanding AI: A Guide to Sensible Governance*, CCIA [described](#) a broad agreement on the requirements for responsible AI development:

- Design for social benefit
- Design to avoid unfair outcomes
- Analyse and minimise risks as you design
- Consider the risks to third parties from AI systems during design, but also the benefits (this can include voluntary collaboration with external bodies such as the UK AI Security Institute (AISI) on pre-deployment testing)
- Use up-to-date safety, security and privacy best practices, potentially including red teaming programmes with external experts
- Monitor and govern identified risks in deployed systems - this can include:
 - automated, scalable guardrails built into AI services; and
 - structured scaling and safety frameworks that set tiered risk thresholds and halt deployment if models fail evaluations.

- Provide appropriate disclosures for deployed AI systems

As that paper noted, while “these principles may be expressed in different ways, any responsible AI framework will incorporate them.” Responsible AI developers have both implemented these measures themselves, but also collaborated to raise standards overall (e.g. the [Frontier Model Forum](#)).

Those high-level principles provide a framework but must be applied in the context of any given application, providing the necessary flexibility to manage risks while maximising the benefits that AI can deliver. In high-risk applications, e.g. medical diagnostics, human supervision and significant disclosure may be appropriate, where in lower-risk applications such as personalisation in entertainment products they would be disproportionate and deter beneficial innovation. However, this does not mean that companies are not still applying the broad approach described above.

Crucially, this approach will often be delivered in the context of technological ecosystems where different parties have different roles to play. Regulators should generally err on the side of allowing companies to play the role consumers expect them to undertake. For example, [Public First research](#) for CCIA has found that consumers generally want app store operators to review for and remove malicious apps and this may increasingly include malicious use of AI. Robust security and governance, delivered at the platform level, can deliver universal guardrails beyond the scope of individual developers, protecting users and supporting deployment by building user trust.

Companies can support digital literacy initiatives, skills and other training that allows users to engage with AI more effectively. Such education and training allows consumers to use AI more safely and therefore trust and benefit more from the services available.

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

Opting out of “AI” is not likely to be a meaningful long-term approach as AI technologies are embedded in a broad range of devices and services, e.g. email spam filtering and maps. If regulators were to require that opting out of AI was possible, it would limit the integration of AI into those devices and services, they would need to be built in such a way they could function without the AI component, and severely constrain ongoing innovation. It is also not clear what upside this would serve as risks are likely to relate to the service and how it functions as a whole, not AI specifically.

AI can be useful as a broad concept, but it necessarily elides a range of techniques and an extremely broad range of potential use cases. Within consumer-facing AI-powered services, granular user controls remain important and are already being implemented by responsible companies.

To the extent that data processing takes place based on consent and not other legal bases such as legitimate interests, the consumer may consent to the terms upon which they use a service as a whole, or to their data being processed to certain ends. There is clearly some

bundling of different data processing techniques permitted or even pre-AI services would be unable to function. Some risks may also be beyond the practical scope for many individual users to understand and are therefore better managed by more sophisticated actors, such as operators of digital ecosystems with a long-term interest in the integrity of the ecosystem.

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?

Rather than rushing to create new rules, regulators should first evaluate whether existing general or sector-specific rules address concerns relating to AI. For example, to the extent that some services pose potential risks to consumers, those might be covered by existing consumer protection laws.

To the extent that new risks are felt to warrant new interventions, the appropriate tools, framework or standards will often be sector-specific and risk-based responses to particular concerns rather than generic AI standards. However there are global higher-level frameworks that can support AI safety. To the extent that regulators can align with those standards, they are likely to reduce overall compliance costs and minimise harmful barriers to AI benefits. These international tools include ISO standards (e.g. ISO/IEC 42001) and existing regulatory work (e.g. the NIST [AI Risk Management Framework](#)).

Any regulation also needs to be considered on a holistic basis so that, for example, competition or access regulations do not create undue risks to other regulatory priorities such as privacy and security. At a more basic level, any requirements need to be technically possible for companies to comply with, which has not always been the case (e.g. the text watermarking obligations in the EU AI Act, which raises significant compliance challenges).

There is also value in regulators not duplicating work that will have been undertaken at other parts of the value chain. Model testing may have been undertaken on a voluntary basis between foundation model developers and the AISI, for example. The best way to ensure regulation is evidence-based may be to adequately resource AISI and similar bodies. Regulation of sector-specific applications developed based on those foundation models can focus on the sector-specific contextual risks associated with that specific use case. It will often not be appropriate or credible for every regulator to address every part of the stack in AI ecosystems that deliver valuable services.

20. What types / modes of guidance or other regulatory tools are most helpful to industry for achieving compliance?

Clear guidance with worked examples and other means to illustrate regulatory expectations can help industry to achieve compliance. This will be most helpful where it aligns with global standards in which industry has been able to provide the technical expertise required to ensure a standard does not create unintended risks or impacts upon AI innovation. That

way guidance does risk unnecessary duplication of compliance efforts which will undermine the dynamism of the sector as services cannot be built for a global market, but must achieve similar regulatory objectives by different means in different jurisdictions. Regulators should also consult upon and carefully review guidance to ensure that it does not inadvertently impose barriers to AI innovation that are disproportionate or exceed legislative intent.

Regulators can also engage positively with regulated companies on a bilateral basis to support them in compliance, particularly where those companies identify opportunities to comply in a way that lowers overall costs, improves outcomes for consumers and/or permits greater innovation. There can be a risk that regulators approach such engagement on a defensive basis, preferring to withhold giving a view and expecting companies to take the risks of subsequent enforcement in the event the new approach is felt not to constitute effective compliance (at a recent CCIA AI Roundtable on AI and data policy, some attendees raised this as a challenge in working with the ICO, for example). This will lead to an ossification of approaches to regulatory compliance which will be particularly costly with the regulation of technologies such as AI where innovation in compliance will need to keep pace with innovation in the underlying technology. Instead regulators should encourage such engagement and work to ensure that opportunities identified in such engagement are communicated to a broad range of regulated companies, diffusing new ideas and promoting best practice.

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

International tools include ISO standards (e.g. ISO/IEC 42001); existing regulatory work (e.g. the NIST [AI Risk Management Framework](#)); and less formal standards (e.g. model cards for disclosures). Those model cards function as a standardised, clear and accessible source of information for consumers, deployers and policymakers covering how a model was designed, its intended use cases and key limitations.

There will also be a wide range of other standards, agreements and collaborative efforts. Regulators should generally look to consult early, before sector- or use case-specific policy options are developed, to identify whether these exist and, if not, why not as the answer may be that the options are simply too nascent and it is preferable to allow companies or groups of companies to develop different options before attempting to standardise either for UK regulated or internationally-agreed standardisation.

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

Consent has important regulatory functions but it should not be misconstrued as either:

- An unrealistic expectation that consumers must understand every part of how goods and services are produced, which is obviously unrealistic in complex global value chains.
- Necessarily separate from the choice to use a particular device or service. In many cases, services are inherently a package deal and users can consent by using those services, or not.

To the extent either of these elements are neglected, consent can either become overwhelming for consumers, forcing them to consider technical distinctions that make little difference to their actual experience, and/or a poor substitute for appropriate design choices to promote user safety. These problems are particularly acute in the context of agentic AI where risks might be novel and a large part of the benefit of the service might be that constant consumer involvement is not required.

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

Generally speaking, these other levers should not be considered with respect to AI specifically, but considering the overall risks of goods or services. Cooling-off periods might be appropriate for consumers making some kinds of ongoing financial commitments, for example, but AI does not particularly make that necessary or unnecessary. Cookies-like requests should generally be seen as a warning, an example of where requirements for consent can achieve little while adding a large overall compliance burden and worsening the consumer experience.

As agentic AI becomes more prevalent, the general principle should be legal continuity: what is illegal without an agent is illegal with an agent and vice-versa. As noted above, while requiring a “human in the loop” might be appropriate in some high risk settings, in lower-risk applications such as personalisation in entertainment products that would be disproportionate and deter beneficial innovation. Requiring repeated consent for automated tasks will cause consumer fatigue, lowering the value of human involvement where it is needed and undermining the benefits of automation.

As noted above, rights to opt out of “AI” are unlikely to be a useful regulatory tool and will hold back beneficial integration in devices and services.

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

As noted above, the best approaches will support the development of global, industry-led approaches to AI safety. This spreads the cost of investment in innovation, enables mutually-beneficial trade and supports UK AI companies to scale. International tools include ISO standards (e.g. ISO/IEC 42001); existing regulatory work (e.g. the NIST [AI Risk Management Framework](#)); and less formal standards (e.g. model cards for disclosures).

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?

Generally speaking, such frameworks should respond to sector-specific risks and regulators should not assume that they are more likely to be needed in the context of AI. While they can be a means to promote positive consumer outcomes, they are necessarily burdensome and in many contexts might be disproportionate and/or counter-productive to the extent they slow investment and innovation that provide consumers with new choices, improving outcomes via innovation and competition.

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

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).

Any changes to liability rules should aim to simplify and reduce the cost of what is already often a highly complex set of rules that can constrain innovation, AI or otherwise. Attempts at reform can worsen outcomes. In the EU, an industry coalition including CCIA [called](#) for the withdrawal of the AI Liability Directive (AILD) arguing that it would add “complexity and uncertainty to an already intricate and novel legal framework impacting AI in Europe.” While some of the concerns related to interactions with other EU laws, the letter noted that the “AILD risks disrupting well-established contractual liability practices that allocate responsibilities along complex supply chains in business-to-business relationships. These practices provide businesses with the flexibility to manage liabilities and negotiate terms tailored to their specific needs, supporting both efficiency and innovation.”

Generally speaking, regulators should look to move cautiously in changing established approaches to liability, where possible make any necessary changes on an evolutionary basis, and make sure they have a thorough understanding of how and why existing approaches function, whether any challenges relating to AI are genuinely novel.

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

Managed risks must exist to the extent that the benefits of AI as a [general purpose technology](#) are potentially very large, diverse and hard to predict. Incurring some degree of risk in the pursuit of those benefits and in the context of iterative improvement will be appropriate particularly for consumers who anticipate particularly large benefits or companies pioneering novel approaches. This will also be a vital process in identifying risks and ensuring that they are appropriately mitigated on an iterative basis over time and for the wider population. Some providers, by contrast, might offer a service that minimises risks by limiting consumer use, with consumers choosing services appropriate to their needs.

Such risks might be particularly likely in the context of delivering on some regulatory objectives, e.g. identifying fraudulent content may carry a necessary risk of false positives, with some consumers losing out as a result. While companies will seek to minimise this, and AI should improve their ability to do so versus less sophisticated approaches, particularly at first there might be a rise in those false positive risks that does not indicate an overall increase in long-term risk to consumers that policymakers should seek to avert. The emphasis should be on accountability and transparency in how those risks are managed, not an unrealistic expectation that they will be avoided entirely.

Industry should manage this through an appropriate governance framework of the sort discussed earlier in this submission. Policymakers should intervene cautiously.

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

Metrics are likely to be highly use case specific, but generally speaking regulators should try to bring nuance to their analysis of trends in AI risks. In particular there may be a healthy trend over time where problems first increase, as services are exposed to more users and are applied in ways designers may not have anticipated, which then subside as those risks are managed by some combination of industry controls and user learning. Regulators have an important positive role to play in helping the public, the media and other stakeholders distinguish between genuine problems and an appropriate management of risk over time.