

WRITTEN EVIDENCE SUBMITTED BY CONNECTED BY DATA

(RAI0013)

Thank you for the opportunity to provide evidence on human rights and regulation of AI. To ensure our response is additive rather than duplicative, our response will focus on:

1. **Decision subject rights** – the rights of people and communities whose lives are affected by automated decisions, even where no personal data about them is processed
2. **Collective rights** – the rights needed to address the fact that many of the most significant AI-related human rights risks and harms occur at group or societal levels
3. **Public participation and collective representation** – the ability of communities to be involved in identifying, challenging and gaining remedies for systemic harms

Our response draws on the special report [Developing a Framework for Collective Data Rights](#)¹, based on an analysis commissioned from the law firm and consultancy AWO, and published by the Centre for International Governance Innovation.

We would also like to draw the Committee's attention to two other reports that discuss relevant questions to this inquiry, namely:

- The Ada Lovelace Institute's 2023 report on [Regulating AI in the UK](#).
- The joint briefing paper '[The Frontier AI Bill and...?](#)' from Demos, the Ada Lovelace Institute and Connected by Data, which focuses on the need for regulation and other activity beyond that focused on frontier AI such as large language models.

About [CONNECTED BY DATA](#)

We campaign to give communities a powerful say in decisions about data and AI to create a just, equitable and sustainable world through collective, democratic and open data governance. We are a non-profit founded in March 2022 with expertise on data, AI and democratic participation.

¹ Tennison, J. (2024) Developing a Framework for Collective Data Rights. *Centre for International Governance Innovation & Connected by Data*.

Human rights issues

1. *How can Artificial Intelligence (AI) affect individual human rights for good or ill, in particular in the areas of:*

- *privacy and data usage*
- *discrimination and bias*
- *effective remedies for violations of human rights?*

The [Interim Report of the Science, Innovation and Technology Committee's inquiry on the Governance of AI](#) detailed twelve challenges arising from AI, many of which are about impacts on human rights, including about bias, privacy, IP and misrepresentation.

As well as impacts that arise from poorly implemented AI (and therefore are open to technological fixes), AI is an accelerant for human rights infringements. The automation offered by AI means that organisations are able to violate human rights at scale and at speed. For example, facial recognition operating across whole crowds in real time, or AI systems that filter and flag social media, are not just invasions of privacy, but have a chilling effect on freedom of association and freedom of expression.

Existing legal and regulatory framework

2. *To what extent does the UK's existing legal framework provide sufficient protections for human rights in relation to AI?*

Traditional privacy protections, such as those under data protection law, are designed to target the use of personal data about identifiable individuals (**data subjects**). However, AI systems are frequently built using data about one group of people (over which they have data protection rights), that then have impacts on another group of people.² For instance, predictive policing tools trained on historical arrest data can disproportionately target neighbourhoods with higher proportions of ethnic minority residents, impacting on the human rights of those affected by the targeted policing (**decision subjects**) without their personal data being used. Similarly, AI systems that use aggregated statistics, satellite imagery, or environmental sensor readings to make predictions and decisions about support for local communities fall outside the scope of current data protection law.

The impacts of AI can be experienced at group and societal levels as well as by individuals³. This is particularly the case when biases in training data are embedded into AI systems that then leads to stereotyping and discrimination. A recent illustrative case is an LSE study which found that large-language-model-generated summaries of women's health records downplayed their needs in comparison to men's⁴. This kind of systemic bias is unlikely to be detected without an aggregate view of the outputs of AI systems. A case brought by an individual is unlikely to pass the threshold for meaningful harm, particularly when the AI system, as in this case, is not the ultimate decision maker. Nevertheless, in aggregate this use of AI is likely to lead to worse outcomes for women as a group.

² Viljoen, S. (2020) A Relational Theory of Data Governance. *Yale Law Journal*, Forthcoming, <https://ssrn.com/abstract=3727562> or <http://dx.doi.org/10.2139/ssrn.3727562>

³ Smuha, N.A. (2021). Beyond the individual: governing AI's societal harm. *Internet Policy Review*, **10**(3). <https://doi.org/10.14763/2021.3.1574>

⁴ Rickman, S. (2025) Evaluating gender bias in large language models in long-term care. *BMC Med Inform Decis Mak* **25**, 274. <https://doi.org/10.1186/s12911-025-03118-0>

Current data-specific regulation and redress mechanisms are primarily designed for individuals to control and challenge misuses of personal data and automated decision making. They do not cover the full range of harms that may arise from AI. While equalities, consumer and competition law may come to bear on some of them, there are specific risks that arise from the use of AI, including biased and inaccurate data, buggy implementations, and impacts on trust and the rule of law.

Even within the narrow scope of AI that is currently regulated, there are real, practical challenges in gaining redress, not least arising from a lack of transparency: people frequently do not know that they have been the subject of AI-based or AI-informed decision making, and researchers and advocacy organisations seldom if ever have access to aggregate statistics that would help them detect when systems are going wrong.

Some of these barriers may be addressed if civil society organisations were able to take legal action under data protection law without needing to identify named affected individuals. However this is not currently possible in the UK as Article 80(2) of the EU GDPR has not been implemented in the UK. As the Open Rights Group says, “Representative actions without authority would be game-changing, as it would allow public interest organisations to litigate against data abuses on behalf of society, without the hurdles of involving and being authorised on an individual basis.”⁵

3. To what extent is the Government's policy approach to deploying AI, expressed in its “AI Opportunities Action Plan”, sufficiently robust in respect of safeguarding human rights?

The AI Opportunities Action Plan contains no meaningful commitment to protecting human rights in the development and deployment of AI systems. Its only explicit reference to rights concerns intellectual property rights over the use of their content in generative AI¹. There is only a cursory mention of privacy, and no recommendations that promote equality, freedom of expression, or effective remedies for people and communities harmed by AI.

Even on its own terms, this is a critical omission. AI opportunities rely on people and organisations adopting AI; adoption is dependent on public trust and confidence; and trust is dependent not only on AI's functionality but on clear accountability, effective regulation, and visible enforcement of rights. Without robust protections, there is a high risk of public resistance and costly remediation of harms identified after deployment, as we have seen with cases like [Robodebt in Australia](#), the [Dutch child benefit scandal](#), and the [Post Office / Horizon scandal in the UK](#).

Possible changes to legal and regulatory framework

4. What would be needed in any future UK legislation to protect human rights?

- *To what extent should the same human rights standards apply to private actors as public bodies when they use AI?*
- *To what extent might different kinds of AI technology require different regulatory approaches?*

Future UK legislation should embed protections that address the rights of decision subjects; recognise the collective nature of many AI impacts; and enable collective participation,

⁵ Open Rights Group (2022), ORG Representative actions under the UK GDPR
<https://www.openrightsgroup.org/publications/org-representative-actions-under-the-uk-gdpr/>

representation, and redress. These principles should apply equally in the public and private sectors, reflecting the reality that the impacts of AI on people's rights often depend more on the nature of the decision being made, and the impacts that it can have on people and communities, than on the organisational form of the decision-maker.

There are different risk levels and types associated with the use of AI in different contexts. For example, facial recognition technology used to unlock phones has a different risk profile to that used to identify criminal suspects. Small automated food delivery vehicles have a different risk profile to self-driving cars, and to military drones.

Closer governance is particularly required around the use of AI in:

- **Public sector and essential services** such as utilities, which people cannot avoid interacting with
- Services that impact people's **wellbeing**, including their mental and financial wellbeing as well as their physical health
- Services that provide **access to opportunities** (such as jobs, training, housing and finance), including automated pricing and selection or vetting algorithms; search engines; and digital marketplaces

Effective governance of AI needs to be able to be context-aware and context-specific. This context-specificity, as well as the fast-moving nature of technology, means it is more effective to regulate AI by defining principles, processes and rights, than it is through specifying lists of safe or unsafe technologies.

While it is possible to set an overall high level framework for AI governance, the details of the deployment of particular AI systems in particular ways and particular contexts need both sector-level and case-by-case consideration, which must also include the voices and concerns of the people and communities who are affected by the use of those systems. Sector-based regulation also requires regulators who are effective and well-resourced to address AI-related risks and harms.

For example, there should be effective general governance processes for the use of AI systems in recruitment, distinct from other purposes to which AI is put, created in consultation with working-age people. It is not currently clear which regulator would develop these. Equally, there should be governance over the specific use of algorithms in recruitment in every organisation, given each will have its own particular context, processes and adopt different AI systems. These should be created in consultation with employees and prospective applicants.

Future UK legislation should therefore include:

1. **Recognition of decision subject rights** – Explicitly extend protections beyond data subjects to all people and communities whose lives are affected by AI decisions, even where no personal data about them is processed.
2. **Collective rights and representation** – Establish mechanisms for collective redress, recognising that many AI harms are only visible in aggregate (e.g. pricing discrimination, systemic bias in recruitment). Enable representative organisations to act

on behalf of affected communities without needing an individual mandate, including the ability to lodge complaints, trigger investigations, and seek remedies.

3. **Public participation in human rights impact assessments** – Require meaningful participation from affected groups in the design, development and deployment of AI. This includes providing transparent information, but also active involvement, for example in the development of proportionate, context-specific impact assessments for AI systems, especially those with potential for systemic or structural harm. Assessments should cover human rights impacts at individual, collective and societal levels, including impacts on equalities, democracy and the environment.

5. Who should be held accountable for breaches of human rights resulting from uses of AI, and on what basis?

- *Where in the process of developing, deploying and using AI technologies should liability arise?*
- *What additional measures, if any, are needed to ensure that individuals have sufficient redress where they have suffered harm because of the use of AI?*

As in health and safety legislation, accountability ultimately rests with the organisations deploying and using AI. It must be possible for those affected – whether individuals or collectively – to hold organisations accountable for the inappropriate deployment of AI products and services. This requires transparency from those deploying AI, including about any customisation, impact assessment, and evaluation of AI systems.

Regulation should also place requirements on other actors in the AI supply chain – such as AI suppliers – to enable those deploying AI to carry out due diligence checks which will give them confidence in its adoption and in taking on the associated liability. Suppliers should be held liable if they misrepresent the functionality of their products and services or mislead AI adopters about its impacts.

6. How might regulation match the pace of AI technology development, such as the emergence of agentic AI, to ensure that human rights are preserved as technology continues to develop?

The speed of AI innovation, and the fact that impacts are highly context-specific, means that regulation cannot rely solely on static rules tied to specific technologies. Instead, regulation must focus on the principles underpinning human-rights-respecting AI; the processes by which AI systems are designed, developed, deployed, and monitored; and the rights of individuals and communities.

To ensure that AI developers and deployers understand potential impacts, as well as to build trust and confidence, regulation should embed participation of people and communities affected as a fundamental principle, and specify the processes through which this is realised such as by having formal requirements for the involvement of affected people and communities in identifying problems to be solved; through impact assessment, design and testing; to ongoing evaluation.

Early, meaningful involvement of affected communities is the most effective way to anticipate and avoid adverse impacts. Those with lived experience of the context in which AI will be deployed can:

- identify risks that designers and deployers may overlook
- suggest mitigations that work in practice
- highlight unintended consequences before they become systemic harms

This approach has two benefits in the face of rapid technological change:

- It allows governance to adapt to emerging risks without waiting for new legislation.
- It builds public trust by demonstrating that AI systems are shaped with, rather than imposed on, the people they affect.

7. How could regulation take account of the international nature of AI? How could it address the potential consequences for human rights in the UK of the malign use of AI by regimes in other countries?

There is an assumption that AI must be governed like large social media platforms, with complex, cross-jurisdictional enforcement challenges. While this is true for certain centralised foundation models and global AI services, most impactful uses of AI are, and will be, deployed locally:

- in specific workplaces, schools, housing associations, local authorities, and NHS Trusts
- using smaller, specialised language models or domain-specific tools
- with decisions tailored to particular populations and contexts

This provides for local governance levers such as procurement rules, sectoral regulation, and domestic human rights law that can be applied without waiting for international consensus.

(Sep 2025)