



House of Lords
House of Commons

Joint Committee on Human Rights

Human Rights and the Regulation of AI

Fourth Report of Session 2026–27

HC 160 / HL Paper 56

Joint Committee on Human Rights

The Joint Committee on Human Rights is appointed by the House of Lords and the House of Commons to consider matters relating to human rights in the United Kingdom (but excluding consideration of individual cases); proposals for remedial orders, draft remedial orders and remedial orders. The Joint Committee has a maximum of six Members appointed by each House, of whom the quorum for any formal proceedings is two from each House.

Current membership

House of Lords

[Lord Alton of Liverpool](#) (Crossbench; Life peer)

[Baroness Chakrabarti](#) (Labour; Life peer)

[Baroness Hamwee](#) (Liberal Democrat; Life peer)

[Lord Murray of Blidworth](#) (Conservative; Life peer)

[Lord Rook](#) (Labour; Life peer)

[Lord Sewell of Sanderstead](#) (Conservative; Life peer)

House of Commons

[Alex Sobel](#) (Labour; Leeds Central and Headingley) (Chair)

[Tom Gordon](#) (Liberal Democrat; Harrogate and Knaresborough)

[Afzal Khan](#) (Labour; Manchester Rusholme)

[Euan Stainbank](#) (Labour; Falkirk)

[Peter Swallow](#) (Labour; Bracknell)

[Sir Desmond Swayne](#) (Conservative; New Forest West)

Powers

The Committee has the power to require the submission of written evidence and documents, to examine witnesses, to meet at any time (except when Parliament is prorogued or dissolved), to adjourn from place to place, to appoint specialist advisers, and to make Reports to both Houses. The Lords Committee has power to agree with the Commons in the appointment of a Chairman.

Publication

This Report, together with formal minutes relating to the report, was Ordered by the House of Lords and House of Commons, on 9 September 2026, to be printed. It was published on 14 September 2026 by authority of the House of Lords and House of Commons. © Parliamentary Copyright House of Commons 2026.

This publication may be reproduced under the terms of the Open Parliament Licence, which is published at www.parliament.uk/copyright.

Committee reports are published on the Committee's website at www.parliament.uk/jchr and in print by Order of the House of Lords and House of Commons

Contacts

All correspondence should be addressed to the Clerk of the Joint Committee on Human Rights, House of Commons, London SW1A 0AA. The Committee's email address is ichr@parliament.uk.

Follow the work of committees on Bluesky [@HOCcommittees.parliament.uk](#), Facebook [UK House of Commons Committees](#), Instagram [@UKCommonsCommittees](#), LinkedIn [House of Commons Committees](#), Reddit [u/UKCommonsCommittees](#), X [@HOCcommitteesUK](#) and YouTube [Commons Committees playlist](#).

Contents

Summary	1
Introduction	3
1 What is artificial intelligence?	5
Definition of artificial intelligence	5
Overview of the AI lifecycle	8
Foundation models	8
2 Human rights issues raised by AI	12
Council of Europe Framework Convention	12
Key risks to human rights	14
Equality and non-discrimination	14
Privacy and data protection	17
The right to an effective remedy	20
Other rights	21
3 Legal and policy framework	25
Main relevant legal provisions	25
Discrimination	25
Data protection	26
Sector-specific requirements	28
Obligations of public bodies	29
Enforcement	30
Redress	31
Role of government bodies	32
Government's position	34
Overall approach	34
Approach to regulation of AI	35

4	Problems with existing protections	38
	Before AI systems are deployed	39
	Design and development	39
	Pre-market evaluation and testing by regulators	41
	When AI systems are deployed	42
	Lack of transparency	42
	The difficulty of identifying the cause of human rights harms	44
	Lack of effective remedy	46
	Limits of human oversight	46
	Gaps in protection and legal uncertainty	47
	Oversight after AI systems are deployed	48
	Uncertain legal liability	48
	Rapid multiplication of risks	49
	Limits of regulatory action	51
	A fragmented regulatory landscape?	52
	Resourcing of regulators	54
	Technical expertise of regulators	54
	Individual redress via regulators	55
	Does regulation stifle or support innovation?	55
	International nature of issue	56
5	How risks to human rights should be addressed	58
	An AI Bill and AI-specific regulation	58
	A risk-based approach	59
	Allocating responsibility where it belongs	62
	Prohibitions	64
	Prior approval for high-risk activities	65
	Obligations of due diligence	66
	Transparency	67
	Automated decision-making	68
	An AI regulator	69
	The AI Security Institute (AISI)	73

Conclusions and recommendations	75
Formal minutes	85
Witnesses	87
Published written evidence	89
List of Reports from the Committee during the current Parliament	93

Summary

Artificial Intelligence systems are being used by businesses, individuals and public authorities in ways that affect every aspect of people's lives in the UK. The technology is developing at a pace that means governments, and human rights protections, risk being left behind. Our report sets out why this is a problem and what must be done to tackle it.

AI systems can also bring real human rights benefits, and have the potential to create economic growth. The approach of the UK government to AI has been to focus on the growth benefits, although it has also taken important steps to reduce some risks, especially around national security.

However, we heard clear evidence that AI presents novel and serious human rights risks. AI systems are opaque, so that it is not clear why they have produced the results they generate—the so-called “black box”. People are also often not aware that AI is being used in a way that affects them and their rights. These factors further increase the likelihood of harm. Because of the scale and speed of AI use, when these harms occur, they may happen at a huge scale and very quickly.

These risks are particularly acute for groups that are already minoritised, such as Black and Minority Ethnic people. We focus on the real damage AI systems may do to people's rights to equality and non-discrimination, to privacy and data protection, and to the right to an effective remedy when people are harmed. The government must act to protect these fundamental rights.

We heard about the human rights impact of AI used to create deepfake sexualised images of women and girls; to profile prisoners in ways which disproportionately affect Black and Minority Ethnic people; to flag up workers for disciplinary action without reasonable cause; and to scan and recognise people's faces in public places without their consent. Many more uses are being created as technology develops.

There are many laws and regulations that may already apply to some AI systems, but overall the legal landscape is patchy and confused. There are many regulators who are charged with tackling some AI harms, but no single body co-ordinates regulation of AI. This leaves gaps which mean people will be affected by AI harms, but will not be able to get redress.

There is also far too much freedom for the large and powerful technology companies who develop AI systems to pass on liability to those who deploy them. These deployers are often less powerful, have less access to resources, and are not in the best position to identify risks or prevent harm occurring. The law does not target those who actually create and shape the systems. Responsibility to prevent harm should sit with those who are best able to do so.

We call on the government to take immediate action. It must bring in a dedicated AI Bill, and set up a new regulator, to create tailored, effective protection for human rights in the UK. The most dangerous uses of AI should be banned, and others should be regulated in a principled and risk-based way, with strong sanctions for companies and organisations who damage human rights through their use of AI systems. None of this need compromise the use of AI to help generate economic growth, where it can be done safely. As one witness told us, “The only growth worth having is one that protects human rights.”

Introduction

1. We launched our inquiry in July 2025 in the context of wide public debate and concern about the growth and diversification of Artificial Intelligence ('AI'). Our aim was to investigate the extent to which existing law and regulation was sufficient to protect human rights in the UK in the context of AI.
2. Important work has been undertaken or is being carried out in related fields by other Committees, including the House of Lords Select Committee on Artificial Intelligence, the House of Lords Justice and Home Affairs Committee, the House of Commons Science, Innovation and Technology Committee, the House of Lords AI Weapons Committee, and others. All touch on human rights impacts and have raised important issues. Because AI is a general purpose technology, we felt that our Committee, with its cross-cutting remit, was especially well placed to consider the wide-ranging impact of AI on human rights.
3. We received over 70 pieces of written evidence, and are grateful to all those who took the time to share their views, experiences and expertise with us. Over the course of ten oral evidence sessions we heard from senior academics, researchers, science communicators, civil society campaigners, policy bodies, legal experts, technology specialists and entrepreneurs, major technology firms, regulators charged with overseeing elements of AI use, and the Minister for AI. We are especially grateful to our Special Adviser, Professor Karen Yeung, Interdisciplinary Professorial Fellow in Law, Ethics and Informatics in the Birmingham Law School and the School for Computer Science at the University of Birmingham, for her advice and guidance.
4. In order to clearly set out the case for change, this report starts by explaining what artificial intelligence is and why we believe it raises vital human rights issues. We then discuss the existing legal and regulatory protections, and explain why we consider they are seriously lacking. Finally, we set out our proposals for comprehensive, tailored law and regulation to address these risks and protect human rights.
5. As our inquiry drew to a close, a new Prime Minister took office. The new administration acted immediately to make clear the importance it places on AI. Kanishka Narayan MP was appointed Minister of State (Minister for Artificial Intelligence) jointly in the Cabinet Office and the Department for Business, Innovation, Science and Trade on 20 July 2026, having previously

been Parliamentary Under-Secretary of State at the then Department for Science, Innovation and Technology, and Minister for AI Opportunities. He will attend Cabinet.¹

6. The Committee considers that this increased prominence of AI within government offers an opportunity to drive economic growth while respecting rights, and makes our report timely. The continuity in this key portfolio also encourages us to believe that the evidence we took, and the conclusions we drew from it, remain relevant and useful to policy-making and public debate. We look forward to continuing to work with the Minister in his expanded role.
7. Finally, since our inquiry finished taking evidence, there have been many high-profile developments which illustrate the importance of the issues we examined. No report can keep pace with the accelerating change created by AI in our economy and society, and we may revisit this area in future.

¹ <https://www.gov.uk/government/news/ministerial-appointments-july-2026>

1 What is artificial intelligence?

Definition of artificial intelligence

8. Although the meaning of the term ‘Artificial Intelligence’ has changed over time, the House of Commons Library recently described it as a shorthand for “technologies that perform the types of cognitive functions typically associated with humans, including reasoning, learning and solving problems.”² Many AI technologies are underpinned by ‘machine learning,’ through which a software algorithm finds patterns in existing data (known as ‘training data’) to create a statistical model (an ‘AI model’) which can then be applied to new data to make predictions or generate other outputs based on those predictions.³ AI systems are machine systems that, broadly speaking, perceive, learn and act.⁴ As a general purpose technology, AI is now widely used across for a diverse range of purposes (or ‘use cases’), ranging from email spam filters through to powerful tools for scientific discovery.⁵
9. AI systems can be broadly conceptualised into two main categories:
 - a. Task-specific machine learning AI systems, which perform a specific task.
 - b. Generative AI systems built on an AI model (‘foundation model’) trained with a large amount of data using self-supervision at scale, that displays significant generality and can competently perform a wide range of distinct tasks and can be integrated into a variety of downstream systems and applications.

2 <https://commonslibrary.parliament.uk/research-briefings/cbp-10003/>
3 <https://post.parliament.uk/research-briefings/post-pb-0057/>; <https://post.parliament.uk/artificial-intelligence-ai-glossary/>; Royal Statistical Society. [AI is Statistics](#). 2 March 2026.
4 United Nations, [Preliminary Report of the Independent International Scientific Panel on AI: Evidence-based assessments of the opportunities, risks and impact of artificial intelligence](#) (2026).
5 Dada et al [“Machine learning for email spam filtering: review, approaches and open research problems”](#) (2019); Wang et al [“Scientific discovery in the age of artificial intelligence”](#) (2023)

10. AI systems are now being incorporated into more complex AI ‘agents’. These are small, specialised pieces of software that can make decisions and operate cooperatively or independently to achieve pre-specified objectives. Agentic AI systems (‘agentic AI’) are composed of several AI agents which are designed to behave and interact autonomously, working together as a single system, to achieve the system’s objectives.⁶
11. A fundamental aspect of current AI systems is their capacity to operate with varying degrees of ‘autonomy’ and adaptiveness. Unlike traditional automated systems, AI systems generate outputs which are not determined in advance by their programming, so there is an unpredictable relationship between input and output.⁷ Some AI models, including foundation models, operate in ways that are so mathematically complex that their outputs are opaque, even to an informed expert. As a result, it is not possible to understand how these AI systems produce outputs (known as the ‘interpretability’ of the system) or to explain why these systems produced a particular output in a specific case (known as the ‘explainability’ of that output).⁸
12. Other significant characteristics of AI systems include their ability to operate automatically without human involvement, the speed at which they can operate, the scale at which they are being deployed,⁹ and their capacity to produce outputs which credibly imitate human behaviour.¹⁰ We heard that

6 <https://www.gov.uk/government/publications/ai-insights/ai-insights-agentic-ai-html>

7 OECD, [Recommendation of the Council on Artificial Intelligence](#) (OECD/LEGAL/0449); OECD, [Explanatory memorandum on the updated OECD definition of an AI system](#), March 2024; European Commission, [Guidelines on the definition of an artificial intelligence system established by the EU AI Act](#), C (2025) 924 final, (14)-(21); Department for Science, Innovation and Technology, [A pro-innovation approach to AI regulation](#), 3 August 2023, section 3.2.1; UK Jurisdiction Taskforce’s [Legal Statement on Liability for AI Harms under the private law of England and Wales](#), July 2026, paragraphs 7–12

8 Explanatory Report to the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, paragraphs 60–61. The concept of ‘opacity’ can refer to (1) opacity as intentional corporate or state secrecy, (2) opacity arising from lack of technical understanding, or (3) opacity arising from the characteristics of machine learning algorithms and the scale required to apply them usefully (Burrell, [“How the machine ‘thinks’: Understanding opacity in machine learning algorithms”](#), 6 January 2016). The term is used here in the third sense.

9 Council of Europe, [A study of the implications of advanced digital technologies \(including AI systems\) for the concept of responsibility within a human rights framework](#) (MSI-AUT(2018)05Rev), 22 May 2019.

10 DeVrio et al. [“AI Automotons: AI Systems Intended to Imitate Humans.”](#) The 2026 ACM Conference on Fairness, Accountability, and Transparency. 2026.

because of these characteristics, AI presents novel risks.¹¹ Professor Kevin Fong told us:

[W]e need a different approach because of the uniqueness of this new technology [...] It is fast; it has a direct impact on populations at scale and very quickly; it has unique properties that we have not seen before in any other technology, and those principally revolve around its ability to play the role of a convincing human actor, even if it has no human component to it.¹²

13. Other witnesses also told us of some of the benefits that could be accrued through AI use, if it is done responsibly.¹³ We heard these included possible benefits to efficiency, consistency of decision making, improved healthcare,¹⁴ increased environmental protection by reducing inefficiencies in transport and shipping; improved human safety by better disaster and conflict prediction; more efficient deployment of scarce resources; enhanced provision of protection and remedies for victims of human rights violations, including children affected by trafficking and exploitation;¹⁵ the use of AI-enabled captioning and sign-language tools to enhance access to public services;¹⁶ enhanced expression and creativity;¹⁷ and improved public participation in governance.¹⁸ We are also aware that AI-enabled surveillance and monitoring systems can be deployed for the purposes that are intended to help prevent crime and enhance public safety, reflected in claims by some English police forces that the use of live facial recognition technology is of great benefit to them in carrying out their role of preserving the peace, apprehending wanted individuals and seeking to prevent violence and disorder.

11 United Nations, [Preliminary Report of the Independent International Scientific Panel on AI: Evidence-based assessments of the opportunities, risks and impact of artificial intelligence](#) (2026). Bengio et al. [International AI Safety Report 2026](#). <https://ico.org.uk/about-the-ico/research-reports-impact-and-evaluation/research-and-reports/technology-and-innovation/tech-horizons-and-ico-tech-futures/ico-tech-futures-agentic-ai/>; Digital Regulation Co-operation Forum, “The Future of Agentic AI”, published 31 March 2026, <https://www.drcf.org.uk/siteassets/drcf/pdf-files/final-drcf-the-future-of-agentic-ai-foresight-paper.pdf?v=423791>

12 [Q27](#)

13 Dr Felipe Romero-Moreno ([RAI0087](#))

14 ACT | The App Association ([RAI0043](#)); Equality and Human Rights Commission ([RAI0075](#))

15 Trilateral Research ([RAI0046](#)); Dr Giulia Gentile (Lecturer at Essex Law School); Dr Matthew Gillett (Senior Lecturer at Essex Law School); Professor Geoff Gilbert (Professor at Essex Law School); Professor Audrey Guinchard (Professor at Essex Law School) ([RAI0047](#))

16 European Center for Not-For-Profit Law (ECNL) ([RAI0008](#))

17 Google ([RAI0084](#))

18 Duality Ltd ([RAI0071](#))

14. What all witnesses agreed on was the ever-growing use and significance of AI, raising fundamental challenges of profound importance. Academics from Essex Human Rights Centre and Essex Law School told us:

The inner functioning, processes, and resources underlying AI raise unprecedented legal questions, and bear intrinsic risks for individuals and the community at large directly impacting fundamental legal entitlements, such as human rights. The regulation of this technology is one of the planetary challenges of this century.¹⁹

Professor Ethan Mollick of Generative AI Labs and Wharton University of Pennsylvania, told us:

there are very few aspects of society that are not being touched [by AI]. It has moved from a peripheral, interesting, useful tool to being central to almost everything that we do.²⁰

Overview of the AI lifecycle

15. A central concept when considering how AI systems are developed and deployed is the 'AI lifecycle'. The phrase refers to the various phases through which the provision of an AI system proceeds.²¹ In general terms, the AI lifecycle begins with planning and design, data collection and processing, development of an AI model (which may include training a model on very large datasets), fine-tuning an AI model, testing, verification and validation, through to launch, deployment, on-going monitoring and maintenance and, finally, retirement and decommissioning.

Foundation models

16. Foundation models, also known as 'general-purpose' AI models, are computer programs that process inputs to perform a range of tasks, such as text synthesis, image manipulation and audio generation and may be used to form the core of larger AI systems.²² Developing a foundation AI model typically requires very significant capital resource, large quantities of relevant data, highly specialised expertise, major compute capability, an abundant and reliable energy supply and major data storage capacity. As a

19 Dr Giulia Gentile (Lecturer at Essex Law School); Dr Matthew Gillett (Senior Lecturer at Essex Law School); Professor Geoff Gilbert (Professor at Essex Law School); Professor Audrey Guinchard (Professor at Essex Law School) ([RAI0047](#))

20 [Q45](#)

21 Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, CETS 225, Vilnius, 5 September 2024.

22 Competition and Markets Authority, AI Foundation Models Initial Report, published 18 September 2023, https://assets.publishing.service.gov.uk/media/650449e86771b90014fdab4c/Full_Non-Confidential_Report_PDF.pdf

result, only a handful of organisations have the resources to build state-of-the-art AI models that serve as the foundation for AI systems designed for specific tasks and purposes.²³

17. Developers typically make their foundation models widely available, although this is entirely at their discretion.²⁴ AI systems built on foundation models can be directly available to consumers as standalone AI systems (such as OpenAI’s GPT, Anthropic’s Claude, Google’s Gemini and Meta’s Llama models) or used as the ‘foundation’ by other businesses to build context-specific AI applications. The launch of foundation models from late 2022 onwards has made AI system development more affordable and accessible, provoking an explosion in their provision and adoption.²⁵ For example, the Ada Lovelace Institute gives examples of private-sector companies that have built AI-enabled products and services on OpenAI’s GPT foundation models,²⁶ including Microsoft’s BingChat (for internet search and retrieval),²⁷ Duolingo Max and Khan Academy’s Khanmingo (for personalised education),²⁸ and Be My Eyes’ Be My AI (for vision assistance).²⁹

AI supply chains

18. Early AI systems were developed by a single organisation which collected its own data, trained its own models and deployed its own infrastructure. With the advent of foundation models, however, a modern AI system is now typically the result of contributions from many organisations and follow a similar pattern. A very large dataset serves as the ‘raw material’ to train a foundation model. An AI developer can then produce a purpose-specific AI-system (for example, a doctor-patient visit summarisation app) by fine-tuning the pre-trained foundation AI model on a specialised (e.g. hospital)

23 <https://www.gov.uk/government/publications/frontier-ai-capabilities-and-risks-discussion-paper/frontier-ai-capabilities-and-risks-discussion-paper>

24 Ian Brown, ‘Allocating accountability in AI supply chains: a UK-centred regulatory perspective.’ Ada Lovelace Institute (2023) p 3. Cobbe et al “[Understanding accountability in algorithmic supply chains](#)” (2023). Schrepel and Pentland “[Competition between AI foundation models: dynamics and policy recommendations](#)” (2025) Models may be made available following a commercial judgment, or for altruistic reasons. There is a debate on the security and other considerations around making models available, which there is not space to address here.

25 Brown, “[Allocating accountability in AI supply chains: a UK-centred regulatory Perspective](#)” (2023). Ada Lovelace Institute.

26 https://www.adalovelaceinstitute.org/evidence-review/foundation-models-public-sector/#_ftn19

27 <https://blogs.microsoft.com/blog/2023/02/07/reinventing-search-with-a-new-ai-powered-microsoft-bing-and-edge-your-copilot-for-the-web/>

28 <https://blog.duolingo.com/duolingo-max/>; <https://blog.khanacademy.org/harnessing-ai-so-that-all-students-benefit-a-nonprofit-approach-for-equal-access/>

29 <https://www.bemyeyes.com/news/introducing-be-my-ai-formerly-virtual-volunteer-for-people-who-are-blind-or-have-low-vision-powered-by-openais-gpt-4/>

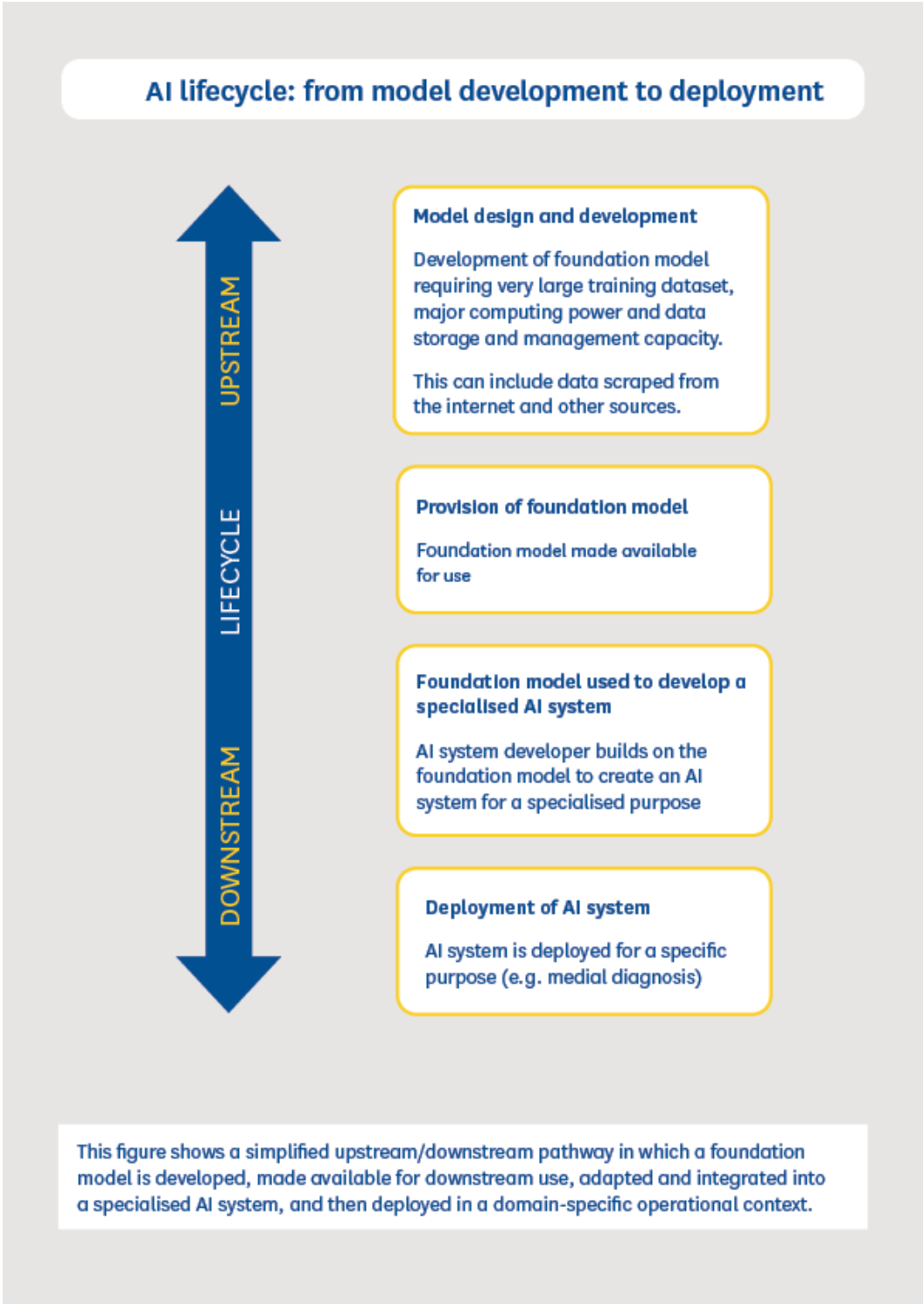
dataset. Multiple actors supply resources and expertise, such as AI models, datasets and other AI system components which are interdependent and make up the AI system's supply chain.³⁰

19. The AI development process is complex, opaque, rapid and dynamic. Across an AI system's supply chain, different lifecycle phases and aspects of its design, development, distribution and deployment are spread across multiple actors. Every AI system will have a different supply chain, with variations depending on the sector, the use case, whether the system is developed in-house or procured, and how the system is distributed (for example, through an application programming interface ('API') or through a hosted platform).³¹ Global, multi-tiered supply chains are a complex web of relationships where multiple actors contribute to enable an AI system to function. This has made it increasingly difficult to identify the source of risk, and the actor primarily responsible for it, if human rights harm subsequently arises from the deployment of an AI system.³²

30 <https://www.adalovelaceinstitute.org/resource/explainer-foundation-models/> Hopkins et al "AI Supply Chains: An Emerging Ecosystem of AI Actors, Products, and Services"(2025); Cobbe et al "Understanding accountability in algorithmic supply chains" (2023).

31 Brown, "Allocating accountability in AI supply chains: a UK-centred regulatory Perspective" (2023). Ada Lovelace Institute. Irene Solaiman (2023) 'The gradient of generative AI release: Methods and considerations.' In Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency, pp. 111-122.

32 Cobbe et al "[Understanding accountability in algorithmic supply chains](#)" (2023).



20.

CONCLUSION

AI supply chains are complex and opaque, involving multiple actors. This makes it difficult to identify the source of human rights risks, and the actors responsible for identifying and addressing those risks.

2 Human rights issues raised by AI

Council of Europe Framework Convention

21. AI has the potential to enhance the protection of human rights, but can also adversely affect them. A useful analysis of the risks to human rights which have provoked the greatest concern is provided by the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law (the Framework Convention),³³ which was agreed following an extensive process of engagement in response to those concerns.³⁴
22. The Framework Convention is the first legally binding international instrument targeting AI, although it has not yet entered into force. Seventeen states and organisations have signed the Convention so far, including the UK, the USA, the EU and many European states.³⁵ It therefore represents a widely accepted international understanding of how AI systems might adversely affect human rights, and the steps that should be taken to ensure appropriate respect for those rights. The Framework Convention aims to ensure that activities within the lifecycle of artificial intelligence systems are fully consistent with human rights.³⁶
23. Parties to the Convention must adopt or maintain measures to:
- respect human dignity and individual autonomy;

33 [Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law](#), CETS 225, Vilnius, 5 September 2024

34 Council of Europe, [Ad hoc Committee on Artificial Intelligence](#) (CAHAI); [Committee on Artificial Intelligence \(CAI\)](#).

35 Council of Europe, [The Framework Convention on Artificial Intelligence](#). Signatories so far are the UK, the European Union, Andorra, Armenia, Bosnia and Herzegovina, Georgia, Iceland, Liechtenstein, Montenegro, Norway, Moldova, San Marino, Switzerland, Ukraine, Canada, Israel, Japan, the USA and Uruguay. The Convention needs 5 ratifications to enter into force (Article 30(3)). On 15 May 2025 the EU ratified the Framework Convention ([European Union ratifies the Council of Europe Framework Convention on Artificial Intelligence](#)).

36 Article 1(1). See also the OECD [Recommendation of the Council on Artificial Intelligence](#) (OECD/LEGAL/0449), paragraph 1.2.

- ensure adequate transparency and oversight;
- ensure accountability and responsibility;
- ensure that activities respect equality and non-discrimination;
- ensure privacy and personal data rights are protected;
- promote reliability; and
- enable safe innovation.³⁷

24. AI Minister Kanishka Narayan told us: “as a country, we want to make sure that AI is grounded in human rights, democracy and law.”³⁸ The then Department for Science, Innovation and Technology (DSIT) wrote to us to affirm the government’s intention to ratify the Convention.³⁹ At the time of writing, however, no timetable or plan has been set out for when the UK will ratify the Convention, or for what steps might need to be taken to achieve this. AI Minister Kanishka Narayan told us “I feel no significant obstacles in us being able to put [the Convention] into effect. It is just that the programme of work required to do so [...] is pretty extensive [...] we are working hard on that.”⁴⁰

25. The Law Society told us the UK government “has an opportunity to be the first state to ratify [the Convention] and continue its leading role in establishing internationally recognised norms in AI regulation.”⁴¹ Justice, the legal reform charity, called for the government to publish and to consult publicly on a ratification plan.⁴²

26. CONCLUSION

The Council of Europe Framework Convention on AI provides a strong basis for both effective regulation of AI within the UK, and collective international action. We urge the government to seize the opportunity to be an early adopter of the Framework Convention.

27. RECOMMENDATION

The UK should set out a timeline for ratification of the Council of Europe Framework Convention detailing what steps will be taken and when. This plan for ratification should be subject to public consultation.

³⁷ [Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law](#), Articles 7–13.

³⁸ [Q101](#)

³⁹ Department of Science Innovation and Technology ([RAI0077](#))

⁴⁰ [Q101](#)

⁴¹ The Law Society of England and Wales ([RAI0021](#))

⁴² Justice ([RAI0073](#))

Key risks to human rights

28. There are many ways in which AI may adversely affect human rights. In particular, witnesses told us of its potential impact on the treatment and conditions of workers involved in the labelling of data and the extraction of raw materials, workers' rights generally,⁴³ and on the environment.⁴⁴
29. For the purposes of this report, however, the Committee has concentrated on three of the key principles identified in the Framework Convention as being fundamental to the protection of human rights in this area: equality and non-discrimination; privacy and personal data protection; and the right to an effective remedy. These elements reflect some of the most tangible and prevalent risks currently presented by a wide variety of AI systems, which we heard are already impacting people in the UK. They were also cited by many of our witnesses as the most salient concerns.⁴⁵

Equality and non-discrimination

30. Article 10 of the Framework Convention requires parties to adopt or maintain measures to ensure that activities within the AI lifecycle respect the principles of equality and non-discrimination, as well as to achieve fair, just and equitable outcomes. The explanatory report refers to “the real and well-documented risk of bias that can constitute unlawful discrimination” arising from activities within the AI lifecycle, and explains that parties to the Convention must consider appropriate solutions to address the different ways in which bias can be incorporated into AI systems.⁴⁶
31. We heard that intentional abuses committed using AI systems, such as deepfake intimate-image-based abuse, can have a disproportionate effect on people from minoritised groups, “compounding [...] existing structural and systemic racism and sexism that impacts victims’ personal and professional lives.”⁴⁷ The Safety not Surveillance Coalition (made up of OpenRights Group, Stopwatch and the Runnymede Trust) told us “AI is having a chilling effect on communities of colour in the UK.”⁴⁸

43 Dr Joe Atkinson (Associate Professor of Employment Law at University of Southampton) ([RAI0080](#))

44 Professor Lorna McGregor (Professor of International Human Rights Law at University of Essex) ([RAI0056](#))

45 European Center for Not-For-Profit Law (ECNL) ([RAI0008](#))

46 Explanatory Report to the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, paragraphs 75–78

47 [Q28](#), Joy ([RAI0001](#))

48 Open Rights Group, StopWatch, Runnymede Trust ([RAI0018](#))

- 32.** Discrimination can also arise unintentionally because of the way AI systems are developed. We heard about how bias is built into AI systems. AI models are trained on large quantities of data as part of their development process. These large datasets often relate to human activity and decision making and therefore inevitably contain “baked in” human-generated biases and can be unrepresentative.⁴⁹ When an AI system is based on an AI model trained with biased data, it can replicate or even exacerbate these biases, leading to discriminatory outcomes that unfairly affect various groups, particularly those which have historically been marginalised.⁵⁰ Unintended bias can also arise from design choices made in the development and configuration of AI systems which reflect the biases of their developers.⁵¹
- 33.** Large technology firms such as Google told us of the extensive efforts they make to prevent and address bias in their AI models.⁵² However, we heard from other witnesses that it is very difficult to achieve this.⁵³ Adjusting data used to train an AI system relies on the decision making of the developer and hence will be subject to an inherent bias.⁵⁴
- 34.** Discrimination can arise when AI systems are deployed because of these “baked in” biases, and also due to choices made by those deploying AI systems. Examples of use cases where biased AI systems could cause, or already have caused, harms and which were cited by our witnesses included tools used for casting in the creative industries,⁵⁵ to support education in

49 [Qq28–33](#); Public Law Project ([RAI0045](#)); Big Brother Watch ([RAI0036](#)); Catholic Bishops’ Conference of England and Wales ([RAI0040](#)); [Q8](#); Dr Malak Benslama-Dabdoub (Lecturer in Law at Royal Holloway University of London) ([RAI0010](#))

50 Department of Science Innovation and Technology ([RAI0077](#)); Fair Vote UK ([RAI0002](#)); Dr Mark Wong (Senior Lecturer and Head of Social and Urban Policy at University of Glasgow) ([RAI0014](#)); Gina Romero (Special Rapporteur for the rights to freedom of assembly and of association at United Nations - OHCHR); Prof. Pete Fussey (Head of Department. Department of Sociology, Social Policy and Criminology at University of Southampton); Dr. Murray Daragh (Senior Lecturer, School of Law at Queen Mary University of London) ([RAI0023](#)); Dr Erin Ferguson (Lecturer in Public Law at University of Aberdeen); Liam Moorhouse (PhD Researcher at University of Aberdeen); Professor Abbe Brown (Professor in Intellectual Property Law at University of Aberdeen); Dr Clare Frances Moran (Lecturer in Law at University of Aberdeen); Professor Irene Couzigou (Professor of Public International Law at University of Aberdeen); Muhammed Nurullah Benli (PhD Researcher at University of Aberdeen); Dr Kevin Allan (Lecturer in Psychology at University of Aberdeen); Andrew Walters (PhD Researcher at University of Aberdeen); Dr Andrew Starkey (Reader in Engineering at University of Aberdeen); Dr Xuan Tung Le (Researcher at University of Aberdeen) ([RAI0048](#)); Glenlead Centre ([RAI0069](#)) Liberty ([RAI0079](#))

51 [Q10](#)

52 [Q68](#)

53 Good Tech Advisory ([RAI0086](#))

54 The Law Society of England and Wales ([RAI0021](#)); [Q19](#); [Q28](#); [Q63](#)

55 Equity ([RAI0055](#))

schools,⁵⁶ and to help assess security categorisation and the risk of violence in prisons.⁵⁷ In each case witnesses told us minoritised groups would be the ones most at risk of harm. The Law Society of England and Wales told us it was already aware of existing cases, on which solicitors were advising, relating to the use of biased data to train AI systems used in public-sector decision making. They told us “this is not a hypothetical problem for the future but one that requires immediate attention to protect human rights.”⁵⁸ AI Minister Kanishka Narayan told us that within government, when AI was used in the public sector, “there was very specific focus on ensuring compliance with the equality duty and avoiding any discriminatory outcomes.”⁵⁹ However as described above, we heard from other witnesses that this was difficult to achieve in practice.

35. Specific examples of AI apparently causing discriminatory outcomes that we heard about include an Uber Eats driver who could not log on at work because the AI facial recognition software used, which was worse at identifying Black people’s faces, did not recognise him;⁶⁰ a system used by Durham Constabulary to assess the risk of recidivism which we were told placed a heavy reliance on individuals’ postcodes and locations, potentially leading to these factors having a significant impact on prosecution decisions;⁶¹ the well-publicised discontinuance by Amazon Ltd of the use of an AI recruitment tool when it was found to discriminate against women applicants;⁶² and use by employers of members of the Communication Workers’ Union of AI tools which it was stated had flagged for investigation workers who had stopped at traffic lights while driving, or had failed to “upsell” products while working in call centres.⁶³
36. We also heard evidence that the degree of certainty which attaches to the conclusions of an AI system can be overstated, risking unfairness or abuse. Javier Ruiz Diaz of Amnesty International UK told us: “Eligibility criteria and risk profiling has an effect on people when they get their benefits cut or they are stopped by the police, but [the system] does not say “we think this is happening”; it says “we know this is happening”.⁶⁴”

56 Dr Ayça Atabey; Dr Kim Sylwander; Professor Sonia Livingstone ([RAI0058](#)); see also 5Rights Foundation ([RAI0074](#))

57 The Howard League for Penal Reform ([RAI0068](#)), Prison Reform Trust ([RAI0041](#))

58 The Law Society of England and Wales ([RAI0021](#))

59 [Q107](#)

60 [Q83](#); See EHRC, [Uber Eats courier wins payout with help of equality watchdog, after facing problematic AI checks](#).

61 [Q28](#)

62 <https://www.reuters.com/article/world/insight-amazon-scraps-secret-ai-recruiting-tool-that-showed-bias-against-women-idUSKCN1MK0AG/>, accessed 25 July 2026

63 Eleonor Duhs (Barrister, Partner, Head of Data & Privacy at Bates Wells LLP) ([RAI0035](#))

64 [Q28](#)

37. The use of AI systems can further entrench and compound discrimination. Once biased systems are deployed, there is a risk that the AI's own outputs create a 'feedback loop' where the AI system produces data which further reinforces its own biased learning.⁶⁵ The then DSIT told us that AI systems trained on incomplete, incorrect, biased, or unrepresentative data can replicate and exacerbate biases, leading to discriminatory outcomes in critical sectors, and can for example potentially worsen pre-existing health inequalities.⁶⁶

38. **CONCLUSION**

AI systems are prone to producing unfairly discriminatory outcomes, which may arise from bias in the datasets on which they are trained as well as from the way they are developed and deployed. This threatens the human rights of people in the UK and elsewhere.

Privacy and data protection

39. Article 11 of the Framework Convention requires parties to adopt or maintain measures to ensure that privacy rights and personal data are protected within the AI lifecycle. The right to privacy is protected by Article 8 of the European Convention on Human Rights, and is given particular expression in data protection law in relation to personal data.⁶⁷ The right to privacy is qualified, which means interferences with that right can be justified if prescribed by law and necessary in a democratic society, meaning that the interference must pursue a legitimate aim in a proportionate manner.⁶⁸
40. Navigating trade-offs between the right to privacy and competing rights and other legitimate public interests (including the interests of public safety, the prevention of disorder or crime, and the protection of health) has produced a considerable body of case law revealing how the courts determine the appropriate balance between competing rights and interests

65 [Qq28-33](#), Amnesty International UK ([RAI0051](#))

66 Department of Science Innovation and Technology ([RAI0077](#))

67 [Explanatory Report to the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law](#), paragraphs 79-83

68 Art 8(1) the European Convention on Human Rights provides that "[e]veryone has the right to respect for his private and family life, his home and his correspondence". However, Art 8(2) of the Convention further provides that "there shall be no interference by a public authority with the exercise of this right except such as is in accordance with the law and is necessary in a democratic society in the interests of national security, public safety or the economic well-being of the country, for the prevention of disorder or crime, for the protection of health or morals, or for the protection of the rights and freedoms of others."

concerning digital and data-driven technologies.⁶⁹ Those cases illustrate that making these trade-offs in a lawful manner requires careful and sometimes finely balanced judgement, suggesting that actors involved in the AI lifecycle may find it difficult to make these trade-offs when developing and deploying AI systems. Dr Sarah Jane Fox, an academic at University of Leicester Law School, told us that “there remains a delicate balance between leveraging AI’s benefits, addressing its inherent restrictions and offering protections to ensure (and even build) public trust, particularly in policing and across law enforcement authorities.”⁷⁰

41. AI systems, particularly those which are built on AI foundation models, including large language models (LLMs), require huge volumes of data to feed into training processes that enable the development of AI system capabilities. This data is gained from a range of sources, and many developers rely on publicly available online sources,⁷¹ which can include personal data. Collecting this data is sometimes referred to as “web scraping”. We heard this had often been done without sufficient consideration of the implications for individuals.⁷²
42. The use of AI by both state and private actors can damage the right to privacy and the right to protection of personal data. We have heard a number of examples of this. AI systems can be configured to collect personal data from individuals without their awareness or cooperation, including biometric data, and to use that data to identify individuals.⁷³ A

69 For example, in *R (Bridges) v Chief Constable of South Wales Police* [2020] EWCA 1058, at [54]-[130], the Court of Appeal held that the South Wales Police Force had breached Article 8 of the ECHR when using live automated facial recognition technology, because there was not a sufficiently clear legal framework governing use of the technology. In *Glukhin v Russia*, App. No. 11519/20, judgment of 4 July 2023, at [58]-[91], the European Court of Rights held that there had been a breach of Article 8 when the applicant was filmed by CCTV cameras, identified by facial recognition technology, and then convicted of an administrative offence for holding a peaceful solo demonstration without prior notification. It was doubtful that the domestic legal provisions were sufficiently clear, and in any case it was disproportionate to use highly intrusive technology to locate and arrest someone for an administrative offence involving the exercise of the right to free expression. A considerable body of case law has been developed by the European Court on Human Rights from adjudicating disputes involving alleged interferences with the right to private life by the European Court of Human Rights. [Fact Sheet ‘New Technologies’, October 2024](#); European Court of Human Rights. [Fact Sheet ‘Personal Data Protection’, February 2024](#).

70 Dr S. J. Fox (Academic (Law)); Associate Director, Institute for Digital Culture at University of Leicester) ([RAI0009](#))

71 Information Commissioner’s Office, [Generative AI first call for evidence: The lawful basis for web scraping to train generative AI models](#) (closed 1 March 2024)

72 Ada Lovelace Institute ([RAI0066](#))

73 UNISON ([RAI0076](#)), Eleonor Duhs (Barrister, Partner, Head of Data & Privacy at Bates Wells LLP) ([RAI0035](#)) Big Brother Watch ([RAI0036](#)); Dr Joe Atkinson (Associate Professor of Employment Law at University of Southampton) ([RAI0080](#))

particularly clear example of risk is offered by facial recognition technology, specifically live facial recognition technology, which has been increasingly used by police forces across England and Wales. We heard that between 1 January 2025 and 29 October 2025, around 3 million people would have had their faces scanned by police facial recognition technology in the UK without their consent.⁷⁴ Ms. Gina Romero, United Nations Special Rapporteur on Freedom of Peaceful Assembly and of Association, told us the use of facial recognition technology posed “expanding risks.”⁷⁵

43. The availability of AI systems also increases the risk that individuals can be re-identified from apparently anonymous data, especially when they can draw on multiple datasets to allow the combination of different potentially identifying factors. We heard this may “allow actors to obtain more private information about people than previously possible,”⁷⁶ further risking the right to privacy. The Ada Lovelace Institute told us that “a new generation of equally invasive inferential [...] technologies claims to infer sensitive internal states like a person’s emotions, intentions, attention or truthfulness—often with low levels of scientific validity.”⁷⁷
44. Within government, safeguards are being applied. AI Minister Kanishka Narayan told us “we have set out new guidelines for AI-ready government datasets [...] focused on how responsible development takes place using those datasets. They require human oversight at key stages to protect privacy rights, a duty to prevent bias, and that people, rather than just machines, are accountable for decisions that matter as well.”⁷⁸

45. **CONCLUSION**

AI systems pose particular risks to privacy and personal data, because of the prevalence of web scraping to collect personal data to train large language models (LLMs), and the potential for AI systems to be used in highly intrusive ways. This threatens the right to privacy of people in the UK and elsewhere.

⁷⁴ [Q28](#)

⁷⁵ Gina Romero (Special Rapporteur for the rights to freedom of assembly and of association at United Nations - OHCHR); Prof. Pete Fussey (Head of Department. Department of Sociology, Social Policy and Criminology at University of Southampton); Dr. Murray Daragh (Senior Lecturer, School of Law at Queen Mary University of London) ([RAI0023](#))

⁷⁶ Ada Lovelace Institute ([RAI0066](#)) UNISON ([RAI0076](#)); Eleonor Duhs (Barrister, Partner, Head of Data & Privacy at Bates Wells LLP) ([RAI0035](#)); Big Brother Watch ([RAI0036](#)) Dr Joe Atkinson (Associate Professor of Employment Law at University of Southampton) ([RAI0080](#)). Ada Lovelace Institute, ‘[An Eye on the Future - A legal framework for the governance of biometric technologies in the UK](#)’ (May 2025).

⁷⁷ Ada Lovelace Institute ([RAI0066](#))

⁷⁸ [Q109](#)

The right to an effective remedy

46. Article 14(1) of the Framework Convention requires parties to adopt or maintain measures to ensure the availability of accessible and effective remedies for violation of human rights arising from the activities within the lifecycle of AI systems.⁷⁹ We heard that when AI systems are deployed to support decision-making that have legal or other significant effects on individuals without the knowledge or awareness of the individual affected, this can present significant risks to the right to an effective remedy. The EHRC told us that “[a] significant challenge for effective remedy is ensuring individuals know that AI has been used and that it can be challenged. Transparency in AI use is essential to support challenge and remedy.”⁸⁰ This concern was emphasised by Ellen Lefley of JUSTICE in her oral evidence, who told us that “You cannot mount redress to an AI-related harm if you do not know that AI has been used in the first place. You do not know what you are asking for. You do not know what you are looking for.”⁸¹
47. We also heard that a lack of transparency by public bodies deploying AI systems was a particular concern, because it adversely affected the ability of affected individuals to challenge decisions by public bodies, affecting the right to an effective remedy.⁸² Even if individuals are made aware of the use of an AI system in these circumstances, we heard that they may not be provided with enough information to understand the basis of the output in order to mount an effective challenge.⁸³
48. Witnesses also expressed concern that the opacity of the outputs, and the complex nature of AI systems, can undermine the capacity of individuals affected by them to challenge those outputs or seek a remedy.⁸⁴ Duality

79 Article 14(2) of the Convention further requires that each Party shall adopt or maintain measures to (a) ensure that relevant information regarding artificial intelligence systems which have the potential to significantly affect human rights and their relevant usage is documented, provided to bodies authorised to access that information and, where appropriate and applicable, made available or communicated to affected persons (b) ensure that the information referred to in subparagraph a is sufficient for the affected persons to contest the decision(s) made or substantially informed by the use of the system, and, where relevant and appropriate, the use of the system itself; and (c) an effective possibility for persons concerned to lodge a complaint to competent authorities.

80 Equality and Human Rights Commission ([RAI0075](#))

81 [Q41](#)

82 The Law Society of England and Wales ([RAI0021](#)); Public Law Project ([RAI0045](#)); [Q30](#)

83 Ada Lovelace Institute ([RAI0066](#))

84 Dr Giulia Gentile (Lecturer at Essex Law School); Dr Matthew Gillett (Senior Lecturer at Essex Law School); Professor Geoff Gilbert (Professor at Essex Law School); Professor Audrey Guinchard (Professor at Essex Law School) ([RAI0047](#)); Chartered Institute of Library and Information Professionals (CILIP) ([RAI0053](#)); Dr Ayça Atabey; Dr Kim Sylwander; Professor Sonia Livingstone ([RAI0058](#)). Minderoo Centre for Technology and Democracy, University of Cambridge ([RAI0067](#)); Mr Burak Batuhan Karakus (Head of Legal Affairs at Human Rights Solidarity) ([RAI0026](#))

Ltd, an UK-based AI company, told us that “[t]he opacity (‘black box’) nature of many AI systems can make it difficult for individuals to seek justice where their rights are violated. If a person is denied a service or penalised based on an AI decision, they may not even know AI was involved, let alone how it reached its decision. This poses a serious challenge to the right to a fair hearing and an effective remedy.”⁸⁵

49. **CONCLUSION**

AI systems pose particular risks to the right to an effective remedy in circumstances where there is insufficient transparency concerning the use of AI systems to support significant decisions about affected persons, and in circumstances where it is difficult for those persons to access mechanisms to challenge those decisions. This threatens the human rights of people in the UK.

Other rights

50. While the rights raised below were not the focus of our inquiry, they illustrate the breadth of the challenge to human rights potentially posed by the use of AI.

Intellectual property rights

51. The inquiry received several submissions which raised the issue of copyrighted and professional artistic works being integrated into datasets for AI systems, and other issues relating to the rights of creators and performers.⁸⁶ Article 1 of the First Protocol to the ECHR protects the right to property, including intellectual property,⁸⁷ and in some circumstances imposes positive obligations on states to ensure that property rights are protected.⁸⁸ There is extensive legislation in the UK to provide such protection, but there is so far limited case law about the extent to which use of protected material to train AI systems might breach intellectual property law.
52. The Independent Society of Musicians told us: “the right to own, control, and benefit from one’s work is essential to sustaining both innovation and cultural production.”⁸⁹ Equity, a creative industries trade union, provided a case study of a member who agreed to the use of her voiceover

85 Duality Ltd ([RAI0071](#))

86 The Independent Society of Musicians ([RAI0004](#)); Equity ([RAI0055](#)); Chartered Institute of Library and Information Professionals (CILIP) ([RAI0053](#)); Leena Sowambur (Founder at Marcomms By Leena) ([RAI0063](#)); British Copyright Council ([RAI0027](#))

87 [Anheuser-Busch v Portugal](#), App. No. 73049/01, judgment of 11 January 2007, at [66]-[72].

88 [Kotov v Russia](#), App. No. 54522/00, judgment of 3 April 2012, at [109]-[115].

89 The Independent Society of Musicians ([RAI0004](#))

performance “for non commercial purposes to help visually impaired readers to access educational content.” According to Equity, the member later found that a cloned version of her voice had been “made available to others via a text-to-speech tool [...] without [her] consent, control or pay,” risking both her commercial livelihood and her professional standing, “as her voice may be used to say things which may damage her reputation and career.”⁹⁰

53.

CONCLUSION

Our inquiry has not focused on intellectual property rights, or the rights of performers. The issues which witnesses have raised clearly merit further investigation, and could require specific measures going beyond the scope of this inquiry. It is important for the UK to protect its rich creative industries and those who contribute to them, and to safeguard performers’ human rights as well as their livelihoods.

Right to a fair trial

54. The increasing use of AI systems in court proceedings could have implications for the right to a fair trial under Article 6 of the ECHR. The Committee heard that although there were potential benefits,⁹¹ there could also be harmful consequences of using misleading, unreliable or biased AI-generated outputs. The former DSIT told us these could lead “people to make poor decisions about their rights or cases [...] or be discouraged from seeking proper legal support.” Commentators have highlighted a risk that the volume and poor quality of AI generated submissions to courts presented by litigants in person, which may appear sophisticated but contain major flaws or “hallucinated” cases with no basis in fact, often described as ‘AI slop’ could also put already stretched court systems under further strain, affecting the operation of the system overall and hence access to justice.⁹² Ellen Lefley of JUSTICE told the Committee there was a risk that “AI for lawyers [...] makes their lives easier, but the people who really need [...] assistance do not get any and the access to justice gap widens.”⁹³ We also heard that the availability of AI might serve as a fig leaf for reduced investment in the court system, thereby indirectly frustrating access to justice.⁹⁴

90 Equity ([RAI0055](#))

91 The Law Society of England and Wales ([RAI0021](#))

92 [In depth: AI slop - law’s quiet epidemic | Law Gazette](#)

93 [Q42](#)

94 [Q42](#)

55.

CONCLUSION

Our inquiry did not focus on the potential effects of AI on the justice system, but witnesses raised some very concerning issues which merit detailed consideration, and might require specific measures going beyond the scope of this inquiry. The system is already under great strain. We note the Justice Select Committee has also recently discussed with ministers the use of AI in connection with the legal system, as part of its ongoing inquiry into Access to Justice.

Wider risks to human rights

56. The Committee heard a range of views regarding the potential future risks which could be anticipated if AI systems were to develop Artificial General Intelligence (AGI).⁹⁵ The long term risks that AI may pose to humanity have caused concern in the sector for several years. In 2023, many senior figures including Professor Geoffrey Hinton, Professor Yoshua Bengio and the CEOs of Google DeepMind and OpenAI signed a statement published by the Centre for AI Safety: “Mitigating the risk of extinction from AI should be a global priority alongside other societal-scale risks such as pandemics and nuclear war.”⁹⁶
57. Several witnesses to our inquiry highlighted similar concerns.⁹⁷ Professor Roman Yampolskiy, Director of the Cybersecurity Laboratory at the University of Louisville, recommended a ban on the creation of uncontrolled general superintelligence, which he characterised as a national security threat.⁹⁸ He told us the risks posed by this technology are such that “the only way to win is not to play the game.”⁹⁹ Other witnesses, however, described any threat from AGI as “remote.”¹⁰⁰
58. In the shorter term, Dr Iulian Serban of LawZero told the Committee that AI systems are already displaying disturbing behaviours, including sycophancy, deception, scheming, power seeking behaviour and a general misalignment with human values and goals.¹⁰¹ We were told of the risks

95 ‘Artificial General Intelligence’ (AGI), broadly speaking, refers to an AI system with general, human-level (or beyond) ability to learn, reason, and apply knowledge across a wide range of tasks and domains. However, there is no clear, accepted definition of AGI in either policy or scientific literature. Xu [‘What is meant by AGI? On the definition of Artificial General Intelligence’](#) (2024)

96 Centre for AI Safety, [AI Extinction Statement Press Release | CAIS](#), published 30 May 2023, accessed 27 August 2026.

97 Andrea Miotti (Founder and CEO at ControlAI); Steven Adler (Machine Learning Researcher and Practitioner formerly at OpenAI ([RAI0031](#)))

98 [Q52](#)

99 [Q50](#)

100 [Q96](#)

101 [Q53](#)

arising from the emergence of agentic AI, which may take unauthorised actions that could have unforeseen and potentially serious systemic consequences.¹⁰² Another risk is suggested by reports of instances of AI chatbots encouraging individuals, especially young or otherwise vulnerable people, to self-harm or suicide, or to acts of violence.¹⁰³ We also received evidence that AI support chatbots may be inadequate in terms of the advice they offer to distressed children, and fail to raise safeguarding alerts.¹⁰⁴ Kanishka Narayan MP, the AI minister, assured us the government was taking self-harm and suicide content online and in the use of AI chatbots extremely seriously and was taking steps to address those risks.¹⁰⁵

59. Our inquiry did not focus on specific ways to address these risks but they are of course of the utmost concern. We will continue to monitor developments in this area closely and may return to the topic in the future.

60. **CONCLUSION**

We welcome the government's stated commitment to ensuring that in the UK, development and use of AI systems is grounded in human rights. AI has the potential to enhance the protection of human rights, but can also give rise to significant interferences with those rights.

61. **RECOMMENDATION**

Human rights principles must inform all aspects of the government's policies and practices relating to AI systems, including its approach to the regulation of AI.

62. **RECOMMENDATION**

Government should encourage the development of AI use cases which can demonstrably improve human rights.

102 [Q13](#) We do not address here questions around where liability would reside for such an action.

103 5Rights Foundation ([RAI0074](#))

104 Dr Ayça Atabey, Dr Kim Sylwander, Professor Sonia Livingstone, Digital Futures for Children centre, LSE ([RAI0058](#))

105 [Q101](#)

3 Legal and policy framework

- 63.** The above chapter sets out some of the ways in which AI may pose a risk to human rights. This chapter examines how existing legal provisions can apply to AI in the UK in order to protect those rights, and the government’s current AI policy.

Main relevant legal provisions

- 64.** Although there are few provisions in UK law specifically concerning the development and deployment of AI systems, existing law provides some protection against interferences with human rights resulting from AI if specific conditions are met. Those laws were generally not formulated with AI in mind, and there is limited case law about how they apply to AI, so there are many questions about their effects that are yet to be resolved.

Discrimination

- 65.** Deployers of AI systems must comply with the Equality Act 2010 in circumstances covered by the Act, such as employment,¹⁰⁶ or the provision of services.¹⁰⁷ In these circumstances, direct or indirect discrimination in relation to any of the protected characteristics covered by the Act is prohibited.¹⁰⁸ For instance, AI systems used in recruitment must not discriminate on the basis of race, sex, age, disability or other protected characteristics.¹⁰⁹

¹⁰⁶ Equality Act 2010, section 39.

¹⁰⁷ Equality Act 2010, section 29.

¹⁰⁸ Equality Act 2010, section 4.

¹⁰⁹ See the guidance published by the Department for Science, Innovation and Technology, [Responsible AI in Recruitment](#), 25 March 2024.

66. The Act only applies to persons concerned with the provision of a service to the public or a section of the public. It does not apply to businesses providing products or services to other businesses, which includes those who provide AI systems to organisational deployers, or to developers further up the AI supply chain.¹¹⁰

Data protection

67. Personal privacy is given legal protection in specific circumstances. When the development or deployment of AI systems involves processing personal data, the person responsible for the processing must comply with data protection law which may impose legal obligations at various stages of the supply chain.¹¹¹ For instance:
- AI developers training AI systems which process personal data must comply with data protection requirements if specific individuals can be identified from the data.¹¹²
 - If the deployment of an AI system will entail the processing of personal data that entails high risks to the rights and freedoms of natural persons, then the proposed deployer has a legal obligation to carry out a data protection impact assessment before so doing.¹¹³
 - Data protection law requires transparency when personal data is being processed, including obligations to be transparent with data subjects about what that processing entails. This obligation applies

110 Equality Act 2010, section 29(1); [Dockers' Labour Club v Race Relations Board](#) [1976] AC 285 at 297D-E; [Johnston v Giving.com](#), Case No 226MC647, 2 November 2022, at [47]-[56]; [\[2024\] EWHC 944 \(KB\)](#) at [16]-[18]. Depending on the circumstances, it may be possible for a person to bring a claim under the Equality Act 2010 against the deployer rather than the developer in respect of discrimination arising from actions taken by the developer. Compare *R (Bridges) v Chief Constable of South Wales Police* [2020] EWCA 1058, at [163]-[202].

111 [Regulation \(EU\) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data \(United Kingdom General Data Protection Regulation\)](#), as it forms part of the law of England and Wales, Scotland and Northern Ireland by virtue of section 3 of the European Union (Withdrawal) Act 2018 (UK GDPR); Data Protection Act 2018

112 UK GDPR, Article 4(1)

113 UK GDPR, Articles 35–36. Data controllers must carry out a data protection impact assessment if “a type of processing in particular using new technologies, and taking into account the nature, scope, context and purposes of the processing, is likely to result in a high risk to the rights and freedoms of natural persons” (UK GDPR, Article 35(1)). See also the ICO’s guidance on [Data Protection Impact Assessments \(DPIAs\)](#).

if the training of an AI model entails the processing of personal data, and if deployment of an AI system entails the processing of personal data.¹¹⁴

68. Additional obligations apply to those deploying AI systems for the purposes of automated decision-making.¹¹⁵

- Data protection law imposes requirements on decisions based solely on automated processing of personal data, where there is no meaningful human involvement in the taking of the decision.¹¹⁶ The Secretary of State has the power to make regulations defining what is or is not to be taken to be “meaningful human involvement” in the taking of a decision.¹¹⁷
- The restriction on automated decision-making applies to processing of special categories of personal data,¹¹⁸ which is only permitted in certain circumstances.¹¹⁹ Other types of personal data may be used for the purposes of automated decision-making as long as safeguards are in place, including that the data subject is told about the decisions, is

114 UK GDPR, Articles 12–14; ICO guidance on Artificial Intelligence: Explaining decisions made with AI: Part I The basics of explaining AI: [Legal framework](#).

115 UK GDPR, Article 22A–22D; Data Protection Act 2018, sections 50A–50D (in relation to processing for law enforcement purposes). RAI0077 (DSIT), paragraphs 49–51.

116 UK GDPR, Article 22A(1); Data Protection Act 2018, section 50A(1) (in relation to processing for law enforcement purposes). As the ICO’s guidance points out, where an AI-assisted decision is not part of a solely automated process because there is meaningful human involvement, processing of personal data will still be subject to the general requirements of the UK GDPR, including the principles of fairness, transparency and accountability (Information Commissioner’s Office, UK GDPR Guidance and Resources: Artificial Intelligence: Explaining decisions made with AI: Part I The basics of explaining AI: [Legal framework](#)).

117 UK GDPR, [Article 22D\(1\)](#); Data Protection Act 2018, section 50D(1)

118 UK GDPR, Article 22B(1); Data Protection Act 2018, section 50B(1). ‘Special categories’ of personal data are defined in Article 9 of the UK GDPR to mean “personal data revealing racial or ethnic origin, political opinions, religious or philosophical beliefs, or trade union membership, and the processing of genetic data, biometric data for the purpose of uniquely identifying a natural person, data concerning health or data concerning a natural person’s sex life or sexual orientation”.

119 If the data subject has given explicit consent (UK GDPR, Article 22B(2); Data Protection Act 2018, section 50B(2)), or if automated decision-making is necessary for a contract or required or authorised by law, provided that the processing is necessary for reasons of substantial public interest, proportionate, and subject to safeguards (UK GDPR, Article 22B(3); see also Data Protection Act 2018, section 50B(3) (in relation to processing for law enforcement purposes), which requires only that the decision is required or authorised by law).

able to make representations, can obtain human intervention, and can contest the decision.¹²⁰ The Secretary of State may make regulations clarifying or adding to the safeguards.¹²¹

Sector-specific requirements

- 69. There are also sector-specific legal obligations which apply to the use of AI in certain circumstances.
- 70. For instance, the Online Safety Act 2023 applies to providers of specific internet services that are regulated by the Act, requiring them to identify, mitigate and manage risks of harm.¹²² The Act does not refer to AI as such, but it covers providers who use AI as part of the regulated service.¹²³ The Crime and Policing Act 2026 gives the Secretary of State power to extend the provisions of the Online Safety Act to AI services.¹²⁴
- 71. Other sector-specific legislation may apply when AI systems are deployed, for instance if AI is used in the provision of financial services,¹²⁵ or incorporated into medical devices.¹²⁶

120 UK GDPR, Article 22C; Data Protection Act 2018, section 50C.

121 UK GDPR, Article 22D; Data Protection Act 2018, section 50D.

122 The Online Safety Act 2023 applies to providers of ‘user-to-user services’ (section 3(1)) and ‘search services’ (section 3(4)). The Act does not directly restrict what content can be made available online, which is a matter for the general criminal law. It does, however, require service providers to comply with specified requirements to identify, mitigate and manage risks of harm from illegal content and content that is harmful to children (sections 7–34). RAI0077 (DSIT), paragraphs 53–55.

123 Ofcom, [Open letter to UK online service providers regarding Generative AI and chatbots](#), 8 November 2024.

124 [Crime and Policing Act 2026](#), section 248 (in force on Royal Assent), inserting new section 216A into the Online Safety Act 2023.

125 Financial Conduct Authority, [AI Update](#), 2024; Bank of England, [DP5/22 - Artificial Intelligence and Machine Learning](#), 11 October 2022.

126 Medicines and Healthcare products Regulatory Agency, [Impact of AI on the regulation of medical products](#), April 2024, page 6

Obligations of public bodies

72. Public bodies are subject to various obligations which will apply when they deploy AI systems. In particular:

- Public authorities must comply with the Human Rights Act 1998, which requires them to ensure their activities are compatible with the European Convention on Human Rights, including when they deploy AI systems.¹²⁷
- They must comply with the public sector equality duty under the Equality Act 2010 when considering whether and how to deploy AI systems.¹²⁸
- They have a particular obligation to take into account statements of government procurement policy when procuring AI systems.¹²⁹
- They must also comply with principles of public law, for instance when a decision made by an AI system constitutes the decision of a public body, or when they decide how to deploy AI systems in carrying out their functions.¹³⁰

127 E.g. in [R \(Bridges\) v Chief Constable of South Wales Police](#) [2020] EWCA 1058, at [54]-[130], the Court of Appeal held that the South Wales Police Force had breached Article 8 of the ECHR when using live automated facial recognition technology, because there was not a sufficiently clear legal framework governing use of the technology.

128 Equality Act 2010, section 149; EHRC guidance on [The Public Sector Equality Duty \(PSED\)](#). In [R \(Bridges\) v Chief Constable of South Wales Police](#) [2020] EWCA 1058, at [163]-[202], the Court of Appeal held that the South Wales Police Force had also breached the public sector equality duty when using live automated facial recognition technology, because they did not seek to satisfy themselves, either directly or by way of independent verification, that everything reasonable had been done to make sure the software did not have a racial or gender bias.

129 Under the Procurement Act 2023, section 13(9), contracting authorities must have regard to the national procurement policy statement published by the Cabinet Office. The [National Procurement Policy Statement](#) of 13 February 2025 says (at page 5) that contracting authorities should apply commercial best practice, including the principles and policies in the Government's Playbook series. [The Digital, Data and Technology Playbook](#) of June 2023 refers to the [Guidelines for AI procurement](#) of 8 June 2020, which set out principles for procuring AI systems such as considering the benefits and risks, establishing appropriate oversight mechanisms, maximising transparency, encouraging explainability and interpretability, and ensuring ethical use throughout the lifecycle of the AI system.

130 See the Law Commission discussion paper, [AI and the Law](#), 31 July 2025, especially pages 16–19; Lord Sales, [AI and Public Law: Automated Decision-Making in Government](#), 5 November 2025; Government Digital Service, [Artificial Intelligence Playbook for the UK Government](#), 10 February 2025, page 62; McGurk and Tomlinson, "Artificial Intelligence and Public Law", 1st edition, Hart 2025.

Enforcement

73. A number of statutory regulators have powers to enforce the law in these areas.¹³¹ In particular:

- The Equality and Human Rights Commission (EHRC) has various enforcement powers in relation to the Equality Act 2010,¹³² and can also institute or intervene in legal proceedings under the Human Rights Act 1998.¹³³
- The Information Commissioner (ICO)¹³⁴ has extensive powers to enforce compliance with data protection law,¹³⁵ including powers to obtain information, carry out investigations, issue warnings and reprimands, require compliance, limit or ban processing, impose fines, and provide advice and opinions.¹³⁶
- Ofcom is required to produce guidance and codes of practice about how it will apply the Online Safety Act 2023.¹³⁷ It also has extensive powers to require information, carry out investigations, issue notices, make decisions, impose penalties, and apply to the court for orders disrupting the business of non-compliant service providers.¹³⁸

131 In 2024 the government [wrote](#) to 13 regulators with responsibilities relating to AI: the Information Commissioner's Office (ICO), the Equality and Human Rights Commission (EHRC), the Office of Communications (Ofcom), the Competition and Markets Authority (CMA), the Financial Conduct Authority (FCA), the Bank of England, the Medicines and Healthcare products Regulatory Agency (MHRA), the Office for Nuclear Regulation (ONR), the Health and Safety Executive (HSE), the Office for Standards in Education, Children's Services and Skills (Ofsted), the Office of Qualifications and Examinations Regulation (Ofqual), the Legal Services Board (LSB), and the Office of Gas and Electricity Markets (Ofgem).

132 Equality Act 2006, sections 20–32. With the exception of the power to institute or intervene in legal proceedings, these powers relate only to breaches of the Equality Act 2010, not to breaches of the Human Rights Act 1998: see the definition of “unlawful” in section 34(1).

133 Equality Act 2006, section 30.

134 Under Part 6 of the [Data \(Use and Access\) Act 2025](#), not yet in force, the office of the Information Commissioner will be abolished and replaced by a body corporate called the Information Commission.

135 UK GDPR, Articles 57–58; Data Protection Act 2018, section 115.

136 The ICO has published an AI and biometrics strategy for 2025–26: [Preventing harm, promoting trust: our AI and biometrics strategy](#), 25 June 2025.

137 See [Online safety regulatory documents and guidance](#).

138 Online Safety Act 2023, Chapter 6.

Redress

74. In some circumstances, the law also enables those who suffer harm to seek financial compensation through the courts for loss caused by use of AI systems.¹³⁹

- If an individual's personal data has been processed in breach of their data protection rights, they are entitled to claim compensation.¹⁴⁰ The victim of discrimination may be able to claim compensation under the Equality Act 2010.¹⁴¹
- If there is a contract between the deployer of an AI system and the person suffering loss, there might be a claim for breach of contract, depending on its terms. Within AI supply chains, in the absence of legislation, the parties to contracts can choose how legal responsibility should be allocated between them.¹⁴²
- If an AI system is deployed in a careless manner causing harm to an individual, there could be liability in negligence under the common law (for instance, if a professional carelessly relied on AI-generated output when providing a service).¹⁴³ But it might not be straightforward to demonstrate who, if anyone, was responsible for causing the harm.¹⁴⁴
- There could be other specific causes of action in some circumstances. For instance, if an AI system (such as a chatbot) makes a defamatory statement, a claim might be available to the person who has been

139 Department of Science Innovation and Technology ([RAI0077](#)) paragraphs 29–30. The legal position is summarised in the UK Jurisdiction Taskforce's [Legal Statement on Liability for AI Harms under the private law of England and Wales](#), July 2026.

140 UK GDPR, Article 82; Data Protection Act 2018, sections 167–169. See also ICO guidance, [Taking your case to court and claiming compensation](#). Claims can be for material or non-material damage, including distress: UK GDPR, Article 82(1); Data Protection Act 2018, section 168(1). Data controllers and processors are not liable if they can prove they were not in any way responsible for the event giving rise to the damage: UK GDPR, Article 82(3).

141 Equality Act 2010, Part 9;

142 See e.g. the [report](#) of the Financial Markets Law Committee on Private Law Issues in AI, 27 October 2025, page 10.

143 UK Jurisdiction Taskforce's [Legal Statement on Liability for AI Harms under the private law of England and Wales](#), July 2026, paragraphs 55–66. Other torts might also be relevant, for instance if an AI chatbot makes a defamatory statement (paragraphs 164–184). There could also be strict liability under the Consumer Protection Act 1987 if AI is incorporated into a defective physical product which cause physical harm (paragraphs 70–80).

144 See the Law Commission discussion paper on [AI and the Law](#), 31 July 2025; the UK Jurisdiction Taskforce's [Legal Statement on Liability for AI Harms under the private law of England and Wales](#), July 2026, especially paragraphs 81–131; RAI0077 (DSIT), paragraph 30.

defamed.¹⁴⁵ If AI is incorporated into a defective physical product which causes physical harm, strict liability could arise under the Consumer Protection Act 1987 without the need to show fault.¹⁴⁶

- If an AI system takes action which would constitute a criminal offence when taken by a person, there could be criminal liability.¹⁴⁷ But if an AI system independently performs the act when it has not been directed to do so, it may be difficult to identify a person with the mental state needed to give rise to a criminal offence.¹⁴⁸

Role of government bodies

- 75.** Parts of government currently engage in activities concerned with AI systems, although they have no formal regulatory powers.
- 76.** The AI Security Institute (AISi) was established in November 2023.¹⁴⁹ It was set up as a research organisation within the former Department for Science, Innovation and Technology but has now moved to the Cabinet Office,¹⁵⁰ to conduct research on some of the most powerful foundation models to help understand their capabilities, test the effectiveness of safety guardrails, and help inform governments about AI safety.¹⁵¹ In February 2025, the AI Safety Institute was renamed the AI Security Institute, with a renewed sole focus on the risks AI poses to national security and serious crime, rather than more general safety concerns, and the removal from its stated agenda of references to algorithmic bias. The Government's stated

145 UK Jurisdiction Taskforce's [Legal Statement on Liability for AI Harms under the private law of England and Wales](#), July 2026, paragraphs 164–184.

146 UK Jurisdiction Taskforce's [Legal Statement on Liability for AI Harms under the private law of England and Wales](#), July 2026, paragraphs 70–80. The Law Commission is carrying out a review of product liability, which will address whether this principle should be extended to pure software that is not part of a tangible product (Law Commission review of product liability, [terms of reference](#), 8 December 2025, paragraph 4(1)). The EU Product Liability Directive ([Directive \(EU\) 2024/2853 of the European Parliament and of the Council of 23 October 2024 on liability for defective products](#)) to be implemented in Northern Ireland under the Windsor Framework by December 2026, updates EU product liability law by clarifying that all types of software, including AI systems, are covered by the law on liability for defective product (Article 4(1) and recital 13).

147 There are offences specifically relating to abuse of AI systems. The Crime and Policing Act 2026 created new offences relating to matters such as child sexual abuse image generators (Sexual Offences Act 2023, section 46A, inserted by section 72 of the 2026 Act) and purported intimate image generators (Sexual Offences Act 2023, section 66I, inserted by section 99 of the 2026 Act). Many general offences may also be committed online, including through use of AI systems. Schedule 7 to the Online Safety Act 2023 lists 'priority offences' under other legislation to which the Act's regulatory provisions apply.

148 See the Law Commission discussion paper [AI and the Law](#), 31 July 2025, pages 11–12.

149 AISi, [About](#), accessed 17 April 2026

150 [Machinery of Government changes: Fact Sheet](#), 22 July 2026 (updated 27 July 2026)

151 AISi, [Our first year](#), accessed 17 April 2026

aim in making this change was to “strengthen [...] protections against the risks AI poses to national security and crime”.¹⁵² AI Minister Kanishka Narayan also told us it may have helped the AISI “attract particular types of talent.”¹⁵³ AISI’s research currently focuses on emerging AI risks with serious security implications.¹⁵⁴ AISI engages collaboratively with developers of the most powerful foundation models (“frontier models”) who allow AISI access to their models to undertake some pre-deployment testing.¹⁵⁵ These arrangements are voluntary. Although AI minister Kanishka Narayan told us that AISI “has some of the world’s most exceptional talent in being able to assess models,”¹⁵⁶ AISI has no regulatory authority and cannot, for example, demand access to foundation models either before or after their deployment or prevent them from being released.¹⁵⁷

- 77.** The Government Digital Service has developed the Algorithmic Transparency Reporting Standard (ATRS),¹⁵⁸ which is mandatory for all government departments, as well as for arm’s-length bodies which provide public or frontline services, or routinely interact with the general public.¹⁵⁹ They are required to publish a record setting out how they use algorithmic tools (including but not limited to tools incorporating AI) which either have a significant influence on a decision-making process with public effect, or directly interact with the general public.¹⁶⁰

152 DSIT press release, [Tackling AI security risks to unleash growth and deliver Plan for Change](#), 14 February 2025

153 [Q104](#)

154 Department of Science Innovation and Technology ([RAI0077](#))

155 These include the potential for AI to help users develop chemical and biological weapons, cause major cyberattacks and create artificial child sexual abuse material, RAI0077, DSIT

156 [Q104](#)

157 [Correspondence from Kanishka Narayan MP, Parliamentary Under Secretary for Science, Innovation and Technology, 1 April 2026](#)

158 [Algorithmic Transparency Recording Standard Hub](#), updated 5 May 2025: “The ATRS is mandatory for all government departments, and for arm’s-length bodies (ALBs) which deliver public or frontline services, or directly interact with the general public.” It is not mandatory for local authorities or for police forces.

159 [Algorithmic Transparency Recording Standard \(ATRS\) Mandatory Scope and Exemptions Policy](#), 17 December 2024, paragraph 3.1. The Government Digital Service is now part of the Department for Digital, Culture, Media and Sport: [Machinery of Government changes: Fact Sheet](#), 22 July 2026 (updated 27 July 2026)

160 [Algorithmic Transparency Recording Standard \(ATRS\) Mandatory Scope and Exemptions Policy](#), 17 December 2024, paragraph 4.2

Government's position

Overall approach

- 78.** At the time of writing, a new Prime Minister has assumed office and made ministerial changes. The Department of Science, Innovation and Technology has been abolished and its functions redistributed, and the previous Department for Business and Trade has become the Department for Business, Innovation, Science and Trade (BIST).¹⁶¹ The former Parliamentary Under-Secretary of State with responsibility for AI, Kanishka Narayan MP, now has the more senior role of Minister of State (Minister for Artificial Intelligence) jointly in the Cabinet Office and BIST, and will attend Cabinet.¹⁶²
- 79.** Our evidence was taken during the previous administration. There have been no policy publications directly focused on AI since the new Prime Minister took office. The government has, however, made statements about matters such as the capabilities of frontier artificial intelligence,¹⁶³ AI security,¹⁶⁴ and protecting children online.¹⁶⁵ The account of the government's position provided below draws on evidence taken, and public statements made, before the new Prime Minister took office, and is presented to highlight the principles that have informed the government's approach to AI in the present Parliament. The AI Minister, as noted above, has a similar, though more senior, role to the one he had when we heard evidence from him. However, we recognise that the government's approach may change under the new Prime Minister.
- 80.** The AI Opportunities Action Plan ('Action Plan') was published by the then DSIT in January 2025.¹⁶⁶ The Action Plan was not itself a statement of government policy—it was commissioned by the government, and produced by the entrepreneur Matt Clifford, who has since joined AI firm Anthropic and has stood down from his role as Chair of the UK Advanced Research and Invention Agency—but the plan found favour and nearly all

161 Written Ministerial Statement, 21 July 2026, <https://questions-statements.parliament.uk/written-statements/detail/2026-07-21/hlws298>

162 <https://www.gov.uk/government/news/ministerial-appointments-july-2026>

163 Written Ministerial Statement, Artificial Intelligence, 7 September 2026, [Written statements - Written questions, answers and statements - UK Parliament](#)

164 [Prime Minister's Questions](#), 9 September 2026

165 Statement, Online Safety, HC Deb, vol 790, col 917 [Online Safety - Hansard - UK Parliament](#)

166 [AI Opportunities Action Plan](#)

its recommendations were accepted by the then administration.¹⁶⁷ The plan focussed on what is needed to drive forward the adoption of AI in the UK. It set out an ambition to “unlock” large datasets to position the UK as an ideal state partner for AI developers and investors.

81. The government has expressed high expectations for the benefits that AI adoption will bring to the UK. The then Secretary of State for Science, Innovation and Technology, the Rt Hon. Liz Kendall MP, set out in January 2026 an ambition to be “the fastest adopting AI country in the G7.”¹⁶⁸ The former DSIT told us that AI is “the defining opportunity of our generation”.¹⁶⁹ The Action Plan states that Britain is the third largest AI market in the world, but is at risk of falling behind the USA and China unless action is taken.¹⁷⁰ It lists three priorities for government:

- Invest in the foundations of AI, i.e. physical infrastructure and skills;
- Increase the adoption of AI in public and private sectors through pilot projects; and
- Improve the UK’s position as a state partner for developers of Frontier AI.

Approach to regulation of AI

82. The former DSIT told us that “the UK’s domestic legal framework is consistent with the principles set out in the [Framework Convention].” It expressed the government’s willingness to amend legislation where needed to mitigate potential AI harms, citing provisions in what is now the Crime and Policing Act 2026 to help combat the production by AI systems of “artificial CSAM, extreme pornography or non-consensual intimate image abuse.” The then Department told us its approach was “targeted, risk-led and designed to help the UK safely and securely prepare [...] for the age of AI” and “safely seize the opportunities AI offers.”¹⁷¹

83. It also told us the government was “committed to proportionate AI regulation which is grounded in evidence and supports growth and innovation.” We heard the government favoured regulation “at the point of

¹⁶⁷ The Government’s response to the plan was published on 13 January 2025. Then Prime Minister, Keir Starmer, said he was “happy to endorse it and take the recommendations forward”. <https://www.gov.uk/government/publications/ai-opportunities-action-plan-government-response/ai-opportunities-action-plan-government-response>, accessed 26 July 2026

¹⁶⁸ DSIT, [AI Opportunities Action Plan: One Year On](#), 29 January 2026

¹⁶⁹ Department of Science Innovation and Technology ([RAI0077](#))

¹⁷⁰ Department for Science, Innovation and Technology, [AI Opportunities Action Plan](#), published 13 January 2025

¹⁷¹ Department of Science Innovation and Technology ([RAI0077](#))

use by our existing expert regulators [...] through sector-specific regulations such as medicines regulation, and cross-cutting legal safeguards including data protection, competition, product safety, and equalities laws.”¹⁷² The then DSIT considered this approach was working well, but also underlined its commitment, made in its response to the AI Action Plan, to “work with regulators to boost their capabilities.” AI Minister Kanishka Narayan MP told us “static forms of regulation, where we have been historically, may not be best adapted to where we are likely to go.”¹⁷³

84. It is notable however that the government, having announced in the 2024 King’s Speech that it would “seek to establish the appropriate legislation to place requirements on those working to develop the most powerful artificial intelligence models,”¹⁷⁴ has not at the time of writing yet brought forward a dedicated AI Bill.
85. The government’s approach to regulating AI was warmly welcomed by the large technology firms from whom we took evidence, with the representative from Meta, owner of Facebook and WhatsApp, describing it as a “thoughtful and sensible approach and in some ways a global model.”¹⁷⁵ However, it was criticised by a wide range of other witnesses, especially those focused on human rights protection, those involved in the legal system, and technical bodies. The approach was described as “uncritical and deregulatory,”¹⁷⁶ and “asleep at the wheel.”¹⁷⁷
86. The AI Action Plan—which, as noted above, was welcomed by the previous administration which set out a plan to implement it—was also widely criticised.¹⁷⁸ Dr Susie Alegre, a barrister at Garden Court Chambers, caught the mood of much of the evidence to the inquiry when she told us the Plan “turns a blind eye to the human rights risks posed by AI.”¹⁷⁹ Witnesses pointed to a lack of rights-based language, an emphasis on “economic growth, competitiveness and trust, which are not proxies for human rights protections,”¹⁸⁰ and what they felt was an excessive focus on retaining a position on the “world AI stage”, which “should not come to the detriment of human rights [or] citizen’s safety and security.”¹⁸¹

172 [Correspondence from Kanishka Naranyan MP, Parliamentary Under Secretary for Science, Innovation and Technology, 1 April 2026](#)

173 [Q99](#)

174 HL Deb, Volume 839, Col 7, 17 July 2024, [King’s Speech - Hansard - UK Parliament](#)

175 [Q89](#)

176 [Q17](#)

177 [Q29](#)

178 Eleonor Duhs (Barrister, Partner, Head of Data & Privacy at Bates Wells LLP) ([RAI0035](#)); Justice ([RAI0073](#))

179 Dr Susie Alegre (Barrister, Associate Tenant, Garden Court Chambers) ([RAI0019](#))

180 Responsible AI UK (RAI UK) ([RAI0049](#))

181 Dr S. J. Fox (Academic (Law); Associate Director, Institute for Digital Culture at University of Leicester) ([RAI0009](#))

87. Among the specific criticisms we heard, witnesses told us the Plan needed to do more to “ensure that AI-based innovation should be human rights compliant,”¹⁸² and to address copyright law, data use and transparency.¹⁸³ Some witnesses criticised a lack of focus on protecting the rights of children within the Plan.¹⁸⁴ The British Standards Institution, which is the UK’s National Standards Body, supported the positive outlook on innovation, but suggested “there might be room for an ‘AI Safety Action Plan’ that focuses more on the risks and other considerations surrounding AI deployment.”¹⁸⁵

88. Later chapters of this report will set out in detail our recommendations for comprehensive and tailored regulation of AI in order to protect human rights.

89. **CONCLUSION**

The AI Action Plan commissioned by the previous administration, while it brings a welcome focus on innovation and economic opportunity, fails to pay sufficient attention to the need to secure the protection of human rights from AI-generated interference throughout the supply chain.

90. **CONCLUSION**

The UK’s leading position as the third largest AI market in the world affords opportunities, not only for economic growth but also for setting strong standards on the responsible adoption and development of new technologies.

91. **RECOMMENDATION**

The Committee is encouraged by the prominence given to AI in the new administration and urges the government to seize the opportunity to regulate AI in a way which puts human rights first.

182 European Center for Not-For-Profit Law (ECNL) ([RAI0008](#))

183 Chartered Institute of Library and Information Professionals (CILIP) ([RAI0053](#))

184 5Rights Foundation ([RAI0074](#))

185 British Standards Institution ([RAI0050](#))

4 Problems with existing protections

- 92.** As set out in the previous chapter, there are many legal provisions which apply or might apply to AI systems.
- 93.** However, we heard evidence that the existing protections are not sufficient to protect human rights. Global Network Initiative, an organisation campaigning on privacy and freedom of expression, told us: “While [the UK’s] existing frameworks already apply to AI in many contexts, they were designed before the emergence of advanced AI models, including foundation models and agentic AI systems with autonomous decision-making capabilities.”¹⁸⁶ Liberty told us that existing legal frameworks are “wholly insufficient for protecting human rights.”¹⁸⁷ UNISON said “the existing legal framework is fundamentally inadequate to address AI-specific risks”.¹⁸⁸
- 94.** A common theme concerned the need to impose mandatory legal obligations on AI developers and providers because the source of human rights harm is often built into AI systems during the development of AI systems.¹⁸⁹ We heard that relying primarily on the mandatory obligations imposed on deployers at the point when AI systems are used cannot provide adequate human rights protection.¹⁹⁰ Michael Birtwistle from the Ada Lovelace Institute told us, “The major gap in our regulatory system is managing risk upstream with those building the tech.”¹⁹¹
- 95.** This chapter considers the various points along the AI lifecycle where human rights risks may arise, and identifies the extent to which existing legal protection adequately addresses them.

186 Global Network Initiative ([RAI0034](#))

187 Liberty ([RAI0079](#))

188 UNISON ([RAI0076](#))

189 UNISON ([RAI0076](#)); Justice ([RAI0073](#)); Trilateral Research ([RAI0046](#)); Dualarity Ltd ([RAI0071](#))

190 The Law Society of England and Wales ([RAI0021](#)) Gina Romero et al ([RAI0023](#)), Amnesty International UK ([RAI0051](#))

191 [Q27](#)

Before AI systems are deployed

Design and development

- 96.** Harm caused by AI systems at the point of use may be a result of things done higher up the AI supply chain.¹⁹² Problems may be built in at the development stage so that, even if the deployer uses the AI system as intended, harmful consequences can occur. As explained in the previous chapter, some existing laws may constrain AI developers to some extent in specific circumstances. But we heard that existing law is not enough on its own to protect against risks to human rights arising from the actions of developers.
- 97.** Data protection law will apply to developers to the extent that their activities involve processing personal data. William Malcolm of the ICO told us, “we believe in the importance of upstream intervention, working with providers of technology to catch technology as it is adapted and built, to get that privacy by design and default at the front end.”¹⁹³
- 98.** However, we heard that it is often not at all obvious when people’s data is being used to train AI systems. Alex Pirlot de Corbion of Privacy International told us of “industry and governments making a mad dash to Hoover loads of data into AI systems [and] models ... it is really hard to know how much of our personal data is being processed by companies and governments.”¹⁹⁴ Javier Ruiz Diaz of Amnesty International UK told us, “we find that current protections around data protection, access and transparency do not seem to work at all. It is pretty much impossible to find out in detail how your data is being processed.”¹⁹⁵ Equity told us that in the creative industries, “it can be very difficult for an artist to establish whether and to what extent their data has been used to train the AI model,” and hence to uphold their rights.¹⁹⁶
- 99.** It will also be difficult in practice for data subjects to rely on data protection law to protect them against privacy intrusions arising during AI development. The Law Society of England and Wales told us that “the use of personal data by AI is opaque and not widely understood as someone may not be aware of their rights being violated.”¹⁹⁷ Even if someone does become aware, Ravi Naik highlighted the difficulty of upholding data rights. Where a correction to personal data is requested by an individual, it may

192 [Q27](#)

193 [Q76](#)

194 [Q28](#)

195 [Q30](#)

196 Equity ([RAI0055](#))

197 The Law Society of England and Wales ([RAI0021](#))

be impossible to tell if the system has actually changed the records it holds: “Are you affecting the intrinsic model or just the output and the way it talks back to you?”¹⁹⁸

- 100.** In addition, there are limits in the extent to which the existing law can protect against discrimination arising from the way an AI system is trained at the development stage. The Equality Act 2010 only applies to persons concerned with the provision of a service to the public or a section of the public. It does not apply between businesses, so that the Act may not apply to developers further up the supply chain.¹⁹⁹
- 101.** The Human Rights Act 1998 will also not provide adequate protection against breaches of human rights arising at the development stage, because it only applies to public authorities. Witnesses told us that because practically all AI development takes place in the private sector, this left a problematic gap in protection for human rights.²⁰⁰
- 102.** Upstream developers may escape legal responsibility by relying on contractual terms in their arrangements with deployers that exclude responsibility for harm even if it can be attributed to actions or omissions of the provider. Louise Hooper of Garden Court Chambers told us that responsibility must not fall solely on AI system deployers: “It is not appropriate for deployers of AI systems to bear the burden of risks created up stream, which they may not properly understand because they have not been explained properly.”²⁰¹
- 103.** We heard that the current law imposes no systematic obligations on AI developers or providers to identify, evaluate, prevent or reduce human rights risks arising from the intended or foreseeable applications of AI systems.

104. CONCLUSION

Taken together, existing UK laws that could apply to AI systems do so primarily at the point of deployment, so that deployers are the primary holders of responsibility for harms arising from AI systems. This means that existing laws do not take due account of the complexity of the AI lifecycle and supply chain and are not directed to the actor that may be best placed to identify and avert the source of risks. They are therefore likely to allow human rights harm that could have been prevented.

198 [Q11](#)

199 Equality Act 2010, section 29(1); [Dockers’ Labour Club v Race Relations Board](#) [1976] AC 285 at 297D-E; [Johnston v Giving.com](#), Case No 226MC647, 2 November 2022, at [47]-[56]; [\[2024\] EWHC 944 \(KB\)](#) at [16]-[18].

200 Dr Erin Ferguson et al ([RAI0048](#)), Dr Giulia Gentile et al ([RAI0047](#)); Professor Francois du Bois (Professor of Law at University of Leicester) ([RAI0072](#))

201 [Q41](#)

Pre-market evaluation and testing by regulators

- 105.** Regulators have power to conduct prior evaluation or testing of AI systems only in limited and specific circumstances: for example, AI-enabled medical devices that are high risk require clinical evaluation by the MHRA.²⁰² The ICO also has power to prevent high-risk processing of personal data at the development stage through the data protection impact assessment process (see Chapter 3),²⁰³ although we heard that compliance with requirements relating to data protection impact assessments is patchy.²⁰⁴
- 106.** AISI has some responsibility for assessing risks arising at the development stage, and representatives from big tech firms spoke highly of AISI's work.²⁰⁵ However, these arrangements are voluntary. AISI has no regulatory authority and cannot, for example, demand access to foundation models either before or after their deployment.²⁰⁶
- 107.** Michael Birtwistle of the Ada Lovelace Institute told us AISI's change in focus, from "safety" to "security" (see section on "Role of government bodies" in Chapter 3), which the then Secretary of State for Science, Innovation and Technology described as bringing a new "focus on serious AI risks with security implications,"²⁰⁷ showed "the deregulatory flavour of the government's attitude towards AI."²⁰⁸ UNISON told us the shift "fundamentally altered the UK's approach to AI risk management, shifting from comprehensive safety considerations to narrower security concerns. In doing so, it fails to recognise that AI can be dangerous even when not criminally misused [...] [excluding algorithmic bias] represents a clear reduction in safeguarding mechanisms and is likely to have wider implications for human rights."²⁰⁹

202 Medicines and Healthcare products Regulatory Agency, [Impact of AI on the regulation of medical products](#), April 2024, page 6

203 The UK GDPR requires data controllers to carry out data protection impact assessments (DPIAs) and consult the ICO when they plan to process data in ways that carry high risks (UK GDPR, Articles 35–36). Data controllers must carry out a DPIA if "a type of processing in particular using new technologies, and taking into account the nature, scope, context and purposes of the processing, is likely to result in a high risk to the rights and freedoms of natural persons" (UK GDPR, Article 35(1)). The controller must consult the ICO prior to processing if the impact assessment indicates that the processing would result in a high risk (UK GDPR, Article 36(1)). If the ICO considers the processing would breach the UK GDPR, it has a range of powers, including gathering information, imposing conditions, or even prohibiting the processing altogether (UK GDPR, Articles 36(2) and 58).

204 Eleonor Duhs (Barrister, Partner, Head of Data & Privacy at Bates Wells LLP) ([RAI0035](#))

205 [Q97](#)

206 [Correspondence from Kanishka Narayan MP, Parliamentary Under Secretary for Science, Innovation and Technology, 1 April 2026](#)

207 [Tackling AI security risks to unleash growth and deliver Plan for Change - GOV.UK](#), dated 14 February 2025, accessed 27 August 2026

208 [Q17](#)

209 UNISON ([RAI0076](#))

- 108. CONCLUSION**
Regulators lack the power to test and evaluate AI systems before they are publicly released, or to prevent their release if they are considered to pose unacceptable risks, except in very limited circumstances.
- 109. CONCLUSION**
Model developers engage with the AISI on a voluntary basis. The AISI has no statutory power.
- 110. CONCLUSION**
No cross-economy regulators, agencies or bodies have systematic powers to prevent AI models or systems being deployed even if they pose significant risks.

When AI systems are deployed

Lack of transparency

- 111.** Article 8 of the Framework Convention says that each Party “shall adopt or maintain measures to ensure that adequate transparency and oversight requirements tailored to the specific contexts and risks are in place in respect of activities within the lifecycle of artificial intelligence systems, including with regard to the identification of content generated by artificial intelligence systems.”²¹⁰ The Explanatory Report notes that parties must strike an appropriate balance between these considerations and other important factors including “privacy, confidentiality (including, for instance, trade secrets), national security, protection of the rights of third parties, public order, [and] judicial independence.”²¹¹
- 112.** However, we heard that this transparency is not at present being achieved in the UK in a comprehensive way. Data protection law imposes obligations on developers and deployers to be transparent about use of AI when they are processing personal data.²¹² The government has adopted transparency requirements in the public sector, in particular under the Algorithmic

210 Council of Europe Treaty Series - No. 225, Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, Vilnius, 5.IX.2024, <https://rm.coe.int/1680afae3c>

211 Council of Europe Treaty Series - No. 225, Explanatory Report to the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, Vilnius, 5.IX.2024, <https://rm.coe.int/1680afae67>

212 UK GDPR, Articles 12–14; ICO guidance on Artificial Intelligence: Explaining decisions made with AI: Part I The basics of explaining AI: [Legal framework](#).

Transparency Reporting Standard (ATRS),²¹³ although it has no statutory basis. However, there are no general transparency requirements concerning AI deployment in UK law.

- 113.** As AI develops, it is increasingly difficult, and may be practically impossible, to tell when AI is being used. Witnesses told us this presented a serious problem. JUSTICE told us, “There is an initial fundamental block to ensuring human rights are practically enforceable under the current regime: knowing if AI has been used to violate your rights in the first place.”²¹⁴ The Ada Lovelace Institute told us “strengthening individual and collective action can only be effective in the context of people having meaningful routes to understanding harm [...] and knowing when their rights have been infringed.”²¹⁵ Kay Firth-Butterfield of GoodTech Advisory told us “AI should always identify itself to users, or users should be told when AI is involved. If you do not do that, you take away a right to remedy.”²¹⁶
- 114.** We also heard that individuals significantly affected by the output of an AI system are not entitled to be provided with a proactive and personalised explanation, nor information on how and why a decision based on an AI-generated output was reached, including specific categories of information such as tailored information on the contribution of the automated or algorithmic tool or system in the decision-making process.²¹⁷
- 115.** We also heard concerns about the transparency of AI use by public decision-makers. JUSTICE told us that “The use of AI by public decision-makers is currently extremely opaque. There is a lack of transparency around whether AI is being used at all, and where it is known that AI is being used it is often unclear how a tool is being used and what its functionality is.”²¹⁸ Similarly, Prof David Leslie told us that “there also needs to be much more transparency about use; the public should be aware when these technologies are being used to measure or observe them.”²¹⁹
- 116.** The effectiveness of the Algorithmic Transparency Reporting Standard (ATRS) (see above, paragraph 77) in ensuring transparency about use of AI systems in the public sector was also a matter of concern for some witnesses. As discussed above, this standard is mandatory for government departments and many Arms-Length Bodies. The EHRC told us: “While this standard is a good starting point for transparency, people should

213 [Algorithmic Transparency Recording Standard Hub](#), updated 5 May 2025: “The ATRS is mandatory for all government departments, and for arm’s-length bodies (ALBs) which deliver public or frontline services, or directly interact with the general public.”

214 Justice ([RAI0073](#))

215 Ada Lovelace Institute ([RAI0066](#))

216 [Q56](#)

217 Public Law Project ([RAI0045](#)); Ada Lovelace Institute ([RAI0066](#))

218 Justice ([RAI0073](#))

219 [Q10](#)

be proactively informed about the use of AI in any significant decision about them, and about how such decisions can be challenged.”²²⁰ Several witnesses raised concerns about gaps in the published information.²²¹ Louise Hooper of Garden Court Chambers told us that the ATRS “has been mandated but we know it is not being used properly. Not all algorithmic use is on there.”²²² The Department for Science, Innovation and Technology told us that “the ATRS is designed as a transparency and communication standard, not an assurance or audit mechanism. ATRS records are reviewed for quality and readability by DSIT as part of the internal clearance and publication process, but accountability for tools sits with the organisation that owns them”.²²³

- 117.** More broadly, some witnesses argued that the lack of mandatory disclosure obligations on AI developers across all sectors was an obstacle to securing end-to-end accountability across the AI supply chain.²²⁴ On the other hand, the ICO told us the “challenges raised [by AI in the context of privacy and data protection] are complex but we do not think they are insurmountable under the current legislation.”²²⁵ In later chapters we set out the changes we would like to see in legislation to address these problems.

118. CONCLUSION

Existing law is deficient by failing to impose legal requirements on AI system deployers to be transparent about their use of AI systems.

The difficulty of identifying the cause of human rights harms

- 119.** We heard that the opacity of AI systems can make it difficult for the existing legal and regulatory regimes to identify and address human rights harms.²²⁶ In written evidence to the Committee, the Department for Science, Innovation and Technology told us:

220 Equality and Human Rights Commission ([RAI0075](#))

221 RAI0045, Public Law Project; Public Law Project, Tracking Automated Government ‘TAG’ Register (9 February 2023) <http://trackautomatedgovernment.org.uk/>; [Q44](#)

222 [Q38](#)

223 [Correspondence from Kanishka Narayan MP, Parliamentary Under Secretary for Science, Innovation and Technology, 1 April 2026](#)

224 [Q7](#)

225 Information Commissioner’s Office (ICO) ([RAI0085](#))

226 [Q53](#); Chartered Institute of Library and Information Professionals (CILIP) ([RAI0053](#)) Dr Ayça Atabey; Dr Kim Sylwander; Professor Sonia Livingstone ([RAI0058](#)) Minderoo Centre for Technology and Democracy, University of Cambridge ([RAI0067](#)) Mr Burak Batuhan Karakus (Head of Legal Affairs at Human Rights Solidarity) ([RAI0026](#)); Eleonor Duhs (Bates Wells LLP) ([RAI0035](#)); Duality Ltd ([RAI0071](#)); Dr Felipe Romero-Moreno ([RAI0087](#))

the opaque nature and technical complexity of some AI systems may make it harder for claimants to [identify] failures and to attribute failures to particular actors in the supply chain against whom a claim for damages could then be brought [...] [it] may be difficult for claimants to prove that output from an AI system was factually caused by some sort of failure to meet appropriate standards on the part of, for example, a system developer or deployer. It remains to be seen how the courts will deal with these sorts of issues which will necessarily be very fact-specific.²²⁷

- 120.** Academics from Essex Human Rights Centre and Essex Law School told us many AI systems “are opaque insofar as they rely on an intricate set of processes and infrastructures that can be deciphered only by experts [...] Sometimes, due to the internal machine-learning processes, it is not possible to truly decipher.”²²⁸
- 121.** This means it is often unclear why an AI system has generated an outcome or decision (sometimes known as the “black box problem”).²²⁹ Witnesses told the Committee that even the designers of an AI system might not be able to explain its processing activities,²³⁰ and that those flagged by AI systems often would not be able to find out why this had happened, contradicting widely held notions of fairness.²³¹ Dr Iulian Serban of frontier AI firm LawZero told us: “large language models—from well-known companies such as OpenAI and Anthropic are brittle black boxes that we do not fully understand. We do not understand how they behave, we cannot predict what they will do and we see over and over again that they do not behave the way we want them to.”²³²
- 122.** This opacity can mean that harm arising from AI systems may elude existing human rights protections. Professor Lorna McGregor of the University of Essex cautioned that the opacity attendant on AI makes it difficult to determine harms as well as to remedy them.²³³ Connected by Data, a campaigning not-for-profit organisation, pointed to research showing that AI-generated summaries of women’s health records downplayed their needs in comparison to men’s, and noted that in such a situation, “a case brought by an individual is unlikely to pass the threshold for meaningful harm [...]”

227 Department of Science Innovation and Technology ([RAI0077](#))

228 Dr Giulia Gentile et al ([RAI0047](#))

229 Dr Giulia Gentile et al ([RAI0047](#)) Professor Francois du Bois (Professor of Law at University of Leicester) ([RAI0072](#))

230 [Q31](#)

231 [Q30](#)

232 [Q53](#)

233 Professor Lorna McGregor (Professor of International Human Rights Law at University of Essex) ([RAI0056](#))

Nevertheless, in aggregate this use of AI is likely to lead to worse outcomes for women as a group.”²³⁴ In later chapters we explain how we think the legal and regulatory framework should change to deal with these concerns.

Lack of effective remedy

- 123.** The opacity of AI processes often also means the information that can be acquired about the cause of harm may be too limited to evidence an effective claim against a deployer even if a claim could in principle be brought.²³⁵ It will often be difficult to identify who should be held legally responsible for the safe functioning of AI systems and where liability lies.²³⁶ Deployers may also not have the expertise or capacity to address harms arising from an aspect of the AI system that was embedded at the development stage.²³⁷
- 124.** Tetyana Krupiy of Newcastle University told us that in the case of a discrimination claim, “it can be very challenging to show that [...] unfavourable treatment is based on the possession of a particular protected characteristic when organisations use artificial intelligence as part of the decision-making process.”²³⁸

Limits of human oversight

- 125.** Deployers may rely on human oversight of AI systems in order to protect against AI harms. There are at present no specific legal obligations relating to human oversight, except in respect of automated decision-making (see chapter 3).
- 126.** Witnesses raised concerns about over-reliance on human involvement as a protection against human rights risks. We were told that humans may not be in a position to question the decision reached by AI, so effective scrutiny may be impossible.²³⁹ JUSTICE commented: “Even highly trained and qualified ‘humans in the loop’ may struggle to scrutinise effectively the AI’s output given that AI tools can process much larger volumes of information than it is possible for a person to do.”²⁴⁰ Humans can also be influenced by AI outputs, which we heard might further reduce the protective influence assumed to derive from human involvement. We were told about the risk of

234 Connected by Data ([RAI0013](#))

235 Ada Lovelace Institute ([RAI0066](#)) Glitch ([RAI0015](#)) Dr Tetyana Krupiy (Newcastle University) ([RAI0012](#))

236 Department of Science Innovation and Technology ([RAI0077](#)); The Law Society of England and Wales ([RAI0021](#)); Dr Erin Ferguson et al ([RAI0048](#))

237 [Q20](#)

238 Dr Tetyana Krupiy (Newcastle University) ([RAI0012](#))

239 Public Law Project ([RAI0045](#))

240 Justice ([RAI0073](#))

“automation bias” when human over-rely on AI outputs without questioning them,²⁴¹ as well as research suggesting that humans become more biased after interacting with biased AI, creating a harmful feedback loop.²⁴²

Gaps in protection and legal uncertainty

- 127.** Additionally, the current law (described in the previous chapter) does not provide comprehensive protection against human rights harms which arise from the deployment of AI system.
- 128.** For instance, the former DSIT noted in written evidence that it is possible for individuals to enforce their ECHR rights under the Human Rights Act 1998 as well as directly in the European Court of Human Rights.²⁴³ But the Human Rights Act 1998 only applies to public authorities. Witnesses told us that because much AI use takes place in the private sector, this left a problematic gap in protection for human rights.²⁴⁴
- 129.** Other legal protections may address certain harms arising from the deployment of AI systems, but they do not provide comprehensive protection. For instance, although the data protection regime is extensive, it is intended to protect the privacy rights of identifiable individuals rather than to address more widespread harms arising from use of AI systems.²⁴⁵ Where deployers are carrying out high-risk activities, they should prepare a data protection impact assessment and notify the ICO (see chapter 3). But as noted above, we heard evidence that compliance with this requirement is patchy.
- 130. CONCLUSION**
The current legal framework does not provide for adequate transparency in relation to the development and deployment of AI systems.

241 Dr Johann Laux (Departmental Research Lecturer in AI, Government and Policy at Oxford Internet Institute, University of Oxford) ([RAI0006](#))

242 UNISON ([RAI0076](#)); Glickman, M., Sharot, T. How human–AI feedback loops alter human perceptual, emotional and social judgements. *Nat Hum Behav* 9, 345–359 (2025). <https://doi.org/10.1038/s41562-024-02077-2>

243 Department of Science Innovation and Technology ([RAI0077](#))

244 Dr Giulia Gentile et al ([RAI0047](#)), Dr Erin Ferguson et al ([RAI0048](#)); Professor Francois du Bois (Professor of Law at University of Leicester) ([RAI0072](#))

245 Privacy International ([RAI0081](#))

131. CONCLUSION

Existing law does not require the provision of notice to individuals that an AI system is in use. In the absence of such knowledge, individuals who have been directly affected by the use of AI systems will be in practice unable to challenge such use or seek an effective remedy for human rights violations. The evidence we heard strongly suggests that many people may be in this situation.

132. CONCLUSION

There are no mandatory obligations on AI developers and providers to provide downstream actors, including deployers, with information concerning the potential risks arising from the deployment of AI systems. This can reduce transparency and traceability across the AI supply chain.

133. CONCLUSION

The ATRS provides some transparency concerning the use of AI systems in the public sector, but is limited in its coverage.

Oversight after AI systems are deployed

Uncertain legal liability

- 134.** We heard that legal uncertainty about the circumstances in which liability might arise can allow harms caused by deployment of AI systems to continue. Public Law Project expressed concern that unclear liability for rights-critical AI uses could result in harmful consequences to the public, particularly marginalised individuals and communities who are often most reliant on public services.²⁴⁶ The former DSIT told us it was assessing what gaps in liability may emerge as AI deployment accelerates, and that “there are limits to the certainty that the government can provide at this stage.”²⁴⁷
- 135.** The Law Society for England and Wales also drew attention to the lack of clarity about where legal responsibility for AI-generated harms may lie.²⁴⁸ The Ada Lovelace Institute told us that “[i]deally, liability risk should be fairly distributed between the relevant actors (upstream developers, downstream developers, deployers), as risk can materialise at different points in the development and deployment process. Currently, legal risks are primarily pushed to downstream actors through the use of contracts

246 Public Law Project ([RAI0045](#))

247 Department of Science Innovation and Technology ([RAI0077](#))

248 The Law Society of England and Wales ([RAI0021](#))

and standard terms & conditions. Consequently, upstream developers (usually LLM or foundation model developers) are largely shielded from risk, even though they have considerable impact on the functioning of the AI product, and its propensity to create human rights risks. If certain actors are shielded from liability risk in this way, they may not be optimally incentivized to invest in precautionary measures to prevent harm from occurring.”²⁴⁹

- 136.** In addition, under the present law, strict liability for producers of defective products only applies to AI systems that are integrated into physical products, and only applies to a narrow set of types of harm.²⁵⁰

Rapid multiplication of risks

- 137.** Under the current framework, accessing redress for AI harms relies on the aggrieved individual taking action via a legal or regulatory route. We heard that this framework may struggle to address the potential for mass harms that can be caused when AI systems go wrong and harm is suffered by thousands or even millions of individuals.
- 138.** We heard from witnesses that because of its scale, AI has the potential to rapidly multiply harms. The Ada Lovelace Institute told us that “Many deployers will be downstream procurers of AI models and systems developed further upstream by AI developers; for example, chatbots or other AI tools being deployed currently within government are usually based on a general-purpose large language model trained by a major AI developer. This means that without careful audit for risks and further finetuning, such systems will inherit all the flaws, biases and risks of their base LLM model. This is known as ‘risk proliferation’.”²⁵¹ Professor Kevin Fong told us of an LLM deployed to help eating disorder patients in the USA which ended up giving patients “very harmful” advice about calorie restriction.”²⁵² Connected by Data told us, “[t]he automation offered by AI means that organisations are able to violate human rights at scale and at speed.”²⁵³ Silkie Carlo of Big Brother Watch highlighted the capacity of Facial Recognition Technology, when used on a wide scale, to rapidly multiply harmful effects on minoritised groups.²⁵⁴

249 Ada Lovelace Institute ([RAI0066](#))

250 UK Jurisdiction Taskforce’s [Legal Statement on Liability for AI Harms under the private law of England and Wales](#), July 2026, paragraphs 70–80; [Product liability – Law Commission](#), published 8 December 2025, accessed 27 August 2026

251 Ada Lovelace Institute ([RAI0066](#))

252 [Q20](#)

253 Connected by Data ([RAI0013](#))

254 [Q29](#)

- 139.** A connected point is that of the extent to which AI systems may apparently ignore or override instructions they have been given. Dr Roman Yampolskiy of the University of Louisville told us in relation to the possibility of future superintelligent systems: “Even if we could completely trust the leaders of large AI companies [...] they are, themselves, not capable of controlling the technology that they are developing.”²⁵⁵
- 140.** Several recent incidents have already given rise to similar concerns. On 21 July 2026, the company OpenAI made a statement in which they described a “new kind of security incident.” This took place when AI models being tested in an internal sandbox found a way to gain internet access using a previously unknown vulnerability, and then hacked into another company, Hugging Face, to find information to help it achieve its test goal. OpenAI has stated that it is taking action to strengthen protections around its tests.²⁵⁶
- 141.** We heard the speed of change in the AI field makes it hard for existing systems to address the risks it poses. The Ada Lovelace Institution told us, “case law is a slow, unresponsive mechanism for addressing the rapidly-developing impacts of emerging technologies and complex interactions with principle-based legal regimes such as human rights law.”²⁵⁷ Similarly, Professor Albert Sanchez-Graells told us, “the current legal framework, built around slow, human-exclusive decision-making, is ill-equipped to handle the speed and scale of AI-driven processes. Procedural rights (e.g., to be heard, to access the file, to be given reasons for decisions) become ineffective when decisions are instant and data-driven.”²⁵⁸
- 142.** There is also the cost of litigation to consider, which witnesses said made litigation inaccessible to many.²⁵⁹ Ravi Naik told us that “the cost of bringing action in this country, particularly against a private actor, is prohibitive for most people. We are talking seven-figure sums to bring action against a private actor just to uphold your rights, let alone seek damages. That has to be a key barrier to accountability in this space.”²⁶⁰ In the next chapter we set out the changes we recommend to law and regulation to address these concerns.

255 [Q46](#)

256 <https://openai.com/index/hugging-face-model-evaluation-security-incident/> accessed 23 July 2026

257 Ada Lovelace Institute ([RAI0066](#))

258 Professor Albert Sanchez-Graells (Professor of Economic Law at University of Bristol Law School) ([RAI0005](#))

259 Eleonor Duhs (Bates Wells LLP) ([RAI0035](#)); [Q11](#); Dr Tetyana Krupiy (Newcastle University) ([RAI0012](#))

260 [Q3](#)

143. CONCLUSION

The recent OpenAI/Hugging Face security incident starkly highlights the societal risks posed by the growing power of AI systems. It illustrates how the newest AI models can engage in activities that would be unlawful if undertaken by a legal person, with the potential to seriously undermine security, including in critical infrastructure. It also suggests that at present society largely relies on private companies to take voluntary action to guard against such incidents. This is unsatisfactory.

144. CONCLUSION

The general-purpose nature of the foundation models which allows them to be embedded into a wide and diverse range of AI systems and contexts means that any vulnerabilities and sources of risk within the model can propagate very rapidly and at scale once deployed. The power of the newest models, and their potential for mass harm, justifies a precautionary approach to regulation.

145. CONCLUSION

Existing law is inadequate to enable those who suffer human rights harm generated by AI systems to secure an effective remedy.

146. CONCLUSION

Even if some existing laws could address some human rights risks posed by AI systems and provide redress to individuals in specific circumstances, their application is often speculative and uncertain, and the cost of individual litigation is in practice prohibitive.

Limits of regulatory action

- 147.** As noted above, the government has stated its preference for a “regulator-led” approach to keep up with AI developments.
- 148.** However, we heard some evidence that existing regulators may struggle to tackle AI risks effectively, despite their broad powers, for several reasons including their remits, resourcing and need for technical expertise. In this section we highlight the problems witnesses identified, and briefly survey the powers of the main regulators in the AI space.

A fragmented regulatory landscape?

- 149.** The current regulatory framework which applies to AI systems in the UK has been described to us as fragmented and inadequate,²⁶¹ with weak enforcement,²⁶² and with a lack of AI expertise at the regulatory level,²⁶³ resulting in a system which does not effectively protect human rights and which risks harm slipping through the cracks between regulators' remits.²⁶⁴ As there is no single body with responsibility for overseeing AI regulation, there is no overall cross-economy monitoring to identify AI incident patterns which could indicate more systemic problems in AI systems, such as inherited bias.²⁶⁵
- 150.** Michael Birtwistle of the Ada Lovelace Foundation told us "AI is regulated in the UK, but only incidentally and not well."²⁶⁶ Justice told us that the "plethora" of overlapping regulators "makes it confusing for individuals, whose rights may have been infringed by the use of AI, to know where to go to make a complaint or commence a claim".²⁶⁷ The Law Society told us that "there are still gaps to protecting people's rights, such as a wider, non-data protection focused right to challenge automated decision-making."²⁶⁸
- 151.** We also heard that existing regulation may not be sufficiently sensitive to the risks of bias and discrimination in AI systems. Dr Mark Wong, Senior Lecturer and Head of Social and Urban Policy at University of Glasgow, told us that "group-based harm and systemic racism are less understood [than individual based harm] and not regulated with adequate, fit-for-purpose safeguarding [...] Relying on existing regulations alone will perpetuate these systemic issues."²⁶⁹
- 152.** The regulators we spoke to did not themselves raise concerns about gaps. The EHRC told us some areas were less well regulated than others: "for example, healthcare is well regulated so, with our concern about equality and human rights, we can work with regulators in the healthcare system to ensure that they are thinking about these issues in relation to AI. Other areas are less well regulated, such as the social security system."²⁷⁰ Andrew Breeze, Director for Online Safety Technology Policy at Ofcom, told the Committee: "Our powers to regulate AI are really at the deployment and use-case level, to the extent that they fall within the sectors that we

261 [Q33](#); Fair Vote UK ([RAI0002](#)); Justice ([RAI0073](#)); [Q62](#)

262 UNISON ([RAI0076](#))

263 UNISON ([RAI0076](#))

264 Duality Ltd ([RAI0071](#))

265 ISAR Global ([RAI0007](#)); See also Chapter 2, The right to non-discrimination

266 [Q17](#)

267 Justice ([RAI0073](#))

268 The Law Society of England and Wales ([RAI0021](#))

269 Dr Mark Wong (University of Glasgow) ([RAI0014](#)); Fair Vote UK ([RAI0002](#))

270 [Q77](#)

have been given powers by Parliament to regulate.”²⁷¹ He also confirmed that misinformation or discriminatory content online would not fall within Ofcom’s remit, except “to the extent that it rises to the level of criminal activity.”²⁷² Ofcom told us there were “edge cases,” such as AI chatbots and AI image generators, which did not fall squarely in their remit.²⁷³ Ofcom told us that they talk often to the ICO and EHRC to identify areas of concern.²⁷⁴ Kanishka Narayan MP, the AI Minister, told us “the AI Security Institute [...] co-ordinates and often speaks to representatives in Ofcom [...] to try to ensure that Ofcom’s regulatory functions are being appropriately deployed upstream.”²⁷⁵

153. The AI Minister, Kanishka Narayan MP, told us the government carried out systematic reviews to identify risks that were not being effectively countered, and took specific action to plug any gaps.²⁷⁶ He also noted that some regulation is cross-sectoral (although as noted above, existing cross-sectoral regulation does not provide comprehensive protection for human rights harm arising from AI systems).

154. An ongoing review by the Law Commission is considering the operation of the existing product liability regime, particularly in relation to digital products and emerging technologies such as AI, to determine what law reform might be required to ensure that the product liability regime is fit for purpose.²⁷⁷

155. CONCLUSION

The current legal framework that applies to AI systems is fragmented and difficult to navigate. It relies primarily on harm-specific and sector-specific legislation leaving gaps in protection which fails to address preventable human rights risks.

156. CONCLUSION

Current product liability law provides strict liability for AI systems incorporated into physical products, but not in relation to stand-alone AI systems. This difference in legal liability is difficult to justify in principle. The government must carefully consider any recommendations from the ongoing Law Commission review, including on this specific point.

271 [Q75](#),

272 [Q80](#)

273 [Q75](#)

274 [Q75](#)

275 [Q100](#)

276 [Q103](#)

277 [Product liability – Law Commission](#), published 8 December 2025, accessed 27 August 2026

Resourcing of regulators

- 157.** We heard some concerns about the extent to which regulators have the resources to deal with the risks linked to AI. These were particularly highlighted in relation to the EHRC. The Ada Lovelace Institute told us UK regulators “are not well resourced or empowered to address AI,”²⁷⁸ and that the EHRC is “critically under-resourced to address AI.”²⁷⁹ The Equality and Human Rights Commission itself wrote to the Committee stating that their ability to undertake work on AI is limited by their resources.²⁸⁰
- 158.** Ravi Naik told us the “ICO is a stretched organisation. I do not think we could expect the ICO to enforce every single complaint [within its remit.]”²⁸¹ However the ICO painted a more positive picture of resources available to them to regulate AI, telling us, “we can always use more to do more, but we work within the envelope we have [...] it is our job to prioritise our resources in areas where we can have the most impact, both in enforcement and upstream regulation.”²⁸²

Technical expertise of regulators

- 159.** Witnesses told us “regulators lack both the AI expertise and enforcement powers necessary to protect human rights effectively.”²⁸³ Given the wide spread of regulators potentially called on to deal with AI matters, we heard that that levels of expertise in AI would vary widely across regulators, potentially making it hard to identify “the real risks” of AI systems.²⁸⁴ Witnesses also called for “building common capacity across the regulators” on AI,²⁸⁵ and for “upskilling and reskilling of regulators in the different sectors.”²⁸⁶ We heard that regulators needed “the best technical talent, who are able to know exactly what is going on, as well as working directly with industry [so] that no loopholes can be found in regulation.”²⁸⁷
- 160.** Confidence among regulators we heard from on this point seemed to vary. The EHRC told us, “we do not have the technical expertise on AI that other regulators have. While proposals from the government have set out shared resources for regulators to draw on to fill this gap, this will

278 [Qq17-27](#)

279 Ada Lovelace Institute ([RAI0066](#))

280 Equality and Human Rights Commission ([RAI0089](#)); the EHRC has also stated that its budget has remained static at £17.1 million since 2016

281 [Q11](#)

282 [Q76](#)

283 UNISON ([RAI0076](#))

284 [Q9](#)

285 [Q9](#)

286 [Q17](#)

287 [Q62](#)

be unlikely to close the capability gap we currently face in the short or medium terms.” Ofcom, however, stated that the biggest challenge they faced in this area was the pace of AI change, but that they were confident they had the resources to meet the challenge, and had focussed on hiring technologists.²⁸⁸

Individual redress via regulators

- 161.** The then DSIT pointed out that regulators offer routes for redress when individuals suffer harm. However, witnesses told us this could not be a sufficient solution. Professor David Leslie told us that because of the asymmetry between consumers, regulators and large corporations, “rather than focusing on what an individual can do to have recourse in these situations [we should] think about what structural conditions can be put in place to ensure that you will not have illegal appropriation of biometric data.”²⁸⁹ Handley Gill Ltd told us, “in practice individuals cannot expect regulators to take action in relation to non-compliance and must instead seek to enforce their rights through the courts, if they have the knowledge, resources and commitment to do so.”²⁹⁰ Ravi Naik told us that if rights cannot be upheld, they are “hollow.”²⁹¹ In later chapters we explain how we think these problems should be addressed through law and regulation.

162. CONCLUSION

Existing regulators cannot effectively address human rights harms arising from AI systems using their present powers and resources.

163. CONCLUSION

Human rights risks produced by AI systems cannot be adequately addressed by relying on individual litigants to bring proceedings.

Does regulation stifle or support innovation?

- 164.** Evidence to our inquiry from the then DSIT set out its “pro-innovation AI regulation strategy, which seeks to embed fairness and transparency while reducing compliance burdens.”²⁹²

288 [Q76](#)

289 [Q11](#)

290 Handley Gill Limited ([RAI0070](#))

291 [Q11](#)

292 Department of Science Innovation and Technology ([RAI0077](#))

- 165.** Several witnesses told us that appropriate regulation could support innovation.²⁹³ We heard regulation could “allow for British companies to really take the lead [to] create, develop and support responsible technology,”²⁹⁴ while some requirements like transparency were essential.²⁹⁵ We heard appropriate regulation could promote “legal certainty and to some extent [enable] early adoption of technology,”²⁹⁶ and that a clear and comprehensive framework “would boost the UK’s economic standing and position us as a leader in AI.”²⁹⁷ Private sector companies told us regulation could be “an enabler of trust,”²⁹⁸ and could support innovation.”²⁹⁹

International nature of issue

- 166.** AI use is growing rapidly across the globe. Witnesses told us this meant that co-ordinated action would be needed to address the risks posed by AI.³⁰⁰ The Law Society told us there was a need to “internationalise human rights standards [for AI] globally”. The Catholic Bishops’ Conference of England and Wales noted the UK’s potential role to influence standards and practice in other countries.³⁰¹
- 167.** Witnesses pointed out that the UK had established a leadership role on AI, and urged government to continue this.³⁰² Professor David Leslie called for the country to “leverage the tremendous amount of intellectual and social capital that we have in this space to retain that leadership role on AI.”³⁰³ Dr Janis Wong of the Law Society said “The UK Government [has] already developed momentum, been part of international conversations and led on the discussions about what it means to incorporate and consider justice and human rights through the lens of technology. [...] [it] should continue to foster those conversations and discussions, and lead on the international

293 Minderoo Centre for Technology and Democracy, University of Cambridge ([RAI0067](#)); Duality Ltd ([RAI0071](#))

294 [Q7](#)

295 [Q23](#)

296 [Q41](#)

297 The Law Society of England and Wales ([RAI0021](#))

298 [Q53](#)

299 Good Tech Advisory ([RAI0086](#))

300 Andrea Miotti and Steven Adler ([RAI0031](#))

301 Catholic Bishops’ Conference of England and Wales ([RAI0040](#))

302 Dr Rishi Gulati (Associate Professor in International Law at University of East Anglia) ([RAI0016](#)); Minderoo Centre for Technology and Democracy, University of Cambridge ([RAI0067](#))

303 [Q5](#)

agenda.”³⁰⁴ Minister Narayan told the Committee that government is focussed on the UK playing a significant role in international co-ordination on AI risk mitigation.³⁰⁵

- 168.** Witnesses also told us that harmonising AI regulation would have a beneficial impact on business by providing certainty. SME Trilateral Research told us “alignment with international and regional approaches will minimise the financial and administrative burden on organisations, especially start-ups and SMEs who work internationally.”³⁰⁶

169. CONCLUSION

To address the global nature of AI harms, and protect human rights, governments need to collaborate to identify risks and promote consistent standards.

170. CONCLUSION

The UK can be justifiably proud of its record of leadership on AI. This has been maintained across administrations from different sides of the political debate.

171. CONCLUSION

An appropriate regulatory regime would not only enhance human rights protections but also help to reinforce the UK’s leadership position in AI.

172. RECOMMENDATION

The UK should continue to play a leading role in international AI summits and seek to be party to international declarations and agreements on AI safety. In particular, the government should use its G20 presidency and other available global leadership opportunities to make the case for enabling or building upon international mechanisms to address AI harms and intervene early where necessary in the development of foundational models.

304 [Q44](#)

305 [Q101](#)

306 Trilateral Research ([RAI0046](#))

5 How risks to human rights should be addressed

- 173.** This chapter sets out our findings on the main steps government should take to address the concerns we identified in our inquiry. These are the creation of AI-specific regulation which applies across the AI lifecycle and supply chain, the creation of an independent AI oversight body, strengthened safeguards for automated decision-making, and a stronger role for the AI Security Institute established on a statutory basis.

An AI Bill and AI-specific regulation

- 174.** As we have explained above, we heard from witnesses that existing law and regulation is not suitable to address the risks raised by AI.³⁰⁷ We heard that the UK should develop AI-specific laws and regulation.³⁰⁸
- 175.** The Law Society of England and Wales highlighted the need for mandatory obligations across the AI lifecycle, stating:
- regulation must require transparency by organisations using AI and tangible risk-assessments from AI producers about how their tools are trained and deployed in line with existing standards around equal treatment. There also needs to be full and regular evaluation of deployment against commonly agreed standards of what good looks like to identify and address any unintended harms.³⁰⁹
- 176.** While we have not sought to prescribe the exact details of an expanded regulatory and legal framework, we did find that there was a high degree of consistency among our witnesses as to the features they felt would create a rights-respecting framework for AI regulation. We briefly set out these key features below.

307 Dr Malak Benslama-Dabdoub (Royal Holloway University of London) ([RAI0010](#)); Duality Ltd ([RAI0071](#)); Dr Felipe Romero-Moreno ([RAI0087](#))

308 Dr S. J. Fox, University of Leicester ([RAI0009](#)); Dr Giulia Gentile et al ([RAI0047](#)) Minderoo Centre for Technology and Democracy, University of Cambridge ([RAI0067](#)); Dr Erin Ferguson et al ([RAI0048](#))

309 The Law Society of England and Wales ([RAI0021](#))

177. CONCLUSION

The UK needs an AI-specific legal framework that applies to all sectors of society and the economy, which comprehensively addresses AI-generated risks and harm across the AI lifecycle to address AI's unique challenges. To ensure that the UK remains at the forefront of regulation, this framework must be proportionate.

178. RECOMMENDATION

The government should introduce a new AI Bill, as promised in the 2024 King's Speech, to place AI regulation on a sound basis and give full effect to the Framework Convention.

A risk-based approach

179. AI systems vary in terms of the risks they present. As the EHRC put it, "AI is not a single thing, and the risks and opportunities are unique to a specific tool or model's purpose, use and utility."³¹⁰ We heard that some AI systems will be naturally lower risk, for example applications used for social event planning, shopping, or to automate simple administrative tasks. In contrast, some witnesses told us that AI system applications in specific contexts may present an unacceptable risk to human rights, warranting outright prohibition. Our witnesses overwhelmingly called for a risk-based approach to regulating AI,³¹¹ "tailored to the severity of potential harms,"³¹² and concentrating on where the greatest risks arise in the supply chain.³¹³

180. The Framework Convention takes a risk-based approach, calling for regulatory measures that are differentiated according to the level of risk associated with a particular class or type of AI systems based on their application, purpose and use-context allocated across the entire AI lifecycle.³¹⁴ This is reflected in Article 1(2), which requires the adoption of measures that are

graduated and differentiated as may be necessary in view of the severity and probability of the occurrence of adverse impacts on human rights, democracy and the rule of law throughout the lifecycle of artificial intelligence systems.

³¹⁰ Equality and Human Rights Commission ([RAI0075](#))

³¹¹ British Standards Institution ([RAI0050](#))

³¹² Professor Subhajit Basu (University of Leeds) ([RAI0032](#)); Professor Belen Olmos Giupponi (Professor of Law at Middlesex University London) ([RAI0062](#)); Justice ([RAI0073](#))

³¹³ British Standards Institution ([RAI0050](#)); Professor Belen Olmos Giupponi (Professor of Law at Middlesex University London) ([RAI0062](#)); Handley Gill Limited ([RAI0070](#)); Duality Ltd ([RAI0071](#)); Professor Subhajit Basu (University of Leeds) ([RAI0032](#)).

³¹⁴ [Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law](#), CETS 225, Vilnius, 5 September 2024

- 181.** Many witnesses pointed to the EU AI Act as a possible model.³¹⁵ This Act takes a risk-based approach. Its legal safeguards against AI-generated harm are proportionately more stringent and demanding the higher the level of risk posed by an AI system.
- 182.** The EU AI Act classifies AI systems and/or practices into several ‘tiers,’ with proportionately more demanding legal obligations that vary in accordance with the severity of the risks they pose to health, safety or fundamental rights: (1) systems deemed to pose an ‘unacceptable’ risk are prohibited; (2) systems classified as ‘high risk’ must comply with ‘essential requirements’ (3) general purpose AI models are subjected to obligations that primarily focus on transparency, intellectual property protection, and the mitigation of “systemic risks”; and (4) AI systems posing a limited risk must comply with specific transparency requirements.
- 183.** We heard considerable support for an approach on these lines. The Ada Lovelace Institute told us it would mean “systems [...] are more likely to be rights-compliant at the point of deployment, helping address risk proliferation.”³¹⁶ Silkie Carlo of Big Brother Watch told us that although the EU Act was far from perfect, “we need something like an EU-style AI Act or a digital Bill of Rights: a specific piece of legislation that takes a risk-based approach to artificial intelligence and [...] can put red lines in.”³¹⁷ The Independent Society of Musicians told us that “the UK must move away from reactive policy-making to a proactive precautionary model that embeds human rights as a baseline for innovation.”³¹⁸
- 184.** Others were critical of the EU AI Act, claiming that it was overly bureaucratic, and risked unnecessarily stifling innovation and growth. Bates Wells LLP told us that the EU AI Act was “overly complex”, citing calls from some in the EU to pause the rollout of the Act.³¹⁹ Dualarity Ltd noted that the Act had been characterised as “overly bureaucratic” especially for smaller companies, and suggested the UK adopt a “UK-flexible way [of regulating] that maintains our attractive environment for AI R&D.”³²⁰
- 185.** We heard that a principles-based approach would be the most suitable basis on which to regulate this rapidly developing sector.³²¹ Witnesses told us this could help overcome the problem of pace, where AI develops more quickly than regulation designed to address its risks. Framework legislation

315 [Regulation \(EU\) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence \(Artificial Intelligence Act\)](#)

316 Ada Lovelace Institute ([RAI0066](#))

317 [Q33](#)

318 The Independent Society of Musicians ([RAI0004](#))

319 Eleonor Duhs (Bates Wells LLP) ([RAI0035](#))

320 Dualarity Ltd ([RAI0071](#))

321 Dr Malak Benslama-Dabdoub (Royal Holloway University of London) ([RAI0010](#)); Dualarity Ltd ([RAI0071](#)); Connected by Data ([RAI0013](#))

could allow for the creation of codes of practice which could more easily be kept up to date than primary legislation,³²² which might “becom[e] soon obsolete as [it] cannot keep up with innovation and this may in turn jeopardise effective human rights protection.”³²³ Alexandria Walden of technology firm Google called for regulation of AI to avoid duplication: “regulating AI is necessary but must be done well [...] We really do not want duplicative laws or to be reinventing the wheel.”³²⁴

186. CONCLUSION

A risk-based approach is needed to address AI harms. The Hugging Face incident illustrates the significant risks that can arise from the newest and most advanced AI models, and legislation should address those risks on a precautionary basis with a view to avoiding harm, by analogy with the regulation of other potentially harmful products such as pharmaceuticals. It should only impose legal obligations where they are necessary and proportionate to prevent unacceptable risks and ensure effective redress for harm. Low risk AI systems should be subject to fewer and less demanding legal requirements. Legislation should not place unjustified burdens on businesses, especially SMEs, to enable UK companies to innovate and for the UK to benefit from the advantages AI can confer.

187. CONCLUSION

To help ensure that regulation can keep pace with technological developments, it is preferable for legislation to take the form of broadly framed principles, with more detailed guidance provided by codes of practice.

188. RECOMMENDATION

The AI Bill we recommend should provide for a regulatory regime that would classify different risk levels, and mandate proportionately more demanding obligations for higher-risk AI systems and AI models.

322 UNISON ([RAI0076](#)); The Law Society of England and Wales ([RAI0021](#))

323 [Q41](#); European Center for Not-For-Profit Law (ECNL) ([RAI0008](#))

324 [Q64](#)

Allocating responsibility where it belongs

- 189.** Many witnesses said that action was needed to ensure that legal responsibility to prevent and mitigate AI-generated harm is fairly, effectively and proportionately allocated across the entire AI lifecycle.³²⁵ Alexandria Walden of Google told us “responsibility [for harm] should look different for different actors, depending on their relation to that harm.”³²⁶ Ellen Lefley told us that this was achievable, drawing an analogy with knives:

It is illegal to manufacture certain types of knives. It can also be illegal to sell certain types of knives. With respect to legal knives that are not banned, it is unlawful to sell them to certain groups of people. There are numerous laws that restrict what people can do with knives. [...]. There are existing examples where things pose risk and we can segment liability across the supply chain. That absolutely needs to be the approach with respect to AI and the risks that it poses.³²⁷

- 190.** As evidence from the then DSIT confirmed, in “contractual situations [...]—with the exception of consumer contracts—parties are generally free to allocate responsibility for harms as they see fit.”³²⁸ We heard that this situation, coupled with “power imbalances,”³²⁹ created a likelihood that better resourced firms upstream will tend to seek to “pass on” the responsibility for mitigating or addressing AI risks through contract terms to smaller or less well resourced firms who are closer to the consumer or point of use.³³⁰
- 191.** In turn, we heard these firms lower down the supply chain are less likely to be able to prevent or address harms, for three reasons: first, because they have fewer resources on which to draw; second, because they operate lower downstream and so are less likely in practice to have access to the components of AI systems, and thirdly, because downstream deployers may lack the expertise and understanding necessary to take appropriate action to address those risks.

325 Chartered Institute of Library and Information Professionals (CILIP) ([RAI0053](#)); Liberty ([RAI0079](#)) Eleonor Duhs (Bates Wells LLP) ([RAI0035](#))

326 [Q72](#)

327 [Q41](#)

328 Department of Science Innovation and Technology ([RAI0077](#))

329 Professor Lorna McGregor (University of Essex) ([RAI0056](#))

330 Ada Lovelace Institute ([RAI0066](#))

192. Witnesses argued strongly for government to introduce regulations which would ensure that legal responsibility for preventing AI harms falls on those most able to do so.³³¹ They called for legally binding due diligence obligations throughout the AI lifecycle, and at every stage of the supply chain including stronger redress mechanisms.³³²

193. CONCLUSION

AI supply chains are complex. Actors which have the greatest capacity to prevent harm are often upstream developers. However at present it is too easy for the risks to be passed down the chain to those who are not in a position to identify and address the source of those risks located upstream.

194. CONCLUSION

Many AI systems developed for a specific purpose can be deployed or adapted for use for different purposes and contexts from those intended by the AI developer. Those deploying the system may lack understanding of the risks arising from the way in which the AI system has been designed and developed, and how its deployment could adversely impact human rights.

195. RECOMMENDATION

The AI Bill we recommend should impose proportionate obligations on actors at all stages of the AI lifecycle and supply chain to ensure that risks are identified, evaluated and effectively addressed at the appropriate stage.

331 Liberty ([RAI0079](#)); Professor Subhajit Basu (University of Leeds) ([RAI0032](#)); Eleonor Duhs (Bates Wells LLP) ([RAI0035](#)); Privacy International ([RAI0081](#)); Public Law Project ([RAI0045](#)); Professor Lorna McGregor (University of Essex) ([RAI0056](#)); Amnesty International UK ([RAI0051](#)); Chartered Institute of Library and Information Professionals (CILIP) ([RAI0053](#)); Dr Ayça Atabey; Dr Kim Sylwander; Professor Sonia Livingstone ([RAI0058](#)); Professor Belen Olmos Giupponi (Middlesex University London) ([RAI0062](#)); Minderoo Centre for Technology and Democracy, University of Cambridge ([RAI0067](#)); Ada Lovelace Institute ([RAI0066](#)); Duality Ltd ([RAI0071](#)); Justice ([RAI0073](#)); UNISON ([RAI0076](#)); European Center for Not-For-Profit Law (ECNL) ([RAI0008](#)); Connected by Data ([RAI0013](#)), Catholic Bishops' Conference of England and Wales ([RAI0040](#)); Information Commissioner's Office ([RAI0003](#))

332 UNISON ([RAI0076](#)); Fair Vote UK ([RAI0002](#)); [Q24](#); Gina Romero et al ([RAI0023](#)); Amnesty International UK ([RAI0051](#)); European Center for Not-For-Profit Law (ECNL) ([RAI0008](#)); Duality Ltd ([RAI0071](#)); Dr Giulia Gentile et al; ([RAI0047](#)); Chartered Institute of Library and Information Professionals (CILIP) ([RAI0053](#)), Ada Lovelace Institute ([RAI0066](#)); Justice ([RAI0073](#)); [Q91](#)

Prohibitions

- 196.** We heard from several witnesses that the UK should introduce ‘red lines’ prohibiting the most egregious AI use cases and practices as part of a risk-based regulatory approach.³³³ This comports with Article 16(4) of the Framework Convention, which provides that a complete ban or moratorium may be required for uses of AI systems assessed as ‘incompatible’ with respect for human rights, the functioning of democracy or the rule of law.
- 197.** As noted above, some AI practices are simply prohibited under the EU AI Act.³³⁴ They include AI systems that use subliminal techniques; exploit people’s vulnerabilities; use social scoring;³³⁵ use profiling to predict the risk that someone will commit a criminal offence; carry out untargeted scraping of images to build facial recognition databases; infer emotions in the workplace or education institutions; use biometric data to deduce personal characteristics; or use real-time biometric identifications systems in public places for law enforcement purposes, except for specified and targeted reasons.
- 198.** There are no equivalent prohibitions in the UK. We heard AI is being used in the public sector for activities which include fraud detection in the benefit system,³³⁶ predictive policing,³³⁷ and using real-time biometric identifications systems in public places.³³⁸ Several witnesses from the NGO sector called for the introduction of ‘red lines’ in the UK, which clearly state when AI should not be used, such as for Facial Recognition Technology,³³⁹ or for “emotional recognition or manipulation” in the workplace.³⁴⁰

333 Amnesty International UK ([RAI0051](#)); [Q33](#); Liberty ([RAI0079](#))

334 [Regulation \(EU\) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence \(Artificial Intelligence Act\)](#), Article 5.

335 Article 5(1)(c) of the EU AI Act prohibits “the placing on the market, the putting into service or the use of AI systems for the evaluation or classification of natural persons or groups of persons over a certain period of time based on their social behaviour or known, inferred or predicted personal or personality characteristics, with the social score leading to either or both of the following: (i) detrimental or unfavourable treatment of certain natural persons or groups of persons in social contexts that are unrelated to the contexts in which the data was originally generated or collected; (ii) detrimental or unfavourable treatment of certain natural persons or groups of persons that is unjustified or disproportionate to their social behaviour or its gravity”.

336 [Universal Credit Advances Model - Fairness Assessment - GOV.UK](#), published 17 July 2025, accessed 27 August 2026

337 Connected by Data ([RAI0013](#)); Dr Mark Wong (University of Glasgow) ([RAI0014](#))

338 Dr S. J. Fox (University of Leicester) ([RAI0009](#)); Glitch ([RAI0015](#)); Big Brother Watch ([RAI0036](#)); Amnesty International UK ([RAI0051](#)); European Center for Not-For-Profit Law (ECNL) ([RAI0008](#))

339 [Q30](#); [Q37](#);

340 Dr Joe Atkinson (University of Southampton) ([RAI0080](#))

199. CONCLUSION

Some uses of AI should be prohibited, because they are incompatible with human rights.

200. RECOMMENDATION

The AI Bill we recommend should prohibit the use of AI in ways that are inconsistent with respect for human rights, such as subliminal techniques, emotional inference or inappropriate use of profiling or biometric data. Because of the very serious human rights risks, this would mean that the development and provision of very powerful AI systems (such as Artificial General Intelligence and Artificial Superintelligence) which risk causing widespread and very serious harm, including the capacity to evade effective human control, would be prohibited.

201. RECOMMENDATION

The detail of such prohibitions should be established following public consultation.

Prior approval for high-risk activities

202. For activities that pose a high risk of harm, we heard that the government should introduce a requirement for a public oversight body to give its approval before such activities can begin. Such a body could also set and monitor strict regulatory requirements that apply on an ongoing basis to such high-risk activities, in order to manage risks before they materialise.³⁴¹

203. For risks that are impossible to meaningfully quantify but potentially catastrophic, witnesses told us a precautionary approach entailing some forms of regulatory intervention may be warranted.³⁴² This could entail periodic reporting, and the provision of information before the activity begins to enable a public oversight body to actively monitor and test for emerging risks and threats.³⁴³

204. RECOMMENDATION

The AI Bill we recommend should include a requirement that AI systems posing a high risk of causing harm to human rights should need prior approval before they can be provided or deployed.

341 [Q26](#); [Qq1-16](#)

342 [Q27](#); [Glitch \(RAI0015\)](#); Amnesty International UK ([RAI0051](#)); Ada Lovelace Institute ([RAI0066](#))

343 The Independent Society of Musicians ([RAI0004](#)); Ada Lovelace Institute ([RAI0066](#)); Liberty ([RAI0079](#))

Obligations of due diligence

- 205.** We heard evidence that both those developing and using AI systems considered high risk should be required both to undertake, and demonstrate that they have taken appropriate due diligence measures, and to demonstrate they have taken due care to prevent human rights harm as appropriate to their role and contribution across the AI lifecycle.³⁴⁴ These include conducting risk management activities, which should include an assessment of the likely impact of the system on the human rights of those affected by its use, to evaluate their seriousness, and to take action to prevent or control those risks, based on the findings of those assessments.³⁴⁵
- 206.** Other submissions also emphasised the need for other measures, including the importance of identifying and assessing the human rights risks associated with design decisions,³⁴⁶ the need for human rights protection measures ‘by design’³⁴⁷ and other technical safeguards, data governance,³⁴⁸ adequate testing,³⁴⁹ and the involvement of those who will be affected in the risk assessment, evaluation and management process.³⁵⁰ We heard that decisions to design, develop and deploy AI systems in the knowledge of their unpredictability “should attract the presumption of responsibility.”³⁵¹
- 207. RECOMMENDATION**
The AI Bill we recommend should place requirements on actors at all stages of the AI supply chain to undertake due diligence obligations necessary to prevent unacceptable human rights risks when designing, developing or deploying high risk AI systems.

344 Privacy International ([RAI0081](#)); The Independent Society of Musicians ([RAI0004](#))
Dr Malak Benslama-Dabdoub (Royal Holloway University of London) ([RAI0010](#)); Joy
([RAI0001](#)); Online Safety Act Network ([RAI0029](#)); Trilateral Research ([RAI0046](#)); The Law
Society of England and Wales ([RAI0021](#))

345 Minderoo Centre for Technology and Democracy, University of Cambridge ([RAI0067](#));
Dr Malak Benslama-Dabdoub (Lecturer in Law at Royal Holloway University of London)
([RAI0010](#)); Mr Burak Batuhan Karakus (Head of Legal Affairs at Human Rights Solidarity)
([RAI0026](#)); Dr Joe Atkinson (University of Southampton) ([RAI0080](#))

346 Professor Subhajit Basu (University of Leeds) ([RAI0032](#)); Privacy International ([RAI0081](#))

347 Privacy International ([RAI0081](#)); Liberty ([RAI0079](#))

348 Privacy International ([RAI0081](#))

349 Professor Subhajit Basu (University of Leeds) ([RAI0032](#))

350 Amnesty International UK ([RAI0051](#))

351 Professor Lorna McGregor (University of Essex) ([RAI0056](#))

208. RECOMMENDATION

We recommend that the due diligence requirements contained in the AI Bill should be differentiated, based on the role of the relevant actor along the AI supply chain, and the nature and seriousness of the risk posed by the relevant AI model or system. To identify the nature and scope of those due diligence obligations, the EU AI Act may provide a useful point of departure, but the provisions should be tailored to the requirements of the UK.

Transparency

- 209.** As discussed above, witnesses told us that deployers of AI systems should be obliged to inform people when they are interacting with an AI system or if an AI system is used to automate or inform legal or other significant decisions about them, or if it uses data they own or have created.³⁵² We heard that without such requirements, accountability and responsibility for the operation of AI systems cannot be secured. As a point of comparison, under the EU AI Act, individuals have a right to know when AI is being used and how it affects them.³⁵³
- 210.** Witnesses told us public authorities should be required to explicitly inform decision subjects, or those affected by a decision or action taken by a public authority, about the use of an algorithmic or automated tool or system when communicating the decision to them. We also heard they should have to provide a proactive and personalised explanation to affected individuals alongside a notification providing information on how and why the decision was reached, including specific categories of information such as tailored information on the contribution of the automated or algorithmic tool or system in the decision-making process.³⁵⁴
- 211.** We also heard evidence about the need for transparency requirements between developers and deployers. Witnesses told us AI suppliers should provide sufficient detail “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.”³⁵⁵

352 Chartered Institute of Library and Information Professionals (CILIP) ([RAI0053](#)); Professor Belen Olmos Giupponi (Middlesex University London) ([RAI0062](#)); Dr Malak Benslama-Dabdoub (Royal Holloway University of London) ([RAI0010](#))

353 EU AI Act, Article 50, in force on 2 August 2026 (see European Commission, [Guidelines on transparency obligations for providers and deployers of certain AI systems](#), 6 August 2026).

354 Public Law Project ([RAI0045](#)); Ada Lovelace Institute ([RAI0066](#))

355 Connected by Data ([RAI0013](#)); Professor Subhjit Basu (University of Leeds) ([RAI0032](#))

- 212.** The Law Society called for the Algorithmic Transparency Recording Standard, which currently applies only to public sector bodies, to be extended to private sector organisations to help them support human rights standards that respect individual rights.³⁵⁶

213. RECOMMENDATION

The AI Bill we recommend should impose mandatory transparency requirements on all actors across the AI lifecycle. The deployment of AI systems which may have significant impacts on individuals, groups and communities should include an obligation to state when AI systems are being used and provide a full and comprehensible explanation of what it is being used for. There should also be an obligation to provide information about the source of data used by the AI system.

Automated decision-making

- 214.** We heard concerns that the protections in the UK GDPR in relation to automated decision-making are insufficiently robust. Chapter 4 summarises the evidence we received about the risks of relying on human involvement as a protection without further safeguards. We were also told that many organisations are unaware of, misunderstand or disregard their legal obligations.³⁵⁷ The ICO’s report on automated recruitment concluded that employers have more work to do to ensure that the use of automated recruitment tools respects people’s information rights.³⁵⁸
- 215.** We heard that individuals may not know they have been subject to a decision made with AI, and therefore may be unable to enforce their rights.³⁵⁹ The Ada Lovelace Institute told us that people who are subject to automated decision-making are usually not given enough information to allow them to understand if the system itself is biased, or whether they were discriminated against in their individual case, and only entitles them to an explanation of the ‘general logic’ of a system rather than a personalised explanation.³⁶⁰ The Public Law Project told us: “New, clearer individual notice requirements are needed to ensure that people are able to meaningfully and effectively make representations and contest the decisions that are made about them.”³⁶¹

356 The Law Society of England and Wales ([RAI0021](#))

357 Eleonor Duhs (Bates Wells LLP) ([RAI0035](#))

358 Information Commissioner’s Office, [Recruitment rewired: an update on the ICO’s work on the fair and responsible use of automation in recruitment](#), 31 March 2026.

359 Glitch ([RAI0015](#))

360 Ada Lovelace Institute ([RAI0066](#))

361 Public Law Project ([RAI0045](#))

216. AI Minister Kanishka Narayan told us that “the existing data protection framework [...] says that we should only consider [automated] processing, [...] where it is proportionate and fair. There are a number of areas [...] where those tests are not likely to be met and, even when they are met, they should be done so only in a way that still allows individuals to be informed, to contest, and to secure human intervention as well.”³⁶²

217. CONCLUSION

Although data protection law provides for specific safeguards when automated decision-making is used, there is a need for greater clarity in the law to ensure that the safeguards operate effectively in practice, and that reliance on human involvement is not misplaced.

218. RECOMMENDATION

The regulation-making powers in the UK GDPR should be used to ensure that there are robust protections against misuse of automated decision-making. The regulations should say that the mere presence of a ‘human in the loop’ is not enough to constitute meaningful human involvement or intervention. It should be necessary to show that the person is sufficiently informed and independent to be able to reach an objective view that has not been improperly influenced by the automated decision. The regulations should also say that safeguards include a requirement to provide information relating to the data subject’s individual circumstances rather than the general operation of the system, and that data subjects must be given enough information to allow them to mount an effective challenge against the decision.

An AI regulator

219. As described above, the UK’s existing regulators cannot fully address the risks we have identified. We heard evidence from several witnesses that there should be a regulator with powers to enforce specific requirements at all stages of the AI lifecycle.³⁶³ Witnesses varied in their views of the form this should take. Some called for an ombudsman-style model or “watchdog” to receive and act on complaints.³⁶⁴ There was also support for a central body to monitor and report to government on AI harms,³⁶⁵

³⁶² [Q110](#)

³⁶³ [Q6](#); Leena Sowambur ([RAI0063](#)) Mr Burak Batuhan Karakus (Human Rights Solidarity) ([RAI0026](#)); Good Tech Advisory ([RAI0086](#))

³⁶⁴ [Q59](#); Ada Lovelace Institute ([RAI0066](#)); Fair Vote UK ([RAI0002](#)); Liberty ([RAI0079](#))

³⁶⁵ Ada Lovelace Institute ([RAI0066](#));

and to have “oversight and audit powers [and] access to the data used to train the AI system and the algorithm to pursue investigations and effective monitoring.”³⁶⁶

- 220.** Many witnesses emphasised the value of a sector-specific regulatory approach. Alexandria Walden of Google told us that the way “AI gets used in [a specific] field should be taken up by that field, because it may look different in the law versus in tax advising or some other field.”³⁶⁷ Ofcom told us that if a co-ordinating body were established, it should delegate to existing regulators on areas that were within their expertise.³⁶⁸ The EHRC told us a “single AI regulator would duplicate and complicate existing regulation”³⁶⁹ and that government should prioritise ensuring that “existing regulators are sufficiently funded and funded to be able to work together.”³⁷⁰ The ICO told us that any proposals for a new oversight body, “should cover well-evidenced gaps that could not be addressed by empowering and resourcing current regulators.”³⁷¹ ACT/The App Association, a body representing tech SMEs warned of the risk of creating overlaps or a lack of consistency with any new regulator: “Parliament must first understand how existing frameworks apply to activities involving AI.”³⁷²
- 221.** Other witnesses highlighted the importance of ensuring that existing regulators were properly supported and resourced to do their jobs properly. Professor Lorna McGregor of the University of Essex, a former EHRC Commissioner, recommended the extension of the remits of existing regulators, accompanied by extra resources.³⁷³ UNISON told us that “enhanced enforcement mechanisms are essential, requiring regulators to have adequate funding and resources along with the technical expertise to examine AI algorithms and datasets.”³⁷⁴ Ellen Lefley of Justice told us that an AI authority with “a duty of vigilance over the sufficiency of existing regulatory schemes and existing law” would be “a good middle ground where you embrace the benefits of a sector-specific regulatory approach, but you also have someone charged with the job of spotting those gaps and filling them.”³⁷⁵

366 Dr Giulia Gentile et al ([RAI0047](#))

367 [Q66](#)

368 [Q77](#)

369 Equality and Human Rights Commission ([RAI0089](#))

370 [Q77](#)

371 Information Commissioner’s Office (ICO) ([RAI0085](#))

372 ACT | The App Association ([RAI0043](#))

373 Professor Lorna McGregor (University of Essex) ([RAI0056](#))

374 UNISON ([RAI0076](#))

375 [Q40](#)

222. One factor in favour of a possible overarching regulator is the fact that, as discussed in previous chapters, there is no single channel for concerns about AI risks or harms to be raised. The burden is on the individual to work out which regulatory body they can contact (if any) and then to pursue a complaint, which may or may not lead to further action. If individuals are ignored by regulatory bodies, there is very little they can do (other than take legal action in relation to their own individual claim).³⁷⁶

223. Public complaints and reports from whistleblowers form a major part of the intelligence that regulators use to guide their actions and investigations.³⁷⁷ We heard that detecting patterns of harm in AI systems at the earliest possible stage in the AI lifecycle is important due to the significant human rights risks posed by AI systems, and the ability for these to proliferate rapidly if left unchecked.³⁷⁸ Some witnesses called for a public database of reported harm.³⁷⁹

224. CONCLUSION

We agree with our witnesses that the present regulatory framework is fragmented and difficult to navigate. Urgent action is needed to close gaps in protection arising from the current framework.

225. CONCLUSION

A single AI oversight body would be able to act as a single point of contact for raising concerns about use of AI, carry out oversight and monitoring of AI harms and risks, and maintain an adverse incident repository to help identify where more systemic issues may lie. We are confident this can be done without creating duplication.

226. RECOMMENDATION

The government should establish an independent AI oversight body on a statutory footing, either as a new body or by expanding the powers of an existing body. The government should consult carefully with existing regulators to avoid unnecessary overlap and promote regulatory consistency and coherence.

³⁷⁶ Eleonor Duhs (Bates Wells LLP) ([RAI0035](#))

³⁷⁷ ICO, [How we handle concerns](#), accessed 15 May 2026; ICO, [Report bad practices as a whistleblower](#), accessed 15 May 2026

³⁷⁸ Ada Lovelace Institute ([RAI0066](#))

³⁷⁹ Fair Vote UK ([RAI0002](#))

227.

RECOMMENDATION

The new oversight body should:

- a.** have the power to facilitate co-ordination with other national regulators in relation to monitoring, investigation and enforcement, knowledge sharing, and capacity building
- b.** make regular reports on its activities to Parliament, identifying opportunities to fill emerging gaps in protection and highlighting trends and sectoral developments in best practice
- c.** have a statutory duty to maintain a public repository of AI incidents and proactively identify patterns of reported harm and any emerging systemic threats, and to consider cases for further investigation
- d.** have the power to publish mandatory codes of practice for the responsible use of AI in the private and public sector, and to sanction developers, deployers and other actors who do not fulfil their responsibilities under these codes of practice
- e.** have the powers necessary to implement and enforce the provisions of the AI Bill we recommend, including:
 - i.** to carry out testing and evaluation of AI systems, including powers to implement a prior testing, auditing and evaluation regime for high-risk uses and applications of AI
 - ii.** to prohibit AI models or systems from being launched or deployed, and to order their withdrawal from the market if unacceptable risks to human rights are identified
 - iii.** to issue codes of practice which would provide greater specificity concerning the due diligence obligations which must be complied with (for example, requiring high risk AI providers to establish and maintain a risk management system, to ensure that practical and acceptable complaints mechanisms are in place and so forth), and to receive reports if certain types of risk are identified
 - iv.** to conduct investigations of allegations concerning breaches of the statutory prohibitions and mandatory codes of conduct, including allegations that AI systems have caused harm to human rights
 - v.** to sanction actors who fail to comply with mandatory codes of practice

- vi. to order that remedies be provided as appropriate in individual cases without requiring the affected individuals to incur prohibitive costs
- vii. to prescribe transparency requirements through its codes of practice, identifying the information actors at each stage of the AI lifecycle must disclose and to whom, including other actors along the AI lifecycle and public sector AI deployers
- viii. where actors in the AI supply chain are located outside the UK, to be able to assess the risks of deploying their systems in the UK, having regard to how those actors are regulated in their home state. If the new AI oversight body considers that deployment in the UK would give rise to unacceptable risks which could not be addressed within the UK, the oversight body should be able to limit or prohibit such deployment.

The AI Security Institute (AISI)

- 228.** We discussed above the voluntary nature of engagement between the AISI and firms developing and deploying AI, and its lack of regulatory authority. We also noted that witnesses felt the AISI’s 2025 shift in focus from “safety” to “security,” which the government said would “strengthen [...] protections against the risks AI poses to national security and crime” (see Chapter 3 above), showed a “deregulatory” approach to AI by government.
- 229.** The UK government has previously stated its intention to place the AISI on a statutory footing,³⁸⁰ but it has not yet done so. The AI Minister told us “the AI Security Institute is unparalleled globally in its access to models pre-deployment” and that he had seen no indication of a need for reform to enhance its access and capabilities, but promised to keep the question under review.³⁸¹
- 230.** Major AI firms have issued voluntary commitments to publish safety frameworks, define ‘intolerable’ risk thresholds that would justify halting the development or deployment of advanced AI systems, and to refrain from deploying such systems if risks could not be reliably reduced below those thresholds and to maintain internal accountability mechanisms with transparency to public authorities and, where feasible, the wider public. We heard evidence from technical experts that many in the AI industry consider

380 [Written questions and answers - Written questions, answers and statements - UK Parliament](#)

381 [Q104](#)

these voluntary commitments inadequate, unreliable and ineffective, particularly given the strong financial incentives driving firms to avoid oversight.³⁸²

231. RECOMMENDATION

The government should place AISI on a statutory basis. AISI should have the power to review new and revised powerful AI models. It should also have a duty to assess, publish findings pre-release, and issue preliminary ‘warning’ statements associated with release of new AI foundation models. To avoid duplication, the government should consider combining the existing functions of the AISI with those of the new regulator.

232. RECOMMENDATION

Developers of powerful AI models must be required to submit new models (and new versions of existing models) to AISI for review, evaluation, and testing. They must also provide technical specifications, including model properties, training data and process, intended use, security and safety testing and risk control measures.

233. RECOMMENDATION

The government must give the AISI a statutory duty to inform the authorities, including the new AI oversight body, of potential risks posed by AI models and systems. This should include a duty to advise the oversight body of the need to prevent AI systems from being launched or deployed, or to order them to be withdrawn from the market, if AISI identifies significant risks to human rights.

382 Andrea Miotti (Founder and CEO at ControlAI); Steven Adler) ([RAI0031](#))

Conclusions and recommendations

What is artificial intelligence?

1. AI supply chains are complex and opaque, involving multiple actors. This makes it difficult to identify the source of human rights risks, and the actors responsible for identifying and addressing those risks. (Conclusion, Paragraph 20)

Human rights issues raised by AI

2. The Council of Europe Framework Convention on AI provides a strong basis for both effective regulation of AI within the UK, and collective international action. We urge the government to seize the opportunity to be an early adopter of the Framework Convention. (Conclusion, Paragraph 26)
3. The UK should set out a timeline for ratification of the Council of Europe Framework Convention detailing what steps will be taken and when. This plan for ratification should be subject to public consultation. (Recommendation, Paragraph 27)
4. AI systems are prone to producing unfairly discriminatory outcomes, which may arise from bias in the datasets on which they are trained as well as from the way they are developed and deployed. This threatens the human rights of people in the UK and elsewhere. (Conclusion, Paragraph 38)
5. AI systems pose particular risks to privacy and personal data, because of the prevalence of web scraping to collect personal data to train large language models (LLMs), and the potential for AI systems to be used in highly intrusive ways. This threatens the right to privacy of people in the UK and elsewhere. (Conclusion, Paragraph 45)
6. AI systems pose particular risks to the right to an effective remedy in circumstances where there is insufficient transparency concerning the use of AI systems to support significant decisions about affected persons, and in

circumstances where it is difficult for those persons to access mechanisms to challenge those decisions. This threatens the human rights of people in the UK. (Conclusion, Paragraph 49)

7. Our inquiry has not focused on intellectual property rights, or the rights of performers. The issues which witnesses have raised clearly merit further investigation, and could require specific measures going beyond the scope of this inquiry. It is important for the UK to protect its rich creative industries and those who contribute to them, and to safeguard performers' human rights as well as their livelihoods. (Conclusion, Paragraph 53)
8. Our inquiry did not focus on the potential effects of AI on the justice system, but witnesses raised some very concerning issues which merit detailed consideration, and might require specific measures going beyond the scope of this inquiry. The system is already under great strain. We note the Justice Select Committee has also recently discussed with ministers the use of AI in connection with the legal system, as part of its ongoing inquiry into Access to Justice. (Conclusion, Paragraph 55)
9. We welcome the government's stated commitment to ensuring that in the UK, development and use of AI systems is grounded in human rights. AI has the potential to enhance the protection of human rights, but can also give rise to significant interferences with those rights. (Conclusion, Paragraph 60)
10. Human rights principles must inform all aspects of the government's policies and practices relating to AI systems, including its approach to the regulation of AI. (Recommendation, Paragraph 61)
11. Government should encourage the development of AI use cases which can demonstrably improve human rights. (Recommendation, Paragraph 62)

Legal and policy framework

12. The AI Action Plan commissioned by the previous administration, while it brings a welcome focus on innovation and economic opportunity, fails to pay sufficient attention to the need to secure the protection of human rights from AI-generated interference throughout the supply chain. (Conclusion, Paragraph 89)
13. The UK's leading position as the third largest AI market in the world affords opportunities, not only for economic growth but also for setting strong standards on the responsible adoption and development of new technologies. (Conclusion, Paragraph 90)

14. The Committee is encouraged by the prominence given to AI in the new administration and urges the government to seize the opportunity to regulate AI in a way which puts human rights first. (Recommendation, Paragraph 91)

Problems with existing protections

15. Taken together, existing UK laws that could apply to AI systems do so primarily at the point of deployment, so that deployers are the primary holders of responsibility for harms arising from AI systems. This means that existing laws do not take due account of the complexity of the AI lifecycle and supply chain and are not directed to the actor that may be best placed to identify and avert the source of risks. They are therefore likely to allow human rights harm that could have been prevented. (Conclusion, Paragraph 104)
16. Regulators lack the power to test and evaluate AI systems before they are publicly released, or to prevent their release if they are considered to pose unacceptable risks, except in very limited circumstances. (Conclusion, Paragraph 108)
17. Model developers engage with the AISI on a voluntary basis. The AISI has no statutory power. (Conclusion, Paragraph 109)
18. No cross-economy regulators, agencies or bodies have systematic powers to prevent AI models or systems being deployed even if they pose significant risks. (Conclusion, Paragraph 110)
19. Existing law is deficient by failing to impose legal requirements on AI system deployers to be transparent about their use of AI systems. (Conclusion, Paragraph 118)
20. The current legal framework does not provide for adequate transparency in relation to the development and deployment of AI systems. (Conclusion, Paragraph 130)
21. Existing law does not require the provision of notice to individuals that an AI system is in use. In the absence of such knowledge, individuals who have been directly affected by the use of AI systems will be in practice unable to challenge such use or seek an effective remedy for human rights violations. The evidence we heard strongly suggests that many people may be in this situation. (Conclusion, Paragraph 131)

22. There are no mandatory obligations on AI developers and providers to provide downstream actors, including deployers, with information concerning the potential risks arising from the deployment of AI systems. This can reduce transparency and traceability across the AI supply chain. (Conclusion, Paragraph 132)
23. The ATRS provides some transparency concerning the use of AI systems in the public sector, but is limited in its coverage. (Conclusion, Paragraph 133)
24. The recent OpenAI/Hugging Face security incident starkly highlights the societal risks posed by the growing power of AI systems. It illustrates how the newest AI models can engage in activities that would be unlawful if undertaken by a legal person, with the potential to seriously undermine security, including in critical infrastructure. It also suggests that at present society largely relies on private companies to take voluntary action to guard against such incidents. This is unsatisfactory. (Conclusion, Paragraph 143)
25. The general-purpose nature of the foundation models which allows them to be embedded into a wide and diverse range of AI systems and contexts means that any vulnerabilities and sources of risk within the model can propagate very rapidly and at scale once deployed. The power of the newest models, and their potential for mass harm, justifies a precautionary approach to regulation. (Conclusion, Paragraph 144)
26. Existing law is inadequate to enable those who suffer human rights harm generated by AI systems to secure an effective remedy. (Conclusion, Paragraph 145)
27. Even if some existing laws could address some human rights risks posed by AI systems and provide redress to individuals in specific circumstances, their application is often speculative and uncertain, and the cost of individual litigation is in practice prohibitive. (Conclusion, Paragraph 146)
28. The current legal framework that applies to AI systems is fragmented and difficult to navigate. It relies primarily on harm-specific and sector-specific legislation leaving gaps in protection which fails to address preventable human rights risks. (Conclusion, Paragraph 155)
29. Current product liability law provides strict liability for AI systems incorporated into physical products, but not in relation to stand-alone AI systems. This difference in legal liability is difficult to justify in principle. The government must carefully consider any recommendations from the ongoing Law Commission review, including on this specific point. (Conclusion, Paragraph 156)
30. Existing regulators cannot effectively address human rights harms arising from AI systems using their present powers and resources. (Conclusion, Paragraph 162)

31. Human rights risks produced by AI systems cannot be adequately addressed by relying on individual litigants to bring proceedings. (Conclusion, Paragraph 163)
32. To address the global nature of AI harms, and protect human rights, governments need to collaborate to identify risks and promote consistent standards. (Conclusion, Paragraph 169)
33. The UK can be justifiably proud of its record of leadership on AI. This has been maintained across administrations from different sides of the political debate. (Conclusion, Paragraph 170)
34. An appropriate regulatory regime would not only enhance human rights protections but also help to reinforce the UK's leadership position in AI. (Conclusion, Paragraph 171)
35. The UK should continue to play a leading role in international AI summits and seek to be party to international declarations and agreements on AI safety. In particular, the government should use its G20 presidency and other available global leadership opportunities to make the case for enabling or building upon international mechanisms to address AI harms and intervene early where necessary in the development of foundational models. (Recommendation, Paragraph 172)

How risks to human rights should be addressed

36. The UK needs an AI-specific legal framework that applies to all sectors of society and the economy, which comprehensively addresses AI-generated risks and harm across the AI lifecycle to address AI's unique challenges. To ensure that the UK remains at the forefront of regulation, this framework must be proportionate. (Conclusion, Paragraph 177)
37. The government should introduce a new AI Bill, as promised in the 2024 King's Speech, to place AI regulation on a sound basis and give full effect to the Framework Convention. (Recommendation, Paragraph 178)
38. A risk-based approach is needed to address AI harms. The Hugging Face incident illustrates the significant risks that can arise from the newest and most advanced AI models, and legislation should address those risks on a precautionary basis with a view to avoiding harm, by analogy with the regulation of other potentially harmful products such as pharmaceuticals. It should only impose legal obligations where they are necessary and proportionate to prevent unacceptable risks and ensure effective redress for harm. Low risk AI systems should be subject to fewer

and less demanding legal requirements. Legislation should not place unjustified burdens on businesses, especially SMEs, to enable UK companies to innovate and for the UK to benefit from the advantages AI can confer. (Conclusion, Paragraph 186)

- 39. To help ensure that regulation can keep pace with technological developments, it is preferable for legislation to take the form of broadly framed principles, with more detailed guidance provided by codes of practice. (Conclusion, Paragraph 187)
- 40. The AI Bill we recommend should provide for a regulatory regime that would classify different risk levels, and mandate proportionately more demanding obligations for higher-risk AI systems and AI models. (Recommendation, Paragraph 188)
- 41. AI supply chains are complex. Actors which have the greatest capacity to prevent harm are often upstream developers. However at present it is too easy for the risks to be passed down the chain to those who are not in a position to identify and address the source of those risks located upstream. (Conclusion, Paragraph 193)
- 42. Many AI systems developed for a specific purpose can be deployed or adapted for use for different purposes and contexts from those intended by the AI developer. Those deploying the system may lack understanding of the risks arising from the way in which the AI system has been designed and developed, and how its deployment could adversely impact human rights. (Conclusion, Paragraph 194)
- 43. The AI Bill we recommend should impose proportionate obligations on actors at all stages of the AI lifecycle and supply chain to ensure that risks are identified, evaluated and effectively addressed at the appropriate stage. (Recommendation, Paragraph 195)
- 44. Some uses of AI should be prohibited, because they are incompatible with human rights. (Conclusion, Paragraph 199)
- 45. The AI Bill we recommend should prohibit the use of AI in ways that are inconsistent with respect for human rights, such as subliminal techniques, emotional inference or inappropriate use of profiling or biometric data. Because of the very serious human rights risks, this would mean that the development and provision of very powerful AI systems (such as Artificial General Intelligence and Artificial Superintelligence) which risk causing widespread and very serious harm, including the capacity to evade effective human control, would be prohibited. (Recommendation, Paragraph 200)
- 46. The detail of such prohibitions should be established following public consultation. (Recommendation, Paragraph 201)

47. The AI Bill we recommend should include a requirement that AI systems posing a high risk of causing harm to human rights should need prior approval before they can be provided or deployed. (Recommendation, Paragraph 204)
48. The AI Bill we recommend should place requirements on actors at all stages of the AI supply chain to undertake due diligence obligations necessary to prevent unacceptable human rights risks when designing, developing or deploying high risk AI systems. (Recommendation, Paragraph 207)
49. We recommend that the due diligence requirements contained in the AI Bill should be differentiated, based on the role of the relevant actor along the AI supply chain, and the nature and seriousness of the risk posed by the relevant AI model or system. To identify the nature and scope of those due diligence obligations, the EU AI Act may provide a useful point of departure, but the provisions should be tailored to the requirements of the UK. (Recommendation, Paragraph 208)
50. The AI Bill we recommend should impose mandatory transparency requirements on all actors across the AI lifecycle. The deployment of AI systems which may have significant impacts on individuals, groups and communities should include an obligation to state when AI systems are being used and provide a full and comprehensible explanation of what it is being used for. There should also be an obligation to provide information about the source of data used by the AI system. (Recommendation, Paragraph 213)
51. Although data protection law provides for specific safeguards when automated decision-making is used, there is a need for greater clarity in the law to ensure that the safeguards operate effectively in practice, and that reliance on human involvement is not misplaced. (Conclusion, Paragraph 217)
52. The regulation-making powers in the UK GDPR should be used to ensure that there are robust protections against misuse of automated decision-making. The regulations should say that the mere presence of a ‘human in the loop’ is not enough to constitute meaningful human involvement or intervention. It should be necessary to show that the person is sufficiently informed and independent to be able to reach an objective view that has not been improperly influenced by the automated decision. The regulations should also say that safeguards include a requirement to provide information relating to the data subject’s individual circumstances rather than the general operation of the system, and that data subjects must be given enough information to allow them to mount an effective challenge against the decision. (Recommendation, Paragraph 218)

- 53.** We agree with our witnesses that the present regulatory framework is fragmented and difficult to navigate. Urgent action is needed to close gaps in protection arising from the current framework. (Conclusion, Paragraph 224)
- 54.** A single AI oversight body would be able to act as a single point of contact for raising concerns about use of AI, carry out oversight and monitoring of AI harms and risks, and maintain an adverse incident repository to help identify where more systemic issues may lie. We are confident this can be done without creating duplication. (Conclusion, Paragraph 225)
- 55.** The government should establish an independent AI oversight body on a statutory footing, either as a new body or by expanding the powers of an existing body. The government should consult carefully with existing regulators to avoid unnecessary overlap and promote regulatory consistency and coherence. (Recommendation, Paragraph 226)
- 56.** The new oversight body should:
- a.** have the power to facilitate co-ordination with other national regulators in relation to monitoring, investigation and enforcement, knowledge sharing, and capacity building
 - b.** make regular reports on its activities to Parliament, identifying opportunities to fill emerging gaps in protection and highlighting trends and sectoral developments in best practice
 - c.** have a statutory duty to maintain a public repository of AI incidents and proactively identify patterns of reported harm and any emerging systemic threats, and to consider cases for further investigation
 - d.** have the power to publish mandatory codes of practice for the responsible use of AI in the private and public sector, and to sanction developers, deployers and other actors who do not fulfil their responsibilities under these codes of practice
 - e.** have the powers necessary to implement and enforce the provisions of the AI Bill we recommend, including:
 - i.** to carry out testing and evaluation of AI systems, including powers to implement a prior testing, auditing and evaluation regime for high-risk uses and applications of AI
 - ii.** to prohibit AI models or systems from being launched or deployed, and to order their withdrawal from the market if unacceptable risks to human rights are identified

- iii. to issue codes of practice which would provide greater specificity concerning the due diligence obligations which must be complied with (for example, requiring high risk AI providers to establish and maintain a risk management system, to ensure that practical and acceptable complaints mechanisms are in place and so forth), and to receive reports if certain types of risk are identified
- iv. to conduct investigations of allegations concerning breaches of the statutory prohibitions and mandatory codes of conduct, including allegations that AI systems have caused harm to human rights
- v. to sanction actors who fail to comply with mandatory codes of practice
- vi. to order that remedies be provided as appropriate in individual cases without requiring the affected individuals to incur prohibitive costs
- vii. to prescribe transparency requirements through its codes of practice, identifying the information actors at each stage of the AI lifecycle must disclose and to whom, including other actors along the AI lifecycle and public sector AI deployers
- viii. where actors in the AI supply chain are located outside the UK, to be able to assess the risks of deploying their systems in the UK, having regard to how those actors are regulated in their home state. If the new AI oversight body considers that deployment in the UK would give rise to unacceptable risks which could not be addressed within the UK, the oversight body should be able to limit or prohibit such deployment. (Recommendation, Paragraph 227)

57. The government should place AISI on a statutory basis. AISI should have the power to review new and revised powerful AI models. It should also have a duty to assess, publish findings pre-release, and issue preliminary ‘warning’ statements associated with release of new AI foundation models. To avoid duplication, the government should consider combining the existing functions of the AISI with those of the new regulator. (Recommendation, Paragraph 231)

58. Developers of powerful AI models must be required to submit new models (and new versions of existing models) to AISI for review, evaluation, and testing. They must also provide technical specifications, including model properties, training data and process, intended use, security and safety testing and risk control measures. (Recommendation, Paragraph 232)

59. The government must give the AISI a statutory duty to inform the authorities, including the new AI oversight body, of potential risks posed by AI models and systems. This should include a duty to advise the oversight body of the need to prevent AI systems from being launched or deployed, or to order them to be withdrawn from the market, if AISI identifies significant risks to human rights. (Recommendation, Paragraph 233)

Formal minutes

Wednesday 9 September 2026

Members present:

Lord Alton of Liverpool, in the Chair

Tom Gordon MP

Baroness Hamwee

Afzal Khan MP

Lord Murray of Blidworth

Lord Rook

Alex Sobel

Euan Stainbank

Peter Swallow MP

Sir Desmond Swayne MP

Human Rights and the Regulation of Artificial Intelligence

Draft Report (*Human Rights and the Regulation of AI*), proposed by the Chair, brought up and read.

Ordered, That the draft Report be read a second time, paragraph by paragraph.

Paragraphs 1 to 233 read and agreed to.

Summary agreed to.

Resolved, That the Report be the Fourth Report of the Committee to both Houses.

Ordered, That the Chair make the Report to the House of Lords and Alex Sobel make the Report to the House of Commons.

Ordered, That embargoed copies of the Report be made available (Standing Order No. 134).

Adjournment

Adjourned until Wednesday 14 October at 2.00pm

Witnesses

The following witnesses gave evidence. Transcripts can be viewed on the [inquiry publications page](#) of the Committee's website.

Wednesday 2 July 2025

Professor David Leslie, Professor of Ethics, Technology, and Society at Digital Environment Research Institute (DERI), Queen Mary University of London, Director of Ethics and Responsible Innovation Research, The Alan Turing Institute; **Ravi Naik**, Solicitor, AWO (a data rights agency), Hon. Professor at UCL Laws, UCL (University College London)

[Q1-16](#)

Wednesday 29 October 2025

Dr Sarah Kiden, Research Fellow, Responsible AI UK/University of Southampton; **Michael Birtwistle**, Associate Director (Law & policy), Ada Lovelace Foundation; **Professor Kevin Fong OBE MBBS MRCP FRCA FFICM**, Broadcaster, The Artificial Human

[Q17-27](#)

tèmitópé lasade-anderson, Executive Director, Glitch; **Javier Ruiz Diaz**, Technology and Human Rights Lead, Amnesty International UK; **Ms. Alex Pirlot de Corbion**, Director of Strategy, Privacy International; **Silkie Carlo**, Director, Big Brother Watch

[Q28-33](#)

Wednesday 17 December 2025

Ellen Lefley, Senior Lawyer, JUSTICE; **Louise Hooper**, Barrister, Garden Court Chambers; **Dr Janis Wong**, Policy Adviser, Data and Technology Law, Law Society

[Q34-44](#)

Professor Ethan Mollick, Co-Director, Generative AI Labs at Wharton, Rowan Fellow, Wharton University of Pennsylvania; **Professor Roman Yampolskiy**, Associate Professor, University of Louisville

[Q45-52](#)

Wednesday 14 January 2026

James Drayson, CEO, Locai Labs; **Kay Firth-Butterfield**, CEO, Good Tech Advisory; **Dr. Iulian Serban**, Senior Director of Research & Development, LawZero

[Q53–63](#)

Wednesday 21 January 2026

Alexandria Walden, Global Head of Human Rights, Google

[Q64–74](#)

Wednesday 4 February 2026

Andrew Breeze, Director for Online Safety Technology Policy, Ofcom; **William Malcolm**, Executive Director of Regulatory Risk & Innovation, ICO; **Dr Mary-Ann Stephenson**, Chair, EHRC

[Q75–88](#)

Wednesday 25 February 2026

Rob Sherman, VP and Deputy Chief Privacy Officer, Policy, Meta; **Ginny Badanes**, General Manager - Tech for Society, Microsoft

[Q89–98](#)

Kanishka Narayan MP, Parliamentary Under Secretary of State for Science, Innovation and Technology, Minister for AI, Department of Science, Innovation and Technology

[Q99–111](#)

Published written evidence

The following written evidence was received and can be viewed on the [inquiry publications page](#) of the Committee's website.

RAI numbers are generated by the evidence processing system and so may not be complete.

1	5Rights Foundation	RAI0074
2	ACT The App Association	RAI0043
3	Ada Lovelace Institute	RAI0066
4	Alegre, Dr Susie (Barrister, Associate - Garden Court Chambers)	RAI0019
5	Amnesty International UK	RAI0051
6	Atabey, Dr Ayca; Dr Kim Sylwander; and Livingstone, Professor Sonia	RAI0058
7	Atkinson, Dr Joe (Associate Professor of Employment Law, University of Southampton)	RAI0080
8	Basu, Professor Subhajit (Professor of Law and Technology, School of Law, University of Leeds)	RAI0032
9	Benslama-Dabdoub, Dr Malak (Lecturer in Law, Royal Holloway University of London)	RAI0010
10	Big Brother Watch	RAI0036
11	Bois, Professor Francois du (Professor of Law, University of Leicester)	RAI0072
12	British Copyright Council	RAI0027
13	British Standards Institution	RAI0050
14	Catholic Bishops' Conference of England and Wales	RAI0040
15	Chartered Institute of Library and Information Professionals (CILIP)	RAI0053
16	Christoph, Joel (PhD Researcher, European University Institute)	RAI0059
17	Connected by Data	RAI0013
18	Dashti, Abdullah	RAI0025

19	Department of Science Innovation and Technology	RAI0077
20	Dualarity Ltd	RAI0071
21	Duhs, Eleonor (Barrister, Partner, Head of Data & Privacy, Bates Wells LLP)	RAI0035
22	Equality and Human Rights Commission	RAI0089
23	Equality and Human Rights Commission	RAI0075
24	Equity	RAI0055
25	European Center for Not-For-Profit Law (ECNL)	RAI0008
26	Fair Vote UK	RAI0002
27	Fox, Dr S. J. (Academic (Law); Associate Director, Institute for Digital Culture, University of Leicester)	RAI0009
28	Ferguson, Dr Erin (Lecturer in Public Law, University of Aberdeen); Moorhouse, Liam (PhD Researcher, University of Aberdeen); Brown, Professor Abbe (Professor in Intellectual Property Law , University of Aberdeen); Moran, Dr Clare Frances (Lecturer in Law, University of Aberdeen); Couzigou, Professor Irene (Professor of Public International Law, University of Aberdeen); Nurullah Benli, Muhammed (PhD Researcher, University of Aberdeen); Allan, Dr Kevin (Lecturer in Psychology, University of Aberdeen); Walters, Andrew (PhD Researcher, University of Aberdeen); Starkey, Dr Andrew (Reader in Engineering, University of Aberdeen); and Le, Dr Xuan Tung (Researcher, University of Aberdeen)	RAI0048
29	Gentile, Dr Giulia (Lecturer, Essex Law School); Gillett, Dr Matthew (Senior Lecturer , Essex Law School); Gilbert, Professor Geoff (Professor, Essex Law School); and Guinchard, Professor Audrey (Professor, Essex Law School)	RAI0047
30	Giupponi, Professor Belen Olmos (Professor of Law, Middlesex University London)	RAI0062
31	Glenlead Centre	RAI0069
32	Glitch	RAI0015
33	Global Network Initiative	RAI0034
34	Good Tech Advisory	RAI0086
35	Google	RAI0084
36	Gulati, Dr Rishi (Associate Professor in International Law, University of East Anglia)	RAI0016
37	Handley Gill Limited	RAI0070

38	Henry Jackson Society	RAI0011
39	ISAR Global	RAI0007
40	Information Commissioner's Office	RAI0003
41	Information Commissioner's Office (ICO)	RAI0085
42	Jorgensen, Dr. Mackenzie (Postdoctoral Research Fellow, Northumbria University); Paul, Dr. Angela (Postdoctoral Research Fellow, Northumbria University); and Oswald, Professor Marion (Professor in Law, Northumbria University)	RAI0020
43	Joy	RAI0001
44	Justice	RAI0082
45	Justice	RAI0073
46	Karakus, Mr Burak Batuhan (Head of Legal Affairs, Human Rights Solidarity)	RAI0026
47	Krupiy, Dr Tetyana (lecturer, Newcastle University)	RAI0012
48	Laux, Dr Johann (Departmental Research Lecturer in AI, Government and Policy , Oxford Internet Institute, University of Oxford)	RAI0006
49	Liberty	RAI0079
50	Loideain, Dr Nora Ni (Senior Lecturer in Law and Director, Information Law & Policy Centre, Institute of Advanced Legal Studies, University of London)	RAI0060
51	McGregor, Professor Lorna (Professor of International Human Rights Law, University of Essex)	RAI0056
52	Minderoo Centre for Technology and Democracy, University of Cambridge	RAI0067
53	Miotti, Andrea (Founder and CEO, ControlAI); and Adler, Steven (Machine Learning Researcher and Practitioner, Formerly OpenAI (co-led the development of dangerous capability evaluations))	RAI0031
54	Music Publishers Association	RAI0039
55	Northern Ireland Human Rights Commission	RAI0038
56	Online Safety Act Network	RAI0029
57	Open Rights Group; StopWatch; and Runnymede Trust	RAI0018
58	Prison Reform Trust	RAI0041
59	Privacy International	RAI0081
60	Public Law Project	RAI0045

61	Responsible AI UK (RAI UK)	RAI0049
62	Romero, Gina (Special Rapporteur for the rights to freedom of assembly and of association, United Nations - OHCHR); Fussey, Prof. Pete (Head of Department. Department of Sociology, Social Policy and Criminology, University of Southampton); and Daragh, Dr. Murray (Senior Lecturer, School of Law, Queen Mary University of London)	RAI0023
63	Romero-Moreno, Dr Felipe	RAI0087
64	Sanchez-Graells, Professor Albert (Professor of Economic Law, University of Bristol Law School)	RAI0005
65	Sowambur, Leena (Founder, Marcomms By Leena)	RAI0063
66	Sutherland, Dr Jessica (Research Fellow, University of Warwick); and Hyams, Professor Keith (Professor of Political Theory and Ethics, University of Warwick)	RAI0022
67	The Howard League for Penal Reform	RAI0068
68	The Independent Society of Musicians	RAI0004
69	The Law Society of England and Wales	RAI0021
70	Trilateral Research	RAI0046
71	Tzanou, Dr Maria (Senior Lecturer in Law , University of Sheffield)	RAI0033
72	UNISON	RAI0076
73	Webb, Professor Philippa (Professor of Public International Law and Director, Oxford Institute of Technology and Justice); and Moynihan, Ms Harriet (Head of Accountability in International Law, Oxford Institute of Technology and Justice)	RAI0042
74	Wong, Dr Mark (Senior Lecturer and Head of Social and Urban Policy, University of Glasgow)	RAI0014

List of Reports from the Committee during the current Parliament

All publications from the Committee are available on the [publications page](#) of the Committee's website.

Session 2026–27

Number	Title	Reference
3rd	Legislative Scrutiny: Northern Ireland Troubles Bill	HC 162
2nd	Human Rights of Children in the Social Care System in England	HC 161
1st	Draft Human Rights Act 1998 (Remedial) Order: Judicial Immunity: Second Report	HC 325

Session 2024–26

Number	Title	Reference
9th	Draft Northern Ireland Troubles (Legacy and Reconciliation) Act 2023 (Remedial) Order 2025: Second Report	HC 1438
8th	Proposal for a Remedial Order to amend the Human Rights Act 1998: Judicial Immunity	HC 1406
7th	Transnational repression in the UK	HC 681
6th	Forced Labour in UK Supply Chains	HC 633
5th	Legislative Scrutiny: Crime and Policing Bill	HC 830
4th	Legislative Scrutiny: Border Security, Asylum and Immigration Bill	HC 789
3rd	Legislative Scrutiny: Mental Health Bill	HC 601
2nd	Accountability for Daesh crimes	HC 612
1st	Proposal for a Draft Northern Ireland Troubles (Legacy and Reconciliation) Act 2023 (Remedial) Order 2024	HC 569

Number	Title	Reference
8th Special	Draft Northern Ireland Troubles (Legacy and Reconciliation) Act 2023 (Remedial) Order 2025: Second Report: Government Response	HC 1697
7th Special	Crime and Policing Bill: Government Response	HC 1439
6th Special	Transnational repression in the UK: Government Response	HC 1405
5th Special	Forced Labour in UK Supply Chains: Government Response	HC 1404
4th Special	Legislative Scrutiny: Border Security, Asylum and Immigration Bill: Government Response	HC 1315
3rd Special	Legislative scrutiny: Mental Health Bill: Government Response	HC 1217
2nd Special	Accountability For Daesh Crimes: Government Response to the Committee's Second Report of Session 2024 - 2025	HC 1211
1st Special	Human rights and the proposal for a "Hillsborough Law": Government Response to the Committees Third Report of Session 2023 - 2024	HC 739