

A Human-Centric Approach Through Rights: The *Missing Dialogues* between Rights, Risks, and Remedies under the EU's Artificial Intelligence

Elif Biber & Herwig C.H. Hofmann

Abstract: A human-centric approach to the regulation of AI must be built on individual rights, which in turn need to be enforceable through adequate remedies. This chapter argues that the EU's Artificial Intelligence Act (AI Act), despite its explicit commitment to a human-centric and rights-based approach to artificial intelligence (AI), fails to establish a meaningful relation between fundamental rights and the risk-based regulatory framework on which it is built. While the AI Act repeatedly invokes the protection of fundamental rights as a core objective, this article demonstrates that such references remain largely abstract and insufficiently translated into enforceable legal mechanisms capable of addressing the concrete societal impacts of AI systems and thereby hampering a truly human-centric orientation of the regulatory regime it built. This is at the heart of enforcement difficulties, weakening the regulatory power of the AI Act. The chapter develops this argument in three steps. First, it situates the AI Act within the broader evolution of the EU's digital *acquis* and its long-standing normative commitment to human dignity, democracy, and the rule of law. It shows that, notwithstanding this tradition, the AI Act adopts a predominantly technocratic governance model that tends to conflate the protection of fundamental rights with technical risk management and product-safety compliance. Secondly, the analysis examines the AI Act's mandatory requirements, focusing on human oversight and the Fundamental Rights Impact Assessment (FRIA). It argues that both instruments, while formally rights-oriented, are structurally ill-equipped to capture the contextual, systemic, and lived impacts of AI, and risk devolving into checklist-based or symbolic safeguards. Thirdly, the chapter analyses the Act's review and conformity-assessment procedures, highlighting how extensive reliance on self-assessment and private actors undermines accountability, weakens access to justice, and obscures the distinctive normative status of fundamental rights. Finally, the chapter advances concrete pathways to bridge this missing dialogue between fundamental rights and the legal architecture intended to secure their effective realisation in practice, underscoring the need to strengthen access to justice. Together, these elements are presented as indispensable to operationalising and giving substantial effect to the EU's human-centric vision of AI governance.

I. Human-Centeredness, Fundamental Rights and AI

European legal frameworks addressing digital technologies place strong emphasis on the centrality of the human factor and fundamental rights in

an era increasingly shaped by automation and artificial intelligence (AI).¹ Although still evolving, the EU's digital legal architecture demonstrates its commitment to European values and fundamental rights, which serve as the guiding constitutional principles for balancing interests through legislation and regulation. By contrast, compared to other leading regulatory models,² the EU explicitly seeks to place individuals and their fundamental rights at the centre of digital governance, in an attempt to reinforce its broader tradition of rights protection and European values. This orientation is not merely a regulatory preference but a normative stance, positioning Europe as an advocate of a rights-based, human-centric vision of digital transformation. The approach to equate the human-centric approach with fundamental rights is well established. The HELG identified it as the attempt to '*ensure that human values are central to the way in which AI systems are developed, deployed, used and monitored, by ensuring respect for fundamental rights...*'³ However, the realisation of the human-centric vision based on rights remains work in progress. Based on the necessity to link a human-centric approach with individual rights, this chapter points at the need for robust review and enforcement mechanisms in order to ensure

-
- 1 The most relevant legal instruments are 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 (General Data Protection Regulation) [2016] OJ L 119/1; Regulation (EU) 2022/2065 of the European Parliament and of the Council of 19 October 2022 on a Single Market for Digital Services (Digital Services Act) [2022] OJ L 277/1; and the 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) [2024] OJ L 1689/1.
 - 2 For a broad, comparative overview, see: Anu Bradford, *Digital Empires: The Global Battle to Regulate Technology* (Oxford University Press 2023).
 - 3 European Commission, 'Ethics Guidelines for Trustworthy AI', High-Level Expert Group on Artificial Intelligence (HLEG-AI), 8 April 2019, available at op.europa.eu/en/publication-detail/-/publication/d3988569-0434-11ea-8clf-01aa75ed71a1 (accessed 13 March 2026). The Guidelines define the human-centric approach as follows '*The human-centric approach to AI strives to ensure that human values are central to the way in which AI systems are developed, deployed, used and monitored, by ensuring respect for fundamental rights, including those set out in the Treaties of the European Union and Charter of Fundamental Rights of the European Union, all of which are united by reference to a common foundation rooted in respect for human dignity, in which the human being enjoy a unique and inalienable moral status. This also entails consideration of the natural environment and of other living beings that are part of the human ecosystem, as well as a sustainable approach enabling the flourishing of future generations to come.*'

realisation of what was articulated in e.g. the Commission's *European Declaration on Digital Rights and Principles for the Digital Decade* (of January 2022)⁴ setting out a comprehensive vision of Europe's digital transition, framed around the principle of 'putting people at the centre.'⁵ It reaffirmed that digital rights and freedoms must be respected and upheld online, with the same degree of protection as they are offline in the analogue world – thereby establishing a clear link between digital governance and the EU's constitutional traditions of human dignity, democracy, and the rule of law.⁶

Nonetheless, while this normative commitment is both ambitious and forward-looking, a substantial, and we might argue, widening gap remains in terms of the concrete implementation and operationalisation of the European human-centric rights-based approach.

Bridging this gap requires not only regulatory innovation and institutional capacity-building, but also procedural innovations that translate abstract values into enforceable legal principles.

This chapter explores the necessary steps to consolidate and concretise the European human-centric, rights-based digital framework, with a focus on how such a model can be rendered both practically effective and globally influential. It first focuses on what we refer to as the dynamics of the *missing dialogue*⁷ between rights, risks and remedies under the EU's Artificial Intelligence Act (AI Act).⁸ It then focuses on the legal pathways to bridge this missing dialogue.

4 European Commission, *European Declaration on Digital Rights and Principles for the Digital Decade* COM(2022) 28 final, 26 January 2022.

5 Elif Biber, 'Art. 2 The Principle of Solidarity and Inclusion' in Rossano Ducato, Alain Strowel and Enguerrand Marique (eds), *The Declaration on European Digital Rights and Principles for the Digital Decade: A Commentary with a Legal Design Perspective* (forthcoming).

6 European Commission, *European Declaration on Digital Rights and Principles for the Digital Decade* COM(2022) 28 final, 26 January 2022, Preamble, para. 3 ('*The digital transformation should not entail the regression of rights. What is illegal offline, is illegal online*').

7 For an earlier discussion on the *missing dialogue* between rights and risks see Elif Biber, *A Rights-Based Inter-Legal Approach to Artificial Intelligence*, (Hart Publishing 2026).

8 See the AI Act in (n 1).

II. A Missing Dialogue between Rights, Risks, and Remedies

The EU's AI Act is primarily grounded in the risk-based regulatory framework, as also underlined in Recital 26 of the AI Act, which claims that proportionate and effective binding rules for AI systems require that 'the type and content of such rules' be tailored to 'the intensity and scope of the risks that AI systems can generate'.⁹ Within the context of the AI Act, the risk-based framework is expressly tied to risks affecting health, safety, and fundamental rights. However, while the Act repeatedly invokes the notion of fundamental rights, its references to the fundamental rights remain largely generic, offering limited guidance on how specific rights are to be addressed within the regulatory scheme and by specific rules.

The engagement with fundamental rights does not appear to run counter to the emphasis on fundamental rights and a human-centric approach that predates the AI Act. For instance, in 2017, the European Council called on the EU to adopt a forward-looking, proactive approach to AI, emphasising the need to act urgently while safeguarding data protection, digital rights, and ethical standards. It further invited the European Commission to develop a distinctively European approach to AI.¹⁰

The Commission subsequently issued the draft and final versions of the Ethics Guidelines for Trustworthy AI in 2018 and 2019, also referring to AI systems as *socio-technical* systems rather than merely as products.¹¹ These Guidelines articulate a European, human-centric vision of AI governance, positioning the protection of human values as a central regulatory objective.¹² They underline the centrality of fundamental rights and human

9 Recital 26 of the AI Act in (n 1).

10 European Council, *European Council meeting (19 October 2017) – Conclusions* EUCO 14/17 (2017) 8 <https://www.consilium.europa.eu/media/21620/19-euco-final-conclusions-en.pdf> accessed 1 September 2025. Responding to this call, the Commission published a Communication in April 2018, which sought to lay the foundations for an ethical and legal framework for AI that reflects the Union's core values and aligns with the EU Charter of Fundamental Rights in European Commission, *Communication on Artificial Intelligence for Europe* COM(2018) 237 final, 25 April 2018; Charter of Fundamental Rights of the EU [2012] OJ C 326/391.

11 *Ethics Guidelines for Trustworthy AI* (n 3) 2 ('... providing guidance on how such principles can be operationalised in socio-technical systems').

12 European Commission, *Ethics Guidelines for Trustworthy AI*, High-Level Expert Group on Artificial Intelligence, 8 April 2019 <https://op.europa.eu/en/publication-detail/-/publication/d3988569-0434-11ea-8c1f-01aa75ed71a1> accessed 1 September 2025.

dignity throughout the development, deployment, use, and monitoring of AI systems, while also recognising the relevance of environmental sustainability and intergenerational equity. However, the AI Act left its practical implications largely underdeveloped with terms such as ‘trustworthy AI’, ‘ethical AI’ or ‘human-centric AI’ lacking contour and ‘the grand formula’ of human-centric, secure, trustworthy and ethical AI all ending up as more or less loose references to fundamental rights.¹³ With Micklitz we thus find that the AI Act’s proclaimed human-centrism suffers from important gaps. Not only is the meaning of the term human-centric underdetermined, as the European digital policy legislation does not explicitly confirm that ‘human control over AI systems is a necessary requirement for protection of human dignity.’ Also, there is a missing link between the human-centric formula and the real-world, engaging with the factual side, the concrete impact of AI systems on society.¹⁴

We argue in this chapter that one reason for the gap between the EU’s human-centric vision and the concrete societal impact of AI systems is the Act’s heavy emphasis on highly complex technical procedures and obligations. This focus leaves only a narrowly defined space for the effective protection of fundamental rights. It limits the extent to which the human-centric approach translates into meaningful safeguards in practice, obscures the dynamics underlying the societal impact of AI, and, hence, results in a lack of *dialogue* between rights and risks. In this chapter, we discuss how a more meaningful dialogue between rights and risks can only take shape once effective legal pathways are established – pathways capable of making the social impact of AI both visible and open to debate. This allows us to link rights, risks and remedies.

1. The Missing Dialogue – an Illustration

According to its Recitals, the protection of fundamental rights is identified as a primary objective of the regulation¹⁵ referring to ‘fundamental rights’

13 Hans-W Micklitz, *The Role of Standards in Future EU Digital Policy Legislation: A Consumer Perspective*, commissioned by ANEC and BEUC, July 2023.

14 *ibid.* In this regard, Micklitz underlines that ‘the call for humans to have the last word requires defining red lines which cannot be crossed in technical standardisation’.

15 See article 1 of the AI Act in (n 1), which states that ‘the purpose of this Regulation is to improve the functioning of the internal market and promote the uptake of human-centric and trustworthy artificial intelligence (AI), while ensuring a high level

more than ninety times, signalling a deliberate and sustained emphasis on rights-based concerns. This emphasis indicates that, even within a predominantly product-safety-oriented regulatory framework, the AI Act treats the protection of fundamental rights not merely as a symbolic reference, but as a normative commitment and a core regulatory aim. The limitations become particularly apparent in the Act's tendency to merge rights protection with safety standards.¹⁶ An absence of meaningful dialogue does not stem from a failure to reference fundamental rights, but rather from a failure to accord them proper engagement within the regulatory framework. However, it is critical to recognise, that as the Ethics Guidelines emphasise, that AI systems are *socio-technical* systems rather than merely as products.¹⁷

By addressing matters predominantly in terms of narrowly defined technical risks, the AI Act underplays its problematic nature with respect to the core requirements of fundamental rights balancing and limitation as expressed under article 52(1) of the EU Charter of Fundamental Rights (CFR). The problem with this risk-based approach lies, in part, in the fact that the risks to a particular AI system's compliance with fundamental rights can seldom be fully identified at the time of its conception or even at deployment.¹⁸ Technological risk evolves continuously as use cases diversify, tools are combined in novel ways, information flows diversify, and data can be recombined with increasing sophistication. The rising proliferation of non-binding standards and guidance from diverse bodies and sources as a remedy has not been able to address the gaps.¹⁹

of protection of health, safety, fundamental rights enshrined in the Charter, including democracy, the rule of law and environmental protection, against the harmful effects of AI systems in the Union and supporting innovation. See also Elif Biber, 'Machines Learning the Rule of Law: EU Proposes the World's First Artificial Intelligence Act' (Verfassungsblog, 17 December 2025), verfassungsblog.de/ai-rol/ (accessed 3 January 2026).

16 *ibid.*, 9. See also theories of fundamental rights in Ernst-Wolfgang Böckenförde, 'Grundrechtstheorie und Grundrechtsinterpretation' (1974) 27 *Neue Juristische Wochenschrift* 1529–1538 (presenting five fundamental rights theories, namely the liberal, the institutional, the democratic-functional, the welfare-state, and the value theory of fundamental rights).

17 *Ethics Guidelines for Trustworthy AI* (n 3) 2 ('... providing guidance on how such principles can be operationalised in socio-technical systems').

18 David Collingridge, *The Social Control of Technology* (St Martin's Press 1980) 19.

19 Lisette Mustert and Cristiana Santos, 'The European Data Protection Board – a (non)consensual and (un)accountable role?' (2025) 59 *Computer Law & Security Review* <https://www.sciencedirect.com/science/article/pii/S2212473X25000896> accessed 11 January 2026.

But serious questions remain about whether the pronounced reliance on technical, highly complex procedures and obligations contributes to a meaningful dialogue between rights and risks. Instead, the risk-based and standard-oriented approach entangles regulatory enforcement, creates legal uncertainty about the nature and extent of obligations, and blurs responsibilities for compliance with requirements set out in the legal framework.

2. The Human Oversight and the Fundamental Rights Impact Assessment Obligations

The Act's tendency to conflate the protection of fundamental rights with technical safety standards also manifests in its mandatory requirements concerning human oversight and fundamental rights impact assessments.

Example One: Human Oversight

An example is the requirement of human oversight, set out in article 14 of the AI Act. This provision requires the implementation of mechanisms for human oversight with the explicit objective of safeguarding individuals' fundamental rights.²⁰ In this context, human oversight is conceived as a means to prevent or mitigate risks to health, safety, or fundamental rights that may arise from the deployment of high-risk AI systems, whether such systems are used in accordance with their intended purpose or under conditions of reasonably foreseeable misuse, particularly where such risks persist despite compliance with other regulatory requirements.²¹ Furthermore, article 14(4)(d) explicitly recognises the significance of individual autonomy by empowering the person tasked with oversight to decide, in specific circumstances, not to deploy a high-risk AI system or to disregard, override, or reverse its output.²² Notwithstanding the questionable technical feasibility of these specific requirements, the provision of this specific technical product-safety-based remedy does not adequately account for the fundamental cognitive differences between human decision-makers and

20 See the requirement of human oversight in Germany in Paul Nemitz and Eike Gräf, 'Artificial Intelligence Must Be Used According to the Law, or Not at All' (Verfassungsblog, 2022) <https://verfassungsblog.de/roa-artificial-intelligence-must-be-used-according-to-the-law/> accessed 15 December 2025.

21 Article 14(2) of the AI Act in (n 1).

22 Article 14/(4)(d) of the AI Act in (n 1).

AI systems.²³ Therefore, we are not only confronted with the questionable remedial power of article 14 AI Act's specific requirements, we are also confronted with a lacking link between the obligations under article 14 and the realisation of the possible limitations to individual rights arising from automated systems using AI.

This observation also arises from a temporary aspect. The faster technological development progresses, and the less foreseeable future technological innovations are, the more 'timeless' the regulatory framework must be to project regulatory power into the uncertain future. The 'half-life' of highly specific technological regulation will be all the shorter, the faster the technological innovation cycle turns. Therefore, to confront the unforeseen and unforeseeable, a return to basics, including regarding the rights to be protected, is preferable to highly prescriptive requirements. The AI Act's article 14 human oversight obligations date back to the early drafting of its regulatory content, nearly a decade ago. At that time, today's possibilities of AI were not understood, as little as we can expect to anticipate the specific risks that a given AI-based technology will pose a decade from now. We do know, however, now and in the future, that fundamental rights must be protected. This amounts to the real futureproofing of law and regulatory content through constitutional steering.

Translating this to human oversight under the article 14 AI Act, this means that, given contemporary technological developments and the increasing complexity of AI models, it is often unrealistic to expect human overseers to meaningfully comprehend, scrutinise, and, if need be, correct the outputs of AI-driven decision-making processes. The systemic nature of such systems renders comprehensive oversight impracticable: human operators are generally unable to monitor the full scope of data collection, mining, and processing activities, nor can they reliably interpret or validate the frequently opaque outputs generated by complex models.²⁴ Meaningful oversight would presuppose an extensive level of transparency and docu-

23 Joseph Weizenbaum, *Computer Power and Human Reason: From Judgment to Calculation* (W H Freeman and Company 1976) 7. See also a debate on this issue in Manuel Alfonseca et al, 'Superintelligence Cannot Be Contained: Lessons from Computability Theory' (2021) 70 *Journal of Artificial Intelligence Research* <https://jair.org/index.php/jair/article/view/12202> accessed 13 September 2025.

24 On the inadequacy of human oversight see Ben Green, 'The Flaws of Policies Requiring Human Oversight of Government Algorithms' (2022) 45 *Computer Law & Security Review* 1–22 (demonstrating the limitations of individuals in performing the required oversight role and advocating for institutional oversight mechanisms). See also Ben Green and Amba Kak, 'The False Comfort of Human Oversight as an

mentation regarding data processing practices that are rarely available in practice today. It would also require human overseers to be able to engage in independent, *de novo* decision-making to assess AI-generated outputs against alternative outcomes and reconstruct the underlying logic and reasoning of automated decision-making processes – an expectation that is often unattainable in real-world settings.

Furthermore, the fundamental cognitive differences between humans and AI systems risk entrenching what the Act and parts of the literature refer to as ‘automation bias’, rather than mitigating it. The AI Act explicitly warns those responsible for oversight to remain ‘aware of the possible tendency of automatically relying or over-relying on the output produced by a high-risk AI system.’²⁵ This warning reflects concern about this phenomenon, namely an ‘undue deference to automated systems by human actors that disregard contradictory information from other sources or do not (thoroughly) search for additional information.’²⁶ However, this characterisation implicitly locates responsibility at the level of individual cognition and judgment, thereby underestimating the structural and organisational conditions in which such deference arises. For instance, in many public and private institutional settings, the bias issue is more plausibly understood as systemic or organisational issues: humans are frequently deprived of the time, resources, technical access, or institutional authority necessary to meaningfully interrogate or reproduce automated decision-making processes. Moreover, organisational incentives may actively discourage deviation from algorithmic outputs, even where doubts exist.

If human oversight is to serve as a genuine safeguard rather than a symbolic one, it must therefore be made *de facto* feasible, not merely formally assigned. While the AI Act correctly acknowledges the risk of uncritical overreliance on automated outputs,²⁷ additional structural and procedural

Antidote to AI Harm’ (Future Tense, 2021) <https://slate.com/technology/2021/06/human-oversight-artificial-intelligence-laws.html> accessed 13 September 2025.

25 Article 14(4)(b) of the AI Act in (n 1).

26 Saar Alon-Barkat and Madalina Busuioc, ‘Human–AI Interactions in Public Sector Decision-Making: “Automation Bias” and “Selective Adherence” to Algorithmic Advice’ (2022) *Journal of Public Administration Research and Theory* 7–8 https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3794660 accessed 13 September 2025.

27 Over-reliance on the outputs of high-risk AI systems can give rise to significant harms. A stark illustration is provided by the Spanish government’s deployment of the VioGén system, an algorithmic tool intended to assess the risk of recidivism in cases of gender-based violence. As documented by AlgorithmWatch, the system generated profoundly unreliable risk assessments. Of the fifteen women killed in

dimensions must be addressed if human oversight is to operate as an effective counterweight rather than a legitimising façade. Otherwise, it is not only foreseeable but likely that the excessive digitalisation, for example, of government services and the automation of decision-making will supplant the uniquely human capacity to exercise discretion informed by contextual sensitivity and measure.²⁸

Example Two: Fundamental Rights Impact Assessments

Another example of a *missing dialogue* between rights and risks can be observed in article 27 of the AI Act, which imposes a mandatory obligation only on deployers – an obligation that may also extend to private entities that provide public services – to conduct a Fundamental Rights Impact Assessment (FRIA).²⁹ Impact assessment tools were initially developed as accountability tools for legislative and administrative procedures.³⁰ Injecting such an obligation into the AI Act's context of specific deployment of AI systems profoundly changes the nature of the tool. FRIA under

incidents of domestic violence in Spain in 2014 – each of whom had previously reported their aggressor – fourteen had been classified by VioGén as facing either a 'low' or 'non-specific' level of risk. See AlgorithmWatch, *Automating Society Report 2020* (2020) 7, <https://automatingsociety.algorithmwatch.org/wp-content/uploads/2020/12/Automating-Society-Report-2020.pdf> (accessed 13 September 2025). For a more detailed account of the system's design and operation, see José Luis González Álvarez and others, 'Integral Monitoring System in Cases of Gender Violence: VioGén System' (2018) 4(1) *Behavior & Law Journal* 29–40.

28 Sofia Ranchordás, 'Empathy in the Digital Administrative State' (2022) 71 *Duke Law Journal* 1341–1389 (examining the role of empathy as a key value in administrative law).

29 It is important to note that FRIA should not be understood in isolation, but rather as part of a broader ecosystem of impact assessment instruments, including Human Rights Impact Assessment (HRIA) and Data Protection Impact Assessment (DPIA) or Privacy Impact Assessment (PIA). The key features of the FRIA under the AI Act are (i) its *ex ante* nature, (ii) its mandatory character, (iii) its broad scope, which encompasses all relevant fundamental rights and not merely data protection issues, and (iv) the fact that it is complemented by the obligation on AI providers to assess fundamental rights impacts within the conformity assessment under article 9 of the AI Act. See also Recital 96 of the AI Act ('Services important for individuals that are of public nature may also be provided by private entities').

30 See e.g. Herwig C.H. Hofmann, *Automated Decision-Making (ADM) in EU Public Law*, in: Herwig C.H. Hofmann, Felix Pflücke (eds.), *Governance of Automated Decision Making and EU law* (Oxford University Press 2024) 1–32; Colin Kirkpatrick and David Parker (eds.), *Regulatory Impact Assessment – Towards Better Regulation?* (Edward Elgar 2007); Claire A. Dunlop, Claudio M. Radaelli (eds.), *Handbook of Regulatory Impact Assessment* (Edward Elgar 2016).

the AI Act constitutes an *ex ante*, rights-based assessment process aimed at identifying, analysing, and addressing the potential impacts of an AI system on fundamental rights before and during its deployment, and it extends beyond data protection considerations to assess how a specific AI use case, situated within its concrete social, institutional, and operational context, may restrict or adversely affect a broad range of fundamental rights and freedoms. In this sense, FRIA has been incorporated into the AI Act to serve as a preventive and iterative governance mechanism, requiring expert-based evaluation, contextual analysis, and reasoned documentation to inform design choices, deployment conditions, and risk management measures throughout the entire lifecycle of AI systems.³¹ Notwithstanding these strengths, the FRIA framework under the AI Act relies predominantly on internal assessments conducted by AI deployers. No obligations to ensure the meaningful participation of rights-holders, affected communities, or civil society actors exist,³² despite this participatory and exploratory screening being a central tenet of impact assessments as they emerged in their original public-law setting. This significantly limits the value of an impact assessment, since it relies on vastly reduced information.

This structural limitation risks reducing FRIA to a predominantly technocratic exercise, insufficiently informed by lived factual experiences, social power asymmetries, and contextual forms of vulnerability. Therefore, indirect, cumulative, or systemic impacts on fundamental rights may remain under-identified or inadequately addressed. As a result, although FRIA purports to strengthen *ex ante* accountability, the necessarily one-sided nature of this internal exercise constrains its ability to fully capture the societal dimensions of AI-related risks and ultimately undermines its potential as a meaningful, human-centric instrument of AI governance.

The concept of conformity assessment procedures, as referred to in the AI Act, applies to high-risk AI systems.³³ The AI Act combines internal assessment through impact assessment activities with external private monitoring. For example, as also examined below, Art 3(20) of the AI Act identifies ‘conformity assessment’ as demonstrating whether the requirements set

31 See a detailed analysis of this process in Alessandro Mantelero, ‘The Fundamental Rights Impact Assessment (FRIA) in the AI Act: Roots, legal obligations and key elements for a model template’ (2024) 54 *Computer Law & Security Review* 1–18.

32 Margot E Kaminski and Gianclaudio Malgieri, ‘Impacted Stakeholder Participation in AI and Data Governance’ (2025) 27(1) *Yale Journal of Law & Technology*, 265.

33 Articles 3(20), 6(1)(b), and 43 of the AI Act n (1). Compare with articles 19 and 43 of the AI Act Proposal of 21 April 2021.

out in various standards and obligations relating to a high-risk AI system 'have been fulfilled'. That can be controlled (Art. 3 (21) AI Act) by a 'conformity assessment body', i.e. a 'body that performs third-party conformity assessment activities, including testing, certification and inspection' and needs to be confirmed by a conformity assessment declaration under article 47 AI Act. Here, self-regulation and self-assessment are combined with a form of private certification. But the problem arises that the amount of intervening norms – from standards to guidelines to assessment criteria risks to give false sense of compliance and risks producing a regulatory haze with lacking clarity of which norms are to be complied with, what is binding and what is only guiding, and whether compliance with the law and its objectives or predominantly compliance with certification procedures is the objective. As a result, the oversight framework assigns a particularly prominent role to private actors, not only with respect to the FRIA but more generally as to compliance with the safety requirements. This raises, on the one hand, questions regarding the legitimacy of delegating such extensive oversight authority to private entities, and on the other hand, concerns relating to their capacity to conduct sufficiently rigorous and comprehensive assessments of complex AI systems. These limitations are especially salient in relation to the evaluation of compliance with fundamental and human rights, an area in which private actors may lack both the mandate and the expertise to ensure adequate protection.³⁴

3. The Missing Dialogue in Enforcement Procedures: Conformity Assessment Procedures and the 'Rights' under the AI Act

This discussion leads to the question of implementation and finally enforcement. The AI Act establishes a number of significant oversight bodies, including the 'notifying authority', 'market surveillance authority', 'conformity assessment body', 'post-marketing monitoring system', 'AI Office', and 'European Artificial Intelligence Board'. Pursuant to article 3(48), the first two of these are classified as 'national competent authorities'. The question of whether the creation of such a multiplicity of bodies constitutes an opti-

34 It should additionally be noted that these actors may commit categorical errors in identifying the right or rights affected, particularly where the legal characterisation of harms requires normative judgment beyond technical assessment. For a detailed analysis see Elif Biber, *A Rights-Based Inter-Legal Approach to Artificial Intelligence* (Hart Publishing 2026).

mal institutional design, and whether the procedural interactions among them are appropriately structured, falls outside the scope of this article.

However, we do consider that article 43(2) and Annex VI of the AI Act, which establishes the mechanisms for assessing compliance with the mandatory requirements, weakens the overall robustness of the regulatory framework. Under these provisions, the majority of compliance reviews are conducted through internal ‘conformity assessment procedures’ performed by the providers themselves – an approach adapted from EU product safety law.³⁵ In practical terms, this regulatory design entrusts providers with the responsibility of self-assessing whether high-risk AI systems satisfy the essential requirements laid down by the AI Act, with the exception of high-risk AI systems deployed in the field of biometrics, which are subject to a conformity assessment carried out by an independent ‘notified body’.³⁶

A similar concern arises with regard to standardisation bodies, as explicitly acknowledged by the European Commission, which underlines that, ‘given the fundamental rights and data protection implications of the European standards and European standardisation deliverables requested, compliance with Union law on fundamental rights and Union data protection law should also be guaranteed’.³⁷ The use of private self-assessment mechanisms, together with the involvement of private certification bodies (‘notified bodies’), creates conditions that are vulnerable to the abuse of power.

35 The conformity-assessment procedure is defined in article 2(12) of Regulation (EC) No 765/2008 as ‘the process demonstrating whether specified requirements relating to a product, process, service, system, person or body have been fulfilled’. See Regulation (EC) No 765/2008 of the European Parliament and of the Council of 9 July 2008 [2008] OJ L 218/30. See also European Commission, *The ‘Blue Guide’ on the Implementation of EU Product Rules 2016* (2016) [https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:52016XC0726\(02\)&from=EN](https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:52016XC0726(02)&from=EN) accessed 13 September 2025, clarifying the conformity-assessment procedure and the CE marking. The AI Act similarly defines the procedure in article 3(20) as ‘the process of demonstrating whether the requirements set out in Chapter III, Section 2 relating to a high-risk AI system have been fulfilled’. The conformity-assessment procedure is linked to the CE marking, defined in article 3(24) of Regulation (EU) 2024/1689 (n 1) as ‘a marking by which a provider indicates that an AI system is in conformity with the requirements (...)’.

36 Recital 125 of the AI Act in (n 1).

37 European Commission, *Commission Implementing Decision of 22 May 2023 on a standardisation request to the European Committee for Standardisation and the European Committee for Electrotechnical Standardisation in support of Union policy on artificial intelligence* C(2023) 3215 final, recital 15 [https://ec.europa.eu/transparency/documents-register/detail?ref=C\(2023\)3215&lang=en](https://ec.europa.eu/transparency/documents-register/detail?ref=C(2023)3215&lang=en) accessed 17 December 2025.

Similar regulatory approaches have been applied in other policy areas, equally subject to critical review.³⁸ Ultimately, the delegation of significant oversight responsibilities to private companies raises legitimate concerns regarding independence and the risk of conflicts of interest.

A further issue that reinforces pre-existing structural imbalances of power concerns the lack of a clearly defined and accessible avenue of redress for individuals affected by AI-driven systems. While the AI Act explicitly positions itself as a regulatory instrument designed to uphold European values and protect fundamental rights amid rapid technological development, there remains a substantive risk that persons harmed by the deployment of prohibited or high-risk AI systems may, in practice, be deprived of effective remedies.

Notably, the AI Act contains only two provisions that explicitly frame entitlements as ‘rights’: ‘*right to lodge a complaint with a market surveillance authority*’ under article 85 and ‘*right to explanation of individual decision making*’ under article 86. article 85 grants both natural and legal persons the ability to submit a complaint where they have grounds to believe that the Regulation has been infringed, providing that ‘*any natural or legal person having grounds to consider that there has been an infringement of the provisions of this Regulation may submit complaints to the relevant market surveillance authority*’.³⁹ Given that the Act is primarily addressed to AI providers and deployers and is modelled on the conceptual architecture of EU product safety law, the formal recognition of such a right constitutes, at least in principle, a noteworthy step towards accountability. However, the practical effectiveness of this mechanism remains an open question. The highly technical and complex nature of the AI Act may significantly impede individuals’ ability to identify, interpret, and substantiate what constitutes an ‘infringement’ of its provisions.⁴⁰ This informational and

38 See the discussion of these approaches in pharmaceutical regulation and the regulation of medical devices: Sylvia Kierkegaard and Patrick Kierkegaard, ‘Danger to Public Health: Medical Devices, Toxicity, Virus and Fraud’ (2013) 29(1) *Computer Law & Security Review* 14 (arguing that the EU has prioritised the health of the single market over the health of the community). See also Hannah van Kolschooten, ‘EU Regulation of Artificial Intelligence: Challenges for Patients’ Rights’ (2022) 59 *Common Market Law Review* 106 (arguing that the Proposal focuses on companies rather than patients).

39 Article 85 of the AI Act in (n 1).

40 This provision should be examined by comparison to the GDPR. Article 57 GDPR assigns to each supervisory authority the task ‘*to handle complaints lodged by a data subject, or by a body, organisation or association in accordance with article 80, and in-*

epistemic asymmetry risks rendering the right to lodge a complaint largely abstract, particularly for those lacking technical expertise or access to legal and computational resources.

In this context, Recital 58 assumes particular significance, as it explicitly acknowledges the potential of certain AI systems to undermine core procedural guarantees and emphasises the need for accountability and effective redress. Although formulated specifically with respect to law enforcement, the considerations can be generalisable in that it links ‘the exercise of important procedural fundamental rights, such as the right to an effective remedy and to a fair trial...’ to the requirement that AI systems be ‘sufficiently transparent, explainable and documented’ in order ‘to avoid adverse impacts, retain public trust and ensure accountability and effective redress.’⁴¹

In addition, article 86 of the AI Act introduces a right to explanation of individual decision-making for persons affected by decisions taken on the basis of outputs generated by certain high-risk AI systems.⁴² Where such decisions produce legal effects or similarly significant impacts on an individual’s health, safety, or fundamental rights, the deployer is required to provide ‘clear and meaningful information’ regarding the role of the AI system in the decision-making process and the main elements underpinning the decision. However, the practical value of this right has been widely questioned in academic discourse.⁴³ Most notably, the obligation to provide explanations is imposed exclusively on the deployer, while the AI provider – who possesses critical knowledge concerning the system’s design, training

investigate, to the extent appropriate, the subject matter of the complaint and inform the complainant of the progress and the outcome of the investigation within a reasonable period, in particular if further investigation or coordination with another supervisory authority is necessary’. In addition, the GDPR explicitly enshrines both the ‘right to lodge a complaint with a supervisory authority’ and the ‘right to an effective remedy against a supervisory authority’ in articles 77 and 78.

41 Recital 58 of the AI Act (n 1).

42 The right to explanation under the GDPR has been extensively discussed in legal literature. See early works in Sandra Wachter, Brent Mittelstadt and Luciano Floridi, ‘Why a Right to Explanation of Automated Decision-Making Does Not Exist in the General Data Protection Regulation’ (2017) 7(2) *International Data Privacy Law* 76; Gianclaudio Malgieri and Giovanni Comandé, ‘Why a Right to Legibility of Automated Decision-Making Exists in the General Data Protection Regulation’ (2017) 7 *International Data Privacy Law* 243.

43 Margot E Kaminski, Gianclaudio Malgieri and Paul De Hert, ‘The Right to Explanation in the AI Act’ (8 March 2025) U of Colorado Law Legal Studies Research Paper No 25–9 <https://ssrn.com/abstract=5194301> accessed 8 January 2026.

data, and operational logic – is excluded from the explanatory framework.⁴⁴ Generally speaking, the knowledge about the information that was used to generate the decision, both in terms of what was taken into account and what was discarded, as well as the criteria under which such information was evaluated, remains a central element. General explanations about ‘main factors’ will generally be of only limited value in terms of the enforcement of rights.⁴⁵

This situation leaves deployers, particularly in the public sector, unable to offer more than ‘box-ticking’ explanations – often just a brief note that a human reviewed the AI’s output. Such explanations are unlikely to enable affected individuals to meaningfully challenge the decision or to assess whether the AI system contributed to discriminatory or otherwise unlawful outcomes. Furthermore, the right applies only to a narrow category of high-risk systems listed in Annex III and is expressly subordinated to existing rights under Union law, significantly curtailing its scope of application. Taken together, these limitations risk rendering the right to explanation largely symbolic, offering weak procedural protection rather than an effective remedy in practice. Such a situation, in fact, further exacerbates the existing lack of meaningful dialogue between rights-based considerations and risk-based assessments.

III. Bridging the Gaps : How to Concretise the European Human-Centric Approach to Artificial Intelligence

We must therefore bridge the gaps that lead to what we describe here as the missing dialogue regarding rights and their protection. Bridging the gaps is essential in order to fulfil the promise inherent in the human-centredness of human self-determination. What this means arises not just from legal thinking in terms of rights, but also of the preconditions of rights protection. Arendt, for example, reminds us that human rights emerge

44 Melanie Fink and Simona Demkova, ‘The Constitutional Foundation of Explanation Rights in EU Digital Regulation: A Contextual Approach’ (November 2025) https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5923283.

45 Herwig C H Hofmann, ‘Assessing Cyber Delegation in EU Public Law’ in Herwig C H Hofmann and Felix Pflücke (eds), *Governance of Automated Decision Making and EU Law* (Oxford University Press 2024) 33–52.

predominantly ‘through the capacity to act in a shared public realm’,⁴⁶ making them not an inner state but a practice within a web of human relationships.⁴⁷ This understanding has profound implications for digital environments,⁴⁸ especially for their regulation in line with standards of human-centeredness. If the rights depend on appearance, contestability, and – we would argue due to the reality of the relation between many users and generally big providers, collective action – then the structures through which digital power operates inevitably acquire legal and political significance. Where not only the architecture of advanced digital technologies systematically restricts the individual and collective capacity to act, but also the law regulating digital actors, through self-regulation and abstract risk terminology, limits individuals’ capacity to remedy the wrong. In this context, human-centeredness will necessarily remain out of reach.

On the other hand, a meaningful bridging of the gaps between risks, rights and remedies can emerge only once effective legal pathways are established – pathways capable of rendering the social impacts of AI visible and open to contestation and debate. These pathways preserve the accessibility of justice by making public participation and the protection of rights explicit, rather than allowing them to be obscured within a ‘legal black box’ of technical obligations produced by the legal system itself.

Furthermore, given that AI systems are still often characterised as black-box systems – producing outputs without transparent or comprehensible explanations⁴⁹ – the role of law becomes critically important in safeguard-

46 Charles Barbour, ‘Between Politics and Law: Hannah Arendt and the Subject of Rights’, in Marco Goldoni, and Christopher McCorkindale (eds) *Hannah Arendt and the Law* (Hart Publishing House 2012) 307.

47 Hannah Arendt, *The Human Condition* (The University of Chicago Press 1958) 204, available at <https://archive.org/details/dli.ernet.528547/page/n1/mode/2up> accessed 13 March 2026 (‘Power preserves the public realm and the space of appearance, and as such it is also the lifeblood of the human artifice, which, unless it is the scene of action and speech, of the web of human affairs and relationships and the stories engendered by them, lacks its ultimate *raison d’être*. Without being talked about by men and without housing them, the world would not be a human artifice but a heap of unrelated things (...)’)

48 See Elif Biber, ‘A Challenge for a *Primavera Digitale*: The Phantom Influence of Artificial Intelligence Systems’, International Journal of Constitutional Law’s Blog, *ICONnect* (New York University School of Law) 3 February 2026, available at: <https://www.iconnectblog.com/a-challenge-for-a-primavera-digitale-the-phantom-influence-of-artificial-intelligence-systems/> (accessed on 3 February 2026).

49 Frank Pasquale, *The Black Box Society: The Secret Algorithms That Control Money and Information* (Harvard University Press 2015).

ing the rule of law, ensuring accountability, and preventing opacity from undermining legal principles. Although the development and evaluation of AI require technical expertise, its far-reaching social consequences make empowering individual engagement indispensable⁵⁰ – an involvement that substantive and procedural legal frameworks must actively enable rather than marginalise in order to make a rights-based human-centredness possible. Courts play an essential role in developing an ideally broadly defined legal framework to new such challenges. One example for this, which has become pathbreaking in many ways, was the clarification of both the socio-technical dimensions and the legal aspects of personal data processing and the human involvement in decision-making processes⁵¹ in the famous CJEU case *Google Spain v AEPD and Mario Costeja González*⁵² – widely discussed in scholarship on digital constitutionalism as an example of the normative power of judicial interpretation in the digital context.⁵³ The case demonstrates how courts can shape the scope and content of fundamental rights in response to technologically mediated forms of decision-making and control, and adapt obligations to an ongoing technological situation and its interactions with societal needs and requirements.

In order to ensure this across the board, however, collective redress and a heightened role of NGOs in representing individual rights is necessary. Developing the procedural right of effective access to justice is therefore necessary to address the potential of societal harms posed by AI systems towards a real-life tool.⁵⁴ As with harm caused by physical persons, AI-

50 Lucile Ter-Minassian, ‘Democratizing AI Governance: Balancing Expertise and Public Participation’ (16 February 2025) arXiv <https://arxiv.org/abs/2502.08651> accessed 6 January 2025.

51 Elif Biber, ‘Between Humans and Machines: Judicial Interpretation of the Automated Decision-Making Practices in the EU’ in Herwig C H Hofmann and Felix Pflücke (eds), *Governance of Automated Decision-Making and EU Law* (Oxford University Press 2024).

52 Case C-131/12 *Google Spain SL and Google Inc v Agencia Española de Protección de Datos (AEPD) and Mario Costeja González* EU:C:2014:317.

53 Giovanni De Gregorio, ‘The Rise of Digital Constitutionalism in the EU’ (2021) 19 *International Journal of Constitutional Law* 41–70; Giovanni De Gregorio, *Digital Constitutionalism in Europe: Reframing Rights and Powers in the Algorithmic Society* (Cambridge University Press 2022); Elif Biber, ‘Digital Constitutionalism in Europe: Exploring a Constellation of Legalities’ in Giovanni De Gregorio, Oreste Pollicino and Peggy Valcke (eds), *The Oxford Handbook of Digital Constitutionalism* (Oxford University 2025).

54 Natalie Smuha, ‘Beyond the Individual: Governing AI’s Societal Harm’ (2021) 10(3) *Internet Policy Review*.

based data processing or decision-making must be subject to meaningful avenues for individual redress to remedy harm and hold relevant actors accountable.⁵⁵ Fundamental rights that individuals enjoy in the physical world must be equally safeguarded in online and digital environments.⁵⁶ This equivalence between the analogue and the digital world does not exclude taking into account and reflecting in procedural provisions and the distinctive features of advanced technological systems, as well as the structural and systemic forms of harm they may generate. The digital environment differs in fundamental ways from the physical one and will therefore require adapted interpretations of existing rights,⁵⁷ and possibly, where that is not sufficient, also the creation of new rights and freedoms tailored to AI's specific architectural conditions. For example, the first preliminary reference to the Court of Justice of the EU (CJEU) concerning the AI Act, submitted in November 2024, centred on the right to an explanation of individual decision-making, which arises from individual freedom, privacy and data rights but was also explicitly articulated in the AI Act itself.⁵⁸

These developments go hand in hand with the increasing role of procedural rights in the CJEU case law, where, for example, the right to access to justice has received increasing prominence in the jurisprudence.⁵⁹ Given the realities of the provision of services and AI embedded services in software is often undertaken by large companies which creates uneven playing fields, ensuring judicial review of individual rights via collective redress is requirement to not only in theory but also in practice realise, develop and uphold human-centredness through the recognition and enforcement of

55 See a recently published blog post on the issue in Vladimir Apraxine, 'Complementing the AI Act with a New Right to Access to Justice: A Discussion on Participatory Governance and Redress in the AI Act' (23 December 2025) <https://www.law.kuleuven.be/ai-summer-school/blogpost/Blogposts/complementing-the-ai-act-with-a-new-right-to-access-to-justice-a-discussion-on-participatory-governance-and-redress-in-the-ai-act> accessed 6 January 2026.

56 Dafna Dror-Shpoliansky and Yuval Shany, 'It's the End of the (Offline) World as We Know It: From Human Rights to Digital Human Rights – A Proposed Typology' (2021) 32(4) *European Journal of International Law* 1249.

57 See e.g. Diana-Urania Galetta, Herwig C.H. Hofmann, *Evolving AI-based Automation – Continuing Relevance of Good Administration*, (2023) 48 *European Law Review* 617–635.

58 CJEU, Case C-806/24 *Yettel Bulgaria* (pending).

59 See Simona Demková, Herwig C.H. Hofmann, *General Principles of Procedural Justice*, in: Katja S. Ziegler, Päivi J. Neuvonen and Violeta Moreno-Lax (eds.), *Research Handbook on General Principles in EU Law*, (Elgar Publishing, Cheltenham 2022), 209–226.

individual rights.⁶⁰ Lessons for the digital field especially arise from ASG-2, which demonstrated that the Court in Grand Chamber, reaffirming the importance of effective judicial protection under article 47 of the Charter, held that national rules render the exercise of the EU right to full compensation practically impossible or excessively difficult, were in violation of EU law, thus requiring collective redress norms in certain conditions. This was confirmed in the *Apple* Case, which was also a competition law case concerning digital markets. The case, brought by two Dutch consumer organisations against Apple concerned the commissions charged on its iOS-based App Store.⁶¹

Both cases underline the particular relevance of representing non-governmental organisations (NGOs) in representing and defending the public interest and the interests of numerous individual rights holders.⁶² Through the publication of detailed analyses on AI governance and regulation, NGOs have brought visibility to well-documented instances of AI-related misuse and abuse. The relevance of NGO engagement is illustrated by the Reclaim Your Face campaign, a Europe-wide civil society initiative advocating for a prohibition on indiscriminate or arbitrarily targeted applications of biometric technologies.⁶³ Similarly, the importance of civil society action is underscored by the landmark *SyRI* judgment, widely regarded as one of the most significant technology-related court decisions globally, which was initiated through the rigorous efforts of human rights activists and non-governmental organisations.⁶⁴ This recognition reflects an understanding of NGOs as institutional intermediaries capable of enhancing accountability where individual access to justice is limited. In light of their demonstrated capacity to act effectively on behalf of individuals and in the public interest, especially in addressing societal harms arising from AI systems, it is imperative that AI governance frameworks incorporate

60 The role of collective redress is well recognized in the enforcement of competition law and the protection of rights therein (see e.g. recently Case C-253/23 ASG 2 *Ausgleichsgesellschaft für die Sägeindustrie Nordrhein-Westfalen GmbH v Land Nordrhein-Westfalen* EU:C:2025:40) also with respect to digital markets (see e.g. Case C-34/24 *Stichting Right to Consumer Justice and Stichting App Stores Claims v Apple Distribution International Ltd and Apple Inc* EU:C:2025:936).

61 *Apple* (n 60) paras 13–21.

62 Center for AI and Digital Policy, *Artificial Intelligence and Democratic Values Index* (2022) <https://caidp.org/reports/aidv-2021/> accessed 8 January 2026.

63 Reclaim Your Face <https://reclaimyourface.eu/> accessed 3 January 2026.

64 *The Hague District Court*, Case C/09/550982/HA ZA 18–388 (2020) paras 6.1 – 6.118.

adequate procedural mechanisms to enable and support such collective representation.

V. Concluding Remarks

The AI Act, despite its explicit and repeated commitment to fundamental rights and human-centric AI governance, currently requires developing a meaningful dialogue between risks, rights and remedies. While the AI Act represents an ambitious and unprecedented attempt to regulate AI through a binding, risk-based framework, using plenty of rights-language, but in with insufficient legal mechanisms through which those rights are made visible, contestable, and enforceable in practice.

The analysis has shown that this missing dialogue is particularly evident in the AI Act's mandatory requirements, review, and enforcement architecture. Human oversight, while formally recognised as a safeguard for fundamental rights, is framed in ways that underestimate the cognitive, organisational, and structural asymmetries between human decision-makers and complex AI systems. Similarly, the FRIA is a potentially powerful governance tool in the hands of public regulators, but its predominantly internal design within the private FRIA under the AI Act limits its capacity to capture the lived, cumulative, and systemic impacts of AI on individuals and communities.

The enforcement model of the AI Act further reinforces this disconnect. By relying heavily on self-assessment and conformity procedures derived from product-safety law, the Regulation entrusts private actors with substantial responsibility for evaluating compliance with requirements that extend far beyond technical safety and deeply into the domain of individual fundamental rights, not just to health and safety but also to data protection, freedom of expression, non-discrimination and many other more. This approach risks conflating conformity with accountability and obscuring the distinct legal status of fundamental rights as values whose limitation and balancing require justification, independent oversight, and effective avenues of redress. In doing so, the AI Act exacerbates power asymmetries and renders access to justice increasingly difficult for those affected by AI-driven decisions.

Against this background, the chapter has argued that linking rights and risks requires constructing explicit legal pathways that translate abstract human-centric commitments into enforceable legal realities. Strengthening

access to justice in the AI context, recognising the institutional role of civil society organisations, and enabling rights-sensitive judicial interpretation emerge as key elements in bridging the gap between what is and what should be possible. These mechanisms do not seek to replace technical regulation but to complement it by embedding democratic participation, public contestation, and normative reasoning within AI governance.

Ultimately, a genuinely human-centric approach to AI cannot be realised through technical compliance alone. It requires the construction of a regulatory regime that acknowledges AI systems as *socio-technical* phenomena with profound societal consequences and that places fundamental rights at the centre of governance, not merely as objectives, but as operative legal standards. As emphasised by the Fundamental Rights Agency, effective protection of fundamental rights requires establishing robust mechanisms, rather than relying predominantly on self-assessment checks conducted by providers.⁶⁵ Without such a recalibration, the promise of a human-centred AI regulation within the AI Act risks remaining aspirational, leaving the dialogue between rights and risks unresolved and the protection of fundamental rights insufficiently secured in the digital age.

65 EU Agency for Fundamental Rights, *Getting the Future Right: Artificial Intelligence and Fundamental Rights* (2020) https://fra.europa.eu/sites/default/files/fra_uploads/fra-2020-artificial-intelligence_en.pdf accessed 8 January 2026.