



LEGAL PROTECTION AGAINST ALGORITHMIC DISCRIMINATION IN EUROPE

Current frameworks
and remaining gaps

In collaboration with
the Interfederal Centre for Equal Opportunities (Unia), Belgium,
the Non-Discrimination Ombudsman (YVV), Finland,
and the Commission for Citizenship and Gender Equality (CIG), Portugal

Raphaële Xenidis

Funded
by the European Union



EUROPEAN UNION

COUNCIL OF EUROPE



CONSEIL DE L'EUROPE

Implemented
by the Council of Europe

LEGAL PROTECTION AGAINST ALGORITHMIC DISCRIMINATION IN EUROPE

Current frameworks
and remaining gaps

Raphaële Xenidis

The report was produced with the financial support of the European Union and the Council of Europe. Its contents are the sole responsibility of the author. The views expressed herein can in no way be taken to reflect the official opinion of either the European Union or the Council of Europe.

The reproduction of extracts (up to 500 words) is authorised, except for commercial purposes, as long as the integrity of the text is preserved, the excerpt is not used out of context, does not provide incomplete information or does not otherwise mislead the reader as to the nature, scope or content of the text. The source text must always be acknowledged as follows: "© Council of Europe, year of the publication". All other requests concerning the reproduction/translation of all or part of the document should be addressed to the Publications and Visual Identity Division, Council of Europe (F-67075 Strasbourg Cedex or publishing@coe.int).

All other correspondence concerning this document should be addressed to the Hate Speech, Hate Crime and Artificial Intelligence Unit of the Council of Europe's Inclusion and Anti-discrimination Programmes Division, Council of Europe, F-67075 Strasbourg Cedex, France
E-mail: anti-discrimination@coe.int

Cover design: Publications
and Visual Identity Division (DPiV),
Council of Europe
Layout: Jouve, Paris.
Photos: Shutterstock

Contents

ABBREVIATIONS	5
ACKNOWLEDGEMENTS	6
INTRODUCTION	7
1. EXISTING LEGAL FRAMEWORKS	9
1.1. Definitions	9
1.2. Algorithmic discrimination in Europe	10
1.3. Use of AI and ADM systems in public administrations and sectors of interest	11
1.4. Existing legal frameworks	15
1.4.1. Governing AI and ADM systems	15
1.4.2. Anti-discrimination instruments	26
1.4.3. Data protection instruments	32
1.4.4. Case law of the European Court of Human Rights and the Court of Justice of the European Union	34
1.5. Key stakeholders	37
2. LEGAL GAPS	41
2.1. Rationales for the adoption of AI and ADM systems	41
2.2. Lack of transparency	42
2.3. Access to justice issues	44
2.4. Enforcement of existing and future provisions	46
2.4.1. Institutional co-operation	46
2.4.2. Conceptual gaps	46
2.4.3. Restrictions on data collection and processing	47
2.5. “De-risking” AI systems	48
2.6. Technical standards, equality bodies and the question of harmonisation	49
SUMMARY – USE OF AI IN BELGIAN PUBLIC ADMINISTRATION AND POLICY RECOMMENDATIONS ON ALGORITHMIC DISCRIMINATION	53
SUMMARY – USE OF AI IN FINNISH PUBLIC ADMINISTRATION AND POLICY RECOMMENDATIONS ON ALGORITHMIC DISCRIMINATION	57
SUMMARY – USE OF AI IN PORTUGUESE PUBLIC ADMINISTRATION AND POLICY RECOMMENDATIONS ON ALGORITHMIC DISCRIMINATION	61
REFERENCES	65

Abbreviations

AI	Artificial intelligence
ADM	Automated decision making
CSOs	Civil society organisations
CAI	Council of Europe Committee on Artificial Intelligence
CEN	European Committee for Standardization
CENELEC	European Committee for Electrotechnical Standardization
FRIA	Fundamental rights impact assessment
GDPR	General Data Protection Regulation
GPAI	General-purpose AI
HUDERIA	Risk and impact assessment of AI systems from the point of view of human rights, democracy and the rule of law
TEU	Treaty on European Union
TFEU	Treaty on the Functioning of the European Union

Acknowledgements

This report was prepared in the context of the “Upholding equality and non-discrimination by Equality bodies regarding the use of artificial intelligence in public administrations” project, which is funded by the European Union via the Technical Support Instrument and co-funded by the Council of Europe. The project is implemented by the Council of Europe in co-operation with the European Commission, the Interfederal Centre for Equal Opportunities (Unia, Belgium), the Non-Discrimination Ombudsman (Finland) and the Commission for Citizenship and Gender Equality (Portugal).

The report was authored by Raphaële Xenidis (Assistant Professor in European Law, Sciences Po Law School, France). The executive summaries of the national contexts were drawn from work authored by Ine van Zeeland (Council of Europe national consultant) for the Belgian section, Emeline Banzuzi (Council of Europe national consultant) for the Finnish section, and Eduardo dos Santos (Council of Europe national consultant) for the Portuguese section.

The Council of Europe wishes to extend its gratitude to the project’s beneficiary institutions and to the European Commission for their sustained engagement throughout the drafting process, and in particular to: Nele Roekens, Nadine Brauns and Thelma Raes (Interfederal Centre for Equal Opportunities, Unia, Belgium); Tiina Valonen and Ville Rantala (Non-Discrimination Ombudsman, YVV, Finland); Carla Peixe, Ana Martinho Fernandes, Alexandra Andrade and Susana Miguel (Commission for Citizenship and Gender Equality, CIG, Portugal), and Massimiliano Santini (Reform and Investment Task Force, Secretariat-General, European Commission) as well as to Menno Ettema, Sara Haapalainen, Ayça Dibekoğlu, and Delfine Gaillard (Anti-discrimination Department, Council of Europe).

The report also benefited from the contributions of national experts from public and private entities in Belgium, Finland, and Portugal who participated in interviews in late 2024 and early 2025. Their perspectives enriched the national sections of the report.





Introduction

The growing integration of artificial intelligence (AI) in diverse applications across a broad range of sectors presents significant challenges to the protection of fundamental rights. Among these, algorithmic discrimination emerges as a particularly pressing concern. Empirical research has demonstrated that algorithmic bias not only reflects but also exacerbates existing social inequalities on a large scale, particularly in domains where algorithmic decision making is prevalent – such as policing, employment, education and insurance. For an extended period, the intersection of non-discrimination and data protection law served as the primary legal foundation for addressing instances of algorithmic discrimination within the European context.

However, recent legal, institutional and political developments have significantly reshaped the governance landscape of AI in Europe. In 2024, both the Council of Europe and the European Union adopted landmark legal instruments: the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law (Framework Convention on AI), 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 (EU AI Act). These instruments establish comprehensive governance frameworks designed to safeguard fundamental rights – including the right to non-discrimination – in the context of AI and automated decision-making (ADM) systems. The adoption of these frameworks is expected to

generate substantial legal and institutional impacts at the national level, particularly given that, by 2026, member states will be required to align their domestic legislation and administrative structures with the obligations set out in the EU AI Act.

The AI Act underscores the necessity for AI technologies and their regulatory frameworks to be developed in alignment with EU values, as enshrined in Article 2 of the Treaty on European Union (TEU), the fundamental rights and freedoms guaranteed by the Treaties, and the Charter of Fundamental Rights pursuant to Article 6 TEU. It explicitly asserts that AI must be a human-centric technology. Similarly, the Council of Europe Framework Convention on AI affirms its objective to ensure that all activities across the AI lifecycle are fully consistent with human rights, democracy and the rule of law.

This report examines how these emerging legal instruments contribute to strengthening protections against algorithmic discrimination in Europe and assesses the lacunae that continue to affect these legal frameworks.

Concurrently, two new EU directives concerning equality bodies aim to enhance the role and capacity of these institutions across Europe (European Union 2024a, 2024b). Equality bodies, along with national human rights institutions and ombudspersons, are anticipated to play a pivotal role in the governance of AI and ADM systems, particularly in ensuring compliance with fundamental rights norms, including the principle of non-discrimination. Hence, this report aims to trace and synthesise those developments in order to equip relevant players, such as equality bodies, with tools to adapt to a changing legal landscape, and where relevant, to intervene in ongoing evolutions.

The report is structured in two main parts:

- ▶ an assessment of existing legal frameworks;
- ▶ an evaluation of remaining lacunae.

The first section offers definitions (1) and examples of the deployment and use of AI and ADM systems by public administrations and in other sectors of interest at European level (2). After explaining how algorithmic discrimination occurs (3), it maps the current legal, procedural and governance frameworks related to anti-discrimination, equality and AI in Europe (4). The section closes by exploring the current and potential role of relevant stakeholders in relation to algorithmic discrimination (5).

The second section charts legal gaps in relation to these frameworks and their enforcement. The scope of this report is primarily limited to AI and ADM systems used by public administrations in Europe, and secondarily, used in private sectors such as large temporary work agencies.



1. Existing legal frameworks

1.1. Definitions

Throughout this report, the terms “AI” and “ADM systems” are used to capture the operations of learning-based and rule-based systems, which can both perpetuate and amplify discrimination. Article 2 of the Council of Europe Framework Convention on AI defines an AI system as “a machine-based system that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations or decisions that may influence physical or virtual environments”. It acknowledges that “[d]ifferent artificial intelligence systems vary in their levels of autonomy and adaptiveness after deployment” (Council of Europe 2024a). This definition echoes the definition offered in Article 3(1) of the EU AI Act.¹ Within the realm of AI techniques, machine learning in particular has become widely used to identify so-called correlations between data points in order to make

1. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence, Article 3(1) defines an AI system as “a machine-based system that is designed to operate with varying levels of autonomy and that may exhibit adaptiveness after deployment, and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments”.

predictions and recommendations or score, classify and rank items or people.² By contrast, some ADM systems lack the level of adaptiveness and autonomy required to qualify as AI systems but are nonetheless important to examine within this report given their discriminatory potential. For example, rule-based systems can misclassify individuals and automate decision-making processes on this basis, with potentially discriminatory consequences. The term “ADM systems” covers a wide range of decision systems, from fully automated systems that do not involve any human input, as defined in, for example, Article 22 of the EU General Data Protection Regulation (GDPR), to semi-automated or decision-support systems, which include both machine-driven and human input and are used very widely in practice.³ This report focuses both on AI and ADM systems.

1.2. Algorithmic discrimination in Europe

Discrimination in algorithmic decision making arises when automated systems unfairly disadvantage certain individuals or groups. There are several key factors that contribute to algorithmic discrimination.

First, bias in training and evaluation data can lead to discriminatory decisions. Not only can skewed data collection and production practices yield unrepresentative datasets, but when AI and data-driven ADM systems are trained on, or process, historical data that reflect prejudices against, inequality between or exclusion of, for example, certain ethnic, gender or socio-economic groups, they also tend to reproduce and/or amplify these biases. For instance, a selection algorithm trained on historical hiring data could unfairly favour male candidates by reproducing past hiring decisions that were skewed by discriminatory stereotypes against women or minority groups.

Second, the design choices made by developers, including the selection of target features and predictors, the type of model, the fairness metrics and thresholds can reflect discriminatory biases. If developers fail to account for systemic inequalities, AI and ADM systems may inadvertently make decisions that disproportionately harm certain groups. For instance, welfare services in Europe increasingly use ADM systems to predict fraud among recipients of social benefits. It has been shown that the very category of “risk” and the predictors operationalised by these ADM systems in reality reflect stereotypes against individuals living in poverty that are deeply entrenched in the practices, forms and institutions of welfare services (Dubois 2021). This problem reflects the pervasiveness of so-called proxy discrimination, namely discrimination based on the algorithmic processing of seemingly neutral proxies that invisibly encode inequalities.

2. According to Binns, “[m]ost uses of machine learning in the public sector are of the supervised variety. Supervision refers to the fact that the learning algorithm has to be shown what a decision-maker wants to predict or classify, unlike unsupervised methods that are designed to discover latent structure in a dataset”. See Binns R. (2020), “Algorithmic decision-making: a guide for lawyers”, *Judicial Review*, Vol. 25, No. 2, pp. 3-4.

3. For a taxonomy of these different ADM systems, see Palmiotto F. (2024), “When is a decision automated? A taxonomy for a fundamental rights analysis”, *German Law Journal*, Vol. 25, No. 2, pp. 210-36, available at: <https://doi.org/10.1017/glj.2023.112>, accessed 6 November 2025.

Third, the lack of transparency and accountability of AI and ADM systems currently in use makes it difficult to monitor or challenge algorithmic discrimination. Many AI and ADM systems, especially those based on complex machine learning models, operate as “black boxes”, meaning that their decision-making processes are not transparent. This lack of transparency makes it difficult to identify when and how discrimination occurs, preventing effective scrutiny and accountability. For instance, the scandal caused in the Netherlands in 2021 by an algorithm used to predict fraud among recipients of childcare benefits showed how difficult it was for those wrongly accused to understand how and why they had been classified as fraudsters, let alone challenge that decision (Amnesty International 2021).

In addition, inadequate testing and monitoring on the part of providers and deployers of AI and ADM systems, both in the public and private sector, reinforces the likelihood that discrimination will arise without notice. Testing and monitoring processes must be put in place before, during and after deployment. The involvement of relevant civil society organisations (CSOs), equality bodies and human rights institutions would ensure the effectiveness of those procedures.

Finally, the speed and scale of AI and ADM systems risk propagating discrimination at a systemic level. The industry’s promise of efficiency and cost saving has encouraged many organisations, including public administrations, to rationalise decision making by deploying AI and ADM systems. Yet, without investing enough resources in preventing and mitigating algorithmic discrimination, for example by training case workers and investing in representative data, proper audit and testing mechanisms, and accountability processes, AI and ADM systems are bound to be discriminatory. Hence, the industry’s narrative, according to which AI and ADM systems will allow the drastic cutting of costs, needs to be challenged: it may well be that ensuring that AI and ADM applications operate lawfully, that is without causing discrimination, through appropriate safeguards and investments, is indeed costly.

1.3. Use of AI and ADM systems in public administrations and sectors of interest

AI and ADM systems are increasingly used both in the private and public sector in Europe. A review of the literature shows examples of application in numerous fields such as the healthcare sector, finance and banking, transport and human resources. Public administrations are progressively adopting or experimenting with AI and ADM systems in fields such as fiscal matters, transport, social security, migration, and justice and policing. Applications can be developed in-house or purchased from external private companies. Besides benign administrative support tasks, the main uses of AI and ADM systems in these fields include fraud detection and surveillance, which present high risks of discrimination. Other widespread applications are chatbots and virtual assistants, which can also manifest discriminatory features but probably with a less harmful impact. That said, the take-up of AI v. ADM systems seems different: while AI applications developed in-house may currently be at experimental stages, ADM systems seem to be deployed more widely.

Challenges related to mapping the use of AI and ADM systems

The inventory exercise conducted in this section comes with an important caveat: mapping usage of AI and ADM systems by public administrations is a difficult task. Research highlights the lack of clear and consistent information on how public administrations use AI and ADM systems. In some countries, like Finland, public administrations are subjected to information obligations and must publish information regarding the use of ADM systems on their websites, and must inform subjects of such use when making decisions. However, databases that systematically record AI and ADM applications are still underdeveloped.⁴ Thus, relevant stakeholders such as equality bodies, national human rights institutions and ombudspersons face difficulties in exercising their mandate to investigate, monitor, identify and challenge cases of algorithmic discrimination, and support victims. These barriers jeopardise the effective application of the right to non-discrimination. Other strategies deployed to obtain information (e.g. freedom of information requests, parliamentary inquiries) also run up against challenges and constraints and do not always allow sufficient insights into the use of AI and ADM systems by public administrations. In France, for example, the freedom of information request made by CSOs to gain clarity on a risk scoring system used by the social welfare administration managing family benefits (CNAF) resulted in the limited release of the source code of former models (as opposed to the current model) and redacted lists of variables used (La Quadrature du Net 2023). Other challenges include the very lack of transparency on the use of AI and ADM systems, which pre-empts questions and inquiries per se: if public administrations do not report the use of such systems, it becomes difficult to know where to enquire in the first place.

Employment

Reported uses include profiling tools used by public employment agencies to predict chances of employment or unemployment, which are deployed to help case workers allocate support resources to jobseekers. Some of these profiling tools have been shown to be discriminatory. In Austria, for example, a prototype developed by the Austrian employment agency called the AMS algorithm aimed at predicting jobseekers' employment prospects to help case workers decide on resource allocation was shown to generate discrimination against *inter alia* single mothers and jobseekers with a migration background (Allhutter et al. 2020). Similar systems have been deployed since 2018 by the French employment agency, France Travail, to assess risks of fraud, predict dropouts and assess recipients' employability (La Quadrature du Net 2024). Similar risks of discrimination against individuals in socially and economically precarious conditions, but also on grounds of sex, race or ethnicity, and disability, have been flagged. Other concerns include the lack of adequate guidance and training offered to case workers who use AI and ADM tools and the ways in which jobseekers' personal data are used in profiling applications. Public employment agencies may also deploy other kinds of tools such as matching algorithms to recommend job vacancies to jobseekers.

4. This will probably change following the implementation of the AI Act.

Law enforcement

AI and ADM systems are also regularly used by law enforcement authorities. Research shows, however, that face recognition technologies can exhibit discriminatory biases.⁵ Even when AI and ADM systems themselves are not reported to be biased, their deployment is criticised for disproportionately targeting and surveilling minorities and carrying out ethnic profiling.⁶ Often, evidence of such uses remains circumstantial and the lack of transparent information about which systems are used and how prevents systematic investigations into the risks they present to fundamental rights, including non-discrimination. As flagged below (see section 1.4.1), the AI Act grants wide exceptions to the police and law enforcement authorities regarding the use of recognition technologies processing images such as faces and other biometric data.

Social and welfare sectors

Various branches of welfare and social security systems in Europe are currently experimenting with, or rolling out, AI and ADM systems (e.g. Austria, Belgium, Denmark, France, Poland, Portugal, Slovenia, Sweden, the Netherlands). In France, the system used by the national agency responsible for the allocation of family benefits aims to predict risks of fraud and errors among recipients to help caseworkers target controls (La Quadrature du Net 2023). Risks of discrimination have been flagged by CSOs and a case has been brought to the French Conseil d'État pointing out *inter alia* discrimination based on sex, family status, age and disability (see Conseil d'État 2024). In the Netherlands, a similar system led to the resignation of the government in 2021 after it was found to have discriminated against recipients on grounds of race, ethnic origin and citizenship.⁷ In the context of education, the Dutch Executive Agency of Education used a risk profiling algorithm to support fraud detection in relation to the receipt of student grants (Algorithm Audit 2024a). The agency recognised that “[s]tudents with a non-European migration background were assigned a higher risk score by a risk profile and were more often manually selected for a home visit”, thus creating indirect discrimination (Algorithm Audit 2024b). In Belgium, the OASIS system was discontinued in 2023 after researchers exposed risks of discrimination on grounds of poverty. Such systems combine data from different administrations (e.g. on tax, employment, family benefits, pensions), sometimes combined with data collected by private companies (e.g. on energy consumption), to profile users and determine risk scores.

5. The CSO Liberty criticised a face recognition application used by the South Wales Police *inter alia* for discrimination on grounds of sex and/or race because it produced a higher rate of positive matches for female faces and/or for black and minority ethnic faces. See the subsequent decision of the UK Court of Appeal in *R (Bridges) v. Chief Constable of South Wales Police* ([2020] EWCA Civ 1058) highlighting that the South Wales Police “Equality Impact Assessment was obviously inadequate and was based on an error of law (failing to recognise the risk of indirect discrimination)” and that its “subsequent approach to assessing possible indirect discrimination arising from the use of AFR is flawed”.
6. This happens for example when police forces keep pictures of erroneous matches for a disproportionate amount of time in their systems, or when over-surveillance of locations largely inhabited by racialised communities feeds back into predictive tools and reinforces mass surveillance of these communities.
7. The SyRI system was also deemed to have disproportionately interfered with end users’ right to privacy because it processed personal data from various government agencies.

Migration and citizenship

AI and ADM systems are also used by certain public administrations in Europe to support decision making in the field of migration, for example regarding decisions on citizenship, asylum or residence. Reported uses include, for example, language identification and assessment, detecting fraud related to identity documents, case management, interacting with migrants including through chatbots, migration forecasting and border surveillance technologies (European Migration Network and Organisation for Economic Co-operation and Development 2022; European Network of National Human Rights Institutions 2024; McGregor and Molnar 2023). Although these technologies are allegedly used to enhance the efficiency of migration management, they can negatively impact migrants' rights by reinforcing discrimination, raising concerns over privacy and data protection, and deterring migrants seeking protection (European Network of National Human Rights Institutions 2014). According to the UN Special Rapporteur on contemporary forms of racism, racial discrimination, xenophobia and related intolerance, "digital borders" encompass a range of governance infrastructures that "increasingly rel[y] upon machine learning, big data, automated algorithmic decision-making systems, predictive analytics and related digital technologies" (UN Special Rapporteur 2021). Technologies used include "identification ... systems, facial recognition systems, ground sensors, aerial video surveillance drones, biometric databases and even visa and asylum decision-making processes" (ibid.).

Tax authorities

Although generally less information is available regarding the use of AI and ADM systems by tax authorities, several European countries are currently experimenting with such techniques, mainly for the purpose of fraud detection. For example, AI and ADM systems are used to flag suspicious files for further investigation. Authorities may be reluctant to reveal information about these systems for fear that users will be able to "game" the system. Risks of discrimination exist as fraud detection systems may disproportionately affect individuals based on protected grounds, for instance socio-economic status, migration background or geographical location, even when these categories are not explicitly used as risk factors in these systems. The above-mentioned SyRI case from the Netherlands included tax authorities suspecting 26 000 families of welfare benefits fraud. The algorithm disproportionately flagged ethnic minorities and non-Dutch nationals.

Private sector

In the private sector, AI and ADM systems are used by personnel services companies. Applications include screening and profiling job applicants, matching job applications to their profiles, or drafting and translating job vacancies. These applications can present risks of discrimination to differing degrees. Similarly, AI and ADM systems are prominently used in finance and commerce, such as in marketing, customer services (e.g. chatbots), dynamic pricing and transaction monitoring. For example, in Finland, a customer was denied online credit by a financial services company, which used a scoring system based on data such as gender, language, age and place

of residence to assess the risk of loan default (Algorithm Watch 2018). The National Non-Discrimination and Equality Tribunal of Finland (2017) concluded that this amounted to direct discrimination based on multiple protected grounds.

1.4. Existing legal frameworks

This section analyses existing regulations, laws, case law, procedures, policies, institutions and so on, and examines how they address discrimination arising from the use of AI and ADM systems.

1.4.1. Governing AI and ADM systems

Council of Europe

At Council of Europe level, several instruments address the issue of algorithmic discrimination or can be utilised to do so.

First, the Council of Europe Framework Convention on AI, which was adopted in 2024 but has yet to be ratified to enter into force, “aim[s] to ensure that activities within the lifecycle of artificial intelligence systems are fully consistent with human rights, democracy and the rule of law”. In its preamble, the Framework Convention on AI acknowledges that AI can “promote ... gender equality and the empowerment of all women and girls”. At the same time, it expresses states parties’ “concerns about the risks of discrimination in digital contexts, particularly those involving artificial intelligence systems, and their potential effect of creating or aggravating inequalities, including those experienced by women and individuals in vulnerable situations, regarding the enjoyment of their human rights and their full, equal and effective participation in economic, social, cultural and political affairs”.

In its Article 10 on equality and non-discrimination, the Framework Convention on AI mandates states parties to “adopt or maintain measures with a view to ensuring that activities within the lifecycle⁸ of artificial intelligence systems respect equality, including gender equality, and the prohibition of discrimination, as provided under applicable international and domestic law”. Parties to the convention must also “adopt or maintain measures aimed at overcoming inequalities to achieve fair, just and equitable outcomes, in line with its applicable domestic and international human rights obligations, in relation to activities within the lifecycle of artificial intelligence systems”. These obligations demand that member states review and, if necessary, reform their legislation to ensure that non-discrimination law captures algorithmic discrimination. As confirmed by point 77 of the explanatory memorandum, the convention also lays out a positive obligation for states to adopt measures

8. See points 14-15 of the explanatory report of the Framework Convention: “This reference to the lifecycle ensures a comprehensive approach towards addressing AI-related risks and adverse impacts on human rights, democracy and the rule of law by capturing all stages of activities relevant to artificial intelligence systems.” Examples of relevant activities include: “(1) planning and design, (2) data collection and processing, (3) development of artificial intelligence systems, including model building and/or fine-tuning existing models for specific tasks, (4) testing, verification and validation, (5) supply/making the systems available for use, (6) deployment, (7) operation and monitoring, and (8) retirement.”

to overcome structural and historical inequalities in relation to activities within the lifecycle of AI systems. Points 72-73 of the explanatory memorandum confirm that the prohibition on algorithmic discrimination laid out in the convention relies on a human rights-based approach to algorithmic discrimination that integrates global and regional human rights frameworks. Points 75-76 of the explanatory memorandum draw up a list of well-known sources and types of algorithmic bias. Among “the different ways through which bias can intentionally or inadvertently be incorporated into artificial intelligence systems at various stages throughout their lifecycle”, the explanatory memorandum cites:

- ▶ development (“due to the conscious or unconscious stereotypes or biases of developers”), modelling (“potential bias built into the model upon which the systems are built”);
- ▶ data (inaccurate or insufficiently representative data sets at training, aggregation or evaluation stages);
- ▶ deployment (“biases introduced when such systems are implemented in real world settings”);
- ▶ interpretation (“automation or confirmation biases, whereby humans may place unjustified trust in machines and technological artefacts or situations where they select information that supports their own views ... ignoring their own potentially contradictory judgment and validating algorithmic outputs without questioning them”);
- ▶ technical bias (“which occurs from problems in applying machine learning that results in additional biases that are not present in the data used to train the system or make decisions”);
- ▶ social bias (“failures to properly account for historical or current inequalities in society in the activities within the lifecycle of artificial intelligence systems”).

To facilitate the enforcement of the obligations contained in the convention, the Council of Europe has designed a “methodology [to] ensure a uniform approach towards identification, analysis and evaluation of risk and assessment of impact of [AI systems] in relation to the enjoyment of human rights, the functioning of democracy and the observance of rule of law”, referred to as the HUDERIA (Council of Europe 2022a). It consists of a context-based risk analysis, a stakeholder engagement process, a risk and impact assessment, and a mitigation plan, and demands iterative review. Even though the approach differs from the risk-based classification adopted under the AI Act (see below the section on the European Union), the HUDERIA methodology introduces similar elements through a “graduated and differentiated approach to measures for risk and impact identification, assessment, prevention and mitigation that takes into account the severity and probability of the occurrence of the adverse impacts on human rights, democracy and the rule of law as well as relevant contextual factors” (Council of Europe 2024b). For example, the determination of risk levels must take into account the scale of potential harms, and their severity, reversibility and probability. “Zero questions”, that is concerning the relevance and adequacy of AI and ADM systems to perform certain tasks, must also be considered with the potential consequence that decisions not to develop or deploy such systems could be taken if risks of human rights violations outweigh potential benefits. The HUDERIA risk and

impact assessment methodology aims to ensure seamless compatibility “with the existing compliance practices followed by the industry” (Council of Europe 2022a) and can foster prevention of algorithmic discrimination by public and private actors. Even though Article 16 of the convention foresees the establishment of a risk and impact management framework, compliance with the HUDERIA is not mandatory to satisfy the obligations of the convention. Although the HUDERIA is not a legally binding instrument and states that it does not offer interpretive guidance in relation to the convention, it will provide a particularly useful tool for identifying, assessing and mitigating risks of algorithmic discrimination.

In addition, diverse recommendations have been published to demand that states parties address algorithmic discrimination in their legislative and procedural frameworks. For example, the Parliamentary Assembly in its Resolution 2343 (2020) on “Preventing discrimination caused by the use of artificial intelligence” called on member states to:

review their anti-discrimination legislation and amend it as necessary, so as to ensure that it covers all cases where direct or indirect discrimination, including discrimination by association, may be caused by the use of AI, and that complainants have full access to justice; in the latter respect, pay particular attention to guaranteeing the presumption of innocence and ensuring that victims of discrimination do not face a disproportionate burden of proof (Council of Europe, Parliamentary Assembly 2020a).

In addition, the resolution requested that member states “ensure that equality bodies are fully empowered to address issues of equality and non-discrimination that arise due to the use of AI”. It also demanded that governments be requested “to notify the parliament before [AI and ADM systems] technology is deployed” and that “the use of such technologies by the authorities ... be systematically recorded in a public register”. It is noticeable that many of these detailed recommendations are not reflected in the recently adopted Framework Convention on AI. The ongoing work of the Steering Committee on Anti-discrimination, Diversity and Inclusion and Gender Equality Commission on a Committee of Ministers recommendation on equality in AI could address these shortcomings (Council of Europe, forthcoming).

Other instruments published by institutions within the Council of Europe can be used to address algorithmic discrimination. For instance, the General Recommendation No. 1 on the Digital Dimension of Violence against Women adopted by the Group of Experts on Action against Violence against Women and Domestic Violence (GREVIO) specifically tackles discriminatory forms of online violence such as online sexual harassment, online and technology-facilitated stalking and the digital dimensions of psychological violence (Council of Europe 2021a).

More transversal recommendations or instruments focused on related topics are worth mentioning as well:

- ▶ Recommendation 2102 (2017) of the Parliamentary Assembly of the Council of Europe on “Technological convergence, artificial intelligence and human rights” (Council of Europe, Parliamentary Assembly 2017);
- ▶ the Declaration of the Committee of Ministers on Risks to Fundamental Rights stemming from Digital Tracking and other Surveillance Technologies (Council of Europe 2013);

- ▶ the Recommendation of the Commissioner for Human Rights on “Unboxing Artificial Intelligence: 10 steps to protect Human Rights” (Council of Europe, Commissioner for Human Rights 2019);
- ▶ the Declaration of the Committee of Ministers on the manipulative capabilities of algorithmic processes (Council of Europe 2019a);
- ▶ Recommendation CM/Rec(2020)1 of the Committee of Ministers to member States on the human rights impacts of algorithmic systems (Council of Europe 2020);
- ▶ Recommendation CM/Rec(2022)13 of the Committee of Ministers to member States on the impacts of digital technologies on freedom of expression (Council of Europe 2022b);
- ▶ the Guidance Note adopted by the Steering Committee for Media and Information Society, “Content moderation. Best practices towards effective legal and procedural frameworks for self-regulatory and co-regulatory mechanisms of content moderation” (Council of Europe 2021b);
- ▶ the Follow-up Recommendation by the Commissioner for Human Rights on “Human rights by design: future-proofing human rights protection in the era of AI” (Council of Europe, Commissioner for Human Rights 2023).

These various instruments offer important guidance regarding compliance with human rights, and in particular non-discrimination, when using AI and ADM systems. However, their non-binding nature and limited enforceability limits their effectiveness in practice.

European Union

At EU level, algorithmic discrimination is addressed by several provisions of the AI Act (European Union 2024c).⁹ The AI Act is intended to complement existing anti-discrimination and data protection frameworks. It adopts a risk-based approach with a number of prohibited practices and high-risk systems that are subjected to specific requirements. This is based on the recognition that, “[a]side from the many beneficial uses of AI, it can also be misused and provide novel and powerful tools for manipulative, exploitative and social control practices ... Such practices are particularly harmful and abusive and should be prohibited because they contradict Union values ... including the right to non-discrimination” (Recital 28).

Other systems are considered high risk because they “may violate the right to dignity and non-discrimination and the values of equality and justice” (Recital 31). A third category of AI systems, considered low risk, are much more loosely regulated. In other terms, “[t]he extent of the adverse impact caused by the AI system on the fundamental rights protected by the Charter is of particular relevance when classifying an AI system as high risk” (Recital 48). That said, Recital 27 of the Act recalls that the seven principles defined by the AI High-Level Expert Group are applicable to all systems regardless of their risk level: these are “human agency and oversight; technical robustness and safety; privacy and data governance; transparency; diversity, non-discrimination and fairness; societal and environmental well-being and

9. See also explanatory memorandum and the accompanying AI Pact available at: <https://digital-strategy.ec.europa.eu/en/policies/ai-pact>, accessed 10 November 2025.

accountability". In this context, "[d]iversity, non-discrimination and fairness means that AI systems are developed and used in a way that includes diverse actors and promotes equal access, gender equality and cultural diversity, while avoiding discriminatory impacts and unfair biases that are prohibited by Union or national law" (ibid.). That principle is meant to apply to every AI system placed on the EU market no matter its risk categorisation.

Prohibited AI systems

The list of prohibited systems, particularly relevant in the context of discrimination, is included in Article 5 of the AI Act. Article 5(a) prohibits "subliminal ... or purposefully manipulative or deceptive techniques [that have the] objective, or the effect of materially distorting the behaviour of a person or a group of persons by appreciably impairing their ability to make an informed decision [and] causes or is reasonably likely to cause [them] significant harm". Article 5(b) bans AI and ADM systems that "exploit any of the vulnerabilities of a natural person or a specific group of persons due to their age, disability or a specific social or economic situation" to distort their behaviour and cause them harm.

Article 5(c) bans social scoring systems that evaluate or classify persons or groups "based on their social behaviour or known, inferred or predicted personal or personality characteristics" with the purpose of treating them unfavourably in unrelated social contexts or in an unjustified or disproportionate manner. Article 5(d) prohibits the use of an "AI system for making risk assessments of natural persons in order to assess or predict the risk of a natural person committing a criminal offence, based solely on the profiling of a natural person or on assessing their personality traits and characteristics". Article 5(e) bans face recognition systems based on scraping of internet or CCTV footage.

Article 5(f) prohibits emotional recognition in the workplace and education institutions. According to Article 5(g) and (h), "the use of biometric categorisation systems that categorise individually natural persons based on their biometric data to deduce or infer their race, political opinions, trade union membership, religious or philosophical beliefs, sex life or sexual orientation" is banned, as is "the use of 'real-time' remote biometric identification systems in publicly accessible spaces for the purposes of law enforcement" unless strictly necessary.

Those prohibitions address important risks of algorithmic discrimination through a blanket ban. Yet, critics have pointed at significant loopholes in relation to the broad exceptions granted to police and law enforcement authorities, particularly in relation to live facial recognition and biometric surveillance in Article 5(g) and (h) (European Digital Rights 2024). The European Commission has issued a set of guidelines on prohibited AI practices, which clarifies the scope of certain provisions of the AI Act. The guidelines state:

even where an AI system is not prohibited by the AI Act, its use could still be prohibited or unlawful based on other primary or secondary Union law (e.g., because of the failure to respect fundamental rights in a given case, such as the lack of a legal basis for the processing of personal data required under data protection law, discrimination prohibited by Union law, etc.) (European Commission 2025a, paragraph 43).

High-risk AI systems

Annex III of the AI Act lists so-called high-risk AI used in biometrics, critical infrastructures (critical digital infrastructure, road traffic, or in the supply of water, gas, heating or electricity), education and vocational training, employment, workers' management and access to self-employment, access to and enjoyment of essential private services and essential public services and benefits, law enforcement, migration, asylum and border control management, and the administration of justice and democratic processes. These high-risk systems are subjected to the specific legal requirements listed below. Article 6(3) foresees that an AI system used in an area listed in Annex III "shall not be considered to be high-risk where it does not pose a significant risk of harm to the health, safety or fundamental rights of natural persons, including by not materially influencing the outcome of decision making". These include systems that "perform a narrow procedural task" or a purely "preparatory task" and systems that are "intended to improve the result of a previously completed human activity" or "to detect decision-making patterns or deviations from prior decision-making patterns and [are] not meant to replace or influence the previously completed human assessment, without proper human review". This provision has the potential to exclude a possibly large range of AI systems used in high-risk areas from the scope of Article 6 and could be misused to avoid compliance with the requirements laid out by the AI Act for high-risk systems (see section 2.5). Yet, such systems are still subjected to certain obligations: providers that consider an AI system as not high risk must nonetheless draw up documentation of the assessment before that system is placed on the market or put into service and must provide that documentation to national competent authorities upon request. They must also register such an AI system in the EU database foreseen under Article 71 of the AI Act.

General-purpose AI models

General-purpose AI (GPAI) models, defined by Article 3(63) as AI models "that displa[y] significant generality and [are] capable of competently performing a wide range of distinct tasks regardless of the way [models are] placed on the market and that can be integrated into a variety of downstream systems or applications"; are specifically regulated under Article 53 of the AI Act. Providers must draw up and present technical documentation and instructions for use, they must comply with the EU's copyright legislation, and they must publish a summary detailing the data used for training purposes. GPAI models of a certain capability and scale are classified as presenting a "systemic risk" under Article 51 of the AI Act and subjected to additional requirements under Article 55 of the AI Act. Providers of such systems must *inter alia* carry out model evaluations and adversarial testing as well as monitor and report serious incidents. In addition, when GPAI is integrated within an AI system, that system is also subjected to the legal requirements applicable to its risk category (unacceptable, high, limited or minimal).

AI systems with limited risks

By contrast, applications presenting more limited risks are only subjected to transparency requirements. For instance, deployers of chatbots must inform users that

they are interacting with an AI system and deployers of GPAI producing text, images, videos or sound must mark the output as artificially generated as per Article 50 of the AI Act. The remaining AI applications are considered as presenting no or minimal risks and are unregulated.

Timeline for compliance

The AI Act is a regulation, so no transposition is necessary on the part of member states: the provisions apply directly. That said, there are a number of milestones necessary for the Act to enter into force. The prohibition of practices under Article 5 entered into force in February 2025 and the Commission has issued guidance on those prohibitions (European Commission 2025a). In 2025, the EU AI Office published a Code of Practice for GPAI models, including those that present systemic risks (European Commission 2025b), and is expected to develop guidelines for the classification of AI systems as high risk and a template for fundamental rights impact assessments for high-risk AI systems (see below the section on fundamental rights impact assessments). By August 2025, member states should have designated the national competent authorities (at least one notifying authority and at least one market surveillance authority) tasked with overseeing the national implementation of the AI Act. In 2026, the EU AI Office should decide on the technical standards elaborated by the EU standardisation bodies, the European Committee for Standardization (CEN) and the European Committee for Electrotechnical Standardization (CENELEC), regarding risk management systems and bias prevention in high-risk AI systems (see below the section on CEN-CENELEC). Finally, the rules on high-risk systems should enter into force in August 2026, or August 2027 for AI systems used as products, or safety components thereof, that are required to undergo a third-party conformity assessment under the Union harmonisation legislation listed in Annex I of the AI Act. For high-risk AI systems already placed on the market or put into service by August 2026 and subject to “significant changes in their designs” or intended to be used by public authorities, the providers and deployers must take the necessary steps to comply with the requirements and obligations of the AI Act by August 2030 (Article 111(2)). All those steps will shape member states’ capacity to prevent and redress algorithmic discrimination.

Data governance

Specific provisions of the AI Act are directly relevant to combating algorithmic discrimination, notably because they lay out specific requirements for high-risk systems. Article 10 on “data and data governance” lays out quality criteria for the training, validation and testing of the high-risk systems listed in Annex III. In particular, Article 10(2), paragraphs (f) and (g) state that providers of AI and ADM systems shall conduct an “examination in view of possible biases that are likely to affect the health and safety of persons, have a negative impact on fundamental rights or lead to discrimination prohibited under Union law, especially where data outputs influence inputs for future operations” and take “appropriate measures to detect, prevent and mitigate possible biases identified according to point (f)”. For these purposes, the sensitive categories of personal data, which the GDPR and the

AI Act generally prohibit using, can be processed based on the exception provided in Article 10(5):

To the extent that it is strictly necessary for the purpose of ensuring bias detection and correction in relation to ... high-risk AI systems [and under certain conditions¹⁰], the providers of such systems may exceptionally process special categories of personal data, subject to appropriate safeguards for the fundamental rights and freedoms of natural persons.

This exception can be useful to identify algorithmic discrimination, yet it does not extend to non-high risk systems.

Risk management system

Article 9 of the AI Act requires providers of high-risk AI systems to put in place a risk management system. It sets out that “[a] risk management system shall be established, implemented, documented and maintained in relation to high-risk AI systems” and that risk management must be understood as “a continuous iterative process planned and run throughout the entire lifecycle”. In particular, providers are asked to identify and analyse “the known and the reasonably foreseeable risks that the high-risk AI system can pose to health, safety or fundamental rights when the high-risk AI system is used in accordance with its intended purpose”. They must estimate and evaluate “the risks that may emerge when the high-risk AI system is used in accordance with its intended purpose, and under conditions of reasonably foreseeable misuse”, as well as “other risks possibly arising, based on the analysis of data gathered from the post-market monitoring system referred to in Article 72”. Finally, they have to adopt “appropriate and targeted risk management measures designed to address the risks identified”. These obligations will be implemented by way of technical standards developed by CEN-CENELEC. Hence, the risk management standard that will be issued by CEN-CENELEC will be a critical device for managing risks to fundamental rights, including the prevention of, and protection against, algorithmic discrimination. Even though not compulsory, compliance with standards will trigger a presumption of conformity of high-risk systems as per Article 40 of the Act and the industry will therefore be incentivised to follow those standards. Nevertheless, it is important to mention that this presumption of compliance is exclusively limited to the requirements of the AI Act for high-risk AI systems, and does not extend to EU fundamental rights law or anti-discrimination law.

10. Conditions apply: “(a) the bias detection and correction cannot be effectively fulfilled by processing other data, including synthetic or anonymised data; (b) the special categories of personal data are subject to technical limitations on the re-use of the personal data, and state-of-the-art security and privacy-preserving measures, including pseudonymisation; (c) the special categories of personal data are subject to measures to ensure that the personal data processed are secured, protected, subject to suitable safeguards, including strict controls and documentation of the access, to avoid misuse and ensure that only authorised persons have access to those personal data with appropriate confidentiality obligations; (d) the special categories of personal data are not to be transmitted, transferred or otherwise accessed by other parties; (e) the special categories of personal data are deleted once the bias has been corrected or the personal data has reached the end of its retention period, whichever comes first; (f) the records of processing activities pursuant to Regulations (EU) 2016/679 and (EU) 2018/1725 and Directive (EU) 2016/680 include the reasons why the processing of special categories of personal data was strictly necessary to detect and correct biases, and why that objective could not be achieved by processing other data.”

Other requirements

These provisions are complemented by several supporting obligations for providers of high-risk AI systems such as Article 11 on technical documentation, Article 12 on record-keeping, Article 13 on transparency, Article 14 on human oversight, Article 15 on accuracy, robustness and cybersecurity, and Article 17 on quality management systems. For instance, Article 15(4) requires providers of high-risk systems “that continue to learn after being placed on the market or put into service ... to eliminate or reduce as far as possible the risk of possibly biased outputs influencing input for future operations (feedback loops), and ... to ensure that any such feedback loops are duly addressed with appropriate mitigation measures”. These obligations can make the enforcement of existing anti-discrimination provisions easier.

EU database

Article 49 paragraphs 1 and 2 mandate providers to register AI systems into the EU database referred to in Article 71 of the Act. In addition, Article 49(3) requires “deployers that are public authorities, Union institutions, bodies, offices or agencies or persons acting on their behalf [to] register themselves, select the system and register its use”. Where AI systems are used in the areas of law enforcement, migration, asylum and border control management, this registration must be made into a non-public section of the EU database that only the European Commission and national authorities can access. According to Article 71, the rest of the database must be “accessible and publicly available in a user-friendly manner”. This public registry will facilitate investigations into the discriminatory impacts of AI and ADM systems, including those deployed by public administrations, as well as potential victims’ task of establishing prima facie evidence of discrimination. Yet, although a good starting point, publicity through an EU-wide database will not suffice to challenge discriminatory usages of AI and ADM systems. In addition, CSOs have expressed concerns over the lack of public access to information pertaining to law enforcement, migration, asylum and border control management contained in the public registry (Access Now 2024).

Fundamental rights impact assessment

Article 27 foresees that certain deployers – mainly deployers that are bodies governed by public law, or are private entities providing public services, but also deployers of AI systems intended to be used to evaluate the creditworthiness of natural persons or establish their credit score and AI systems intended to be used for risk assessment and pricing in relation to natural persons in the case of life and health insurance – have to conduct a fundamental rights impact assessment (FRIA) before using high-risk AI systems. These consist of a description of the deployer’s processes; the period of time within which and the frequency with which the AI system will be used; the categories of natural persons and groups likely to be affected by use of the AI system in the specific context; the specific risks of harm likely to impact those groups; human oversight measures taken; and how those risks are addressed, including through internal governance and complaint mechanisms. The FRIA will be submitted to market surveillance authorities and a summary must be registered

in the EU database (Annex IV, section C, point 4). Access to FRIAs will facilitate the challenging of discriminatory AI and ADM systems for individuals and equality bodies. In addition, as per Recital 96:

deployers of high-risk AI system, in particular when AI systems are used in the public sector, could involve relevant stakeholders, including the representatives of groups of persons likely to be affected by the AI system, independent experts, and civil society organisations in conducting such impact assessments and designing measures to be taken in the case of materialisation of the risks.

FRIAs are complemented by Article 60, which establishes a framework in which providers (with or without deployers) can test high-risk AI systems in real-world conditions. While serious incidents must be reported to the national market authority (Articles 60(7) and 73) and immediately addressed by mitigation measures, the extent to which relevant stakeholders like equality bodies can access testing results and mitigation measures is unclear and may depend on national context. On the one hand, Article 71(4) provides that “information [relating to testing] registered in accordance with Article 60 shall be accessible only to market surveillance authorities and the Commission, unless the prospective provider or provider has given consent for also making the information accessible to the public”. In fact, it excludes the testing information registered under Article 60(4)(c) from the accessibility and publicity requirements. On the other hand, Article 77(1) states:

National public authorities or bodies which supervise or enforce the respect of obligations under Union law protecting fundamental rights, including the right to non-discrimination, in relation to the use of high-risk AI systems referred to in Annex III *shall have the power to request and access any documentation created or maintained under this Regulation* in accessible language and format when access to that documentation is necessary for effectively fulfilling their mandates within the limits of their jurisdiction. (emphasis added).

Authorities protecting fundamental rights (Article 77 bodies)

In fact, Article 77 AI Act on “Powers of authorities protecting fundamental rights” is particularly important because it lays out the enforcement framework. As explained above, Article 77(1) grants fundamental rights supervision authorities designated by member states the power to request and access any documentation created or maintained under the AI Act when necessary. In so doing, it provides a range of authorities protecting fundamental rights with rights to access documentation drawn up under the AI Act, such as risk management plans, impact assessments or any other documentation that is necessary for effectively fulfilling their mandate. The European Commission has issued a broad interpretation of the list of relevant authorities. It is not limited to bodies usually understood as human rights structures such as equality bodies, national human rights institutions and ombudspersons, but also includes data protection authorities, consumer protection authorities, child protection authorities, labour law authorities, media supervisory authorities and authorities in charge of ensuring electoral integrity. Authorities granted access rights under Article 77 must either supervise or enforce relevant Union legislation (or both). The Commission deems that Article 77 does not intend to cover national courts or other judicial authorities, acting in their judicial capacity, which already

have access rights under Article 47 of the EU Charter of Fundamental Rights. National public authorities or bodies referred to in Article 77(1) should be announced publicly and notified to the European Commission by member states. Given the new powers entrusted to equality bodies (see below the section on role of equality bodies), these institutions could be called on to play a major role with regard to enforcing the AI Act in relation to algorithmic discrimination.

In addition, Article 77(3) provides that:

Where the documentation referred to in paragraph 1 is insufficient to ascertain whether an infringement of obligations under Union law protecting fundamental rights has occurred, the public authority or body referred to in paragraph 1 may make a reasoned request to the market surveillance authority, to organise testing of the high-risk AI system through technical means. The market surveillance authority shall organise the testing with the close involvement of the requesting public authority or body within a reasonable time following the request.

Besides, Article 73(7) AI Act provides that designated national authorities protecting fundamental rights must be informed by a market surveillance authority that has received a notification related to a serious incident, including discrimination, in relation to a high-risk AI system. As per Article 79(2), those authorities must be informed by a market surveillance authority, which must fully co-operate with them, where the latter identify a risk to fundamental rights, including discrimination, in relation to an AI system presenting a risk. The relevant operators must also co-operate with them, as necessary. According to Article 82(1), national authorities protecting fundamental rights must be consulted by a market surveillance authority that finds that a compliant high-risk AI system presents a risk to fundamental rights, such as discrimination, where the affected fundamental right is relevant for the mandate of such an authority or body.

AI liability

In 2022, the European Commission proposed a specific AI liability directive that aimed to address important barriers to access to justice in the context of AI. It proposed to harmonise liability rules across the EU for non-contractual fault-based civil law claims for damages related to AI harms. However, the proposed directive was withdrawn in 2025, making it all the more urgent to address AI liability issues at national and European level, including through other means. The draft AI liability directive recognised that “the specific characteristics of AI ... pose a problem for existing liability rules” and proposed to “eas[e] the burden of proof in a very targeted and proportionate manner through the use of disclosure and rebuttable presumptions” to facilitate access to justice in cases of AI-induced harms, including algorithmic discrimination (European Commission 2022a). In particular, Article 3 of the withdrawn proposal stated that a court or tribunal could order providers or deployers of high-risk AI systems that are suspected of having caused damage to disclose evidence when necessary in judicial proceedings. Not complying with a request to disclose or preserve evidence would have triggered a rebuttable presumption of non-compliance on the part of the defendant. Explanatory memorandum for the Article 4 of the withdrawn proposal recognised that “[i]t can be challenging for claimants to establish a causal link between non-compliance and the output produced by an AI system that gave

rise to a given damage”, such as discrimination (ibid). Hence, non-compliance with a duty of care pursuant to the AI Act or other EU rules was to be understood as a fault and would have triggered a rebuttable presumption of a causal link between the harm and the fault. For example, if the provider of a high-risk AI system had not adequately put in place a risk management system under the AI Act, a court could have presumed that there was a causal link between this breach of the duty of care and the harm produced by that system, such as discrimination. These rules could have facilitated the establishment of a *prima facie* case of algorithmic discrimination in court where the defendant had not complied with its obligations under the AI Act. However, three conditions would have had to be met: the claimant would have had to show a breach of the duty of care, it would have had to be reasonably likely that such a fault had influenced the harmful output of the AI system, and the claimant would have had to show that the output produced by the AI system gave rise to the damage. For individual victims of discrimination, establishing these elements would have already posed significant hurdles. Hence, reflections on adjusting existing rules on presumptions of discrimination and the shift of the burden of proof must take place in the context of AI and ADM systems.

Other EU legislation and guidelines

Other pieces of legislation can be used to tackle algorithmic discrimination at EU level. For instance, Article 34 of the Digital Services Act on risk assessment by providers of very large online platforms and of very large online search engines includes risks of discrimination (European Union 2022). So does Article 35 on the mitigation of those risks.¹¹ Additionally, the European Commission has recently launched a public consultation on the Digital Fairness Act, which will aim to strengthen protection and digital fairness for consumers.

The EU has also adopted important policies on the governance of AI and ADM systems that can be used by equality bodies, for example, in making policy recommendations to decision makers or raising awareness. For example, the European Declaration on Digital Rights and Principles for the Digital Decade adopted in 2022 offers a digital bill of rights (European Commission 2022b). It expresses a commitment to “ensuring that algorithmic systems are based on suitable datasets to avoid unlawful discrimination and enable human supervision of outcomes affecting people”. The 2020 White Paper of the European Commission (2020) is also concerned with algorithmic discrimination.

1.4.2. Anti-discrimination instruments

While AI-specific rules are important new tools for equality bodies and other organisations combating discrimination, activating anti-discrimination law and data

11. On how equality bodies can address AI-driven online discrimination, including hate speech, see for example: Equinet (2018), “Extending the agenda. Equality bodies addressing hate speech”, available at: <https://equineteurope.org/extending-the-agenda-equality-bodies-addressing-hate-speech/>, accessed 7 November 2025; Facing Facts Network/CEJI (2022), “Current activities and gaps in hate speech responses: a mapping report for the Facing Facts Network”, available at: www.facingfacts.eu/hate-speech-report/, accessed 7 November 2025.

protection law together with these new provisions will be key to tackling algorithmic discrimination. AI-specific legal frameworks complement, but by no means replace, equality law. The AI Act's requirements for high-risk AI systems, in particular, offer a set of technical safeguards, which support – but cannot ensure – compliance with anti-discrimination law. In this context, equality bodies will play an important role in applying existing equality law frameworks to algorithmic discrimination. For instance, they can offer informed advice on the concrete application of rules pertaining to presumptions of discrimination and the shift of the burden of proof in cases of information asymmetries. They can also play an important role in relation to legal qualifications of algorithmic harms in terms of direct, indirect and intersectional discrimination, or advise on questions such as how to apply exceptions and objective justifications in this context.

Council of Europe

At Council of Europe level, several instruments ban discrimination and are applicable to algorithmic discrimination.

Most importantly, Article 14 of the European Convention on Human Rights laying out the prohibition of discrimination states that “[t]he enjoyment of the rights and freedoms set forth in [the] Convention shall be secured without discrimination on any ground such as sex, race, colour, language, religion, political or other opinion, national or social origin, association with a national minority, property, birth or other status”. This non-exhaustive list of prohibited forms of discrimination can be used to tackle the discriminatory impacts of AI and ADM systems whenever one of the fundamental rights listed in the Convention has been breached.

This is complemented by Article 1(1) of Protocol 12 to the Convention, which lays out an independent prohibition against discrimination: “The enjoyment of any right set forth by law shall be secured without discrimination on any ground such as sex, race, colour, language, religion, political or other opinion, national or social origin, association with a national minority, property, birth or other status”. In addition, Article 1(2) states that “No one shall be discriminated against by any public authority on any ground such as those mentioned in paragraph 1”. This prohibition against discrimination applies – regardless of whether any other fundamental right is breached – to public administrations that deploy discriminatory AI and ADM systems in countries that have ratified Protocol 12. This anti-discrimination framework is complemented, in relation to work, by Article 20 of the European Social Charter on the right to equal opportunities and equal treatment in matters of employment and occupation without discrimination on the grounds of sex.

Other relevant instruments include the Convention on Preventing and Combating Violence against Women and Domestic Violence (the Istanbul Convention) that, in Article 1(1)(b), aims to “contribute to the elimination of all forms of discrimination against women and promote substantive equality between women and men, including by empowering women” and that construes violence against women as “a violation of human rights and a form of discrimination against women”. In turn, Article 4 of the Framework Convention for the Protection of National Minorities states that “any discrimination based on belonging to a national minority shall be prohibited” and

this includes “acts of discrimination, hostility or violence as a result of [a person’s] ethnic, cultural, linguistic or religious identity” (Article 6). In parallel, Article 7(2) of the European Charter for Regional or Minority Languages protects persons against discrimination based on “the use of a regional or minority language”.

Beyond these foundational instruments, other important documents can be drawn on in combating algorithmic discrimination. For example, the Council of Europe Gender Equality Strategy 2024-2029 recognises the “specific impact of AI on gender equality and women’s rights” and urges “member States [to] address algorithmic gender-based and intersectional discrimination through human rights-based and multifaceted gender equality and non-discrimination strategies [and] implement[ing] newly developed standards in the area of artificial intelligence and gender equality” (Council of Europe 2024c). This creates a mandate for member states to review their legislation to address shortcomings in relation to algorithmic discrimination.

Committee of Ministers Recommendation CM/Rec(2019)1 to member States on preventing and combating sexism recognises that “the internet has provided a new dimension for the expression and transmission of sexism, especially of sexist hate speech, to a large audience, even though the roots of sexism do not lie in technology but in persistent gender inequalities” (Council of Europe 2019b). It states that “artificial intelligence poses specific challenges in relation to gender equality and gender stereotypes” and that “[t]he use of algorithms can transmit and strengthen existing gender stereotypes and therefore may contribute to the perpetuation of sexism”. It acknowledges the role of AI and ADM systems in “deepen[ing] the scrutiny to which women’s bodies, speech and activism are subjected”, especially online. The Recommendation calls on member states to “[i]ntegrate a gender equality perspective in all policies, programmes and research in relation to artificial intelligence to avoid the potential risks of technology perpetuating sexism and gender stereotypes and examine how artificial intelligence could help to close gender gaps and eliminate sexism”. It recommends “increas[ing] the participation of women and girls in the information and technology area” and demands that the “design of data-driven instruments and algorithms ... factor in gender-based dynamics”. In addition, “[t]ransparency around these issues should be improved and awareness raised about the potential gender bias in big data; solutions to improve accountability should be offered”. All in all, this Recommendation represents an important legal instrument to combat algorithmic discrimination.

Committee of Ministers Recommendation CM/Rec(2022)17 to member States on protecting the rights of migrant, refugee and asylum-seeking women and girls contains a section on AI, ADM and data protection. The Recommendation provides that “any design, development and application of artificial intelligence and automated decision-making systems by the public or private sectors or by service providers and contractors should be non-discriminatory, consistent with privacy principles, transparent and have clear governance mechanisms” in the context of border and migration management (Council of Europe 2022c). It encourages states to ensure that human rights impact assessments are conducted before the introduction of AI and ADM systems in the field of migration. The recommendation also demands that CSOs be involved in discussions on the development and deployment of new technologies affecting migrant, refugee and asylum-seeking women and girls (ibid.).

The Guidelines of the Committee of Ministers of the Council of Europe on upholding equality and protecting against discrimination and hate during the Covid-19 pandemic and similar crises in the future contain several provisions addressing the discriminatory potential of digitalisation, AI and contact tracing technologies on “vulnerable groups”, though the focus is on digital exclusion and lack of access (Council of Europe 2021c). Recommendation CM/Rec(2022)16 on combating hate speech and Recommendation CM/Rec(2024)4 on combating hate crime call on member states to “protect human rights and fundamental freedoms in the digital environment”, underline the role and responsibilities of internet intermediaries in disseminating hate speech, highlight the importance of digital evidence and data in investigating hate crimes, and acknowledge the role of “extremist movements operating primarily through digital channels and online communications” in hate crimes (Council of Europe 2022d). Yet, these recommendations do not fully address the role of algorithmic recommender systems on social media in fostering extreme polarisation online and thus in fuelling discriminatory stereotypes and prejudices among users.

The Council of Europe Convention on the Protection of Children against Sexual Exploitation and Sexual Abuse (Lanzarote Convention) recognises “the increased use by both children and perpetrators of information and communication technologies” and the role of digital technologies in facilitating child abuse (Council of Europe 2007). Committee of Ministers Recommendation CM/Rec(2010)5 of the on measures to combat discrimination on grounds of sexual orientation or gender identity, adopted in 2010, can be used transversally to address algorithmic discrimination, but does not contain any provision specifically addressing the discriminatory impact of AI and ADM systems, though mounting evidence points to the effects of those technologies in fuelling prejudices and exclusion both online and offline (Council of Europe 2010). The same is true of existing instruments prohibiting discrimination against Roma and travellers, for example Recommendation CM/Rec(2024)1 of the Committee of Ministers to member States on equality of Roma and Traveller women and girls (Council of Europe 2024d).

The General Policy Recommendations of the European Commission against Racism and Intolerance also provide helpful legal resources to address algorithmic discrimination, in particular through the lens of the field-specific application of AI and ADM systems:

- ▶ Recommendation No. 6 on Combating the dissemination of racist, xenophobic and antisemitic material via the Internet can help tackle racist content, hate speech, cyber-harassment and extremism fuelled by algorithmic recommender systems (European Commission against Racism and Intolerance 2000);
- ▶ Recommendation No. 7 on National legislation to combat racism and racial discrimination offers interesting procedural resources to tackle algorithmic discrimination at structural level:

The law should provide that organisations such as associations, trade unions and other legal entities which have, according to the criteria laid down by the national law, a legitimate interest in combating racism and racial discrimination, are entitled to bring civil cases, intervene in administrative cases or make criminal complaints, even if a specific victim is not referred to (European Commission against Racism and Intolerance 2017: 8)

...

[This is] essential for addressing those cases of discrimination where it is difficult to identify such a victim or cases which affect an indeterminate number of victims (ibid.: 22);

- Recommendation No. 8 on Combating racism while fighting terrorism invites member states to:

ensur[e] that no discrimination ensues from legislation and regulations – or their implementation – notably governing the following areas:

- checks carried out by law enforcement officials within the countries and by border control personnel
- administrative and pre-trial detention
- ...
- fair trial, criminal procedure
- protection of personal data
- protection of private and family life
- expulsion, extradition, deportation and the principle of non-refoulement
- issuing of visas
- residence and work permits and family reunification
- acquisition and revocation of citizenship

(European Commission against Racism and Intolerance 2004).

As AI and ADM systems are increasingly deployed in these fields, in particular, to facilitate law enforcement, the recommendation provides a useful framework to address their discriminatory potential even though it itself does not directly address these technologies;

- Recommendation No. 10 on combating racism and racial discrimination in and through school education applies where AI and ADM systems are also deployed, for example in selection processes (European Commission against Racism and Intolerance 2006);
- Recommendation No. 11 on combating racism and racial discrimination in policing, a sector where AI and ADM systems are deployed through, for example, face recognition and face matching technologies but also crime prediction and surveillance systems, is relevant (European Commission against Racism and Intolerance 2007);
- Recommendation No. 14 on combating racism and racial discrimination in employment provides useful guidelines in a field where CV-screening and hiring support tools are increasingly deployed (European Commission against Racism and Intolerance 2012);
- Recommendation No. 15 on Combating Hate Speech demands that member states “ensure that the scope of [hate speech related] offences is defined in a manner that permits their application to keep pace with technological developments” (European Commission against Racism and Intolerance 2015).

This framework is complemented by other recommendations focused on specific forms of discrimination such as racism, xenophobia, anti-gypsyism, antisemitism, anti-Muslim hatred, and discrimination against irregularly present migrants and lesbian, gay, bisexual, transgender and intersex (LGBTI) persons.¹²

European Union

At European Union level, the principles of non-discrimination and gender equality are guaranteed in Article 2 of the TEU on EU values and Article 3 TEU on the internal market, as well as Article 8 and 10 of the Treaty on the Functioning of the European Union (TFEU) on mainstreaming gender equality and non-discrimination, Article 19 TFEU mandating the EU to adopt anti-discrimination legislation and Article 157 TFEU on equal pay. Article 21(1) of the EU Charter on Fundamental Rights provides that “[a]ny discrimination based on any ground such as sex, race, colour, ethnic or social origin, genetic features, language, religion or belief, political or any other opinion, membership of a national minority, property, birth, disability, age or sexual orientation shall be prohibited”. Article 23 of the Charter guarantees that “[e]quality between men and women must be ensured in all areas, including employment, work and pay”.

Several directives lay out a dense web of rules against discrimination, including:

- ▶ Directive 2000/43/EC implementing the principle of equal treatment between persons irrespective of racial or ethnic origin;
- ▶ Directive 2000/78/EC establishing a general framework for equal treatment in employment and occupation;
- ▶ Directive 2004/113/EC implementing the principle of equal treatment between men and women in the access to and supply of goods and services;
- ▶ Directive 2006/54/EC on the implementation of the principle of equal opportunities and equal treatment of men and women in matters of employment and occupation (recast);
- ▶ Directive 2010/41/EC on the application of the principle of equal treatment between men and women engaged in an activity in a self-employed capacity;
- ▶ Directive 79/7/EEC on the progressive implementation of the principle of equal treatment for men and women in matters of social security.

These are applicable to algorithmic discrimination provided discriminatory AI and ADM systems are deployed in their material scope of application. That scope, however, displays gaps, for instance in relation to discrimination on grounds of age, sexual orientation, disability and religion or belief in the purchase of goods and services.

European Union Directives on equality bodies

Two directives have recently reformed the mandate of equality bodies, key players in combating algorithmic discrimination:

12. See all recommendations at: www.coe.int/en/web/european-commission-against-racism-and-intolerance/ecri-standards, accessed 9 November 2025.

- ▶ Directive 2024/1499 on standards for equality bodies in the field of equal treatment between persons irrespective of their racial or ethnic origin, equal treatment in matters of employment and occupation between persons irrespective of their religion or belief, disability, age or sexual orientation, equal treatment between women and men in matters of social security and in the access to and supply of goods and services;
- ▶ Directive 2024/1500 on standards for equality bodies in the field of equal treatment and equal opportunities between women and men in matters of employment and occupation.

By June 2026, EU member states must adapt their national legislation to the provisions of the two EU directives on equality bodies. In relation to automated systems and AI, these directives foresee that “equality bodies should be equipped with appropriate human and technical resources. Those resources should, in particular, enable equality bodies to use automated systems for their work on the one hand and to assess such systems as regards their compliance with non-discrimination rules on the other hand”.

1.4.3. Data protection instruments

Data protection laws are another piece in the legislative puzzle that equality bodies and other organisations combating inequality in Europe need to activate to tackle algorithmic discrimination.

Council of Europe

At Council of Europe level, the legislative framework on data protection can be utilised to tackle algorithmic discrimination. The Modernised Convention for the Protection of Individuals with Regard to the Processing of Personal Data (Convention 108+), and in particular Article 6 restricting the use of special categories of data, call on member states to put in place “safeguards ... guard[ing] against the risks that the processing of sensitive data may present for the interests, rights and fundamental freedoms of the data subject, notably a risk of discrimination” (Council of Europe 2018). Article 11 of the Framework Convention on AI on privacy and personal data protection extends these restrictions to “activities within the lifecycle of artificial intelligence systems”, in particular by asking member states to ensure that “privacy rights of individuals and their personal data are protected, including through applicable domestic and international laws, standards and frameworks” and that “effective guarantees and safeguards have been put in place for individuals, in accordance with applicable domestic and international legal obligations”. Upon entry into force, Convention 108+ will also apply to the areas of national security and defence, meaning that its scope is broader than that of the GDPR.

This framework is complemented by the Guidelines on artificial intelligence and data protection issued in 2019 by the Consultative Committee of the Convention for the protection of individuals with regard to automatic processing of personal data (Convention 108) (Council of Europe 2019c). It highlights that “[i]n all phases of the processing, including data collection, AI developers, manufacturers and service providers should adopt a human rights by-design approach and avoid any potential

biases, including unintentional or hidden, and the risk of discrimination or other adverse impacts on the human rights and fundamental freedoms of data subjects". It also encourages "[c]ooperation ... between data protection supervisory authorities and other bodies having competence related to AI, such as: consumer protection; competition; anti-discrimination; sector regulators and media regulatory authorities".

In addition, the Guidelines on facial recognition adopted in 2021 provide that "[t]he use of facial recognition for the sole purpose of determining a person's skin colour, religious or other beliefs, sex, racial or ethnic origin, age, health or social condition should be prohibited unless appropriate safeguards are provided for by law to avoid any risk of discrimination" (Council of Europe 2021d). It also offers guidelines for developers, manufacturers and service providers regarding the representativeness of datasets. Grounding its recommendations in Article 5 of Convention 108+ on data accuracy, the guidelines state that developers, manufacturers and users "have to avoid mislabelling, thereby sufficiently testing their systems and identifying and eliminating disparities in accuracy, notably with regard to demographic variations in skin colour, age and gender, and thus avoid unintended discrimination". The guidelines also request that law enforcement authorities "while considering the deployment of facial recognition technologies in uncontrolled environments ... address the risk to various fundamental rights, including the rights to data protection, privacy, freedom of expression, freedom of assembly and freedom of movement, or the prohibition of discrimination, depending on the potential uses in different locations" notably through data protection impact assessments.

European Union

At European Union level, several provisions of the GDPR can play an important role in relation to addressing discrimination in AI and ADM systems (European Union 2016a). First and foremost, Article 9 on sensitive categories of personal data prevents, in principle, providers and deployers from using such data as variables (e.g. labels, input variables, risk factors) in AI and ADM systems. That said, as stated above, Article 10(5) of the AI Act allows such use, in interaction with Article 9(1)(g) GDPR, when the aim is to detect and mitigate algorithmic discrimination. Indeed, Recital 70 of the AI Act clarifies:

In order to protect the right of others from the discrimination that might result from the bias in AI systems, the providers should, exceptionally, to the extent that it is strictly necessary for the purpose of ensuring bias detection and correction in relation to the high-risk AI systems, subject to appropriate safeguards for the fundamental rights and freedoms of natural persons ... be able to process also special categories of personal data, as a matter of substantial public interest within the meaning of [Article 9(2)(g) GDPR.]

Besides, Article 22 GDPR confers a right not to be subjected to fully automated decisions. In particular, "[t]he data subject shall have the right not to be subject to a decision based solely on automated processing, including profiling, which produces legal effects concerning him or her or similarly significantly affects him or her". This provision has been used to challenge discriminatory credit scoring in the SCHUFA Holding and Dun & Bradstreet Austria cases (see section 1.4.4). Although Recital 71 of the GDPR grants data subjects a right "to obtain an explanation of the decision reached after ... assessment" by a fully ADM system, for a long time, legal uncertainty

surrounded the question of whether the GDPR offered a right to explanation and the possible nature and contours of such a right (Wachter, Mittelstadt and Floridi 2017; Malgieri and Comandé 2017; Selbst and Powles 2017). GDPR entitle data subjects to be provided with “information necessary to ensure fair and transparent processing” (Articles 13(2) and 14(2)), including regarding “the existence of automated decision-making, including profiling, referred to in Article 22(1) and (4) and, at least in those cases, meaningful information about the logic involved, as well as the significance and the envisaged consequences of such processing for the data subject” (Articles 13(2)(f), 14(2)(g) and 15(1)(h)). In the *Dun & Bradstreet Austria* decision (see section 1.4.4), the Court of Justice of the European Union confirmed the existence and extent of the right to explanation contained in Article 15(1)(h) GDPR, which conditions the access to meaningful information about how a decision was made, and offers a basis to challenge automated decisions based on Article 22(3) GDPR. Article 22(1) GDPR states that data subjects have the right “not to be subject to a decision based solely on automated processing, including profiling, which produces legal effects concerning [them] or similarly significantly affects [them]”. However, this right does not apply where automated processing or profiling is “necessary for entering into, or performance of, a contract between the data subject and a data controller”, or in case the data subject has given “explicit consent” as per Article 22(2)(a)(c) GDPR. In this case, Article 22(3) GDPR requires the data controller to implement “suitable measures to safeguard the data subject’s rights and freedoms and legitimate interests, at least the right to obtain human intervention on the part of the controller, to express his or her point of view and to contest the decision”. The right to explanation recently carved out by the Court of Justice can facilitate access to justice at least in cases of algorithmic discrimination based on profiling or fully ADM processes.

In the field of criminal law, Directive 2016/680 on the protection of natural persons with regard to the processing of personal data by competent authorities for the purposes of the prevention, investigation, detection or prosecution of criminal offences or the execution of criminal penalties, and on the free movement of such data is applicable (European Union 2016b). It provides more limited guarantees that may not suffice to challenge algorithmic discrimination in those areas of deployment.

1.4.4. Case law of the European Court of Human Rights and the Court of Justice of the European Union

European Court of Human Rights

In its decision in *Glukhin v. Russia* in 2023, the European Court of Human Rights assessed the legality of the use by the police of live face recognition technologies (European Court of Human Rights 2023). In this case, an individual was arrested following the posting on his social media of photos and a video of a peaceful solo demonstration he held in the Moscow Metro, which the police discovered during routine monitoring of the internet. To identify him, the police used video surveillance images from CCTV cameras that film public spaces combined with biometric data, in this case face recognition technology. The Court concluded there was a violation of the applicant’s freedom of expression under Article 10 and of his right to private life under Article 8 of the European Convention on Human Rights. In particular, the

Court stated “that the use of highly intrusive facial recognition technology in the context of the applicant exercising his Convention right to freedom of expression is incompatible with the ideals and values of a democratic society governed by the rule of law, which the Convention was designed to maintain and promote”. By taking into account “the difficulty for the applicant to prove his allegations because the domestic law did not provide for an official record or notification of the use of facial recognition technology”, the Court also acknowledges that the lack of transparency over the use of such technologies makes it difficult to adduce evidence and mitigate these hurdles (ibid., paragraph 72).

Court of Justice of the European Union

A request for a preliminary ruling by Belgium on its broad implementation of the EU Passenger Name Record Directive, including on how the automation of the data processing system should be implemented by national authorities, gave rise to a decision in C-817/19 *Ligue des droits humains ASBL v. Conseil des ministres* (Court of Justice of the European Union 2022). The case focused on the tension between automated data processing through machine learning algorithms and non-discrimination. The Court of Justice underlined safeguards on automated data processing, barring the use of self-learning AI algorithms to prevent discriminatory results. Moreover, pre-determined criteria used to assess terrorist threats cannot be based on “a person’s race or ethnic origin, political opinions, religion or philosophical beliefs, trade union membership, health, sexual life or sexual orientation”(ibid.).

In its decision in C-634/21 *SCHUFA Holding* the Court of Justice classified profiling in the form of risk scoring in relation to credits as a form of an individual automated decision falling under Article 22 GDPR (Court of Justice of the European Union 2023). In the context of the dispute giving rise to this decision, the applicant was refused a loan by a financial institution based on the risk score provided by SCHUFA Holding, a German credit scoring company. The question was “whether the establishment of [such] a probability value ... constitutes automated individual decision-making within the meaning of Article 22(1) of the GDPR”, in other terms “a decision based solely on automated processing, including profiling, which produces legal effects concerning the data subject or similarly significantly affects him or her” (ibid.: paragraphs 20, 27). Although the risk score provided by SCHUFA is transmitted to a third party, which is ultimately responsible for making the contractual decision, the Court argued that that third party “draws strongly on that value” (ibid.: paragraph 48). It also suggested that the notion of “decision” in Article 22 GDPR “is broad enough to encompass the result of calculating a person’s creditworthiness in the form of a probability value concerning that person’s ability to meet payment commitments in the future” (ibid.: paragraph 46). Hence, where “the probability value established by a credit information agency and communicated to a bank plays a determining role in the granting of credit, the establishment of that value must be qualified in itself as a decision producing vis-à-vis a data subject ‘legal effects concerning him or her or similarly significantly [affecting] him or her’ within the meaning of Article 22(1) of the GDPR” (ibid.: paragraph 50). This is a landmark decision that offers a legal ground to challenge discrimination based on profiling and other decisions based overwhelmingly on automated risk scores.

The Court of Justice further developed the legal remedies available to victims of algorithmic discrimination in its *Dun & Bradstreet Austria* decision (Court of Justice of the European Union 2025). In this case, the applicant was refused a mobile phone contract by an Austrian telecommunications company on the basis of an automated credit assessment. The algorithmic risk score provided by Dun & Bradstreet, an undertaking specialising in the provision of credit assessments, indicated that she did not have sufficient financial creditworthiness to be granted this contract, which would have required a monthly payment of €10 (*ibid.*: paragraphs 2, 16). The applicant seized the Austrian data protection authority, which ordered the credit scoring company to disclose “meaningful information about the logic involved in the automated decision-making based on personal data” to the applicant (*ibid.*: paragraph 17). Invoking trade secrets, the credit assessment company initially refused to provide the applicant with information other than her risk score (*ibid.*). It had only stated that “certain socio-demographic data concerning [the applicant] had been ‘given equal weighting’” to establish that risk score, which the Austrian court deemed insufficient. In an attempt to enforce that decision, a domestic court appointed an expert who considered that fulfilling the obligation to provide “meaningful information” would require Dun & Bradstreet Austria to disclose to the applicant which personal data (e.g. date of birth, address, sex) were processed to generate “factors” and “the specific value” attributed to the applicant for each factor, “the mathematical formula” on which the score was based, “the precise intervals within which the same value is attributed to different data for the same factor” as well as “a list of scoring for the period covering the six months preceding and the six months following the establishment of [the applicant]’s score, as obtained using the same calculation rule” for purposes of accuracy verification (*ibid.*: paragraphs 23-25). In the domestic court’s view, the possibility for applicants to verify the accuracy of profiling is crucial for enforcement purposes.¹³ In the dispute at stake, for example, “the information provided to [the applicant], including, *inter alia*, the score obtained, showed [her] to have very good credit standing, [whereas] the actual profiling led to her being regarded as not creditworthy” (*ibid.*: paragraph 27). Access to meaningful information therefore conditions applicants’ ability to contest discriminatory profiling or automated decisions in the context of Article 22(3) GDPR (see section 1.4.3.) (*ibid.*: paragraph 55).¹⁴

The Court of Justice explained that “the right to obtain ‘meaningful information about the logic involved’ in automated decision-making, within the meaning of [Article 15(1)(h)], must be understood as a right to an explanation of the procedure and principles actually applied in order to use, by automated means, the personal data of the data subject with a view to obtaining a specific result, such as a credit

13. The domestic court argued that “[i]n the event that Article 15(1)(h) of the GDPR does not guarantee this, the right of access to the data subject’s personal data and other information provided for therein would be rendered meaningless and useless, especially since each controller could in that case be able to provide incorrect information” (*ibid.*, paragraph 29).

14. The Court stated that “in the specific context of the adoption of a decision based solely on automated processing, the main purpose of the data subject’s right to obtain the information provided for in Article 15(1)(h) of the GDPR is to enable him or her effectively to exercise the rights conferred on him or her by Article 22(3) of that regulation, namely the right to express his or her point of view on that decision and to contest it” (*ibid.*: paragraph 55).

profile”, which “must be provided by means of relevant information and in a concise, transparent, intelligible and easily accessible form” (ibid.: paragraphs 58-59). While the Court did not list relevant elements, it indicated that “the mere communication of a complex mathematical formula, such as an algorithm, or ... the detailed description of all the steps in automated decision-making” would not satisfy the requirements of Article 15(1)(h) because “none of those would constitute a sufficiently concise and intelligible explanation” (ibid.: paragraph 59). In other terms:

the “meaningful information about the logic involved” in automated decision-making, within the meaning of Article 15(1)(h) of the GDPR, must describe the procedure and principles actually applied in such a way that the data subject can understand which of his or her personal data have been used in what way in the automated decision-making at issue, with the complexity of the operations to be carried out in the context of automated decision-making not being capable of relieving the controller of the duty to provide an explanation (ibid.: paragraph 61).

For the Court, it may be “sufficiently transparent and intelligible to inform the data subject of the extent to which a variation in the personal data taken into account would have led to a different result” (ibid.). In the context of anti-discrimination law, this right to explanation is essential as it could enable applicants to compare themselves with real or hypothetical comparators, an essential building block of the discrimination test, which has however become highly inaccessible in the context of algorithmic opacity and epistemic fragmentation (see section 1.4.3.). Finally, where trade secrets might be compromised, “the controller is required to provide the allegedly protected information to the competent supervisory authority or court, which must balance the rights and interests at issue with a view to determining the extent of the data subject’s right of access provided for in Article 15 of the GDPR” (ibid.: paragraph 76).

In France, 15 organisations recently challenged a decision of the CNAF (Caisse nationale des allocations familiales, a branch of the social security) to use a risk scoring algorithm to decide which beneficiaries to investigate for fraud with the Conseil d’État. The ADM system was shown to generate discriminatory effects *inter alia* for women and people living in poverty. The litigants asked the Conseil d’État to refer questions to the Court of Justice, including in relation to algorithmic discrimination. This case should be monitored because it will provide the Court of Justice with an opportunity to consider algorithmic discrimination through the lens of direct and indirect discrimination.

1.5. Key stakeholders

European AI Office

Several institutions will play a role in combating algorithmic discrimination. The newly created European AI Office¹⁵ (established within the European Commission) will support the implementation of the AI Act and will investigate possible infringements. Under Article 90, the European AI Office will be responsible, in particular, for

15. See <https://digital-strategy.ec.europa.eu/en/policies/ai-office>, accessed 10 November 2025.

dealing with alerts triggered by the Scientific Panel of Independent Experts (Article 68 AI Act) on GPAI presenting systemic risks. Akin to a market surveillance authority, the AI Office will also have the powers to monitor and supervise the compliance of AI systems that embed a GPAI model (Article 75(1) AI Act). It will also co-operate with market surveillance authorities to monitor, investigate and assess compliance of GPAI used in high-risk areas (Article 75(2) AI Act). As per Article 75(3) AI Act, in case “a market surveillance authority is unable to conclude its investigation of the high-risk AI system because of its inability to access certain information related to the general-purpose AI model despite having made all appropriate efforts to obtain that information, it may submit a reasoned request to the AI Office, by which access to that information shall be enforced”. In light of these provisions, it is likely that equality bodies designated under Article 77 AI Act will co-operate with the AI Office directly or indirectly in the framework of their co-operation with market surveillance authorities when addressing risks to, or infringements of, anti-discrimination rights. Indeed, as explained under section 1.4.1, Article 73(7) AI Act requires the relevant market surveillance authorities to inform the national public authorities or bodies referred to in Article 77(1) AI Act in case of a serious incident understood as “the infringement of obligations under Union law intended to protect fundamental rights” under Article 3(49)(c) AI Act. Article 79(2) AI Act also requires “[w]here risks to fundamental rights are identified” market surveillance authorities to “inform and fully cooperate with the relevant national public authorities or bodies referred to in Article 77(1)”. Consultation with Article 77(1) authorities is also mandated under Article 82(1) AI Act.

European Artificial Intelligence Board

Article 65 of the EU AI Act also establishes a European Artificial Intelligence Board.¹⁶ The AI Board is composed of representatives of member states and grants observer status to the European Data Protection Supervisor. The AI Board plays a key role in the enforcement of the AI Act, in particular through helping “coordinate and ensure cooperation between EU member states, aiming for consistent implementation and application of the AI Act across the Union” (European Commission 2025c). It provides a co-operation platform for national competent authorities responsible for enforcing the AI Act and facilitates the exchange of technical and regulatory expertise as well as best practices. It is also tasked with creating sub-groups to facilitate co-operation in enforcement matters, in particular two standing sub-groups to provide a platform for co-operation and exchange among market surveillance authorities and notifying authorities, but also to provide advice on AI policy topics.

AI Advisory Forum

Article 67 AI Act establishes an AI Advisory Forum¹⁷ responsible for providing technical expertise and advising and supporting the work of the AI Board and the European Commission. The Advisory Forum is composed of “a balanced selection of stakeholders, including industry, start-ups, [small and medium-sized enterprises]

16. See <https://artificialintelligenceact.eu/article/65>, accessed 10 November 2025.

17. See <https://artificialintelligenceact.eu/article/67/>, accessed 10 November 2025.

SMEs, civil society and academia” with balanced representation of commercial and non-commercial interests as well as SMEs and other undertakings.

Scientific Panel of Independent Experts

Article 68 AI Act establishes a Scientific Panel of Independent Experts that aims to support the enforcement activities under the Act. Its membership consists of independent experts selected by the Commission on the basis of up-to-date scientific or technical expertise in the field of AI necessary for the tasks of:

- ▶ supporting the implementation and enforcement of the AI Act, especially with regard to GPAI models (e.g. by alerting the AI Office of possible systemic risks, contributing to the development of tools and methodologies for evaluating the capabilities of GPAI, providing advice on the classification of general-purpose AI models with or without systemic risk, and contributing to the development of tools and templates);
- ▶ supporting the work of market surveillance authorities, at their request;
- ▶ supporting cross-border market surveillance activities;
- ▶ supporting the AI Office in carrying out its duties in the context of the Union safeguard procedure.

CEN-CENELEC

The [CEN-CENELEC Joint Technical Committee 21](#) (JTC 21) is conducting work on standards for AI based on a standardisation request by the European Commission.¹⁸ Currently, the JTC 21 is working on producing a harmonised standard translating the requirement, established by Article 9 AI Act, for providers of high-risk AI systems to have a risk management system in place.

Role of equality bodies

Equality bodies have a key role to play. Algorithmic discrimination often takes place in ways that are difficult to detect and prove. This presents a significant barrier to access to justice for those who are affected. Individuals who experience algorithmic discrimination may face challenges in understanding how or why they were treated unfairly, making it difficult to challenge these decisions. Equality bodies thus play a crucial role in supporting victims and promoting accountability. In this context, the recently adopted equality bodies directives (Directive 2024/1500 and Directive 2024/1499) strengthen the role and powers of equality bodies in relation to their independence, resources, and role in facilitating access to justice. They also confirm the role of equality bodies in addressing AI-driven discrimination:

Devoting attention to the opportunities and risks presented by the use of automated systems, including artificial intelligence, is key. In particular, equality bodies should be equipped with appropriate human and technical resources. Those resources should, in particular, enable equality bodies to use automated systems for their work on the one

18. European Commission's implementing decision of 22 May 2023 containing standardisation request C(2023)3215.

hand and to assess such systems as regards their compliance with non-discrimination rules on the other hand. Where the equality body is part of a multi-mandate body, the resources necessary to carry out its equality mandate should be ensured (European Union 2024b: Recital 21; European Union 2024a: Recital 22).

In this context, equality bodies can play a critical role in supporting victims, raising awareness, conducting research, and advocating for better regulation of AI and ADM systems:

- ▶ acting *ex officio* and handling complaints from victims of algorithmic discrimination: equality bodies can have *ex officio* powers of investigation, and they serve as an accessible point of contact for potential victims of algorithmic discrimination. They can handle complaints from victims, investigate alleged instances of algorithmic bias, and provide advice and support on how to pursue legal action;
- ▶ providing legal assistance and taking cases to court: this could include deciding cases involving algorithmic discrimination, advising individuals on how to challenge discriminatory decisions driven by AI or ADM systems or, in some cases, representing potential victims or acting on their behalf in court;
- ▶ raising awareness and informing victims and the wider public: equality bodies can play a key role in educating the public on the potential harms of algorithmic discrimination, including by providing information about how to challenge unfair automated decisions and what rights individuals have in relation to AI-driven and ADM processes;
- ▶ collecting data and carrying out research on discrimination linked to AI and ADM systems: another critical role of equality bodies is to collect data on the incidence of algorithmic discrimination. By systematically tracking complaints related to AI and ADM systems and carrying out research, data collection and monitoring, equality bodies can identify patterns of discrimination and highlight systemic issues that may require regulatory intervention;
- ▶ making recommendations to policy makers and legislators: such data can also provide concrete evidence of the need for legal reforms (e.g. in relation to registration and information obligations). Equality bodies can offer evidence-based recommendations to policy makers and legislators on how to regulate AI and ADM systems based on their monitoring activities, particularly from the perspective of enforcement and related difficulties;
- ▶ engaging with stakeholders to promote equality in AI and ADM systems: equality bodies can engage with a wide range of stakeholders – including employers, service providers, public institutions and CSOs – to encourage the adoption of good equality practices in relation to AI and ADM systems. This could involve helping organisations develop and implement equality plans that minimise the risk of algorithmic bias, providing guidance on how to test AI and ADM systems for discrimination.



2. Legal gaps

This section identifies shortcomings and needs in the current legal, procedural and governance frameworks in light of discriminatory risks arising from the use of AI and ADM systems in public administrations and other sectors falling within the scope of this study. It also summarises the main hurdles, gaps and needs encountered by equality bodies and other relevant stakeholders in their monitoring and supervision functions.

2.1. Rationales for the adoption of AI and ADM systems

Section I has shown that the take-up of AI and ADM technologies is increasingly widespread. Yet several problems arise in relation to their adoption. On the one hand, the narratives promoting the adoption of these systems are problematic and heighten legal risks. AI and ADM systems are often presented as “neutral” or at least “less biased” than human decision makers. Their deployment could be justified by a willingness to rationalise decision-making processes. Since these systems are known to be widely biased, such a narrative can create the illusion that the roll out of AI and ADM reduces or even prevents discrimination. This creates risks that breaches of anti-discrimination legislation are overlooked.

Moreover, AI and ADM systems are often promoted as means to increase cost efficiency in decision-making processes, not least in public administrations that are subject to enhanced public scrutiny in the context of spending cuts and austerity politics. This rationale leads to the adoption of AI and ADM for specific usages, such as fraud detection or risk assessment, which are presented as drivers of spending efficiency. Such a narrative obfuscates at least two important questions:

- ▶ should such technologies be adopted for alternative usages, for example to enhance access to rights and public services, which studies regularly show to be underclaimed? (Défenseur des Droits 2024);
- ▶ can such systems really allow the lowering of public spending while maintaining the same quality of public services, once all costs related to the safeguard of fundamental rights have been internalised?

For example, preventing risks of discrimination from materialising through adequate testing, monitoring and audit procedures, training case workers that handle algorithmic recommendations to ensure effective human oversight, and setting up effective information and redress mechanisms for users of public services is costly. Such costs should arguably be fully internalised by public administrations when decisions to implement such systems are made. Such deliberation should also prompt public decision makers to ask “zero questions”: when the deployment of an AI or ADM system appears too risky or costly once prevention measures and fundamental rights safeguards have been adequately implemented, this should lead to a decision not to deploy that system.

2.2. Lack of transparency

As detailed in the relevant sections below, there is a generalised lack of systematic, clear and accessible information on the development, experimentation and deployment of AI and ADM systems in public administrations and beyond. The use of AI and ADM systems in public administrations is not systematically mapped despite ongoing plans or efforts to create databases of such uses (e.g. the Dutch Algorithms Register).¹⁹ Cases of algorithmic discrimination are still very few, and are often framed around issues of data protection rather than discrimination and equality. The rare existing cases have often emerged in courts with the support of equality bodies and CSOs. Yet, the lack of information around when, why and where AI and ADM systems are used by public administrations – and hence the lack of legal reporting obligations applicable so far – makes it difficult for these actors to effectively and systematically monitor, assess and challenge cases of algorithmic discrimination despite their investigative powers. Indeed, to effectively use these investigative powers, they first and foremost have to be aware of the existence of AI or ADM applications. Instead, equality bodies often have to rely on media investigations, specialist CSOs or informal networks to identify potential cases of algorithmic discrimination. While national experts report that requests for information (e.g. under freedom of information legislation) are regularly sent to public authorities for the sake of creating public records of usages of AI and ADM systems in public administrations, their legitimacy

19. See <https://algoritmes.overheid.nl/en>, accessed 10 November 2025.

is sometimes called into question, or they may remain unanswered when they have no binding character, especially when those requests touch on sensitive applications such as digital surveillance technologies. In addition, applicants often cannot verify the information received. Other channels for obtaining information are sometimes used, such as questions to the government raised in collaboration with members of national parliaments. However, the effectiveness of both freedom of information requests and questions in parliament is undermined by the very lack of transparency on the use of AI and ADM applications. Some equality bodies offer positive examples in this regard: the Finnish Non-Discrimination Ombudsman, for example, can send mandatory requests for information and clarification to public administrations. Such powers can facilitate access to information and transparency with regard to the use of AI and ADM systems and their potential discriminatory effects.

To some extent, the lack of information will be addressed by the new AI regulations. For example, the EU AI database foreseen under Article 71 AI Act will feature public information about high-risk AI applications.²⁰ In addition, Article 14(a)(b) of the Framework Convention on AI demanding 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” allows these persons “to contest the decision(s) made or substantially informed by the use of the system”. These provisions are further strengthened by Articles 8 and 9 of the convention on transparency and oversight and on accountability and responsibility. Such information and transparency requirements can enable public scrutiny. In the same vein, Article 16 of the convention requires the adoption and maintenance of a risk and impact management framework (supported by the HUDERIA standardised impact assessment), and Article 16(2)(f) of the convention demands that “risks, actual and potential impacts, and the risk management approach” be documented. Having access to such information could help bodies in charge of protecting fundamental rights to prevent and address algorithmic discrimination. In turn, the AI Act foresees that the authorities designated under its Article 77 have access to technical documentation and risk assessment systems that are based on the AI Act and fundamental rights impact assessments when mandatory under Article 27 AI Act (limited to certain deployers that offer public or essential private services).

That said, such information and transparency obligations could have been cast in broader terms. At the moment, this only concerns high-risk AI systems, leaving out systems described as non-high risk and ADM systems that are not considered to fall within the definition of AI. In 2020, the Parliamentary Assembly of the Council of Europe demanded that governments be requested “to notify the parliament before [AI and ADM] technology is deployed” and that “the use of such technologies by the authorities ... be systematically recorded in a public register” (Council of Europe, Parliamentary Assembly 2020b). This is not directly reflected in the recently adopted

20. High-risk AI systems, including biometrics, used in the areas of law enforcement, migration, asylum and border control management, are registered in a non-public section of the database to which the European Commission and specific national market surveillance authorities (e.g. data protection agencies) have access.

Framework Convention on AI, which does not require the creation of a public register. In turn, the AI Act does not foresee notifications to national parliaments or EU institutions upon deployment of AI or ADM systems by public authorities.

2.3. Access to justice issues

Access to justice is made difficult due to multiple information and power asymmetries. As explained above, one major finding of the research conducted at national level shows that the lack of clear and consistent mapping of usages of AI and ADM systems by public administrations jeopardises the right to non-discrimination and its application. Not knowing that a given administrative decision has been made with (the support of) an AI or ADM system prevents subjects of such decisions from even suspecting algorithmic discrimination. Even where suspicion arises, the lack of generalised obligations for public administrations to provide information about the type of system used, the decision factors implemented and the way it is used in the decision-making process, prevents victims from adducing sufficient evidence to establish presumptions of discrimination in courts.²¹

To be sure, European non-discrimination laws foresee a sharing of the burden of proof between applicants and defendants to ease applicants' burden when claiming discrimination. EU anti-discrimination law foresees that "when persons who consider themselves wronged because the principle of equal treatment has not been applied to them establish, before a court or other competent authority, facts from which it may be presumed that there has been direct or indirect discrimination, it shall be for the respondent to prove that there has been no breach of the principle of equal treatment".²² Similarly, the European Court of Human Rights shifts the burden to prove that there has been no discrimination onto defendants when an applicant presents credible and consistent facts from which it may be presumed that discrimination has taken place: "[P]roof may follow from the coexistence of sufficiently strong, clear and concordant inferences or of similar unrebutted presumptions of fact" (European Court of Human Rights 2007). In principle, these rules facilitate applicants' task by only requiring them to establish a presumption of discrimination before the burden of proof can shift to the defendant. It is then for the defendant to rebut the presumption of discrimination.

However, these rules insufficiently address power and information asymmetries that manifest when a decision is made with (the support of) an AI or ADM system, and may fall short in cases of algorithmic discrimination. On the one hand, national experts report that courts sometimes set a high evidentiary threshold to trigger a shift the burden of proof onto defendants. On the other hand, algorithmic opacity complicates the task of individual applicants attempting to establish a presumption of algorithmic discrimination (see section 2.2. above). Information and power asymmetries also take the form of what Milano and Prunkl (2024) call "epistemic fragmentation", namely "a structural characteristic of algorithmically-mediated

21. To some extent, this difficulty can be mitigated through activating the right to explanation provided for in the GDPR, though the procedural hurdles and costs could be high for individual victims.

22. See Article 8(1) of Directive 2000/43/EC. This provision is present in all anti-discrimination directives: 2000/78/EC, 2004/113/EC and 2006/54/EC.

environments that isolate individuals” thus “hinder[ing] people and communities from meaningfully sharing and comparing their experiences”. Epistemic fragmentation thus “mak[es] it more difficult to ... identify and conceptualise emerging harms in these environments”, in particular because European non-discrimination law relies on victims’ ability to show that they have been subjected to differential treatment or a particular disadvantage compared to other persons in a similar situation (ibid.). Together, algorithmic opacity and epistemic fragmentation make it very difficult to challenge discriminatory AI and ADM legally.

Thus, access to justice is compromised in practice. To give a concrete example, the ADM systems used by several social security administrations to detect fraud in Europe create a category of *de facto* “suspicious individuals” that is highly discriminatory but very difficult for those individuals to challenge, with dire consequences on their lives (suspension of social benefits, claims for undue payments, etc.). Though the right to explanation recently confirmed by the Court of Justice in the context of Article 15(1)(h) and Article 22 of the GDPR in the case of *Dun & Bradstreet Austria* partly address the problem (see section 1.4.4.), power and information asymmetries remain, which must be addressed by easing the burden of proof even further for applicants, through empowering equality bodies and relevant stakeholders to support potential victims and through increased prevention.²³ This issue was underlined as early as 2020 by the Parliamentary Assembly of the Council of Europe, which highlighted the need to “pay particular attention to guaranteeing the presumption of innocence and ensuring that victims of discrimination do not face a disproportionate burden of proof” (Council of Europe, Parliamentary Assembly 2020b). It demanded that member states “ensure that equality bodies are fully empowered to address issues of equality and non-discrimination that arise due to the use of AI”.

In this perspective, the powers of some European equality bodies to promote protection against discrimination at the collective and structural level can help mitigate these shortcomings. Indeed, some equality bodies can contribute to preventing algorithmic discrimination through supervision, investigation and sanctions powers. Additionally, some equality bodies can monitor the use of AI and ADM systems and challenge discriminatory outputs in court in the name of the public interest, without any individual victim needing to be identified beforehand (see Court of Justice of the European Union 2008). While that is not the case for all European equality bodies, such powers can mitigate the difficulty of having to rely on individual victims reporting discriminatory harms before investigations can be launched. However, to be effective such powers have to be backed by robust information rights: equality bodies must be entitled to request and obtain relevant and meaningful information regarding how AI and ADM systems are used. Such powers should be extended to situations where AI and ADM systems are used to support and assist decision making, as opposed to being restricted to situations of full automation only, to reflect the fact that hybrid decision making is the most widespread form of use. The transposition of the new directives on standards for equality bodies by June 2026 will be a good opportunity for reviewing the powers of equality bodies in light of these new

23. This is all the more important since the AI liability directive has been withdrawn. The directive proposed to facilitate the establishment of presumptions and the reversal of the burden of proof, though to a limited extent.

challenges. In addition, alternative procedural routes such as collective complaint mechanisms (e.g. *actio popularis*, class actions) could mitigate access to justice issues.

2.4. Enforcement of existing and future provisions

2.4.1. Institutional co-operation

The enforcement system established under the AI Act is complex and does not leave much space for individual complaints. Though market surveillance authorities can receive complaints, there is no obligation for them to investigate on that basis. In this perspective, the role of equality bodies becomes even more important. The enforcement system established by the AI Act foresees multi-stakeholder co-operation between equality bodies, other fundamental rights supervision authorities designated under Article 77, market surveillance authorities and other relevant institutions. In addition, equality bodies will co-operate with other fundamental rights institutions designated under Article 77 AI Act. According to some of the institutions interviewed for this report, however, there could be significant challenges in this regard. In particular, co-operation methodologies must be developed and implemented, and institutional communication channels ensured, in order to promote an approach based on institutional complementarity and intersectionality. Furthermore, competent authorities, including single contact points and market surveillance authorities, need to have an adequate understanding of the risks of discrimination presented by algorithmic technologies, and of legal and institutional ways of addressing discrimination.

2.4.2. Conceptual gaps

Other enforcement problems arise in relation to the concepts and procedures enshrined in European anti-discrimination law. Concretely, the functioning of AI, and in particular machine learning, is based on correlations, not causation. This enhances the risk of so-called proxy discrimination, which judges could be reluctant to treat as direct discrimination even when such proxies are known to correlate with protected grounds.²⁴ For example, in certain contexts, individuals' place of residence is known to strongly correlate to race and ethnic origin. At European level, residence does not feature in the list of protected grounds enshrined in non-discrimination law. Hence, such algorithmic proxy discrimination is likely to be treated as indirect discrimination, except in cases where the place of residence is considered to be "inextricably" or "inherently linked" with a protected ground such as race or ethnic origin.

In this regard, EU secondary law prohibiting discrimination could prove too limited as it only does so in relation to a "closed" list of protected grounds: sex, race or ethnic origin, disability, sexual orientation, religion or belief and age. Algorithmic harms

24. For more information about proxy discrimination, see Bartoletti I. and Xenidis R. (2023), "Study on the impact of artificial intelligence systems, their potential for promoting equality, including gender equality, and the risks they may cause in relation to non-discrimination", available at: <https://edoc.coe.int/en/artificial-intelligence/11649-study-on-the-impact-of-artificial-intelligence-systems-their-potential-for-promoting-equality-including-gender-equality-and-the-risks-they-may-cause-in-relation-to-non-discrimination.html>, accessed 10 November 2025.

related to other social characteristics would only be regarded as directly discriminatory if they are considered to be “inherently linked” to one of these protected grounds. Such harms could however be treated as indirect discrimination if they affect a group related to one of these protected grounds. Other algorithmic harms would fall outside the scope of legal protection provided they are not captured by specific national legislative frameworks, which can go beyond EU law.

By contrast, so-called “open-ended” lists such as the protected criteria enshrined in Article 21(1) of the European Charter of Fundamental Rights and Article 14 of the European Convention on Human Rights could offer more flexibility in tackling algorithmic proxy discrimination. This is because they allow judges to find discrimination based “on any ground such as”, respectively, “sex, race, colour, ethnic or social origin, genetic features, language, religion or belief, political or any other opinion, membership of a national minority, property, birth, disability, age or sexual orientation” and “sex, race, colour, language, religion, political or other opinion, national or social origin, association with a national minority, property, birth or other status”.

The question of how to qualify algorithmic proxy discrimination is important. Whereas this qualification does not make a difference under the Convention’s anti-discrimination law due to the existence of a unitary regime of open-ended justifications for both direct and indirect discrimination, the problem is different under EU anti-discrimination law. There, direct discrimination benefits from a stricter justification regime based on closed exceptions. Given the lack of transparency, opacity and complexity of AI and ADM systems, an open-ended justification regime that requires a contextual analysis may not be the most appropriate. Yet, substantiating direct discrimination is often problematic in the algorithmic context. As argued elsewhere, beyond the problems of justification, the very concepts of direct and indirect discrimination are not entirely well suited to capturing algorithmic discrimination, which mixes characteristics of both notions (Xenidis 2021).

The high level of differentiation performed by AI and ADM systems also enhances risks of intersectional discrimination, especially when users combine several risk categories that accrue to their risk score (e.g. gender and migration history in examples of biased unemployment predictions). Yet, intersectional discrimination is still not fully recognised in the EU anti-discrimination directives (save for the recent Equal Pay Transparency Directive 2023/970) and is not uniformly prohibited in national legislation despite recent evolutions. Hence, member states will need to review their anti-discrimination legislation and amend it as necessary, so as to ensure that it covers all cases where direct or indirect discrimination, including discrimination by association and intersectional discrimination, may be caused by the use of AI and ADM systems.

2.4.3. Restrictions on data collection and processing

Finally, existing legal restrictions on the collection of equality data curtail possibilities to effectively investigate algorithmic discrimination and promote equality in the use of AI and ADM systems. Advocating for more adequate equality data collection procedures may be necessary to effectively address discrimination in AI. Further enforcement issues relate to the use of sensitive categories of personal data. Article

10(5) AI Act exceptionally allows processing of sensitive categories of personal data (defined in Article 9 GDPR) such as ethnicity or sexual orientation in order to ensure bias detection, yet this possibility does not extend to AI systems that fall outside the high-risk category. This could present important limitations when it comes to detecting, preventing and mitigating discriminatory biases systems regarded as low risk to which non-discrimination obligations apply transversally. Whether equality bodies and other fundamental rights supervision authorities will be able to avail of the exceptions in Article 9 GDPR and Article 10(5) AI Act to gauge the discriminatory impacts of given AI systems is also uncertain. Algorithmic discrimination may be more difficult to detect and monitor when sensitive categories of personal data – which overlap with many (but not all) of the categories protected under EU and national non-discrimination law – cannot be used for testing, auditing and assessment purposes.

2.5. “De-risking” AI systems

The risk-based classification of the AI Act has been subjected to criticism. The list of high-risk applications provided in Annex III of the Act has been described as overly rigid in the context of rapidly evolving technologies. In addition, the AI Act offers no clear guidance for deciding whether an application falling within the sectors listed in Annex III is high risk or not. This is particularly problematic in light of Article 6(3) AI Act, which provides for a number of exceptions. For example, providers can declare that systems that would normally fall within the list of high-risk sectors in fact pose no significant risk of harm to the health, safety or fundamental rights of natural persons. AI systems are not considered high risk when such systems are used to “perform a narrow procedural task” or a “preparatory task to an assessment”,²⁵ to “improve the result of a previously completed human activity” or to “detect decision-making patterns or deviations from prior decision-making patterns ... [where such systems are] not meant to replace or influence the previously completed human assessment, without proper human review”. Providers that consider their systems as not high risk must only “document its assessment before that system is placed on the market or put into service” (Article 6(4)) and register the system in accordance with Article 49(2).

There is no systematic control at this point as national competent authorities can, but have no systematic obligation to, request the assessments performed. Hence, this creates a potential loophole that could potentially be used widely to “de-risk” AI systems in order to avoid complying with the requirements and safeguards applying to high-risk systems under the AI Act.²⁶ Here, the approach of the Framework Convention on AI, which is rights based as opposed to risk based, could prove particularly useful in extending fundamental rights safeguards to “derisked” AI and ADM applications. Another potential gap that could be exploited in both sets of regulations relates to the definition of AI. As both texts apply to AI as opposed to, for instance, simple rule-based systems, certain ADM systems could be excluded if they

25. That said, under Article 6(3)(d), an AI system used for profiling natural persons in the areas listed in Annex III is always considered high risk.

26. Those requirements and safeguards are listed under section on high-risk AI systems.

are considered to fall outside the scope of the AI definition (European Commission 2025d). This is another loophole that providers of simpler ADM systems may be able to use to evade the obligations of the AI Act.

2.6. Technical standards, equality bodies and the question of harmonisation

The AI harmonised standards that will be issued by CEN-CENELEC will be a critical device for supporting compliance with fundamental rights, including the prevention of, and protection against, algorithmic discrimination. The key question is, however, which risks of discrimination will the standards require providers to identify, evaluate and address? In particular, which forms of discrimination will providers have to take into account and how will they have to address them to be considered compliant with the AI Act? Several problems arise here. The first is that of the power entrusted to the private organisations involved in the co-regulation process. Most of the organisations involved in the JTC 21 of CEN-CENELEC, in charge of drawing up these standards, are industry players that have the resources to fund such participation. CSOs and organisations with a legitimate interest can participate to some extent (Equinet is for instance currently involved in the standard-making process), but often have limited financial means to do so. This set-up also means that expertise on the ground may be unequally distributed, with an adequate level of technical expertise and more limited legal expertise, especially with regard to human rights. The draft standards are not public, which makes it difficult to monitor and assess the effectiveness and flaws of the enforcement strategies that are currently under discussion. This pre-empts attempts to influence these discussions to address existing limitations.

The second problem is that of the compatibility and complementarity of this approach with that of EU fundamental rights law, including non-discrimination law. On the one hand, the level of risk mitigation that the standards will prescribe for compliance purposes is still unclear. This is a key question because Article 9(5) AI Act foresees that “[t]he risk management measures ... shall be such that the relevant residual risk associated with each hazard, as well as the overall residual risk of the high-risk AI systems is judged to be acceptable”. In relation to bias and discrimination, acceptable residual risks under the AI Act will not guarantee compliance with EU fundamental rights law, which adopts a different approach centred on actual harms. In fact, Article 52(1) of the EU Charter of Fundamental Rights prescribes that “[a]ny limitation on the exercise of the rights and freedoms recognised by this Charter must be provided for by law and respect the essence of those rights and freedoms” and that “[s]ubject to the principle of proportionality, limitations may be made only if they are necessary and genuinely meet objectives of general interest recognised by the Union or the need to protect the rights and freedoms of others”. The scope of permissible interferences – constrained by the essence of fundamental rights and conditioned on necessity and the general interest – appears to be in tension with the residual risk approach adopted in the AI Act. At the very least, the relationship between these concepts is ambiguous and will require clarification, potentially by the Court of Justice.

On the other hand, the standardisation approach adopted by the EU, combined with the (at least partial) maximum harmonisation framework enshrined in the EU's new AI regulation, sits in tension with the minimum harmonisation approach of the EU non-discrimination directives. While EU secondary law prohibits six grounds of discrimination (race and ethnic origin, sex and gender, disability, sexual orientation, religion or belief, and age), it also allows member states to go beyond these minimum requirements and to prohibit discrimination more widely through so-called "more favourable provisions", as long as they comply with the EU Treaties. Yet, it is unclear which forms of bias providers of high-risk systems will have to test for. For example, if bias testing and mitigation is only required in relation to the six grounds protected at EU level, and providers complying with these standards are then presumed to comply with the AI Act, this will allow them to sell their products within the whole internal market. Even though deployers may still be liable for discrimination under EU and national anti-discrimination laws requiring wider protection, questions arise in at least two regards.

First, more favourable national provisions prohibiting discrimination are explicitly allowed under EU discrimination law and should not be cast as obstacles to free movement. Such an approach would also breach Article 21(1) of the EU Charter of Fundamental Rights, which prohibits discrimination "based on any ground such as sex, race, colour, ethnic or social origin, genetic features, language, religion or belief, political or any other opinion, membership of a national minority, property, birth, disability, age or sexual orientation". Another scenario would be to include the wider definition of discrimination enshrined in the Charter. Yet, the tensions with more protective national legislation, expressly authorised under EU discrimination law, would remain. At the same time, should the standards make no mention of the prohibited grounds of discrimination that providers of high-risk AI systems have to account for when testing for bias or assessing the impact of their system on fundamental rights, there is a significant risk that such procedures will end up being largely ineffective.

The second question relates to the influence of such compliance mechanisms on the shift of the burden of proof in discrimination cases. Here it is important to recall that the presumption of compliance foreseen in Article 40 AI Act only applies to the requirements of the Act for high-risk systems (see the section on risk management system), and not to anti-discrimination or other fundamental rights law. In this regard, Article 82(1) AI Act does acknowledge that "although a high-risk AI system complies with this Regulation, it [may] nevertheless present a risk to the health or safety of persons, to fundamental rights".

Hence, it must be emphasised that the AI Act's (partial) maximum harmonisation approach should not be interpreted to override or diminish existing anti-discrimination obligations established under EU primary and secondary law.

What does that mean for equality bodies and relevant authorities protecting fundamental rights such as those designated under Article 77 AI Act? First, assessing the compliance of risk management systems and fundamental rights impact assessments with anti-discrimination law may be difficult if those do not provide information regarding the grounds of discrimination that have been tested for. The compelling question is whether equality bodies and relevant stakeholders enforcing fundamental

rights will have sufficiently precise information to assess how discrimination has been prevented and/or risks thereof mitigated in relation to legally protected groups in the context of high-risk AI systems, and how these groups have been defined during risk assessment procedures. Second, the legal dissonance highlighted above creates a gap between providers' (and, as the case may be, deployers') obligations to prevent and address risks of discriminatory bias in the context of high-risk systems, and the mandate of equality bodies and other relevant authorities. Most equality bodies are mandated to address discrimination based on all grounds protected under national law, which frequently extends beyond the grounds protected under EU secondary anti-discrimination law. Depending on the approach favoured in the new standards or in their application (e.g. lowest common denominator or unspecified types of bias), equality bodies will have to be particularly vigilant in relation to grounds of discrimination that are only prohibited under national law. Indeed, risks of algorithmic discrimination may not have been tested for by providers in relation to these nationally protected grounds.

Besides, standardisation presents several challenges for equality bodies. It may be important for equality bodies to monitor and influence how new harmonised AI standards will be defined by CEN-CENELEC because compliance with those standards will trigger a presumption of compliance with the essential requirements of the AI Act for providers of high-risk AI systems. Several difficulties arise, however. First, even though Equinet is involved in some sub-groups of the JTC 21 responsible for drafting those AI standards, the degree of publicity and accessibility of agreed standards remain unclear. In the case *Public.Resource.Org and Right to Know v. Commission and Others* the Court of Justice decided that standards must be made publicly accessible to EU citizens (Court of Justice of the European Union 2024). Since the dispute concerned access to standards related to toy safety, the question remains whether this transparency requirement will apply similarly to AI standards, and whether they will be made accessible free of charge.

Second, technical expertise may be required to understand and critically evaluate whether AI standards sufficiently ensure compliance with fundamental rights and non-discrimination legislation. According to a report by Equinet, the standards "will determine what information and level of detail are included in [the] documentation" produced under the essential requirements for high-risk AI systems under the AI Act and thus "the type of risk management and post-market monitoring to be carried out by providers to anticipate, identify, and mitigate risks to fundamental rights" (Mittelstadt 2025).

Hence, equality bodies should be particularly attentive to the technical choices implemented by the new AI standards:

- ▶ Which fairness metrics are used to measure performance? Which thresholds are used to quantify risks of discriminatory bias? Which de-biasing measures and fairness methods are used? Do they align with European anti-discrimination legislation and case law?
- ▶ How are acceptable levels of "residual risk" of discrimination defined?
- ▶ How are systems' intended purpose, expected level of accuracy and foreseeable misuses defined?

- ▶ How are discriminatory risks and impacts measured and mitigated with regard to groups protected under European anti-discrimination legislation?
- ▶ How are particular groups defined in technical documentation, risk management and impact assessments?
- ▶ Is intersectional discrimination adequately accounted for?
- ▶ Do technical group definitions cover the entire population protected under prohibited grounds of discrimination?

Being attentive to those questions is important to prevent ethics- and fundamental rights-washing.

This is a challenge for equality bodies that may have to seek new expertise, know-how and resources to be able to address these questions. Establishing robust co-operation templates with relevant authorities (market surveillance authorities, single points of contact, the AI Office, etc.) and bodies possessing complementary expertise (data protection agencies, ombudspersons, consumer protection offices, etc.) as well as relevant third parties (researchers, CSOs, etc.) can help equality bodies face the challenge of expertise and resources.

Summary – Use of AI in Belgian public administration and policy recommendations on algorithmic discrimination



The summary is based on the As-is and gap analysis report, in which the Belgian section was authored by Ine van Zeeland. This is a confidential report reviewing how AI and ADM systems are used in the public sector in Belgium, assessing AI-related legal and policy frameworks from equality and non-discrimination perspectives and providing recommendations to authorities and decision makers.

Use of AI in public administration

Belgian public administration uses several artificial intelligence (AI) applications and has various ongoing AI projects, for example in the police, finance, employment and education sectors. Due to the lack of a comprehensive public registry of these systems, it is not possible to fully and coherently map the use of AI and automated decision-making (ADM) systems across public services in Belgium. The Federal Public Service (FPS) Policy and Support (BoSA) hosts an AI Observatory that has launched efforts to collect information on the use of AI or ADM systems in public services in Belgium. Similar efforts to collect information are conducted or are being planned by other bodies such as the FARI Institute for Brussels and FPS Social Security.

Legislation and relevant case law

The EU AI Act together with European and Belgian anti-discrimination and data protection legislation form a legal framework for addressing algorithmic discrimination. Belgium's case law on AI and ADM and discrimination is relatively nascent. The As-is and gap analysis report identified a few cases relevant to the topic, such as the decision of the Court of Justice in C-817/19 *Ligue des droits humains ASBL v. Conseil*

des ministres,²⁷ which ruled on the tension between automated data processing through machine learning algorithms and non-discrimination.

In relation to data protection, the GBA-APD Litigation Chamber has ruled twice on cases related to ADM systems with potential relevance to anti-discrimination; one case involved credit scoring by a car-sharing company²⁸ and the other was on the development of an algorithmic model based on payment card data.²⁹ However, the Litigation Chamber decisions in both cases pertained to GDPR requirements that had no direct bearing on discrimination. Overall, there may be undetected Court cases of unequal treatment by AI and ADM systems. Indeed, in 2022, the Brussels Court of First Instance ruled in a case regarding indirect discrimination through ADM by scan cars deployed by Parking.³⁰

In addition to EU *acquis*, further opportunities to prevent and address algorithmic discrimination are provided in the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, which opened for signature in September 2024, in which Article 10 on equality and non-discrimination prohibits discrimination and requires states to “maintain measures aimed at overcoming inequalities to achieve fair, just and equitable outcomes”.

Gaps in legal and policy frameworks

Lack of transparency

A critical issue in Belgium is the lack of transparency and reporting on the use of AI and ADM systems, as currently there is no formal obligation for public sector organisations to disclose or register their use of AI and ADM systems. The absence of a comprehensive and publicly accessible database presents risks for both oversight and accountability, in addition creating challenges for equality bodies to monitor risks of discrimination or to support individuals who may risk being discriminated against by such systems. Research in the area consistently finds that cases of algorithmic discrimination are underreported and under-investigated.

Lack of awareness and harmonised approach to fundamental rights

Despite various trainings and ethical guidelines, systematic knowledge of risks related to discrimination is limited in many public services. Overall, public administration lacks a standardised approach to assess risks related to fundamental rights and discrimination in the use of AI in public services. There is a lack of a harmonised approach to include the perspectives of affected groups or individuals, and fundamental rights experts in design, development, and deployment of AI and ADM systems.

27. C-817/19 *Ligue des droits humains ASBL v. Conseil des ministres* EU:C:2022:491.

28. Décision 168/2023 du 19 décembre 2023, Chambre Contentieuse GBA-APD.

29. Décision quant au fond 46/2024 du 15 mars 2024, Chambre Contentieuse GBA-APD.

30. Case 21/6495/C, available at: www.unia.be/fr/legislation-et-jurisprudence/jurisprudence/tribunal-de-premiere-instance-de-bruxelles-le-2-mai-20221, accessed 10 November 2025. The plaintiffs in this case were a person with a disability and the Collectif accessibilité Wallonie Bruxelles (CAWaB), supported by the interfederal equality body, Unia.

Fragmented landscape for oversight

Belgium has an institutionally complex oversight landscape, where mandates are distributed across federal, regional and community levels. This creates a challenge for harmonised oversight and legal interpretations, and the risk of inconsistent application of the AI Act across public services in Belgium.

Resource and expertise gaps

Equality bodies and supervisory authorities in Belgium face significant challenges in addressing AI- and ADM-driven discrimination due to limited public financial, human and technical resources.

Key recommendations

- ▶ **Set up a public registry of AI and ADM systems:** Introduce a legal obligation to report the use of AI and ADM systems in the public sector to a national registry or repository with a certain level of public access, in addition to the Europe-wide registry, to ensure transparency and efficient oversight.
- ▶ **Set up advanced co-ordination mechanisms between national competent authorities at various levels:** Facilitate and organise formal co-operation agreements to guarantee harmonised oversight and decision making by the respective authorities and create guidelines for mutual involvement to ensure Article 77 AI Act authorities and national human rights structures are included in investigations proactively.
- ▶ **Form a network to share promising practices and enhance synergies:** Facilitate and support public authorities in setting up a structural interfederal co-operation network for practical advice and guidance and learning from best practices. This can address the fragmentation of relevant knowledge and capacities needed to develop or procure AI and ADM systems that incorporate fundamental rights protections by design, or at least allow for impact monitoring throughout the system life cycle. Such a network should include anti-discrimination, fundamental rights and personal data protection experts.
- ▶ **Ensure procedural guarantees to make redress more effective in cases involving AI or ADM systems, including amending non-discrimination law where necessary:** Introduce a shifting of the burden of proof in law by introducing a presumption of algorithmic bias that would help individuals who otherwise face near-impossible evidentiary hurdles. Where necessary, non-discrimination law should be amended to support this presumption. At the same time, it should be made easier for public interest organisations to initiate collective redress procedures in contexts such as employment, health-care and social security, as is common in consumer-focused regulation and complaint procedures.
- ▶ **Conduct and publish summaries of fundamental rights impact assessments:** Introduce a harmonised approach to fundamental rights impact assessments (FRIAs). These should be applied across public sector organisations and include the perspectives of affected groups and equality bodies, and the summaries of FRIAs should be published.

- ▶ **Include anti-discrimination in public procurement of AI and ADM systems:** Ensure that in public procurement of AI and ADM, anti-discrimination requirements are integrated to improve due diligence.
- ▶ **Raise public awareness of AI use and increase AI literacy:** Public awareness of AI use in public services should be increased through awareness campaigns and AI literacy efforts, including in education.
- ▶ **Guarantee availability of sufficient public resources and expertise for equality bodies and supervisory authorities on AI use:** These are indispensable for equality bodies and supervisory authorities to fulfil their mandate under the AI Act and anti-discrimination legislation.
- ▶ **Train public sector staff on fundamental rights impacts of AI and ADM:** Training must go beyond ethics to address the fundamental rights impacts of AI, with particular attention to discrimination risks.
- ▶ **Strengthen collaboration between equality bodies, other public institutions and civil society:** Reinforce and institutionalise collaboration, for example through “communities of practice” that lower barriers for affected individuals to seek support.

Summary – Use of AI in Finnish public administration and policy recommendations on algorithmic discrimination



The summary is based on the As-is and gap analysis report, in which the Finnish section was authored by Emeline Banzuzi. This is a confidential report reviewing how AI and ADM systems are used in the public sector in Finland, assessing AI-related legal and policy frameworks from equality and non-discrimination perspectives and providing recommendations to authorities and decision makers.

Established use of automated decision making, emerging use of artificial intelligence in public administration

In Finland, the use of automated decision-making (ADM) systems by public administrations is long established, and is deployed to efficiently tackle a vast number of administrative decisions, particularly in areas such as migration, policing, employment and recruitment. By contrast, the use of artificial intelligence (AI) remains limited, and is mostly limited to the use of generative AI (GenAI) for efficiency purposes, such as using Microsoft Copilot to support document drafting or note-taking. Several public authorities are nonetheless exploring or piloting AI applications, including the Finnish Immigration Service (Migri), which is considering AI use particularly for speech-to-text translation; the Finnish Police, which uses AI to read licence plates; and the City of Helsinki, which is piloting “Helsinki GPT”, an AI tool for automatic categorisation and prioritisation of customer feedback and generation of answers.

Legal and institutional framework

The EU AI Act together with European and Finnish anti-discrimination and data protection legislation form a legal framework for addressing algorithmic discrimination.

Finnish authorities are preparing to adapt their legal and institutional frameworks to the obligations introduced by the AI Act. There are two ongoing legislative proposals to implement the AI Act in Finland. The first seeks to, for example, regulate the oversight of certain AI systems and the enforcement of fines as well as designate in national legislation the competent authorities and a single point of contact as required by the AI Act. The second aims to regulate, for example, the establishment and operation of “AI sandboxes”. Other institutional preparatory work includes establishing co-ordination mechanisms and clarifying oversight roles.

In Finland, the use of ADM is regulated separately within its Administrative Act, which allows for the use of ADM only in circumstances where no case-specific judgment for the decision is required or where an official resolves any aspect that requires case-specific judgment. ADM must be based on predetermined rules by a natural person.

Non-discrimination and equality legislation in Finland is relatively comprehensive, as it reinforces the promotion of equality rather than focusing solely on preventing discrimination. The duty to promote equality requires public authorities to take active and anticipatory measures to ensure equal treatment, and this obligation also applies to the design, development and use of AI and ADM systems. In addition to EU *acquis*, further opportunities to prevent and address algorithmic discrimination are provided in the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, opened for signature in September 2024 and signed by the EU, in which Article 10 on equality and non-discrimination prohibits discrimination and requires states to “maintain measures aimed at overcoming inequalities to achieve fair, just and equitable outcomes”.

Notably, the only well-known judgment concerning discrimination and ADM in Finland is a case concerning a financial services company in which the National Non-Discrimination and Equality Tribunal of Finland ordered the company to cease using a statistical credit-scoring method that relied on prohibited grounds of discrimination, namely place of residence, age and gender.³¹

Gaps in legal and policy frameworks

Need for a standardised approach to assessing equality and non-discrimination impacts
The report finds that equality and non-discrimination impact assessments are not systematically conducted on AI or ADM systems, despite the legal obligation to promote equality and assess equality impacts applying to public authorities. However, a national assessment framework developed under the Finnish Government’s analysis, assessment and research activities³² provides an existing tool for assessing AI systems. A common methodology for equality and fundamental rights impact assessments (FRIAs) is needed to ensure that equality considerations are consistently integrated throughout the lifecycle of AI and ADM systems. The Non-Discrimination Ombudsman has urged the government to ensure that public authorities have an

31. Decision 216/2017 of the National Non-Discrimination and Equality Tribunal of Finland.

32. Ojanen A. et al. (2022), “Promoting equality in the use of Artificial Intelligence – An assessment framework for non-discriminatory AI”, Policy Brief 2022:25.

effective and systematically applied method for assessing the impacts on equality, including in public procurement.³³

Narrow awareness of discrimination risks

While public authorities have received training and developed ethical guidelines, they would benefit from additional resources, further targeted training, and a deeper understanding of algorithmic discrimination and non-discrimination laws as part of fundamental rights obligations.

Limited resources and lack of human rights-based technical expertise for equality bodies and oversight authorities

Equality bodies and other authorities responsible for fundamental rights supervision are expected to play an enhanced role under Article 77 AI Act, gaining powers to access documentation and request testing of high-risk systems. However, their ability to fulfil this expanded role may be constrained by limited human, technical and financial resources even as the use of AI systems is increasing, with public awareness and AI literacy still limited. Additionally, significant knowledge on fundamental rights protection, including the identification and assessment of different forms of discrimination, is required among all actors with supervisory tasks under the AI Act. Resources should also be increased to ensure an effective preventive mechanism to safeguard equality before discrimination risks materialise.

Lack of a public and harmonised registry of AI and ADM systems

There is no single public national database compiling information on the use of AI or ADM systems in Finnish public administrations. However, ADM use is shaped by the transparency obligations of the Administrative Act: authorities must publish general information on use of ADM on their websites and inform individuals whenever decisions concerning them have been made through ADM. Although agencies must publish information about such use on their own websites, these disclosures remain fragmented.

Fragmented oversight and institutional co-ordination

Oversight responsibilities are currently distributed across several authorities acting as national competent authorities, such as market surveillance authorities and fundamental rights authorities. While this division ensures coverage across sectors, it also risks inconsistent interpretations and enforcement gaps in ensuring fundamental rights protection, especially if co-operation mechanisms are not formalised and authorities are not provided with sufficient non-discrimination expertise.

33. "Non-Discrimination Ombudsman's recommendations to the governmental programme 2023-2027" (Finnish orig. *Yhdenvertaisuusvaltuutetun suosituksset hallitusohjelmaan 2023-2027*).

Key recommendations

- ▶ **Ensure adequate awareness and implementation of non-discrimination law, including promotional duties of equality:** This should be achieved by ensuring sufficient non-discrimination expertise, providing training, and improving awareness and understanding of these duties among competent authorities. Sufficient non-discrimination expertise and guidance should be provided also to those involved or overseeing procurement and grants related to the purchase and use of AI.
- ▶ **Introduce a standardised framework for equality and fundamental rights impact assessments in the lifecycle of AI systems:** Such assessments should be applied across public authorities, with the participation of equality bodies, civil society and affected groups. Summaries of assessments should be published to enhance accountability.
- ▶ **Increase resource allocations for equality bodies and supervisory authorities:** National competent authorities responsible for enforcing the AI Act should be provided with sufficient resources to effectively carry out their supervisory duties and to participate meaningfully in the national AI governance framework.
- ▶ **Establish a national public registry of AI and ADM systems used by public authorities, complementing the EU-wide database foreseen under the AI Act:** Such a registry should increase transparency, support oversight, and support equality bodies and the public to monitor risks of discrimination.
- ▶ **Standardise co-ordination and co-operation mechanisms among competent authorities, including on protection and promotion of equality:** This should include equipping competent authorities with sufficient equality and non-discrimination expertise to ensure coherent implementation of the AI Act. It is important to strengthen co-operation between the Non-discrimination Ombudsman and the Data Protection Ombudsman on algorithmic discrimination, as cases in this area may overlap with the right to non-discrimination and with data subject rights.

Summary – Use of AI in Portuguese public administration and policy recommendations on algorithmic discrimination



The summary is based on the As-is and gap analysis report, in which the Portuguese section was authored by Eduardo dos Santos. This is a confidential report reviewing how AI and ADM systems are used in the public sector in Portugal, assessing AI-related legal and policy frameworks from equality and non-discrimination perspectives and providing recommendations to authorities and decision makers.

Use of artificial intelligence in public administration

In Portugal, the use of artificial intelligence (AI) is nascent but gradually expanding through pilot initiatives and digitalisation efforts. For example, the national public procurement portal lists 119 AI-related contracts since 2009, with nearly half awarded in 2023 and 2024 alone.³⁴ AI is applied in sectors such as justice and healthcare, and several public services have launched AI chatbots for virtual assistance of services.

Legal and institutional framework

The EU AI Act together with European and Portuguese anti-discrimination and data protection legislation form a legal framework for addressing algorithmic discrimination. In addition to EU *acquis*, further opportunities to prevent and address algorithmic discrimination are provided in the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, opened for signature in September 2024, in which Article 10 on equality and non-discrimination prohibits discrimination and requires states to “maintain measures aimed at overcoming inequalities to achieve fair, just and equitable outcomes”.

34. Portal Base, available at: www.base.gov.pt, accessed 10 November 2025.

In Portugal, a notable court case involving algorithmic decision making occurred in 2021, when the public–private airline TAP used an algorithm to select up to 500 employees for dismissal based on parameters such as productivity, absenteeism, experience, contribution, cost and qualifications. Workers criticised the lack of transparency, as the weighting of each criterion was undisclosed. In 2023, the Lisbon Labour Court, in a case of 23 employees, ruled the collective dismissal unlawful.

The Portuguese Government has adopted several strategic instruments to foster digitalisation and responsible AI use, such as the National Artificial Intelligence Strategy³⁵ published in 2019 as part of the National Digital Competences Initiative (INCoDe.2030), and National Digital Strategy³⁶ to 2030. However, the As-is and gap analysis report notes that the proliferation of strategies, plans and agendas over the years may be counterproductive if not properly co-ordinated and integrated.

In 2024, Portugal established the Council for Digital in Public Administration with a dedicated Artificial Intelligence Technical Working Group to monitor the implementation of the AI Act. Another central role is played by the former Agency for Administrative Modernisation, restructured as the Agency for Technological Reform of the State in 2025, which ensures the responsible use of digital and AI systems in public administration. Public entities intending to acquire goods or services in this area, including AI systems, should inform the agency and may be subject to its assessment. The agency has developed a guide on AI in public administration and a risk assessment tool, covering certain ethical, equality and inclusivity concerns.

Gaps in legal and policy frameworks

Limited transparency and absence of a national AI and automated decision-making registry

There is no comprehensive public registry of the AI or automated decision-making (ADM) systems used in Portuguese public administration. Currently, no public mechanism allows citizens or oversight bodies to identify which AI or ADM systems are being used by public entities. This lack of transparency can make it difficult to assess the potential discrimination risks and harms of such systems.

Exclusion of fundamental rights and equality bodies from the AI Act governance framework

Portugal's list of appointed authorities under Article 77 AI Act includes a total of 14 entities: the general inspectorates for finance, defence, justice, internal administration and education, as well as sectoral regulators such as the National Communications Authority, Food and Economic Safety Authority, and Authority for Labour Conditions. No fundamental rights institutions, equality bodies or the Ombudsperson are included

35. National Artificial Intelligence Strategy, available at: www.incode2030.gov.pt/en/estrategia-nacional-de-inteligencia-artificial-en, accessed 10 November 2025.

36. Portugal Digital Strategy – The key to simplification, available at: <https://digital.gov.pt/documentos/portugal-digital-strategy>, accessed 10 November 2025.

as entities “protecting fundamental rights, including the right to non-discrimination” pursuant to Article 77. The absence of such bodies raises concerns regarding how potential cases of algorithmic discrimination will be addressed according to EU and national law and how individuals facing discrimination will be able to consistently seek redress under the AI Act’s governance framework.

Fragmented oversight and co-ordination mechanisms

Responsibilities for AI governance are dispersed among several authorities and no structured communication channels exist between these entities and institutions with human rights or equality mandates, such as the Ombudsperson or equality bodies such as the Commission for Citizenship and Gender Equality (CIG). The absence of a co-operation mechanism makes it difficult to identify and address potential overlaps or gaps in supervision, particularly regarding fundamental rights and non-discrimination obligations.

Limited awareness, resources and technical capacity

In Portugal, equality and human rights considerations are not yet systematically integrated into AI design or procurement processes and into the use of AI in public administration. Equality bodies and other competent institutions have limited ways to engage sufficient AI technical expertise and lack the resources required to engage in technical evaluations of AI systems to support individuals who may be affected by discrimination. The Commission for Citizenship and Gender Equality and other relevant bodies face resource constraints that limit their ability to address emerging challenges, as well as limited powers (i.e. lack of investigative and litigative mandates), which restrict their capacity to address issues of algorithmic discrimination.

Key recommendations

- ▶ **Establish a public registry of AI and ADM systems used in public administration:** Introducing an obligation for public authorities to report and document AI systems in a centralised registry would enhance transparency, support oversight by supervisory and equality bodies, and increase public trust.
- ▶ **Enhance the investigative and enforcement powers of equality bodies and include fundamental rights and equality bodies as national Article 77 bodies:** In the backdrop of the transposition of the EU directives on standards for equality bodies,³⁷ this would strengthen these institutions to enable proactive monitoring, engagement with AI governance, and effective enforcement of anti-discrimination obligations under both European and national law.
- ▶ **Until fundamental rights and equality bodies are appointed as Article 77 bodies, clarify institutional mandates and strengthen inter-agency co-ordination between authorities:** Defining clear responsibilities and formal co-operation mechanisms between equality, data protection and consumer

37. Directive EU 2024/1499 and Directive EU 2024/1500.

protection authorities would ensure coherent exchange of information and implementation of the AI Act.

- ▶ **Create a harmonised framework for assessing equality, gender equality and fundamental rights impacts, building on the national AI-related strategies, the Agency for Administrative Modernisation's guide and the AI Act:** This framework would develop a standardised national approach to guide public authorities in evaluating the risks of algorithmic discrimination and ensure consistent application of safeguards in the public sector.
- ▶ **Provide training and capacity-building programmes on fundamental rights for public officials and competent authorities under the AI Act:** Targeted training and capacity-building efforts should be implemented for authorities to identify and address algorithmic discrimination, including gender-based discrimination.
- ▶ **Incorporate anti-discrimination and gender equality principles more prominently into public procurement of AI and ADM systems:** This includes continuing the oversight of AI systems with regular audits while they are deployed and used.
- ▶ **Promote AI literacy and public awareness campaigns:** Awareness campaigns and educational materials to inform the public about how AI is used in public administration, its potential impacts and available redress mechanisms need to be developed. Strengthening AI literacy would enhance transparency and accountability in the use of these systems.
- ▶ **Ensure adequate funding and technical AI expertise for equality bodies:** Equality bodies need to be provided with sufficient funding and specialised technical AI expertise to address algorithmic discrimination.
- ▶ Explore how the government's ongoing initiative to develop a new generation of omnichannel, human-centred public services under the gov.pt platform would enable citizens to submit complaints and referrals related to AI-driven discrimination.

References

Access Now (2024), *Joint statement – A dangerous precedent: how the EU AI Act fails migrants and people on the move*, available at: www.accessnow.org/press-release/joint-statement-ai-act-fails-migrants-and-people-on-the-move, accessed 10 November 2025.

Algorithm Audit (2024a), *Preventing prejudice. Recommendations for risk profiling in the College Grant Control process: a quantitative and qualitative analysis*, available at: https://algorithmaudit.eu/algoprudence/cases/aa202401_preventing-prejudice/, accessed 10 November 2025.

Algorithm Audit (2024b), *DUO control process biased towards students with a non-European migration background*, available at: https://algorithmaudit.eu/events/press_room/#DUO_CBS, accessed 15 November 2025.

Algorithm Watch (2018), *Finnish credit score ruling raises questions about discrimination and how to avoid it*, available at: [Finnish Credit Score Ruling raises Questions about Discrimination and how to avoid it - AlgorithmWatch](#), accessed 15 November 2025.

Allhutter D. et al. (2020), "Algorithmic profiling of job seekers in Austria: how austerity politics are made effective", *Frontiers in Big Data*, Vol. 3, No. 5.

Amnesty International (2021), *Dutch childcare benefit scandal an urgent wake-up call to ban racist algorithms*, available at: www.amnesty.org/en/latest/news/2021/10/xenophobic-machines-dutch-child-benefit-scandal, accessed 10 November 2025.

Conseil d'État (2024), *Requête introductive d'instance*, 15 October.

Council of Europe (2007), *Convention 201 on the Protection of Children against Sexual Exploitation and Sexual Abuse* (Lanzarote, 25 October 2007), available at: <https://rm.coe.int/1680084822>, accessed 10 November 2025.

Council of Europe (2010), *Recommendation CM/Rec(2010)5 of the Committee of Ministers on measures to combat discrimination on grounds of sexual orientation or gender identity*, available at: <https://search.coe.int/cm?i=09000016805cf40a>, accessed 10 November 2025.

Council of Europe (2013), *Declaration by the Committee of Ministers on risks to fundamental rights stemming from digital tracking and other surveillance technologies*, available at: <https://search.coe.int/cm?i=09000016805c8011>, accessed 10 November 2025.

Council of Europe (2018), *Convention 108+ for the protection of individuals with regard to the processing of personal data*, available at: <https://rm.coe.int/convention-108-convention-for-the-protection-of-individuals-with-regar/16808b36f1>, accessed 10 November 2025.

Council of Europe (2019a), *Declaration by the Committee of Ministers on the manipulative capabilities of algorithmic processes*, Decl(13/02/2019)1, available at: <https://search.coe.int/cm?i=090000168092dd4b>, accessed 10 November 2025.

Council of Europe (2019b), Recommendation CM/Rec(2019)1 of the Committee of Ministers to member States on preventing and combating sexism, available at: <https://rm.coe.int/168093b26a>, accessed 10 November 2025.

Council of Europe (2019c), “Guidelines on artificial intelligence and data protection”, Consultative Committee of the Convention for the protection of individuals with regard to automatic processing of personal data (Convention 108), available at: <https://rm.coe.int/guidelines-on-artificial-intelligence-and-data-protection/168091f9d8>, accessed 10 November 2025.

Council of Europe (2020), Recommendation CM/Rec(2020)1 of the Committee of Ministers to member States on the human rights impacts of algorithmic systems, available at: <https://search.coe.int/cm?i=09000016809e1154>, accessed 10 November 2025.

Council of Europe (2021a), *General Recommendation No. 1 on the digital dimension of violence against women*, Group of Experts on Action against Violence against Women and Domestic Violence (GREVIO), available at: <https://edoc.coe.int/fr/violence-l-gard-des-femmes/10643-grevio-general-recommendation-no-1-on-the-digital-dimension-of-violence-against-women.html>, accessed 10 November 2025.

Council of Europe (2021b), “Content moderation: best practices towards effective legal and procedural frameworks for self-regulatory and co-regulatory mechanisms of content moderation”, Guidance Note, Steering Committee for Media and Information Society (CDMSI), available at: <https://rm.coe.int/content-moderation-en/1680a2cc18>, accessed 10 November 2025.

Council of Europe (2021c), “Guidelines on upholding equality and protecting against discrimination and hate during the Covid-19 pandemic and similar crises in the future”, Committee of Ministers, CM(2021)37-add1rev, available at: <https://edoc.coe.int/en/living-together-diversity-and-freedom-in-europe/9745-guidelines-of-the-committee-of-ministers-of-the-council-of-europe-on-upholding-equality-and-protecting-against-discrimination-and-hate-during-the-covid-19-pandemic-and-similar-crises-in-the-future.html>, accessed 10 November 2025.

Council of Europe (2021d), “Guidelines on facial recognition”, Consultative Committee of the Convention for the protection of individuals with regard to automatic processing of personal data (Convention 108), available at: <https://edoc.coe.int/en/artificial-intelligence/9753-guidelines-on-facial-recognition.html>, accessed 10 November 2025.

Council of Europe (2022a), “Outline of Huderia risk and impact assessment methodology”, Committee on Artificial Intelligence, available at: <https://rm.coe.int/cai-bu-2022-03-outline-of-huderia-risk-and-impact-assessment-methodolo/1680a81e14>, accessed 10 November 2025.

Council of Europe (2022b), Recommendation CM/Rec(2022)13 of the Committee of Ministers to member States on the impacts of digital technologies on freedom of expression, available at: <https://search.coe.int/cm?i=0900001680a61729>, accessed 10 November 2025.

Council of Europe (2022c), Recommendation CM/Rec(2022)17 of the Committee of Ministers to member States on protecting the rights of migrant, refugee and asylum-seeking women and girls, available at: <https://rm.coe.int/>

[prems-092222-gbr-2573-recommandation-cm-rec-2022-17-a5-bat-web-1-1680a6ef9a#:~:text=Member%20States%20should%20ensure%20that,the%20same%20conditions%20as%20nationals](https://rm.coe.int/prems-092222-gbr-2573-recommandation-cm-rec-2022-17-a5-bat-web-1-1680a6ef9a#:~:text=Member%20States%20should%20ensure%20that,the%20same%20conditions%20as%20nationals), accessed 10 November 2025.

Council of Europe (2022d), Recommendation CM/Rec(2022)16 of the Committee of Ministers on combating hate speech, available at: <https://search.coe.int/cm?i=0900001680a67955>, accessed 10 November 2025.

Council of Europe, Recommendation CM/Rec(2024)4 of the Committee of Ministers on combating hate crime, available at: <https://search.coe.int/cm?i=0900001680af9736>, accessed 10 November 2025.

Council of Europe (2024a), Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law, Vilnius.

Council of Europe (2024b), “Methodology for the risk and impact assessment of artificial intelligence systems from the point of view of human rights, democracy and the rule of law (Huderia Methodology)”, Committee on Artificial Intelligence, available at: <https://rm.coe.int/cai-2024-16rev2-methodology-for-the-risk-and-impact-assessment-of-arti/1680b2a09f>, accessed 10 November 2025.

Council of Europe (2024c), Gender Equality Strategy 2024-2029, available at: <https://rm.coe.int/prems-073024-gbr-2573-gender-equality-strategy-2024-29-txt-web-a5-2756/1680afc66a>, accessed 10 November 2025.

Council of Europe (2024d), Recommendation CM/Rec(2024)1 of the Committee of Ministers to member States on equality of Roma and Traveller women and girls, available at: <https://search.coe.int/cm?i=0900001680af27e4>, accessed 10 November 2025.

Council of Europe (forthcoming), Draft Recommendation of the Committee of Ministers to member States on the impact of artificial intelligence systems, their potential for promoting equality – including gender equality – and the risks they may cause in relation to non-discrimination, Committee of Experts on Artificial Intelligence, Equality and Discrimination (GEC/ADI-AI).

Council of Europe, Commissioner for Human Rights (2019), “Unboxing AI: 10 steps to protect human rights”, available at: <https://book.coe.int/fr/commissaire-aux-droits-de-l-homme/9664-unboxing-artificial-intelligence-10-steps-to-protect-human-rights.html>, accessed 10 November 2025.

Council of Europe, Commissioner for Human Rights (2023), “Human rights by design: future-proofing human rights protection in the era of AI”, Follow-up Recommendation, available at: <https://book.coe.int/fr/commissaire-aux-droits-de-l-homme/11759-pdf-human-rights-by-design-future-proofing-human-rights-protection-in-the-era-of-ai.html>, accessed 10 November 2025.

Council of Europe, Parliamentary Assembly (2017), Recommendation 2102 (2017) on “Technological convergence, artificial intelligence and human rights”, available at: <https://pace.coe.int/en/files/23726/html>, accessed 10 November 2025.

Council of Europe, Parliamentary Assembly (2020a), Resolution 2343 on “Preventing discrimination caused by the use of artificial intelligence”, available at: <https://pace.coe.int/en/files/28807/html>, accessed 10 November 2025.

Council of Europe, Parliamentary Assembly (2020b), Recommendation 2183 (2020) on “Preventing discrimination caused by the use of artificial intelligence”, available at: <https://pace.coe.int/en/files/28809/html>, accessed 10 November 2025.

Court of Justice of the European Union (2008), C-54/07 *Centrum voor gelijkheid van kansen en voor racismebestrijding contre Firma Feryn NV* EU:C:2008:397.

Court of Justice of the European Union (2022), C-817/19 *Ligue des droits humains ASBL v. Conseil des ministres* EU:C:2022:491.

Court of Justice of the European Union (2023), C-634/21 *SCHUFA Holding* (Scoring) EU:C:2023:957.

Court of Justice of the European Union (2024), C-588/21 P *Public.Resource.Org and Right to Know v. Commission and Others* EU:C:2024:201.

Court of Justice of the European Union (2025), C-203/22 *Dun & Bradstreet Austria* EU:C:2025:117.

Défenseur des Droits (2024), “Algorithms, AI systems and public services: what rights do users have?”, available at: www.defenseurdesdroits.fr/our-last-reports-and-studies-316, accessed 10 November 2025.

Dubois V. (2021), *Contrôler les assistés : Genèses et usages d’un mot d’ordre*, Raisons d’agir, Paris.

European Commission (2020), “White Paper: On artificial intelligence – A European approach to excellence and trust”, COM(2020)65 final.

European Commission (2022a), “Proposal for a Directive of the European Parliament and of the Council on adapting non-contractual civil liability rules to artificial intelligence (AI liability directive)”, COM/2022/496 final.

European Commission (2022b), “European Declaration on Digital Rights and Principles for the Digital Decade”, COM(2022) 28 final.

European Commission (2025a), “Commission Guidelines on prohibited artificial intelligence practices established by Regulation (EU) 2024/1689 (AI Act)”, available at: <https://ec.europa.eu/newsroom/dae/redirection/document/112367>, accessed 10 November 2025.

European Commission (2025b), *The General-Purpose AI Code of Practice*, available at: <https://digital-strategy.ec.europa.eu/en/policies/contents-code-gpai>, accessed 10 November 2025.

European Commission (2025c), *AI Board*, <https://digital-strategy.ec.europa.eu/en/policies/ai-board>, accessed 10 November 2025.

European Commission (2025d), “Commission Guidelines on the definition of an artificial intelligence system established by Regulation (EU) 2024/1689 (AI Act)”, C(2025) 924 final.

European Commission against Racism and Intolerance (2000), “ECRI General Policy Recommendation No. 6 on combating the dissemination of racist, xenophobic and antisemitic material via the internet”, available at: <https://>

rm.coe.int/ecri-general-policy-recommendation-no-6-on-combating-the-dissemination/16808b5a8d, accessed 10 November 2025.

European Commission against Racism and Intolerance (2004), “ECRI General Policy Recommendation No. 8 on combating racism while fighting terrorism”, available at: <https://rm.coe.int/ecri-general-policy-recommendation-no-8-on-combating-racism-while-figh/16808b5abc>, accessed 10 November 2025.

European Commission against Racism and Intolerance (2006), “ECRI General Policy Recommendation No. 10 on combating racism and racial discrimination in and through school education”, available at: <https://rm.coe.int/ecri-general-policy-recommendation-no-10-on-combating-racism-and-racia/16808b5ad5>, accessed 10 November 2025.

European Commission against Racism and Intolerance (2007), “ECRI General Policy Recommendation No. 11 on combating racism and racial discrimination in policing”, available at: <https://rm.coe.int/ecri-general-policy-recommendation-no-11-on-combating-racism-and-racia/16808b5adf>, accessed 10 November 2025.

European Commission against Racism and Intolerance (2012), “ECRI General Policy Recommendation No. 14 on combating racism and racial discrimination in employment”, available at: <https://rm.coe.int/ecri-general-policy-recommendation-no-14-on-combating-racism-and-racia/16808b5afc>, accessed 10 November 2025.

European Commission against Racism and Intolerance (2015), “ECRI General Policy Recommendation No. 15 on combating hate speech”, available at: <https://rm.coe.int/ecri-general-policy-recommendation-no-15-on-combating-hate-speech/16808b5b01>, accessed 10 November 2025.

European Commission against Racism and Intolerance (2017), “ECRI General Policy Recommendation No. 7 (revised) on national legislation to combat racism and racial discrimination”, available at: www.coe.int/en/web/european-commission-against-racism-and-intolerance/recommendation-no.7, accessed 10 November 2025.

European Court of Human Rights (2007), Application No. 57325/00, *D.H. and Others v. the Czech Republic* (Grand Chamber, 13 November 2007), paragraph 178.

European Court of Human Rights (2023), Application 11519/20, *Glukhin v. Russia* (Third Section, 2023).

European Digital Rights (2024), *How to fight biometric mass surveillance after the AI Act: a legal and practical guide*, available at: <https://edri.org/our-work/how-to-fight-biometric-mass-surveillance-after-the-ai-act-a-legal-and-practical-guide>, accessed 10 November 2025.

European Migration Network and Organisation for Economic Co-operation and Development (2022), “The use of digitalisation and artificial intelligence in migration management”, available at: www.oecd.org/en/topics/sub-issues/migration-policies-returns-and-attracting-talent/migration-policy-debates-and-data-briefs.html, accessed 10 November 2025.

European Network of National Human Rights Institutions (2024), “Technologies, migration, and human rights: the role of European NHRIs”, ENNHRI scoping paper,

available at: <https://ennhri.org/publications-statements/#artificial-intelligence-publications>, accessed 10 November 2025.

European Union (2016a), Regulation 2016/679 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC.

European Union (2016b), Directive 2016/680 on the protection of natural persons with regard to the processing of personal data by competent authorities for the purposes of the prevention, investigation, detection or prosecution of criminal offences or the execution of criminal penalties, and on the free movement of such data.

European Union (2022), Regulation (EU) 2022/2065 of the European Parliament and of the Council of 19 October 2022 on a Single Market For Digital Services.

European Union (2024a), Council Directive 2024/1499 on standards for equality bodies in the field of equal treatment between persons irrespective of their racial or ethnic origin, equal treatment in matters of employment and occupation between persons irrespective of their religion or belief, disability, age or sexual orientation, equal treatment between women and men in matters of social security and in the access to and supply of goods and services.

European Union (2024b), Council Directive 2024/1500 on standards for equality bodies in the field of equal treatment and equal opportunities between women and men in matters of employment and occupation.

European Union (2024c), Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act), available at: https://eur-lex.europa.eu/legal-content/EN/TXT/HTML/?uri=OJ:L_202401689, accessed 10 November 2025.

Lanzarote Committee (2024), Declaration on protecting children against sexual exploitation and sexual abuse facilitated by emerging technologies, available at: <https://rm.coe.int/declaration-on-protecting-children-against-sexual-exploitation-and-sex/1680b25a78>, accessed 10 November 2025.

La Quadrature du Net (2023), *Notation des allocataires : l'indécence des pratiques de la CAF désormais indéniable*, available at: www.laquadrature.net/2023/11/27/notation-des-allocataires-lindcence-des-pratiques-de-la-caf-desormais-indeniable, accessed 10 November 2025.

La Quadrature du Net (2024), *À France Travail, l'essor du contrôle algorithmique*, 25 June, available at: www.laquadrature.net/2024/06/25/a-france-travail-lessor-du-controle-algorithmique, accessed 10 November 2025.

Malgieri G. and Comandé G. (2017), "Why a right to legibility of automated decision-making exists in the General Data Protection Regulation", *International Data Privacy Law*, Vol. 7, Issue 4, p. 243-65.

McGregor L. and Molnar P. (2023), "Digital border governance: a human rights based approach", University of Essex and the United Nations Human Rights Office of the High Commissioner.

Milano S. and Prunkl C. (2024), "Algorithmic profiling as a source of hermeneutical injustice", *Philosophical Studies*, Vol. 182, pp. 185-203.

Mittelstadt B. (2025), "How to use the Artificial Intelligence Act to investigate AI bias and discrimination: a guide for Equality Bodies", European Network of Equality Bodies (Equinet), Brussels.

National Non-Discrimination and Equality Tribunal of Finland (2017), Decision 216/2017.

Selbts A. D. and Powles J. (2017) "Meaningful information and the right to explanation", *International Data Privacy Law*, Vol. 7, No. 4, pp. 233-42.

UN Special Rapporteur (2021), *Racial and Xenophobic discrimination and the use of digital technologies in border and immigration enforcement – Report of the Special Rapporteur on contemporary forms of racism, racial discrimination, xenophobia and related intolerance*, Achiume E. Tendayi, UN Special Rapporteur on contemporary forms of racism, racial discrimination, xenophobia and related intolerance, A/HRC/48/76.

Wachter S., Mittelstadt B. and Floridi L. (2017), "Why a right to explanation of automated decision-making does not exist in the General Data Protection Regulation", *International Data Privacy Law*, Vol. 7, Issue 2, pp. 76-99.

Xenidis R. (2021), "Tuning EU equality law to algorithmic discrimination: three pathways to resilience", *Maastricht Journal of European and Comparative Law*, Vol. 27, Issue 6, pp. 736-58.

The growing integration of artificial intelligence (AI) in diverse applications across a broad range of sectors presents significant challenges to the protection of fundamental rights. Among these, algorithmic discrimination emerges as a particularly pressing concern. Empirical research has demonstrated that algorithmic bias not only reflects but also exacerbates existing social inequalities on a large scale, particularly in domains where algorithmic decision making is prevalent – such as policing, employment, education and insurance. In 2024, both the European Union and the Council of Europe adopted landmark legal instruments: the EU AI Act (Regulation 2024/1689) and the Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law. This report examines how these emerging legal instruments contribute to strengthening protections against algorithmic discrimination in Europe and assesses the lacunae that continue to affect these legal frameworks. The report also aims to trace and synthesise these latest developments to equip relevant players, such as equality bodies, with tools to adapt to a changing legal landscape, and where relevant, to intervene in ongoing evolutions.



YHDENVERTAISUUSVALTUUTETTU
NON-DISCRIMINATION OMBUDSMAN



COMISSÃO PARA A CIDADANIA
E A IGUALDADE DE GÉNERO

PREMS 163825

ENG

The Member States of the European Union have decided to link together their know-how, resources and destinies. Together, they have built a zone of stability, democracy and sustainable development whilst maintaining cultural diversity, tolerance and individual freedoms. The European Union is committed to sharing its achievements and its values with countries and peoples beyond its borders.

www.europa.eu

The Council of Europe is the continent's leading human rights organisation. It comprises 46 member states, including all members of the European Union. All Council of Europe member states have signed up to the European Convention on Human Rights, a treaty designed to protect human rights, democracy and the rule of law. The European Court of Human Rights oversees the implementation of the Convention in the member states.

www.coe.int



EUROPEAN UNION

COUNCIL OF EUROPE



CONSEIL DE L'EUROPE