



European Center for
Not-for-Profit Law

ECNL Position Statement on the EU AI Act

23 July 2021



ECNL Position Statement on the EU AI Act

Contents

APPROACH OF THE AI ACT	3
RISK LEVELS	6
1. Prohibitions (art. 5)	6
2. High-risk systems (art. 6-7)	8
3. Low-risk systems ("certain AI systems") – art. 52	11
4. No-risk systems (all other systems)	12
5. Governance mechanisms	12
THE CORPORATE RESPONSIBILITY TO RESPECT HUMAN RIGHTS	14
STAKEHOLDER ENGAGEMENT & CIVIL SOCIETY PARTICIPATION	15
HUMAN RIGHTS IMPLICATIONS OF MARGINALIZED AND AT-RISK GROUPS	17
BIBLIOGRAPHY	19

Acknowledgement: This paper was drafted by ECNL's Marlena Wisniak, with support from Vanja Skoric and Francesca Fanucci. We are grateful for and recognize the contributions of our colleagues in this field. We would especially like to thank Sarah Chander, Richard Wingfield, Imogen Parker, Lilian Edwards, Angela Mueller, and Iverna McGowan who have provided invaluable feedback on very short notice. We look forward to continuing our thinking and advocacy related to the EU AI Act going forward with more representatives from civil society, academia, and affected communities.



APPROACH OF THE AI ACT

The European Center for Not-for-Profit Law (ECNL) strongly supports rights-based regulation of artificial intelligence (AI) systems and welcomes the European Commission's initiative to draft a proposal for an EU-wide AI Act. We are also pleased that the legislative process has provided external stakeholders, especially civil society organisations and affected groups, opportunities to contribute to the development of the AI Act. We encourage the European Commission – and going forward the European Parliament – to expand and strengthen the participatory process.

That said, ECNL is deeply concerned about the current approach of the AI Act. In its current state, the AI Act misses an opportunity to effectively protect the rights of persons and communities being subjected to AI systems, placing business and operational interests as well as harmonization of the internal market for AI products above people's fundamental rights. This is despite the Act's specific objective of ensuring that such AI systems respect existing law on fundamental rights and Union values. Indeed, while a superficial reading of the Act suggests that it aims to protect European values and fundamental rights (e.g. 3.5), a deeper analysis exposes that the underlying motives are instead technological innovation, economic development, national security and counter-terrorism, border control, and criminal justice.

Recognising that AI systems can adversely impact – and at times be incompatible with – fundamental rights, the AI Act nonetheless stresses that the development of AI is necessary for sectors as diverse as environmental, health, finance, mobility, home affairs, and culture, both in the private and public sectors (para. (1.1)). This narrative suggests that the Act is promoting an uptake of AI systems. These systems can indeed have beneficial impacts, but ECNL is concerned that as the Act does not sufficiently address the severe power imbalance that exists between those who develop and deploy AI systems, and the communities that are subjected to them. This imbalance is especially acute for historically marginalized and under-represented groups, such as racialized groups, women and gender non-conforming persons, religious minorities, LGBTQIA+, disabled persons, migrants and refugees, children and the elderly, and persons of lower socio-economic status, among others. When considering potential opportunities that can arise from AI systems, it is therefore important to begin with a power analysis and centre the needs of the most at-risk communities. With this in mind, it is crucial to consider the following elements:



This work is licensed under a [Creative Commons Attribution-ShareAlike 4.0 International License](https://creativecommons.org/licenses/by-sa/4.0/).

- I. Who will benefit from these systems (specifically, which demographic groups and/or sectors) and who will be harmed?
- II. Is the root cause of a (social, economic, political or other) issue effectively being addressed by deploying the AI system, or are we merely offering performative and superficial solutions?

In reality, there are no systems that only present opportunities or risks from a strictly binary perspective, but instead there are systems that provide different opportunities or risks depending on the targeted population, context, and situation in which they're deployed. In its current form, the AI Act falls short of addressing this concern.

The overarching approach of the regulation does not adequately protect human rights and fundamental freedoms, especially since only a few AI systems are subject to (inadequate) legal requirements, while the vast majority of AI systems are under no impact assessment nor regulation at all. At its core, the AI Act is rooted in EU free market policy, preventing market fragmentation and ensuring product safety regulation. Promoting neoliberal and industry-friendly narratives, the AI Act actually supports the acceleration of AI systems by preventing Member-States from further regulating them. As noted by scholars Veale and Zuiderveen Borgesius, the material scope of the AI Act “appears to rule out the possibility that [it] is a general ‘minimum harmonisation’ instrument, setting a horizontal regulatory floor.”[1] They caution that “the Act may contribute to deregulation more than it raises the regulatory bar.”[1] This is especially worrisome, given that the regulation has the potential to set legal and policy standards on a global level, promoting “trustworthy AI” instead of rights-respecting AI and therefore focusing on perception rather than on effective rights protection.

ECNL strongly supports establishing minimum transparency and human rights safeguards for all AI systems, irrespective of their level of risk taking into consideration parallel requirements under the GDPR and other relevant laws (e.g. Digital Services Act, Digital Markets Act, European Convention on Human Rights, EU Charter of Fundamental Rights, etc.). Once such level is properly assessed, higher safeguards should be applied whenever the risk is high or higher, also depending on the context and area of application. Importantly, obligations should fall on both AI providers and users.



This work is licensed under a [Creative Commons Attribution-ShareAlike 4.0 International License](https://creativecommons.org/licenses/by-sa/4.0/).

- I. The scope and list of prohibited AI systems is too narrow and fails to include other AI systems that are incompatible with human rights.
- II. The European Commission's proposal for the AI Act proposes very limited oversight and safeguards to AI systems determined as "high risk", while merely establishing weak transparency requirements for "certain AI systems." As a result, all AI systems that do not fall under the "high risk" category remain effectively unregulated. Yet, the level of risk often depends on the application of the AI system, and differs greatly for specific demographic and geographic groups subjected to the technology as well as depending on the context of its application. Finally, it unduly puts the burden of proof on the person subjected to the AI system, as opposed to those developing or deploying it. This would create an uneven playing field and incentivize classifying AI systems as low risk to avoid further regulation. It also runs against an established principle of EU non-discrimination law that the victim is not expected to prove the wrong (see for e.g. EU Directive for equal treatment in employment and occupation[2]).

There are compelling arguments that a thorough, inclusive and transparent human rights impact assessment (HRIA) must be the starting point for all subsequent regulatory actions of any AI system, which is in line with the EU's proclaimed risk-based approach. In addition to HRIAs, ECNL encourages taking a sector-specific approach when it comes to further regulatory requirements, with an emphasis on potential users and developers. This was reflected in the recent online consultation conducted by the Council of Europe Ad-Hoc Committee on Artificial Intelligence (CAHAI). The vast majority of different stakeholders responded with overwhelming support for regulating all AI systems, irrespective of their risk level, and a strong preference for a human rights-based approach, with HRIAs being the preferred choice of governance mechanism.[3]

ECNL is highly concerned that the *obligations* generally pertain to AI providers only, failing to consider those pertaining to AI users. ECNL urges policymakers to develop parallel obligations on users of AI systems, given that they are best positioned to understand the context in which the systems are deployed, and importantly, the impacts that the use will have on affected communities. ECNL strongly recommends that AI users conduct human rights due diligence, including human rights impact assessments, before deploying the AI systems, and continuously thereafter. Importantly, this should be done in close consultation with affected groups, especially marginalized and at-risk ones.



Similarly, the *rights of redress* under the AI Act generally pertain to AI providers only. The rights of affected communities and persons subjected to AI systems are not protected, and these stakeholders have no access to remedy in case of harm. This is exemplified by the fact that transparency requirements regulate the relationships between AI providers and users, as opposed to any direct responsibility towards people subjected to the systems or affected by them (article 13). ECNL strongly recommends that an effective right to redress for affected groups be added to the AI Act, with meaningful support (including adequate resources) to stakeholders so that they can fully exercise this right. Similarly, requirements apply to providers only, and not to the deployers or users who merely need to get instructions from providers (article 14.4).

RISK LEVELS

1. Prohibitions (art. 5)

The EU AI Act includes a few important and welcome prohibitions of AI practices whose use is considered unacceptable, since they contravene Union values and violates fundamental rights. However, the objective of the prohibitions is marred by their narrow scope and broad exceptions and derogations. ECNL recommends (1) expanding the list of prohibited AI practices in line with the European Data Protection Board and the European Data Protection Supervisor's demands; (2) removing the condition to prove "physical or psychological harm"; and (3) narrowing down the scope of exceptions.

- I. The definition of remote biometric identification in publicly accessible spaces is overly narrow and the standards are difficult to meet. Article 5(1)d) proposes a prohibition for a few specific uses of biometric technologies when deployed by law enforcement. One notable example of a prohibited application, which we fully support, is using real-time facial recognition against people protesting. Unfortunately, the prohibition only applies to "real-time" uses, and does not cover other harmful use cases due to "post" - remote biometric identification. This leaves out other dangerous systems that should be prohibited but are instead allowed to be placed on the market. In alignment with many other civil society organisations,[4] the European Data Protection Board and the European Data Protection Supervisor, ECNL calls for the



following prohibitions: (1) “a general prohibition on any use of AI for automated recognition of human features in publicly accessible spaces, such as recognition of faces, gait, fingerprints, DNA, voice, keystrokes and other biometric or behavioural signals, in any context; [(2)] a prohibition on AI systems using biometrics to categorise individuals into clusters based on ethnicity, gender, political or sexual orientation” among others”; (3) a prohibition of emotion recognition systems and the use of AI for social scoring[5] and (4) a prohibition of risk assessment tools for criminal justice and asylum.

- II. Furthermore, the threshold that some AI systems or applications must meet to fall under the scope of the prohibition makes it difficult, if not practically impossible, to achieve. Specifically, the requirement of individual “physical or psychological harm” is very difficult to prove in the case of AI systems, where collective and indirect harms are ubiquitous. What’s more, the definition of public spaces includes spaces (both private and public) that are accessible to the public, excluding online spaces (articles 5.2, 5.3, 5.4). ECNL is concerned about the risks for activists and civil society organisations operating in digital spaces, and the human rights implications of using AI-driven biometrics systems to identify people whose faces appear on the internet.
- III. The broad exceptions to the prohibition of remote biometric identification at best fail to capture all the risks of these technologies, and at worst legitimise their development and use (articles 5.1d, 5.2, 5.3). ECNL agrees with the European Commission that exceptions should be “exhaustively listed and narrowly defined situations” and for “substantial public interest”, but believes that the current proposal falls short of such goals. The AI Act authorises the use of facial recognition systems for the search of missing children, to prevent terrorism or to predict crime under a few conditions (notably, the need to get a judicial authorization, which we support.) Such an approach is highly dangerous. The exception for missing children perpetuates a ‘techno-solutionist’ approach in the absence of any publicly available information that facial recognition is indeed effective in this case. Even more worrisome is the fact that remote biometric identification, as authorized under the AI Act’s exceptions, generally aims to identify, surveil, and possibly detain individuals under the justification of security. Racialized persons, political dissidents, and activists are often disproportionately targeted and are at risk of great physical harm, from arbitrary detention to potentially torture or



death. At the same, by requiring that groups such as migrants or refugees consent to AI-driven biometric systems for processing their request or accessing essential services, they risk excluding people from accessing life-saving opportunities in the context of justice or border control, should they not consent.

- IV. Importantly, these exceptions risk exacerbating existing racial and social inequity. AI-driven surveillance technologies in the hands of powerful actors such as judicial bodies or law enforcement officials have the potential to do great harm, with minorities and racialized groups, human rights defenders, activists and journalists bearing the most significant risk. These risks are heightened by the fact that the systems can only be deployed in specific areas, for example where there is an “indication of threat of presence of an alleged perpetrator.” While this limitation is good at first sight, it will likely lead to deploying biometric identification systems in areas that already over-policed, and where the residents are predominantly poor, migrants, and persons of colour. Other exceptions include derogations for international and bilateral agreements (article 4), for example in the context of national security or counter-terrorism. Ultimately, these exceptions enable the uptake of these systems without any publicly available evidence that they’re effective to combat crime or terrorism, at the expense of fundamental rights and open civic space.

2. High-risk systems (art. 6-7)

Grounded in product safety legislation, the AI Act classifies AI systems on the basis of their *intended* purpose (article 5.2.3). While ECNL agrees that this is an important criterion, it is only one of many relevant factors. Importantly, it excludes unintended (or hidden) purposes, such as collateral harm and misuse or abuse. In practice, it will be difficult to determine the true and underlying purpose of an AI system, given the incentives of providers to frame their intent in a way that limits risk.

- I. We believe that criteria for determining the risk level should include, at a minimum, those related to product design (including intent); severity of impact; due diligence mechanisms; causal link; and potential for remedy.[6] The context in which AI systems are deployed is also critical, as areas such as law enforcement,



migration and border control, and access to justice should de-facto be considered high-risk.

- II. We are deeply concerned that there is currently no provision nor clearly identified procedure allowing for adding new categories to annex III related to the list of high-risk uses of AI systems. This should be amended to permit future additions and follow a thorough human rights impact assessment, where affected communities and civil society are included in the revision process.
- III. We generally support the added requirements for AI providers of high-risk systems, although we do not think that they're robust enough to effectively protect human rights. In any case, we urge the European Parliament to expand these requirements to all AI systems proportionate to their risk level. We welcome the establishment of public registers (article 60) and recommend expanding them to include information about who is deploying them and for what purpose, in line with other civil society organisations.[7] We also would like to see support for civil society organisations and affected groups to access and understand these databases.
- IV. Regarding the requirements for high data quality and other testing, validation, and accuracy standards (inter alia articles 10 and 44), we are disappointed that these apply to high-risk systems only. Moreover, ECNL recommends to include information about who will determine the level of acceptable accuracy, robustness and cybersecurity, who determines evaluation metrics, and what role historically at-risk, marginalized and affected communities will play in establishing these norms. We are also concerned that these standards will promote a narrative that de-biasing or improving the quality of data prevent risks of discrimination, when ultimately the problem lies in how the use of these systems can exacerbate existing discrimination and inequality.
- V. Placement of high-risk systems on the market is subject to conformity assessments. Given that these are self-assessments, there is a high risk of conflict of interest and incitement to lower standards or to easily approve the process. To ensure public interest, ECNL recommends that an external audit of certification be required as conforming with the goals of the regulation, instead of the developers of the AI system themselves. Such conformity assessments should be conducted on an ongoing basis anytime that they are deployed in a



new high-risk (geographic, social or political) context or application, in addition to when systems are substantially modified as required by article 43.4.

- VI. ECNL strongly recommends removing the derogation from the conformity procedure for “exceptional reasons of public security” (article 47). This derogation is especially alarming, as the use of technology for counter-terrorism or national security purposes is already notoriously opaque and under-regulated, and the risks of human rights abuses and adverse impacts are elevated in such contexts. If anything, we recommend conducting *enhanced* assessments (as consistent with the United Nations Guiding Principles for Business and Human Rights), but certainly not a derogation thereof.
- VII. Conformity assessments do not follow a rights-based and community-driven methodology. Human rights impact assessments (HRIAs) do, which is why ECNL believes they are the type of assessment that best prevents adverse impacts on fundamental rights. Following strict and high standards outlined in the United Nations Guiding Principles for Business and Human Rights (UNGPs), HRIAs are based on stakeholder engagement and transparency principles. Accordingly, the conformity assessment should be based on input from affected communities and stakeholder groups, including civil society. This should begin – and where necessary end – with the questions: (i) is the purpose of the technology a legitimate one?; (ii) if so, is the technology effective in achieving that purpose?; and (iii) even if it is effective, is it proportionate, i.e., is there no other less intrusive way to achieve the same result? The results of the conformity assessment should be made available and accessible. Importantly, users who deploy AI systems should also be subject to assessment requirements, given the wide scope of possible applications of AI systems and corresponding human rights risks. HRIAs are best suited for this purpose, where understanding impacts within specific contexts is critical.
- VIII. The AI Act has limited notification requirements to national authorities for serious incidents or malfunctioning of high-risk systems (article 29.4). Affected individuals or groups have the right to information under international human rights law, but the present draft law inhibits their ability to realise this right, as it does not allow for them to contact national authorities to report any incidents that impact them. This leads to an over-reliance on the good faith of AI providers, who have a strong conflict of interest and are



incentivized to under-report. ECNL recommends including an obligation to make these reports publicly available.

- IX. Most regrettably, articles 16 to 29, which outline obligations for high-risk systems to AI providers and users, do not include any requirements to consult with or notify civil society organisations and affected communities. ECNL strongly supports requiring stakeholder engagement and notification to external stakeholders in corrective actions processes (article 21), duty of information (article 22), and obligations of importers and distributors (article 26 and 27), among others. ECNL believes that users, who deploy the AI systems in specific contexts and applications, have a particular duty to consult with affected groups given the significant implications that these systems can have on their rights.

3. Low-risk systems (“certain AI systems”) – art. 52

AI systems considered as low-risk are barely regulated in the AI Act, which merely imposes minimal transparency obligations for providers of a few technologies.

- I. As mentioned above when discussing prohibitions and high-risk systems, some systems listed as low-risk are shockingly misclassified. Emotion recognition technology; biometric categorisation for the purpose of predicting ethnicity, gender, political or sexual orientation; and risk assessments for criminal justice and asylum should be prohibited entirely. The use of bots (at least in some contexts) and deepfakes should be considered high-risk (article 52(2) and 52(3)).
- II. These transparency measures are even more inadequate given the broad exceptions in the AI Act that further reduce their effectiveness. This is particularly worrisome in the context of criminal justice, where there is already lots of opacity and discrimination of racialized persons and religious minorities. High-risk contexts and applications such as predictive policing and sentencing should demand more transparency, not less. Moreover, given their severe human rights impacts (e.g. right to life, liberty, and security), they should be prohibited or at the very least considered high risk.



4. No-risk systems (all other systems)

The vast majority of AI systems are left unregulated in the AI Act. This is even more problematic given that the AI Act precludes governments from further regulating them at the national level (see above).

- I. AI providers are merely encouraged to adopt voluntary codes of conduct (art. 69), which have long been criticized as ineffective. They are at best performative, and at worst legitimise an uptake of AI under the promise of ‘good conduct’. In any case, AI providers should instead adopt human rights policies and implement the UNGPs (see below section on corporate accountability). AI providers are merely *encouraged* to *voluntarily* apply the mandatory requirements for high-risk AI systems, which is highly insufficient to protect fundamental rights. As mentioned throughout this document, minimum legal requirements should apply to all AI systems, irrespective of their risk level, and enhanced obligations should be proportionate to their risk level.
- II. AI providers are finally “*encouraged* to apply on a *voluntary* basis additional requirements related to, for example (...) stakeholders’ participation in the design and development of AI systems” (art. 69.2). Meaningful stakeholder participation, including external stakeholders such as civil society organisations, should be mandatory in the context of human rights due diligence by AI providers and users.

5. Governance mechanisms

Overall, the governance mechanisms outlined in the AI Act are inadequate to effectively prevent and remedy harm. Resources are for their part not well allocated.

- I. The European Data Protection Board, which will be responsible for supervising the Union bodies that fall within the scope of the AI Act, would also be the most effective authority to oversee enforcement of the AI Act. Yet this competence is assigned to a newly established European Artificial Intelligence Board (the ‘Board’), composed of representatives of Members



This work is licensed under a [Creative Commons Attribution-ShareAlike 4.0 International License](https://creativecommons.org/licenses/by-sa/4.0/).

States and the Commission (art. 56 and 57). ECNL recommends expanding membership eligibility to include human rights experts, researchers and affected communities to ensure multi-stakeholder representation. Resources should also be allocated to these stakeholders to ensure their meaningful participation.

- II. Regarding ex-post enforcement, ECNL welcomes the fact that public authorities should have the “powers and resources to intervene in case AI systems generate unexpected risks, which warrant rapid action” (para. (5.2.6); art. 65–68). This includes withdrawing a product when there’s a risk, despite it being in compliance with the regulation. ECNL sees this as good example of the principle of precaution and hopes that adequate implementation will follow. We also support putting the burden of proof of demonstrating that a product is no longer a risk on AI providers. In any case, ECNL encourages expanding this provision to explicitly include the possibility of affected stakeholders and civil society organisations to sound the alarm on adverse impacts and call for the removal of a system from the market.
- III. Relatedly, requirements around withdrawing a product and notifying authorities at the national level seem to ignore the need to inform the general public. While it’s appropriate to not publicly share the details of an ongoing investigation, there is a strong public interest in communicating that a particular AI system is being investigated and/or has been temporarily suspended, and when an AI system has been removed from the market. This is also a necessary element to enable future access to remedy. While the AI Act does not grant a direct right to remedy, affected individuals and groups still need to be informed to seek remedy on other legal grounds. Yet too often, adversely impacted individuals and groups are not even aware of a violation to begin with. Civil society organisations and academics would also benefit from such knowledge for their research, organising and advocacy efforts.



THE CORPORATE RESPONSIBILITY TO RESPECT HUMAN RIGHTS

AI providers, the vast majority of which are private sector actors, are tasked with carrying out most of the requirements in the AI Act, yet there is no mention of companies' responsibility to respect human rights in their activities and supply chains.

- I. ECNL is pleased that the AI act applies to multinational enterprises that are located in a third country where the output of the system is used in the Union (article 1.1.) We would recommend extending this application to companies' value chains, which at the moment are merely required to cooperate with providers and users in complying with the AI Act (para. (60), art. 24-28).
- II. Overall, the AI Act fails to take a human-rights based approach to corporate responsibility, as consistent with the UNGPs, recent country-level laws on mandatory human rights due diligence, and the European Parliament's legislative initiative on mandatory human rights due diligence for supply chains.[8] The AI Act misses an important opportunity to require (or at the very least recommend) AI providers to have in place human rights policies, conduct human rights due diligence, and establish operational grievance mechanisms, in close consultation with affected stakeholders.
- III. ECNL is alarmed about the disproportionate role that standardization bodies like CEN and CENELEC have, and their power to adopt standards related to the AI Act. Given that AI providers will *de facto* follow these standards when conducting conformity assessments, external stakeholders including civil society, academics and affected communities should participate in the development of standards. As Veale and Zuiderveen Borgesius have pointed out, "[n]otified bodies [article 33] checking a provider's self-assessment may play a small role, but there are few situations where they are required." [1] As a result, AI providers will essentially operate unchecked. We refer to our section above on conformity assessments on further analysis on how these assessments fall short of preventing adverse human rights, and why we advocate for HRIAs instead.



STAKEHOLDER ENGAGEMENT & CIVIL SOCIETY PARTICIPATION

The AI Act misses an important opportunity to enable civic participation and require meaningful stakeholder engagement, especially of at-risk and marginalized groups. Ultimately, “the Act lacks a bottom-up force to hold regulators to account for weak enforcement.” [1]

- I. Affected individuals and communities do not have standing to claim redress under the AI Act, which ECNL views as a major shortcoming of the regulation. Only those with obligations under the AI Act can challenge regulators' decisions. In other words, the AI Act creates no legal right to sue a provider or user for failures to comply with the obligations therein. As the European Digital Rights Initiative (EDRI) rightfully warns, this has a dual effect of, first, stripping individuals whose fundamental rights have been impacted from their right to seek remedy, and second, increasing the power of AI providers (which are generally private companies) to shape rules for how public authorities should use AI systems.[7] The only right to contest decisions that is granted to “parties having a legitimate interest in that decision” is an appeal against decisions of notified bodies (article 45). As mentioned above, however, notified bodies have a very limited role in overall compliance and oversight. ECNL also recommends adding an explicit right of civil society organisations and external stakeholders to appeal these decisions and consider them as having a legitimate interest.
- II. Stakeholder engagement obligations related to operationalising the requirement in the AI Act are strongly inadequate, and mostly absent altogether. Meaningful engagement with a wide range of stakeholders, including unions and worker representatives, is needed on an ongoing basis and during re-assessments of conformity (art. 43.4). As mentioned above, notified bodies (article 33) and the market surveillance authority who verify the conformity assessment in limited cases have only minimal importance and responsibility, in practice, given the power given to standardisation bodies.[1] AI system providers can generally choose the notified body, which raises important conflict of interest questions (article 43.1). This is especially problematic given that European Standardisation



Organisations do not have a strong track-record on stakeholder engagement. ECNL urges these organisations to reform their processes to include representatives of civil society and incorporate feedback from at-risk and marginalised groups. Unfortunately, the AI Act makes no mention whatsoever about the need to include affected communities in neither the verification, nor the standardisation process. As noted by Veale and Zuiderveen Borgesius, “It is unclear whether limited existing efforts to include stakeholder representation will enable the deep and meaningful engagement needed from affected communities. The vast majority will have absolutely no experience of standardisation, and may lack EU-level representation.”[1] More generally, there needs to be deeper discussion about who, practically, will be part of notification and standardization bodies, and what (if any) role will civil society play. Will they have sufficient human rights expertise, and will they include the lived experiences of affected communities and marginalized and vulnerable groups? All information related to the bodies, and their assessments, should be made publicly available and accessible. Moreover, users should consult with the market surveillance authority and affected groups before deploying the AI system.

- III. The AI Board has the capacity to invite external experts (article 57(4)). While ECNL encourages the inclusion of experts, we are concerned that, unless specific efforts are made to include affected communities, ‘experts’ will mostly be white, cis male and representatives of industry and academia. Affected communities, especially historically marginalized and vulnerable groups, have lived experienced and understanding of risks on the ground, thus providing valuable input and often neglected expertise. Similarly, civil society organisations are well positioned to inform on risks to fundamental rights and should be referred to as experts. In this respect, sufficient resources should be dedicated to supporting meaningful participation of civil society. The AI Act foresees the need to provide financial and human resources to national competent authorities (article 59(4)). We recommend adding a similar provision focusing on affected groups, especially marginalized ones, and civil society organisations. We also recommend extending the requirement to dedicate adequate resources to all areas where stakeholder engagement is needed, including when setting standards, conducting risk assessments, and enforcing the regulation.



HUMAN RIGHTS IMPLICATIONS OF MARGINALIZED AND AT-RISK GROUPS

AI systems disproportionately impact already marginalized and at-risk groups, further exacerbating existing inequality. The section below outlines a few important shortcomings of the AI Act from the perspective of marginalized groups. Importantly, any analysis should be intersectional at its heart, i.e. acknowledging that persons with intersecting forms of identity face elevated (often unique) harms. This reinforces the importance of including – and centring – affected communities, as they are the best positioned to inform on risks and harms of AI systems in their communities.

- I. To ensure that women, trans people, and gender non-binary persons are protected against harmful consequences of AI systems, the AI Act should include the human rights impacts on gender beyond “sex” only and consider gender as non-binary. Overall, the AI Act makes no mention of the specific risks of AI systems to trans people and gender non-binary persons. What’s more, the AI Act does not address the risk of harm to LGBTQIA+ communities, and loses an opportunity to ban dangerous technology such as automated recognition of gender and sexual orientation.[9]
- II. The use of some AI systems such as polygraphs, emotion recognition technology, and risk assessment tools for the purpose of migration, asylum or border control management by public authorities are considered as high-risk (annex III s. 7). However, these systems are widely inaccurate and effective, relying on racist and pseudoscientific technology, and should be prohibited. They are often incompatible with or have severe human rights impacts on migrants, refugees and asylum seekers’ rights, such as their right to life, the prohibition of arbitrary detention, right to movement, and non-discrimination, among others.
- III. The use of some AI systems such as polygraphs, emotion recognition technology, risk assessment tools and predictive police for the purpose of law enforcement and criminal justice are considered as high-risk (annex III s. 6). However, these systems are widely inaccurate and effective, relying on racist and pseudoscientific technology, and should be banned. The AI Act



acknowledges the risk of “potential biases, errors, and opacity” in these systems as they can harm racialised people. However, it fails to consider how these systems will exacerbate structural racial inequality regardless of the capacity to “de-bias” systems. As cautioned by EDRI, “[b]y relying on technical checks for bias as a response to discrimination, the proposal risks reinforcing a harmful suggestion that removing bias from such systems is even possible.”[7]

- IV. The AI Act does not give proper attention to accessibility for persons with disabilities. It does not analyse the specific risks that persons with disabilities face when AI systems are deployed in different contexts. Conversely, it does not require accessibility to technologies necessary for realising their human rights, only having voluntary recommendation for non-high-risk AI providers.*
- V. The use of AI systems for social services risks exacerbating existing social and economic inequality, disproportionately harming persons of lower socio-economic status. Examples of applications include determining whether and how benefits and services (e.g. welfare, education, healthcare, etc.) should be allocated. The AI Act does not sufficiently address the risks to people’s right to social protection (and other economic and social rights) and non-discrimination, failing to consider the intersection between poverty, race and gender (despite acknowledging the problem at para. (37) and (35)). The AI Act rightfully recognises how the use AI systems in the workplace can harm already vulnerable workers by “perpetuat[ing] historical patterns of discrimination, for example against women, certain age groups, persons with disabilities, or persons of certain racial or ethnic origins or sexual orientation” (para. (36)). It should go further by prohibiting the use of some systems (e.g. emotion recognition in the workplace) and calling out the risks and harm related to worker surveillance, especially given the power differentials that exist between employers and workers.

* Question and comment made by Mher Hakobyan from the European Disability Forum during the ECNL workshop “AI regulation in Europe – Opportunities for civil society to engage in policymaking” on June 30, 2021.



BIBLIOGRAPHY

- [1] M. Veale and F. Z. Borgesius, “Demystifying the Draft EU Artificial Intelligence Act.” SocArXiv, Jul. 05, 2021. doi: [10.31235/osf.io/38p5f](https://doi.org/10.31235/osf.io/38p5f)
- [2] *Council Directive 2000/78/EC of 27 November 2000 establishing a general framework for equal treatment in employment and occupation*, vol. OJ L. 2000. Accessed: Jul. 23, 2021. [Online]. Available: <http://data.europa.eu/eli/dir/2000/78/oj/eng>
- [3] “Consultation on the elements of a legal framework on AI,” *Artificial Intelligence*. <https://www.coe.int/en/web/artificial-intelligence/cahai-multi-stakeholder-consultation> (accessed Jul. 19, 2021).
- [4] “Ban Biometric Surveillance,” Access Now. <https://www.accessnow.org/ban-biometric-surveillance/> (accessed Jul. 13, 2021).
- [5] “EDPB & EDPS call for ban on use of AI for automated recognition of human features in publicly accessible spaces, and some other uses of AI that can lead to unfair discrimination | European Data Protection Board.” https://edpb.europa.eu/news/news/2021/edpb-edps-call-ban-use-ai-automated-recognition-human-features-publicly-accessible_en (accessed Jul. 13, 2021).
- [6] Marlena Wisniak, European Center for Not-for-Profit Law, “Evaluating the risk of AI systems to human rights from a tier-based approach | ECNL,” Mar. 23, 2021. <https://ecnl.org/news/evaluating-risk-ai-systems-human-rights-tier-based-approach> (accessed Jul. 13, 2021).
- [7] “EU’s AI law needs major changes to prevent discrimination and mass surveillance,” *European Digital Rights (EDRI)*. <https://edri.org/our-work/eus-ai-law-needs-major-changes-to-prevent-discrimination-and-mass-surveillance/> (accessed Jul. 13, 2021).
- [8] “MEPs: Companies must no longer cause harm to people and planet with impunity | News | European Parliament,” Oct. 03, 2021. <https://www.europarl.europa.eu/news/en/press-room/20210304IPR99216/meps-companies-must-no-longer-cause-harm-to-people-and-planet-with-impunity> (accessed Jul. 14, 2021).
- [9] Access Now, Reclaim Your Face, “Ban automated recognition of gender and sexual orientation,” Apr. 21, 2021. <https://campaigns.allout.org/ban-AGSR> (accessed Jul. 19, 2021).



This work is licensed under a [Creative Commons Attribution-ShareAlike 4.0 International License](https://creativecommons.org/licenses/by-sa/4.0/).

European Center for Not-for-Profit Law Stichting
5 Riviervismarkt, 2513 AM, The Hague, Netherlands
www.ecnl.org twitter.com/enablingNGOlaw



European Center for
Not-for-Profit Law