

Artificial Intelligence Act

2021/0106(COD) - 14/06/2023 - Text adopted by Parliament, partial vote at 1st reading/single reading

The European Parliament adopted, by 499 votes to 28 with 93 abstentions, amendments to the proposal for a Regulation of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (artificial intelligence act) and amending certain Union legislative acts.

The matter was referred back to the committee responsible for inter-institutional negotiations.

Purpose

The regulation lays down a uniform legal framework in particular for the development, the placing on the market, the putting into service and the use of artificial intelligence in conformity with Union values. Its aim is to **promote the uptake of human centric and trustworthy artificial intelligence** and to ensure a high level of protection of health, safety, fundamental rights, democracy and rule of law and the environment from harmful effects of artificial intelligence systems in the Union.

Supporting innovation

To boost AI innovation and support SMEs, Members added exemptions for research activities and AI components provided under open-source licenses. The new law promotes so-called **regulatory sandboxes**, or real-life environments, established by public authorities to test AI before it is deployed.

General principles applicable to all AI systems

All operators covered by the Regulation should make every effort to develop and use AI or general purpose AI systems in accordance with the following general principles: (i) ‘human agency and oversight’; (ii) ‘technical robustness and safety’; (iii) ‘privacy and data governance’; (iv) ‘transparency’; (v) ‘diversity, non-discrimination and fairness’; and (vi) ‘social and environmental well-being’.

AI literacy

When implementing the regulation, the Union and the Member States should promote measures for the development of a sufficient level of AI literacy, across sectors and taking into account the different needs of groups of providers, deployers and affected persons concerned, including through education and training, skilling and reskilling programmes and while ensuring proper gender and age balance, in view of allowing a democratic control of AI systems.

Prohibition of AI practices

AI systems posing an unacceptable level of risk to personal safety will be prohibited. Members expanded the list to include bans on intrusive and discriminatory uses of AI, such as:

- systems that use **subliminal techniques or deliberately manipulative or deceptive techniques**, with the aim of substantially distorting a person’s or a group of persons’ behaviour by appreciably impairing the person’s ability to make an informed decision, thereby causing the person to take a decision that that person would not have otherwise taken in a manner that causes or is likely to cause that person, another person or group of persons significant harm;

- systems which exploit the possible vulnerabilities of a given person or group of persons, in particular known or predictable personality traits or the social or economic situation, age, physical or mental capacity of that person or group of persons, with the aim or effect of substantially altering that person's behaviour;
- placing on the market, putting into service or use of biometric categorisation systems that categorise natural persons according to **sensitive or protected attributes** (e.g. gender, race, ethnic origin, citizenship status, religion, political orientation), or characteristics or based on the inference of those attributes or characteristics;
- **systems used for social rating** (classifying people according to their social behaviour or personality characteristics);
- the use of 'real-time' remote biometric identification systems in publicly accessible areas;
- **predictive policing** systems (based on profiling, location or past criminal behaviour);
- emotion recognition systems in law enforcement, border management, the workplace, and educational institutions; and
- untargeted scraping of **facial images** from the internet or CCTV footage to create facial recognition databases (violating human rights and right to privacy).
- emotion recognition systems used in law enforcement, border management, the workplace and educational establishments;
- **'post' remote biometric identification systems**, with the only exception of law enforcement for the prosecution of serious crimes and only after judicial authorisation.

The following have been added to the list of high-risk systems:

- systems intended to be used as security components in the management and operation of the supply of water, gas, heating, electricity and critical digital infrastructures;
- systems intended to be used to assess the appropriate level of education of an individual and which substantially influence the level of education and vocational training from which that individual will benefit or to which he or she will have access;
- systems intended to be used to monitor and detect prohibited behaviour in students during tests in the context of, or within, education and training institutions;
- systems intended to be used to make or substantially influence decisions on the eligibility of natural persons for health and life insurance;
- systems intended to evaluate and classify emergency calls from individuals;
- AI systems intended to be used by public authorities in the **management of migration, asylum** and border controls to process, control and verify data for the purpose of detecting, recognising and identifying natural persons;
- systems intended to be used to **influence the outcome of an election or referendum** or the voting behaviour of individuals in the exercise of their vote in elections or referendums;

- AI systems used in **recommender systems** operated by major social media platforms.

Obligations for general purpose AI

Generative AI systems based on such models, like ChatGPT, would have to comply with **transparency requirements** (disclosing that the content was AI-generated, also helping distinguish so-called deep-fake images from real ones) and ensure safeguards against generating illegal content. Detailed summaries of the copyrighted data used for their training would also have to be made publicly available.

AI Office

The proposal establishes the AI Office, which should be an **independent body** of the Union. It is proposed that it should be based in Brussels. Its tasks should include the following:

- support, advise and cooperate with Member States, national supervisory authorities, the Commission and other Union institutions, bodies, offices or agencies on the implementation of this Regulation;
- monitor and ensure the effective and consistent application of the Regulation;
- contribute to the coordination between the national supervisory authorities responsible for the application of the Regulation;
- mediate in discussions on serious disagreements which may arise between competent authorities concerning the application of the Regulation;
- coordinate joint investigations.

The IA office should be accountable to the European Parliament and the Council, act independently and ensure a high level of transparency.

Right to lodge a complaint with a national supervisory authority

Every natural persons or groups of natural persons will have the right to lodge a complaint with a national supervisory authority if they consider that the AI system relating to him or her infringes this Regulation. Lastly, Members want to boost citizens' right to file complaints about AI systems and receive explanations of decisions based on high-risk AI systems that significantly impact their fundamental rights and socio-economic well-being.