

Brussels, 21.4.2021
SWD(2021) 85 final

COMMISSION STAFF WORKING DOCUMENT
EXECUTIVE SUMMARY OF THE IMPACT ASSESSMENT REPORT

Accompanying 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**

{COM(2021) 206 final} - {SEC(2021) 167 final} - {SWD(2021) 84 final}

Executive Summary Sheet
Impact assessment on a Regulatory Framework for Artificial Intelligence
A. Need for action
What is the problem and why is it a problem at EU level?
Artificial Intelligence (AI) is as an emerging general-purpose technology: a highly powerful family of computer programming techniques. The uptake of AI systems has a strong potential to bring societal benefits, economic growth and enhance EU innovation and global competitiveness. However, in certain cases, the use of AI systems can create problems. The specific characteristics of certain AI systems may create new risks related to (1) safety and security and (2) fundamental rights, and accelerate the probability or intensity of the existing risks. AI systems also (3) make it hard for enforcement authorities to verify compliance with and enforce the existing rules. This set of issues in turn leads to (4) legal uncertainty for companies, (5) potentially slower uptake of AI technologies, due to the lack of trust, by businesses and citizens as well as (6) regulatory responses by national authorities to mitigate possible externalities risking to fragment the internal market.
What should be achieved?
The regulatory framework aims to address those problems in order to ensure the proper functioning of the single market by creating conditions for the development and use of trustworthy AI in the Union. Specific objectives are: (1) ensuring that AI systems placed on the market and used are safe and respect the existing law on fundamental rights and Union values; (2) ensuring legal certainty to facilitate investment and innovation in AI; (3) enhancing governance and effective enforcement of the existing law on fundamental rights and safety requirements applicable to AI systems; and (4) facilitating the development of a single market for lawful, safe and trustworthy AI systems and prevent market fragmentation.
What is the value added of action at the EU level (subsidiarity)?
The cross-border nature of large-scale data and datasets which AI applications often rely on means that the objectives of the initiative cannot be achieved effectively by Member States alone. The European regulatory framework for trustworthy AI aims to establish harmonised rules on the development, placement on the market and use of products and services embedding AI technology or stand-alone AI applications in the Union. Its purpose is to ensure a level playing field and to protect all European citizens, while strengthening Europe's competitiveness and industrial basis in AI. EU action on AI will boost the internal market and has a significant potential to provide European industry with a competitive edge at a global level, based on economies of scale that cannot not be achieved by individual Member States alone.
B. Solutions
What are the various options to achieve the objectives? Is there a preferred option or not? If not, why?
The following options have been considered: Option 1: EU legislative instrument setting up a voluntary labelling scheme; Option 2: a sectoral, "ad-hoc" approach; Option 3: a horizontal EU legislative instrument establishing mandatory requirements for high-risk AI applications; Option 3+: the same as 3 but with voluntary codes of conduct for non-high-risk AI applications; and Option 4: a horizontal EU legislative instrument establishing mandatory requirements for all AI applications. The preferred option is Option 3+ since it offers proportionate safeguards against the risks posed by AI, while limiting the administrative and compliance costs to a minimum. The specific question of liability for AI applications will be addressed through distinct future rules and thus not covered by the options.
What are different stakeholders' views? Who supports which option?
Businesses, public authorities, academics and non-governmental organisations all agree that legislative gaps exist or that new legislation is needed, although the majority among businesses is smaller. Industry and public authorities agree with limiting the mandatory requirements to high-risk AI applications. Citizens and civil society are more likely to disagree with limiting mandatory requirements to high-risk

applications.
C. Impacts of the preferred option
What are the benefits of the preferred option (if any, otherwise of main ones)?
For citizens, the preferred option will mitigate risks to their safety and fundamental rights. For providers of AI, it will create legal certainty and ensure that no obstacle to the cross-border provision of AI-related services and products emerge. For companies using AI, it will promote trust among their customers. For national public administrations, it will promote public trust in the use of AI and strengthen enforcement mechanisms (by introducing a European coordination mechanism, providing for appropriate capacities, and facilitating audits of the AI systems with new requirements for documentation, traceability and transparency).
What are the costs of the preferred option (if any, otherwise of main ones)?
Businesses or public authorities that develop or use AI applications that constitute high risk for the safety or fundamental rights of citizens would have to comply with specific horizontal requirements and obligations, which will be put in place through technical harmonised standards. The total aggregate cost of compliance is estimated to be between €100 million and €500 million by 2025, corresponding to up to 4-5% of investment in high-risk AI (which is estimated to be between 5% and 15% of all AI applications). Verification costs could amount to another 2-5% of investment in high-risk AI. Businesses or public authorities that develop or use any AI applications not classified as high risk would not have to incur any costs. However, they could choose to adhere to voluntary codes of conduct to follow suitable requirements, and to ensure that their AI is trustworthy. In these cases, costs could be as high as for high-risk applications at most, but most probably lower.
What are the impacts on SMEs and competitiveness?
SMEs will benefit more from a higher overall level of trust in AI than large companies who can also rely on their brand image. SMEs developing applications classified as high risk would have to bear similar costs as large companies. Indeed, due to the high scalability of digital technologies, small and medium enterprises can have an enormous reach despite their small size, potentially impacting millions of individuals. Thus, when it comes to high-risk applications, excluding AI supplying SMEs from the application of the regulatory framework could seriously undermine the objective of increasing trust. However, the framework will envisage specific measures, including regulatory sandboxes or assistance through the Digital Innovation Hubs, to support SMEs in their compliance with the new rules, taking into account their special needs.
Will there be significant impacts on national budgets and administrations?
Member States would have to designate supervisory authorities in charge of implementing the legislative requirements. Their supervisory function could build on existing arrangements, for example regarding conformity assessment bodies or market surveillance, but would require sufficient technological expertise and resources. Depending on the pre-existing structure in each Member States, this could amount to 1 to 25 Full Time Equivalents per Member States.
Will there be other significant impacts?
The preferred option would significantly mitigate risks to fundamental rights of citizens as well as broader Union values, and will enhance the safety of certain products and services embedding AI technology or stand-alone AI applications.
Proportionality?
The proposal is proportionate and necessary to achieve the objectives, as it follows a risk-based approach and imposes regulatory burdens only when the AI systems are likely to pose high risks to fundamental rights or safety. Where this is not the case, only minimal transparency obligations are imposed, in particular in terms of information provision to flag the use of an AI system when interacting with humans or use of deep fakes if not used for legitimate purposes. Harmonised standards and supporting guidance and compliance tools will aim to help providers and users to comply with requirements and minimise their costs.

D. Follow up
When will the policy be reviewed?
The Commission will publish a report evaluating and reviewing the framework five years following the date on which it becomes applicable.