

An innovation-friendly approach based on European values for the AI Act

- Joint Non-paper by IT, FR and DE -

- We acknowledge the need for a comprehensive regulation of AI systems and from the beginning welcomed the Commission proposal for an AI Act in this regard. The AI Act will provide EU citizens with protection and confidence in the AI products distributed on the single market.
- This new regulation will complement the comprehensive legal toolbox already applicable in the EU, for instance on data privacy with GDPR, or with Digital Services Act or the Terrorist Content Online regulation.
- The EU intends to position itself at the forefront of the AI revolution. This requires a regulatory framework which fosters innovation and competition, so that European players can emerge and carry our voice and values in the global race of AI.
- In this context, we reiterate our common commitment for a balanced and innovation-friendly and a coherent risk-based approach of the AI Act, reducing unnecessary administrative burdens on Companies that would hinder Europe's ability to innovate, that will foster contestability, openness and competition on digital markets.
- We welcome the efforts from the Spanish Presidency to find a compromise with the European Parliament and Commission to reach a satisfactory solution for all parties and stakeholders.
- Together we underline that the AI Act regulates the application of AI and not the technology as such. This risk-based approach is necessary and meant to preserve innovation and safety at the same time.
- Legal certainty, clarity and predictability are of utmost importance.
- Special attention should be paid to definitions and distinctions. We should continue to follow a thorough discussion on this topic. Definitions should be clear and precise. To this regard we strongly underline and welcome the efforts of the Spanish presidency.
- We suggest a distinction between models and general purpose AI systems that can be available for specific applications.
- We believe that regulation on general purpose AI systems seems more in line with the risk-based approach. The inherent risks lie in the application of AI

systems rather than in the technology itself. European standards can support this approach following the new legislative framework.

- When it comes to foundation models we oppose instoring un-tested norms and suggest to instore to build in the meantime on mandatory self-regulation through codes of conduct. They could follow principles defined at the G7 level through the Hiroshima process and the approach of Article 69 of the draft AI Act, and would ensure the necessary transparency and flow of information in the value chain as well as the security of the foundation models against abuse.
- We are however opposed to a two-tier approach for foundation models.
- To implement our proposed approach, developers of foundation models would have to define model cards.
- Defining model cards and making them available for each foundation model constitutes the mandatory element of this self-regulation.
- The model cards must address some level of transparency and security
- The model cards shall include the relevant information to understand the functioning of the model, its capabilities and its limits and will be based on best practices within the developers community. For example, as we observe today in the industry: number of parameters, intended use and potential limitations, results of studies on biases, red-teaming for security assessment.
- An AI governance body could help to develop guidelines and could check the application of model cards.
- This system would ensure that companies have an easy way to report any noticed infringement of the code of conduct by a model developer to the AI governance body. Any suspected violation in the interest of transparency should be made public by the authority.
- No sanctions would be applied initially. However, after an observation period of a defined duration, if breaches of the codes of conduct concerning transparency requirements are repeatedly observed and reported without being corrected by the model developers, a sanction system could then be set up following a proper analysis and impact assessment of the identified failures and how to best address them.
- European standards could also be an important tool in this context as this also creates the adaptive capacity to take into account future developments. Further standardization mandates could be foreseen in this regard