



GSMA-ETNO position paper on European Commission proposal for an Artificial Intelligence Act

GSMA and ETNO welcome the initiative of the European Commission to propose a Regulation on Artificial Intelligence, and notably the central role of the risk-based approach, which keeps citizens at the centre and is based on a robust ethical framework. Nevertheless, the risk-based approach requires more precision in order to achieve the intended results. It is encouraging to see that the legal focus is put on the use of AI and its impacts to society, rather than the technology itself.

AI applications will be a key driver of innovation for the European telecoms industry, notably as part of the shift towards 5G, virtualised networks. 5G and fibre connectivity will accelerate the digitisation of services and industrial processes, enabling the rapid expansion of the Internet of Things (IoT). The massive amounts of data generated by IoT connections and devices will open up new growth opportunities for data analytics and AI services in Europe. High-class, secure

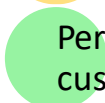
connectivity will then drive IoT, and IoT will in turn fuel European AI. Together, they can form a truly powerful virtuous circle that our industry is committed to nurturing. Digital network providers themselves increasingly deploy AI solutions in various areas, typically to improve efficiency in and reduce the carbon footprint of network operations (e.g., network security, predictive maintenance and energy efficiency), improve cybersecurity, enhance customer experience, and enable better product and service development. AI applications in the telecoms industry include, but are not limited to, network planning optimisation, intelligent pricing, B2B sales optimisation, and customer service. These use cases have to be enabled and the regulation should not disincentivise these initiatives.

As much as the telecoms industry is at the heart of many technological innovations already shaping tomorrow's digital society, operators are also leading the way in transparency and inclusivity in this digital age. Indeed, as advances in AI, IoT and 5G give rise to intelligent connectivity, maintaining public trust and skills are crucial to enable these technologies to deliver a better connected living for everyone. A digital society where everybody can embrace new tools will provide fertile ground for continued innovation across all sectors and throughout the global economy, bolstered by strong values and principles, including the telecoms industry's [Digital Declaration](#) and [AI Ethics Principles](#).

AI Opportunities for the Telecoms Industry



Core Business optimisation



Personalised and improved customer interaction



AI-Driven Mobile Data Insights

The GSMA and ETNO believe it is important to ensure a balance between **legal certainty** for providers and users of AI on the one hand, and guaranteeing that the Regulation will be **sufficiently flexible so as to be future-proof**. Crucial to this is the risk-based approach, which forms the crux of the Commission proposal, ensuring that any new requirements are proportionate and strictly targeted at AI applications that pose a high risk. It is for this reason we argue that AI systems deployed in the operation of electronic communications networks **should not be considered as high-risk within Annex III of the regulation**.

As noted in the explanatory note accompanying the European Commission's proposal, an AI system is classified as high-risk based on the function it performs, the specific purpose and modalities for which the system is used and consideration of the risks posed to individuals **health, safety and fundamental rights**. Furthermore, concerning Annex III, the list of critical infrastructure is intended to enumerate "limited number of AI systems whose risks have already materialised or are likely to materialise in the near future". The function of AI systems and services in telecommunications infrastructure **is solely to manage and optimize operators' networks in a way that enhances resilience and security** (e.g. predictive maintenance, trouble ticketing management, optimising energy use). As there is no direct engagement with end users, these AI systems do not give rise to risks to health, safety and fundamental rights that are addressed under this regulation. It was for these reasons that the European Commission did not include ECN under Annex III following the Impact Assessment process.

Furthermore, the Slovenian Presidency compromise text bases the justification of classifying AI systems used to control or as safety components of digital infrastructure as high-risk under Annex III on account of the "increasing digitalisation of all sectors of the economic and public life" (recital 34, Slovenian compromise text). Notwithstanding that said inclusion has occurred in the absence of an Impact Assessment by the European Commission or co-legislators, the text has been added ostensibly to guarantee the resilience of digital infrastructure. However, the resilience, safety and security of electronic communications networks is already ensured by EU and national frameworks, for example:

- The European Electronic Communications Code (EECC) demands that providers of public electronic communications be required to take measures to **safeguard the security of their networks and services**, respectively, and to **prevent or minimise the impact of security incident**;
- The proposed NIS-2 Directive will introduce requirements for **critical and important entities** to manage the risks posed to the **security of network and information systems** which those entities use in the provision of their services (e.g. risk analysis, incident handling, crisis management, business continuity);
- The proposed Directive on the Resilience of Critical Entities will set out obligations for critical entities aimed at enhancing their resilience and improving their ability to provide those services in the internal market (e.g. incident prevention and mitigation, resilience plans).

Therefore, considering that the application of AI systems and services in digital infrastructure and ECN poses no risk to the **health, safety and fundamental rights** of individuals and that the security and resilience of digital infrastructure and ECN is already governed under existing

national and EU frameworks, there is no justification to include them as part of the list of High-Risk AI systems under Annex III.

Annex I - III: Legal Certainty

Providers and users of AI must be provided with more legal certainty and predictability on the techniques and definitions in scope, which is not the case if the Regulation allows for the annex to be expanded continuously. Since the European Commission is empowered to adopt delegated acts to amend the list of techniques and approaches in Annex I and the list of high-risk applications in Annex III, to which the definition of AI in article 3 and indeed the scope of Regulation are intrinsically linked, strict rules should be foreseen for the periodic review of the Annex to provide legal certainty for providers and users of AI. As such, **any adjustment of the annex must be proportionate** and be based on a regular and institutionalised exchange with industry, should not hinder market entry and innovation, and must take into account evolutionary path of the technology: delegated acts should be drafted on the basis of broad stakeholder consultation, announced in a public rolling plan, and undergo due process including public consultation before adoption.

Article 3: Definition of Artificial Intelligence (and Annex I)

While we welcome the Commission's effort of providing a precise definition of AI, we also caution that the definition of AI for the purposes of the Regulation as it currently stands – namely to apply different regimes, obligations and protections to AI applications of a high-risk versus medium- and low-risk nature – is too broad. **An overly broad or open-ended definition risks to erroneously include all software and create a disproportionate burden for developers of technology that is not strictly AI.** For example, statistical methods and Bayesian estimation (see Annex 1.c) are a part of machine- and deep-learning as well as normal data processing. In our view, they are covered without being specifically mentioned or already regulated under GDPR. If Bayesian and statistical methods were to be explicitly mentioned, their use should only be considered as AI if they are used to extract decisions from data in an automated way. GSMA and ETNO therefore welcome the new definition of AI systems proposed by the Slovenian Presidency, which introduces the **concept of inference**, crucial for distinguishing artificial intelligence systems from traditional ones. A reference to the ability of AI systems to identify patterns through multiple inferences/correlations could further improve the definition. Another missing element in the definition, which characterizes some AI systems, is the capacity to make **autonomous decisions**.

Title III, Chapter 1. High-risk AI: scope of application

We support the more targeted approach of the proposal to define high-risk applications compared to previous proposals in the White Paper, which intended to identify certain sectors as high-risk. This now creates more legal certainty and gives an opportunity to create “best practices” of implementation of trustworthy AI practices in narrow domains and high-risk products that incorporate AI systems, which are already subject to conformity assessment obligations and which the Commission lists in Annex II.

Nonetheless, clarifications will be necessary regarding **applications** considered to be high-risk. While we appreciate the risk-based approach towards different levels of risk (high, low, no risk), a more precise classification of use cases within the risk areas of Annex III is missing. Context and intended purpose of an AI-based system should be taken into account to

determine risks of harm to the health and safety, or a risk of adverse impact on fundamental rights. In the current proposal, these risks have to be taken into consideration with regards to expanding annex III to other areas (Article 7.2). **It should also be used for considering whether a specific application within a high risk area – as outlined in Annex III – should be treated as high risk.**

The proposal establishes a mere ex-ante presumption/classification of what constitutes high-risk but the potential risk level of AI systems often depends on the system's context/use, for instance the use and purpose of chatbot: a chatbot used for communication in HR, which is merely providing **information to potential applicants**, has to be treated differently to a bot that takes **autonomous recruitment decisions** for example. This is not possible with the current static approach towards high-risk applications under Annex III. By refining the notions of harm and intended use, the scope would be more precise and targeted. Implementing the envisaged obligations by the Commission for such a "high-risk" scenario would likely lead to chilling effects on the use of AI systems, as the regulatory burden would outweigh the benefits.

Furthermore, certain parameters assessing whether a given context or use of an AI system is to be considered high-risk could be introduced, such as the following:

- **Decision:** the outcome of the AI system should be taken into account when assessing if the context and use is indeed high-risk: AI based systems taking autonomous decisions about for instance which customer service agent an inquiry from a customer should be redirected to, do not pose a risk to safety or fundamental rights.
- **Methods:** it should be taken into consideration which methods are used, as this also has an impact on the real risk posed – self-learning systems and machine learning could pose a greater risk than AI systems based on predictive statistics.
- **Transparency:** the less transparency, the more difficult it is to monitor or control the AI system, and as such the risk of using the system is higher.

Such parameters can be used to precisely define not only high-risk applications of AI, but also those contexts and use-cases where the risk is especially relevant.

Title III, Chapters 2 and 3: Obligations for providers and users of high-risk AI (art. 6-28 and Annex III)

It is positive that the list of high-risk applications in Annex III are clearly enumerated so as to ensure that providers and users of AI are concretely aware of the obligations incumbent on them. As such, we welcome the Commission's intention to **differentiate between the provider on the one hand, and user, distributor and importer on the other** of AI, as the former group will have to meet more obligations to mitigate risks. In many instances, only the provider is in the position to know whether the updates are required and in a position to provide them.

This is particularly the case for General Purpose AI systems developed by providers, which, as noted in draft compromise text proposed by the Slovenian Presidency of the Council until Title IV, "are able to perform generally applicable functions such as image/speech recognition, audio/video generation, pattern detection, question answering, translation etc.". These AI systems pose a wide range of established and anticipated risks, which are recognised by the

international research community¹. Such risks encompass, for example, training corpora that include harmful language and overrepresentation of social identities. To mitigate such harm, secure innovation of trustworthy AI solutions in Europe, and set global ethical standards, the AI Act must assign obligations for compliance for providers of such models if they are to be deployed in Europe. In particular, those providers should be under an obligation to assess and mitigate the ethical challenges and risks that could arise as a result of the use of a general-purpose AI system.

Nonetheless, while we agree with the overall approach of the proposed Regulation, the **obligations for providers and users of high-risk AI could result in being costly and excessive**, with the consequence that innovation, development and investment are all chilled, or certain market actors fail to fully comply, both of which would be detrimental to the stated aims of the legislative intervention, in particular the requirement to keep technical documentation for over 10 years, including pictures, the definition of residual risks in terms of risk management, and the record keeping of logs.

In addition to the administrative burden brought by such obligations and its risk to innovation, there is an **increased risk for professional secrets and IPR**, in case of breach of confidentiality from regulators, market surveillance authorities or notified bodies. As such, requests by national authorities to prove compliance with Chapter II (arts. 8 to 16) should be based on founded risks or concerns of non-compliance, otherwise developers could be subject to arbitrary requests leading to unjustified burden.

AI value chains are complex, and it is important to avoid that the regulation creates disadvantages for the European AI industry. Therefore, we propose where possible the use of existing documentation systems for managing risks. Due to the close links between data protection and AI, an extension of the GDPR risk management system with AI aspects should be considered instead of or as an alternative for an additional one for AI only. **At a minimum, requirements in the various regulations** (e.g. GDPR, AI Regulation, future EU mandatory environmental and human rights due diligence) **must be congruent and aligned**. Where possible, risk management for the use of AI should be integrated into existing risk management systems, such as for data protection, to avoid red-tape and leverage efficiencies.

Considering distributors already check whether an AI system has EC conformity and the required documentation, distributors should not be burdened with NSA's obligations; at most, **distributors should be responsible to ask for a certification of compliance, not assessing compliance themselves**.

Further clarification would be necessary as to what constitutes a substantial modification of high-risk AI systems, also with regard to AI systems that continue to learn after being placed on the market. In addition, there is no clear distinction made between AI that is continuously learning, and controlled AI and we note that Article 16 implies a considerable administrative burden, which moreover requires multiple notifications to "national competent authorities". The objective of retraining and changes to the algorithm that could be interpreted as substantial modifications are normally aimed to improve the accuracy of an AI system and thereby reduce risks. A new conformity assessment should only be required if those modifications impact the risks posed by the AI system in a negative way.

¹ Weidinger, L. et al. Ethical and social risks of harm from Language Models. DeepMind. ARXIV. December 2021. <https://deepmind.com/research/publications/2021/ethical-and-social-risks-of-harm-from-language-models>

Some requirements would benefit from greater clarification to ensure that the objectives of the Regulation are met. Article 12 should define in greater detail the standards data logging should meet and whether new, additional standards are needed to explain AI decisions. In particular, for highly complex and/or black box models such as deep neural networks, the defined minimal requirements are not sufficient to make the decision-making of the systems traceable. Article 15.4 correctly mandates that systems should be checked to ensure there are no vulnerabilities to adversarial / poisoning attacks. This should be expanded to protect systems against **exploratory attacks**, which can be aimed at revealing training data.

Title III, Chapter 5: Harmonised standards, conformity assessment and governance design

More legal certainty needs to be provided for providers and users of AI, but in addition, clear, predictable and timely working plans for standardisation efforts need to be made available. Today there are no agreed benchmarks for the conformity assessment. Harmonised standards and common specifications can be developed over time as the industry and society gains knowledge and experience. **The development of these standards should be mandated to European SDOs rather than developed by implementing acts**, prepared or defined by the Commission as stated in art. 41. There should be a transition period foreseen, to make sure AI development and uptake is not unduly hampered due to the lack of those common benchmarks.

Furthermore, the industry needs a process design that focuses on clear and easy-to-handle requirements. It is important that, sector by sector, the relevant institutions are operational so as not to create additional delays in market access. We see a **risk that the sum and overlap of requirements and obligations, including in parts vaguely defined high risk applications, and the associated legal uncertainty in operationalisation, creates a complexity and compliance burden** that inhibits the development of AI applications in the area of high risk in the EU. As such, we recommend that lessons from data protection should be considered where overlapping competences lead to slow and contradicting decisions.

Access to training, validation and testing datasets in case of high-risk AI (art. 10 and 64)

Concerning access to data by authorities and the data retention obligation, it is not always possible for many AI applications to retain all of the system training data as this represents an extremely high volume of data. Furthermore, this data is often complex, and not always accessible, if it is generated deep in an AI-driven model, or generated by machine learning, and where the operator has less clear visibility. **A time limit for data retention**, defined on the basis of the retraining periodicity of the systems (variable from system to system), which is consistent with the timing of notification of incidents and the normal supervisory activity of the authority, should be envisaged. Furthermore, consideration must be given to the protection of such data when yielded to competent authorities, as this could represent a major risk to trade secrets and intellectual property rights.

In addition, Article 10.3 requires that high quality data sets are “free of errors and complete”. It is not clear when a data set can be considered complete, nor is it feasible to conclude that any data set is “free of errors”. Furthermore, it should be clarified if “free of errors and complete” refers to data sets or the “Training, validation and testing” processes. Regarding point 10.5, bias needs to be specified in the legislative proposal because it can follow various meanings and interpretations. Data can not only be biased but can have other deficits as well. Therefore, **data quality should be specified in general**. The entire process of data processing

(i.e. pre-processing, feature selection/engineering) should be documented, not only the aggregation of data. If pre-trained models are used, their source should be clearly marked.

A general consideration must be given to the fact that data-sets are not static, but dynamic, and as such it is important to consider how meta-sets are structured and used.

Title V: Measures in support of innovation

We welcome the emphasis of the proposal on the support of innovation, in particular regulatory sandboxes. This is a positive development, but the AI Act could be more ambitious still in favour of innovation. The proposed regulation only establishes rules and regulatory oversight mechanisms for regulatory sandboxes but does not require Member States to establish any regulatory sandbox. It would be desirable to **mandate Member States to include the establishment of regulatory sandboxes in their national AI strategy**.

Enforcement and Implementation

There lies a risk that implementation of the AI regulation will be fragmented across the Member States. As such, cooperation amongst Member States is needed to ensure a harmonised approach. We see a risk that the sum and overlap of requirements, including for vaguely defined high risk applications, could lead to legal uncertainty, complexity and compliance burden. For example, as the legislation grants powers on AI compliance to national competent authorities, we question whether Member States will apply the same risk assessment procedures. As such, there should be a **requirement that the guidelines to complement the “implementation of requirements” should be adopted 3 years before entry into force of the legislation**.

Concerning **penalties** (Article 71), the intended fines are high. Therefore, the requirements and designation process of high-risk systems has to be extremely clear to minimize the risk of significant financial damage, particularly in combination with equally high fines related to data protection.

Additional Areas of Concern

AI Board

We note the lack of business representatives, researchers and consumers in proposed EC AI board and recommend their inclusion.

Low risk AI applications

While we welcome codes of conduct for low-risk AI applications and encourage their usage, they should not use the stringent legal requirements required for high-risk AI.

Policy contacts

Paul Alix

Manager EU Affairs, GSMA
palix@gsma.com

Ross Creelman

Public Policy Manager, ETNO
creelman@etno.eu