



GSMA Response to the Draft Ethics Guidelines for Trustworthy AI

January 2019

Introduction

The GSMA supports the European Commission's endeavour to maximise the benefits of artificial intelligence (AI) while minimising the risks to individuals and communities, and appreciates the opportunity to comment on the Commission's new draft ethics guidelines. We support the view that the development and deployment of AI systems should respect fundamental human rights and applicable regulation, as well as principles and values ensuring an 'ethical purpose'. A growing number of GSMA members have already committed to responsible development of AI technologies.

The general approach set out by the High-Level Expert Group on Artificial Intelligence (HLEG) in the consultation document is clear and well-considered. Only AI that is proportionate, trustworthy and robust has the potential to achieve mass market adoption. By emphasising trust, we are confident that EU technology providers will be strong players in an AI-driven world economy.

We need only observe the global reaction to the General Data Protection Regulation (GDPR) to see how the exercise of 'soft power' based on a strong rules-based framework can indeed help to shape global markets and strengthen the EU economy. We see no reason why AI ethics should be any different, and the GSMA is strongly committed to working with the European Commission and multi-stakeholder groups to deliver on this ambition.

As a general comment on the Draft Ethics Guidelines for Trustworthy AI, they should be focused as much as possible on AI, and the relationship between AI and human rights impact assessments and data protection impact assessments should be made more clear. The document repeats what is already set out elsewhere, and the reader would be given better guidance if it referred to more comprehensive documents on the general approach, such as the OECD Due Diligence Guidance for Responsible Business Conduct. Through such clarification, readers outside of Europe would also more readily accept the universality of the framework applied. With these changes, the document could focus more narrowly on good practice and safeguarding fundamental rights that are unique to AI, thereby increasing its impact on practical business.

Our detailed comments on the guidelines are primarily intended to ensure consistency with existing regulatory frameworks and to advise where measures proposed by the HLEG are either disproportionate or technically unfeasible. In all cases, we have suggested alternative wording that

should align the guidelines with industry best practice and ensure a clearer link with AI developments currently underway across the telecommunications ecosystem.

We also observe that there are many beneficial applications of AI, for example in fraud prevention, mobile network optimisation and improved IT security, that should be encouraged rather than impeded by the guidelines. Our view is that AI technologies are implicitly useful tools whose application requires an appropriate level of human oversight and application of law, such as the GDPR and ePrivacy Regulation.

Even though the guidelines are voluntary and of the ‘soft law’ nature, it is vital that they do not introduce new terminology or rules pertaining to areas that are well established in law — from human rights to privacy and data protection. We are concerned that the guidelines may be based on certain misconceptions of EU data protection law, especially the GDPR, and would suggest a reexamination of these interpretations before final publication. The main issues include excluding data protection and privacy from fundamental rights, incorrect terminology (PII vs. personal data), incorrect applicability of data subjects’ rights (only erasure and portability mentioned as being interchangeable rights), and inaccurate reflection of the existing data protection obligations under EU law (e.g., legal grounds, transparency, automated decision-making).

Section A: Rationale and Foresight of the Guidelines

- Endorsement mechanism [p. 2]

The introduction of a mechanism under which stakeholders will be able to ‘formally endorse’ (p. 2) the guidelines raises questions regarding its practicality: What are the consequences of an endorsement? Would this (fully or partly) replace self-regulatory initiatives such as codes of conduct or self-binding guidelines? Would signatories thereby fall under specific external governance/auditing? And would choosing not to sign these guidelines create a false impression that a stakeholder does not support ethical considerations regarding AI? Lastly, it appears difficult to achieve broad endorsement of the guidelines in the form of a ‘take-it-or-leave-it’ approach, where some guidance might be considered acceptable by those unwilling to accept the guidelines in total. While the intention to regularly update and evolve the guidelines by treating them as a ‘living document’ (page iv) is understandable, it might also lower stakeholders’ willingness to endorse them formally. It is also important to note that governments and policymakers can likewise develop, deploy or use AI and thus also qualify as stakeholders.

- Trustworthy AI [p. 2]

We agree with the assessment that “no legal vacuum currently exists, as Europe already has regulation in place that applies to AI” (p. 2), not least due to the technology’s cross-sectoral nature. While the guidelines are not intended “as a substitute to any form of policy-making or regulation” (p. 3), the aforementioned conclusion nevertheless must be taken into account for the HLEG’s second deliverable, i.e., the AI Policy and Investment Recommendations, due in May 2019.

In this context, we suggest a footnote clarifying that due to fast technological developments, the existing legal framework may need to be further developed and adapted to new requirements, such as with regard to cybersecurity and information security. When it comes to competition law, the authorities should be equipped with the necessary tools to intervene in cases of market abuse related to exclusive access to data and platforms and to address emerging issues such as algorithmic pricing.

Section B, Chapter I: Respecting Fundamental Rights, Principles and Values - Ethical Purpose

The GSMA agrees with the fundamental rights, high-level principles and correlating values identified in the consultation document. At the same time, the terminology and content of the guidelines should be fully aligned with the legal terminology of such concepts as human rights and strive to avoid extended interpretation of these well-established areas of law.

However, this response will focus more on the guidelines' potential implementation issues. Predictability on how to implement and monitor conformity with the guidelines is of paramount importance to mobile operators. In order to achieve the intended results, the GSMA proposes the following:

Providers of AI technology should be empowered to contextualise and make adjustments to suit various use cases.

AI is not legally defined in EU law. A definition should:

1. Exclude software systems based on traditional and determined algorithms which are clearly not based on AI.
2. Focus specifically on AI algorithms that require human supervision only when the purpose may constitute a risk to individuals' fundamental rights.
3. Capture the fact that the AI algorithm takes decisions as a consequence of the application of advanced analytical techniques (machine learning, deep learning and natural language processing) in combination with automation advanced feedback loops to solve problems.
4. Introduce a risk-based approach related to ethical issues:
 - a. Benign AI algorithms should not be submitted to ethical scrutiny. For example, AI algorithms that act as recommendation engines for audio-visual content, speech recognition or translation should not be submitted for ethics scrutiny (only GDPR rules).
 - b. Application of AI algorithms that may have legal or security effects on individuals should not be subject to ethics guidelines that duplicate or contradict their existing regulatory requirements via horizontal privacy/security laws.
 - c. AI algorithms that may take lethal decisions, i.e. AI algorithms for weapons (LAWS) should be excluded.

The GSMA considers the use of AI to fall into two broad categories: technology-focused AI and commercially focused AI. For the former, AI is used to assist with fault detection, predictive maintenance and network planning and optimisation, all of which enables operators to make more

efficient use of their physical assets. Use cases in this category often do not involve processing of personal data and have little direct impact on the fundamental rights of individuals. AI is also used for commercial purposes such as pricing promotions, predictive care, smart retail and through the deployment of virtual assistants (such as Tobi, the Orange-Deutsche Telekom chatbot on the Djingo smart speaker and the Telefónica Aura virtual assistant) and more.

The HLEG should establish at the outset that a one-size-fits-all approach is not appropriate, and that a determination of how ethics guidelines apply should be made on a case-by-case basis. For instance, evaluation of the ethical purpose will vary considerably between the use of AI in relation to virtual assistants and that of any of the areas mentioned in Chapter B.I.5 (critical concerns). The context in which AI is applied must always be kept in mind.

- From Fundamental rights to Principles and Values [p.5]

The proposed ethical principles represent a widely accepted approach to the development of AI, and they echo many of the principles released by GSMA member companies. We would also like to highlight that robust data governance mechanisms are essential for any business that focuses on data, including those that pursue AI solutions. The designation of a Data Protection Officer or Chief Privacy Officer, the adoption of strong policies and procedures, and a culture of compliance are precursors to the implementation of an ethical approach to AI.

The GDPR should provide confidence that personal data is processed according to the regulation and that data subjects' rights are respected. In addition, public authorities should ensure that the population has a basic understanding of what AI entails by encouraging educational institutions to teach this topic.

The HLEG proposes that 'informed consent' is a value needed to operationalise the principle of autonomy. In this context, the HLEG should note that current legislation does not require consent from individuals interacting with AI systems under all circumstances (cf. the principle of explicability). Both private and public sectors should be able to process personal data based on legal grounds other than consent, including when implementing AI technology. Consent is not necessarily a precondition for a human-centric or a privacy-friendly AI; instead, the balancing of interests required by the GDPR represents the proper approach. The GDPR allows for individual control in appropriate circumstances. For example, Article 22 allows data subjects to object to decision-making based solely on automated processing *when the decisions produce legal effects concerning the data subject or similarly significantly affects him or her*.

While consent can be one solution to guarantee accountability and transparency towards users, it is not the only one. According to the GDPR, processing of personal data (including for the purpose of offering an AI-based service) is permissible when it is justified by one or more of six different legal bases (including consent), such as processing necessary for the performance of a contract or for legitimate interest. In addition, the principle of compatible further processing (Article 6(4) GDPR) allows companies to use personal data for purposes other than the initial basis without the need for an additional legal basis. Consent is thus not the sole value and solution to enhance explainability. Voluntary approaches such as a one-pager that explains in simple terms the purpose for which personal data is being collected can enhance transparency

significantly. Therefore, the notion of informed consent is given too much prominence in these guidelines, creating a misleading perception that it is the only and best requirement to preserve autonomy and explainability.

The GSMA supports the recommendation that consent of the data subject should be obtained in many circumstances, e.g., for the use of facial recognition technology, in line with the current privacy and data protection laws. At the same time, it would not be appropriate to provide notice and require consent when facial recognition technology is used in the course of a criminal investigation, for example.

An important example for the mobile industry relates to mobile network optimisation, which customers have a right to expect as part of network service delivery. This will be a key area for the application of AI. Again, the guidelines should not create new rules and interpretations of the existing legislation (in this case regarding consent).

- The Principle of Beneficence: “Do Good” [p.8]

The GSMA supports the principle of explicability, while noting that some uses of AI technology will require explanation to data subjects and the public generally, and others will only require a business to explain elements of its technology to a regulator or expert body. Determinations regarding the level of explicability and transparency should be made according to the level of risk to individuals presented by an AI solution.

In keeping with this approach, a business’s use of internal review boards or consultation with independent experts should be considered good practice.

- The Principle of Non maleficence: “Do no Harm” [p.9]

Reference should be made to the ‘risk-based approach’ that underpins the GDPR, e.g., through the use of privacy by design, data protection impact assessments and the evaluation of data breaches to determine whether data subjects need to be notified, etc. Some of these concepts could easily be grafted onto the AI ethics guidelines to avoid reinventing the wheel. Some work has been done around the classification of harm in the privacy context, which could be leveraged for the AI context, although harms that would not directly affect individuals would need to be considered. This approach would also recognize that responsibility and flexibility may be more effective than regulation. In addition to the privacy principles, one could add that AI should contribute positively to the UN Sustainable Development Goals.

We refer the HLEG to the UN Guiding Principles Reporting Framework, which enshrines the duty of states to protect human rights and of corporate entities to respect human rights.¹

The second sentence of the section on ‘Do no Harm’ should therefore be amended to:

“AI implementations should *respect* the dignity, integrity, liberty, privacy safety and security of human beings in society and at work.”

- Principle of Justice [p. 10]

Instead of stressing “that AI systems must provide users with effective redress if harm occurs,” the guidelines should emphasize that ultimately humans are responsible. Operators of AI should

¹ <https://www.ungpreporting.org/framework-guidance/>

know and make clear who is responsible for which AI system or feature.

- The Principle of Explicability: “Operate transparently” [p. 10]

The guidelines need to be clear if explicability should be required for all AI systems or only for those that can potentially have a negative impact on their users if the wrong decision is taken. The principle of explicability should be proportionate to the level of harm that the AI system can cause.

Critical Concerns raised by AI

It is helpful to consider concerns, and we understand how different points of view can coexist. Here are some initial thoughts of the GSMA members:

- Identification without consent

It would be worthy to differentiate between ‘identification’ as a goal and as a side-effect. AI allows for easier identification of individuals (e.g. via facial recognition). The emphasis should thus be on not abusing this functionality. Use of AI-enabled identification and surveillance processes should follow current legal practices. The emphasis here should be on reliable anonymisation/de-identification methods and adhering to the GDPR.

- Covert AI systems [p. 11]

We generally support the recommendation that people should always know whether they are interacting with a human being or a machine. However, companies should have the flexibility to implement this requirement in the best way.

- Information on process, purpose and methodology of the scoring [p.12]

For complex systems, the GSMA has doubts that it will be practical to understand the logic developed by AI in order to explain to customers or even IT specialists how a decision is made by AI. In practice, decision-making in complex systems that do not use AI can be similarly difficult to understand, so the GSMA does not consider this to be a unique attribute of AI systems, although the self-learning nature of AI can be a barrier at scale. The GSMA would advocate that there is a focus on the learning process for AI systems, including strong human oversight on matters such as data set selection, target-setting and verification of results to ensure there is robustness and fairness in the automated processing implemented by AI.

Additionally, this situation needs to be aligned with the privacy right to know the logic behind automated processing.

- Tensions between individuals and society

The GSMA agrees with the remark in Chapter I.4 (p.8), that “tensions may arise between the principles when considered from the point of view of an individual compared with the point of view of society, and vice versa.” There will be AI solutions that are good for society (e.g. for the environment, disaster relief, healthcare, optimisation of public transportation, national security, immigration control) but that individuals may not perceive as bringing immediate personal

benefit. For instance, societal changes brought by an increased use of AI or production/business processes may lead to temporary unemployment, which may raise a negative perception of the use of AI in the population. Public authorities need to consider how to manage such societal changes, in partnership with the private sector and civil society. Regulators should be open to innovation and innovative AI solutions, and only intervene when the legal and human rights of individuals are at risk. Reliance on the EU Treaties and Charter, as well as existing and new case law where the aforementioned conflict of interest has been addressed, shall become significant.

It is of the utmost importance that private operators should not be expected to make a determination on where the correct balance lies between fundamental rights. Such determinations should only be made by actors with a clear public mandate to decide where the appropriate balance should lie — usually judicial authorities or elected officials.

As it is impossible to foresee all intended or unintended consequences, even with technical and nontechnical methods in place, as discussed in Chapter II.2, it would be good to recommend establishing an AI ethical committee at the governmental level. This could work like the ethic committees established in each Member State in the area of clinical trials according to Directive 2001/20/EC. The purpose would be to provide guidance and create debate about new uses of AI that impact people and societies at large.

Section B, Chapter II: Realising Trustworthy AI

The implementation and realisation of trustworthy AI is critical for achieving the desired outcomes of the guidelines. Clarity that allows for predictability, understanding and policing of the guidelines will pave the way.

- Implementation of the EU's Rights-Based Approach to Ethics [p.5]
While the EU's focus is correctly based on the EU Treaties and Charter of Fundamental rights, the guidelines should ensure that ethics are considered in relation to how organisations comply with the law, or how they should act where the law does not specify or address the specific context. In other words, the guidelines should focus on how the EU's rights-based approach to ethics should be implemented.

Ethical considerations and guidelines should not contradict legal requirements. Legal instruments, such as the The Universal Declaration of Human Rights (UDHR), the EU Charter of Fundamental Rights and the EU GDPR provide both terminology and basic requirements that will need to be reflected in the document. If EU legislation is changed, then such changes will need to be reflected in the document. It should also be noted that the same rights and protections should apply online as well as offline.

It should be made clear that these guidelines do not call for new requirements in law, but instead encourage good practice in developing and applying AI while appropriately safeguarding fundamental rights. The guidelines should enable organisations to identify and weigh good and bad outcomes to determine the best course of action. As the law will generally reflect the needs of a range of stakeholders, the European rights-based approach provides the natural point of

departure for guiding the principles and values that help to understand what ‘good’ and ‘bad’ practices may be.

- Accountability [as part of ‘Requirements of Trustworthy AI, p. 14]
‘Accountability’ as described in the consultation document seems to focus only on redress and remediation. However, accountability goes beyond that. There are already multi-stakeholder efforts in place that encourage good practice, including multi-stakeholder efforts such as the Partnership on AI.
- Data Governance [p.14]
The GSMA emphasizes that the GDPR provides a robust and comprehensive framework for the processing of personal data involved in AI solutions. GDPR provisions, which tailor rules to the sensitivity of data and how it is used, and include data subject rights, are sufficient to address data governance and privacy concerns related to AI. Existing privacy principles are relevant here, e.g., the quality of the data sets (adequate, not excessive) and avoiding bias (fairness, impact assessment, privacy by design). As with the GDPR, this requirement should not become a disproportionate burden when implemented.

The guidelines should ensure responsible approaches to data selection and training to avoid bias and discrimination. They should also encourage the necessary steps to ensure reliable AI performance.

- Design for all [p.15]
We are concerned that this principle lacks specificity: If something is intended to be prohibited or restricted, then the harms should be quite clearly articulated. There are unlawful forms of distinction (e.g., racial or gender-based discrimination) and lawful forms (e.g., differential pricing, age-restricted items). Is the intention to also limit the lawful distinctions? We ask the HLEG to provide additional clarity on lawful/unlawful differentiations in the context of the guidelines.
- Governance of AI Autonomy (Human oversight) [p.15]
Existing privacy principles may be helpful to consider here, such as GDPR Article 22 (referenced above). Companies should have the flexibility to decide how to best operationalise this requirement in a proportionate manner.

In our view, it is not necessary to guarantee human control at all levels (e.g., where AI is deployed deep in the network for fault detection). Human control may in some cases only be necessary in setting the outcomes, whereas in other cases, where there is a significant impact on individuals, human control is essential. The ‘conservative approach’ and ‘human oversight’ principles should help with this objective.

Human control and a stop-button failsafe may not be necessary precautions for all types of self-learning AI approaches. Indeed, if we mandate these types of control, even for AI that is deployed deep in our networks and has no human interaction or customer-facing element, we could deprive AI of its greatest potential: to solve problems (like cancer treatment, reversing global warming etc.) that humans have not been able to.

We request that the HLEG give more detailed thought to some of these subsidiary questions before including human control and/or stop buttons as a requirement in this section of the guidelines. A contextual understanding of AI, where different use cases are permitted with differing levels of human control, appears to us to be the optimum outcome.

- Respect for (& Enhancement of) Human Autonomy [p.16].

The notion that an AI system would result in abuse should lead to an obligation to re-assess the requirements for Trustworthy AI as described in chapter III.

We note that AI services are already deployed successfully in recommendation systems, as used in e-commerce sites and media consumption, as well as services such as search engines. The suggestion to allow the user to specify preferences and limits for system intervention is something we believe is covered already under the GDPR as part of the consent requirements as well as ePrivacy regulations. In practice, as AI solutions become more sophisticated, we have concerns that it will be difficult to provide fine user controls that influence the outcomes of AI solutions. We think it would be better to focus this section on the following ethical practices:

- That AI not be designed to deceive users, for example by the creation of virtual users, customers, reviews, etc., which manipulate the behaviour of users based on peer views;
- That AI be designed to be fair, e.g., not to suppress bad reviews of a product or service, or bias results in such a way so as to purely maximise profits at the expense of customer requirements or interests; and
- That bias (racial, gender, age, etc.) be knowingly engineered out of AI systems both when the AI system is originally designed and as learning adapts.

Earlier we noted that there are legitimate applications for AI such as fraud prevention and network optimisation that will benefit mobile users greatly. We think it would negatively affect users if there was the ability to opt out of such beneficial AI services.

- Robustness — Resilience to Attack [p. 17]

The requirements described in the guidelines regarding resilience and robustness apply to AI systems as well as to any ICT system (e.g., IoT systems). Having said that, the guidelines would benefit from further consideration of the precautions that can be taken to raise the security level of AI systems. The highest security requirements should apply in AI development and application. All security features such as notification of security vulnerabilities, emergency stop buttons or security updates should be aimed towards a clear attribution of responsibility. Besides the risk of weak spots being exploited by hackers, the self-learning capabilities of corrupted AI systems raise the risk of damage exponentially. For security in the development and application of AI, this specifically means that ensuring IT security is a key requirement for product safety of AI applications or products that implement AI applications. This correlation must always be considered by developers and industrial users (i.e., security by design). Mandatory risk assessments analogous to the data protection impact assessment of the GDPR could contribute to highly sensitive AI applications, as in healthcare. The current regulatory focus

on operators of critical IT infrastructures as in the ICT, healthcare or energy sectors, is no longer sufficient because critical issues arise increasingly on an ad-hoc basis. This would for example be the case in a multitude of connected, self-driving vehicles.

- Respect for Privacy [p.17]

The GDPR is principles-based, and this enables it to accommodate new technologies including AI. The GDPR is also based on the identification of risk of harm and on the concept of accountability so that organisations are encouraged to adopt technological and operational measures to control risk, including privacy by design, data-privacy impact assessments, the appointment of a Data Protection Officer, good record-keeping and being able to demonstrate compliance. To the extent that an AI deployment makes use of personal data, it is already regulated by the GDPR. The draft guidelines should therefore acknowledge this explicitly and avoid the duplication of requirements, which could cause uncertainty. Most of the harms discussed in the draft guidelines are in fact privacy harms. Privacy and data protection are separate issues. Some of the principles of the GDPR would have to be extended to accommodate a range of harms, such as harm to groups of individuals, harm to society, harm to the environment, and organisational responses to the GDPR would have to be adapted accordingly. However, the core mechanisms of basing requirements on flexible principles and the identification of risk of harm are already in the GDPR and could be extended to or replicated in the AI context. The requirements proposed in the draft guidelines are also very similar to privacy requirements (e.g., transparency, explicability, right to know when data is being collected) and do not need to be restated or reinvented just because AI is processing the personal data.

The GDPR also recognizes that pseudonymisation and encryption, which facilitate beneficial uses of data while reducing the risk of harm to individuals' privacy are good practices. We therefore suggest these safeguards as additional technical methods in the guidelines. Pseudonymisation has an advantage vis à vis anonymous data, namely that the necessary 'identifiers' remain intact for big data applications, to be able to merge large amounts of data from various sources. The technique thereby eliminates the direct link between the data and the data subject, while the pseudonym used as an identifier allows to repeatedly merge data from different sources over a period of time. This is a key requirement for valuable data-driven services, also in the field of AI. Private companies can thus play a role through the introduction of technical measures which reduce the need for difficult tradeoffs. In the context of the ePrivacy reform, for example, pseudonymisation could be deployed to ensure that electronic communications data can be used without impinging on fundamental rights.

In addition, the challenges related to AI and the principles of data minimisation and purpose limitation should be emphasised. AI will also inherently challenge the principle of transparency. If absolute transparency is a condition, it will rule out deep learning.

- Transparency [p.18]

Again, much can be learned from EU data protection law, where processing of data is intended for a specified purpose (purpose limitation). Although "informed consent should be sought" sounds appealing, it is not always the most appropriate way to protect people, and it is

extremely difficult in practice. The guidelines should not reinvent the wheel vis à vis existing legislation. Overreliance on informed consent could lead to people agreeing to everything or, on the other hand, create an insurmountable barrier to the good outcomes that AI could achieve.

However, if a consumer's data is being used to make decisions about that person through the use of AI, they should be informed that this is happening. Where personal data is being used, an individual already has the right to opt out under the GDPR, so there is no need for an additional requirement.

The context of AI use is very relevant here. If AI is used within a communications network to improve energy efficiency or routing, there should be no need for a consumer to have a right to opt in or out of such use.

Regarding accountability, the guidelines should be clear on the kind of accountability intended. If this merely implies that organisations will be held liable under the law, then the issues to be explored are: Who is liable and under what circumstances? But, most important, the guidelines are not the right legal instrument to determine liability. If they are a call to organisations to hold themselves to a high ethical standard regardless of the law, then this should be the starting point for a chapter about how to enable ethical decision-making in practice.

The section talks about explaining how the system makes decisions, rather than explaining the system's decision. This should be more clear. Neural networks are difficult to explain, but their decisions can, to some extent, be explained.

Technical methods

- Ethics & Rule of law by design (X-by-design) [p.19]

A clear distinction should be made between safety-critical or ethics-critical systems and non-critical systems when demanding failsafe shutdown mechanisms and robustness vs adversarial examples. Demanding this for all AI systems (e.g., a music recommender) does not seem practical and will hinder innovation and commercial adoption of AI.

- Architectures for Trustworthy AI [p.19]

Integrating ethical goals and requirements at sense level for adaptive and learning systems is certainly the preferred way, but not the only way. Unwanted actions could also be filtered out. The latter is probably much easier.

- Testing & validating [p.20]

Testing within predictable bounds should be made offline, before deployment. It is even more important to monitor in the real world continuously.

Section B, Chapter III: Assessing Trustworthy AI

The GSMA considers the menu of potential assessment questions posed by the HLEG to be helpful to entities developing and using AI technologies in a manner consistent with human rights. It is important that any approach to assessment be flexible, reflecting the different types of AI solutions companies pursue, including those that have little direct impact on individual rights.

- Ethics in autonomous systems and time-scales

The GSMA and its members believe it to be instructive to look at ethics and intelligent autonomy based on time scales. For example, do we allow AI systems to take full control and decision autonomy below a certain time limit beyond the human capability (e.g., below 500 milliseconds)? What situations do we allow this to happen in? What about the time scales in which humans are able to intervene in automated decisions, or what if they do want to be in complete control? What about time scales where autonomous decisions may become reversible or iterative as more information becomes available? Just because there is sufficient time to reverse a decision, there may still be domains where AI should not to make such decisions without human oversight. However, below the 1-minute threshold, there may be many situations where it could be crucial to let the machine take control. This aspect has not been addressed at all, but it is very important. For example, network management functions such as beamforming in 5G networks — aimed at increasing spectrum efficiency — will require autonomous systems to make decisions in fractions of a second in order to ensure uninterrupted connectivity.

Under existing data protection rules, it is already incumbent upon organisations to consider points of ethics and fairness in order to avoid harm. These require significant assessments, processes and record keeping. As mentioned before, any operational or practical guidelines developed in the context of AI should be aligned to the greatest extent possible in order to minimise duplication.

- Technical robustness is most likely to lead to positive outcomes if underpinned by sustainable business models

For AI systems to be technically robust and reliable, sustainable business models are key. In the case of the mobile industry, intangible assets (data, insights, analysis, services) can be transferred from a mobile operator for use by a demand-side agency under mutually beneficial terms that enable an ongoing relationship between both partners. This aspect of sustainability allows for robust, repeatable and replicable use of mobile big data across different geographies and use cases, underpinned by secure funding that enables continuity in supply and analysis of the data. For more information please see

<https://www.gsma.com/betterfuture/resources/sustainable-business-models-report>

- Assessment List

The GSMA supports the objective of this chapter to provide a practical checklist of questions to ensure the development of trustworthy/human-centred AI from an early stage of the development cycle.

However, our concern with this section is that while the questions here are appropriate considerations, they lack the technical detail and specificity that would make them useful or practical for AI product developers or engineers, particularly within a small organisation. In general, there is too much repetition of the early sections of the guidelines and not enough

effort to provide a clear structure for AI developers that can be easily understood in a variety of languages and for different AI use cases. The list is currently too broad and covers many requirements that are not specific to AI. The work would benefit from being even more specific about new challenges.

The next steps could also benefit from including some less sensitive use cases of AI. The current four are all use cases where everyone agrees on the high risks towards safety, trust and ethics. It would be good to see how the assessment list would look for less risky use cases, such as customer service applications, marketing or similar.