

Analysis of EU AI Office stakeholder consultations: defining AI systems and prohibited applications

Final study report

Written by
CEPS - Centre for European Policy Studies
Place du Congrès 1, B- 1000 Brussels
Tel. +32 2 229 39 11
www.ceps.eu



EUROPEAN COMMISSION

Directorate-General for Communications Networks, Content and Technology (CNECT)

Artificial Intelligence Office

Contact: Martin Ulbrich

E-mail: martin.ulbrichec.europa.eu

European Commission

B-1049 Brussels

LEGAL NOTICE

The information and views set out in this document are those of the authors and do not necessarily reflect the official opinion of the Commission. The Commission does not guarantee the accuracy of the data included in this document. Neither the Commission nor any person acting on the Commission's behalf may be held responsible for the use which may be made of the information contained therein.

© European Union, 2025

Reproduction is authorised provided the source is acknowledged.



The Commission's reuse policy is implemented by Commission Decision 2011/833/EU of 12 December 2011 on the reuse of Commission documents (OJ L 330, 14.12.2011, p. 39 – <https://eur-lex.europa.eu/eli/dec/2011/833/oj>).

Unless otherwise noted (e.g. in individual copyright notices), the reuse of this document is authorised under the Creative Commons Attribution 4.0 International (CC BY 4.0) license (<https://creativecommons.org/licenses/by/4.0/>). This means that reuse is allowed, provided appropriate credit is given and any changes are indicated.”

More information on the European Union is available on the Internet (<http://www.europa.eu>).

Luxembourg: Publications Office of the European Union, **2025**

KK-01-25-028-EN-N

ISBN 978-92-68-25693-0

DOI 10.2759/6218665

© European Union, **2025**

Table of contents

1.0	EXECUTIVE SUMMARY	7
2.0	DEMOGRAPHIC DATA OF RESPONDENTS	8
3.0	SECTION 1: QUESTIONS IN RELATION TO THE DEFINITION OF AN AI SYSTEM	16
3.1	Question 1: Elements of the definition of an AI system	16
3.2	Question 2: Simple software systems out of scope of the definition of an AI system ..	23
4.0	SECTION 2. QUESTIONS IN RELATION TO THE PROHIBITIONS (ARTICLE 5 AI ACT)	24
4.1	QUESTIONS IN RELATION TO HARMFUL SUBLIMINAL, MANIPULATIVE OR DECEPTIVE PRACTICE	24
4.2	QUESTIONS IN RELATION TO HARMFUL EXPLOITATION OF VULNERABILITIES	29
4.3	QUESTIONS IN RELATION TO UNACCEPTABLE SOCIAL SCORING PRACTICES	32
4.4	QUESTIONS IN RELATION TO INDIVIDUAL CRIME RISK ASSESSMENT AND PREDICTION	36
4.5	QUESTIONS IN RELATION TO UNTARGETED SCRAPING OF FACIAL IMAGES	40
4.6	QUESTIONS IN RELATION TO EMOTION RECOGNITION	43
4.7	QUESTIONS IN RELATION TO BIOMETRIC CATEGORISATION	47
4.8	QUESTIONS IN RELATION TO REAL-TIME REMOTE BIOMETRIC IDENTIFICATION.....	51
4.9	QUESTION IN RELATION TO INTERPLAY WITH OTHER UNION LEGISLATION.....	55
5.0	References.....	57

Table of Figures

Figure 1: Distribution of respondent types	9
Figure 2: Distribution of the geography of respondents	10
Figure 3: Country of headquarters/residence	11
Figure 4: Distribution of the geography of respondents	11
Figure 5: Distribution of organisation size.....	12
Figure 6: Stakeholder distribution	13
Figure 7: Sectoral distribution	14
Figure 8: Clarification Needs #1	16
Figure 9: Clarification Needs #2	17
Figure 10: Clarification Needs #3	18
Figure 11: Clarification Needs #4	19
Figure 12: Clarification Needs #5	20
Figure 13: Clarification Needs #6	20
Figure 14: Clarification Needs #7	21
Figure 15: Do you know examples of AI systems	26
Figure 16: Examples of AI systems needing clarification.....	27

1.0 EXECUTIVE SUMMARY

This report analyses the results of stakeholder consultations conducted by the EU AI Office regarding two critical aspects of AI regulation: **the definition of AI systems and prohibited AI applications** (European Commission, 2024). This report synthesises stakeholder feedback to 88 *questions* and informs the development of clear, practical, and effective AI regulatory frameworks. The report was drafted by the Centre for European Policy Studies (CEPS).

The respondent demographics showed a strong industry and technical stakeholder bias, with over 47.2% representing industry organisations, AI providers, and deployers, while ordinary citizens comprised only 5.74%. This skew suggests the consultation may not fully represent the interests of those most affected by AI systems. The consultation also shows notable geographic concentration in major EU economies.

The consultation revealed **significant concerns about the clarity and scope of the definition of "AI system" in the EU AI Act**, with stakeholders strongly emphasising the need for **more precise definitions of terms like "adaptiveness," "inference," and "autonomy."** A major concern was the potential for current definitions to **inadvertently include traditional software** systems that do not exhibit true AI characteristics.

Regarding prohibited practices, stakeholders expressed **particular concern about emotion recognition systems, biometric categorisation, and social scoring practices**. There was strong consensus that clearer guidelines are needed to **distinguish between legitimate uses (such as in healthcare and safety) and prohibited applications** that could infringe on fundamental rights.

The consultation highlighted significant **ambiguity around what constitutes "manipulation" and "significant harm" in AI systems**. Stakeholders called for specific examples and clearer thresholds to help organisations understand when their AI applications might cross into prohibited territory, particularly in marketing and consumer interaction contexts.

Privacy and data protection emerged as critical concerns, especially regarding facial recognition and biometric data collection. Respondents emphasised the need for better alignment between the AI Act (2024) and existing regulations like GDPR (2016), calling for clear guidelines on how these frameworks interact.

The responses revealed particular concern about AI systems' potential to perpetuate **discrimination and bias, especially in law enforcement, employment, and financial services**. Stakeholders urged for robust safeguards and oversight mechanisms to prevent the misuse of AI technologies against marginalised communities.

Real-time biometric identification emerged as a contentious area, with stakeholders seeking clearer definitions of terms like "real-time," "publicly accessible spaces," and "law enforcement

purposes." There was strong emphasis on the need for strict limitations and oversight to prevent mass surveillance.

The consultation identified **significant challenges for small and medium-sized enterprises (SMEs) in complying with the AI Act**. Respondents called for practical guidance and support mechanisms to help smaller organisations navigate the regulatory requirements without stifling innovation.

Based on these findings, the EU AI Office needs to prioritise developing detailed guidelines that clearly define technical terms and provide concrete examples of prohibited practices. These guidelines should include specific use cases illustrating the boundaries between acceptable and unacceptable AI applications.

The EU AI Office should also focus on creating a balanced regulatory framework that protects fundamental rights while enabling beneficial AI innovation. This includes developing clear compliance pathways for SMEs, establishing robust oversight mechanisms for high-risk applications, and ensuring better alignment with existing EU regulations.

The report presents a comprehensive analysis of responses to each of the 88 questions of the stakeholder consultation, organised into nine key sections. The analysis begins with demographic data about respondents (Section 2), followed by detailed examinations of stakeholder views on AI system definitions (Section 3), prohibited practices (Section 4) focusing specifically on harmful subliminal, manipulative or deceptive practices (4.A), harmful exploitation of vulnerabilities (4.B), unacceptable social scoring practices (4.C), individual crime risk assessment and prediction (4.D), untargeted scraping of facial images (4.E), emotion recognition (4.F), biometric categorisation (4.G), real-time remote biometric identification (4.H) and the interplay with other EU legislation (4.I). Each section systematically presents quantitative response distributions alongside qualitative analysis of stakeholder concerns, recommendations, and specific examples provided.

The results are produced under a contract with the European Commission. The opinions expressed are those of the Centre for European Policy Studies only and do not represent the European Commission's official position.

2.0 DEMOGRAPHIC DATA OF RESPONDENTS

Question 0.1: "1. Do you represent one or more organisations (e.g., industry organisation or civil society organisation) or act in your personal capacity (e.g., independent expert)?"

Out of the total responses, a substantial majority, approximately 80%, indicated that they represent one or more organisations, highlighting the collective nature of the input received. This suggests that the perspectives shared are likely influenced by organisational agendas, priorities, and experiences, which may shape the overall discourse. In contrast, a smaller

segment of respondents, less than 20%, stated that they are acting in a personal capacity. Independent experts and individuals may provide insights based on personal experiences rather than organisational affiliations and their contributions could offer unique viewpoints that diverge from the collective stance of organisations, potentially enriching the discussion with diverse, more grassroots, perspectives.

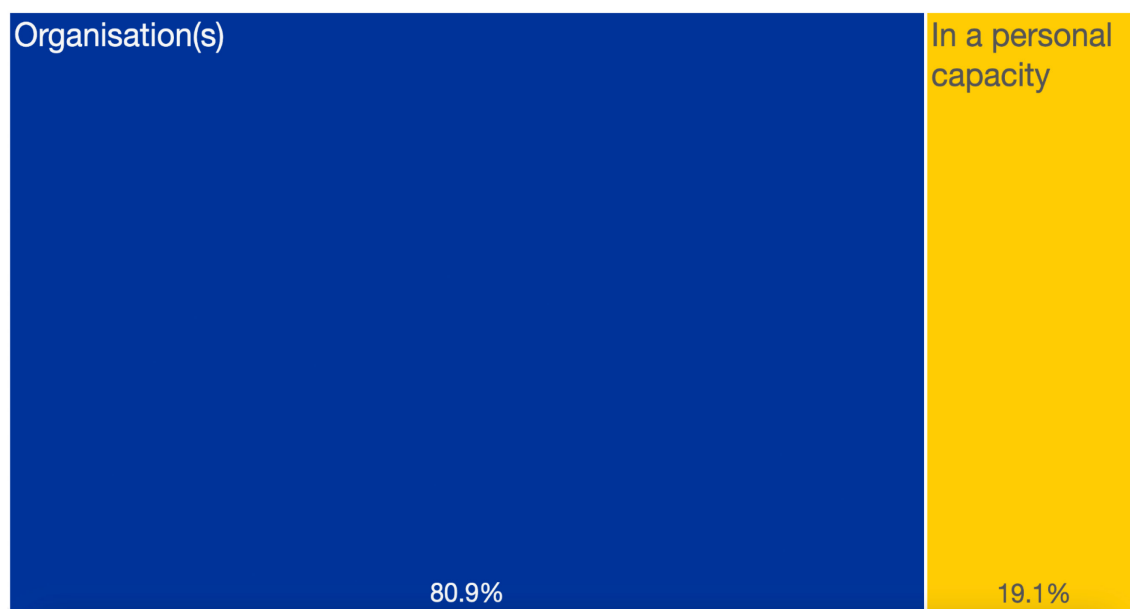


Figure 1: Distribution of respondent types

Question 0.2: "If you are organisation(s), please specify the name(s):"

The responses came from a diverse ecosystem of organisations, with major technology companies (such as Microsoft,) alongside digital rights advocacy groups (Electronic Frontier Foundation, Access Now), financial institutions (Alior Bank, French Banking Federation), healthcare organisations (MedTech Europe), consumer rights groups (BEUC- The European Consumer Organisation), and public policy bodies (CNIL). This mix of commercial and non-profit entities, spanning from AI developers to rights advocates, suggests a broad engagement of this consultation across multiple sectors, with particularly strong representation from technology and advocacy organisations.

Questions 0.3; 0.4; 0.5 and 0.6 are standard administrative questions (affiliation, first name, last name, email) used to establish the identity and contact information of survey participants but are not relevant to this report and therefore not analysed.

Question 0.7: "Are you headquartered/residing in the EU?"

A significant majority of respondents affirming their presence within the EU. Specifically, 321 respondents (approximately 84%) indicated "Yes," while 54 respondents (about 14%) answered "No." Only about 2% answered "Other". This strong EU representation shows that the target audience of EU stakeholders is reached. The consultation collected valuable insights

from those directly affected by EU policies. It also introduces geographic bias and might limit the survey's global applicability.

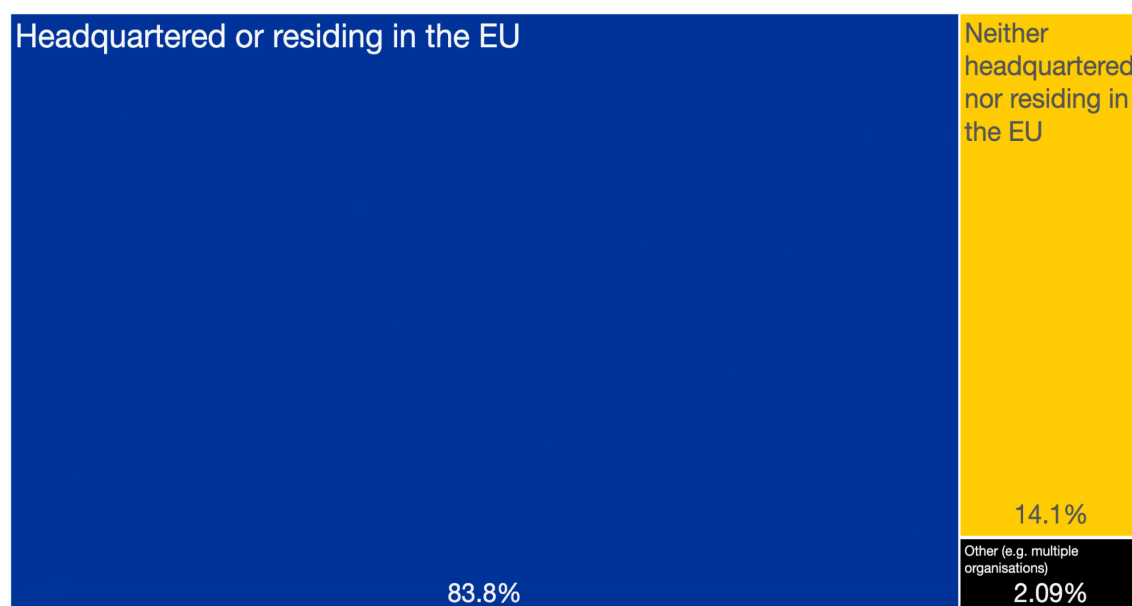


Figure 2: Distribution of the geography of respondents

Question 0.8: "Headquarter / Country of residence"

The consultation captured responses from 22 EU Member States, representing a broad coverage of the European Union, with only five member states not represented: Bulgaria, Cyprus, Latvia, Lithuania, and Malta. Looking at the distribution¹ of EU Member States respondents, Germany leads with the highest participation (72 respondents), accounting for 18.8% of the total, Belgium follows with 58 respondents (15.1%) and France ranks third with 39 respondents (10.2%). In Non-EU Countries the United States provided 25 respondents (6.53% of the total) and the United Kingdom contributed 10 respondents (2.61%). This distribution shows a significant representation from the major EU economies, while also including valuable input from key non-EU partners like the US and UK.

¹ We reclassified responses as 'Headquarter / Country of residence = Yes' when initially answered as a 'No' but the country was Belgium, Germany, or Ireland, and as non-EU (No) if it was Japan or the United Kingdom but initially answered as a 'Yes'.

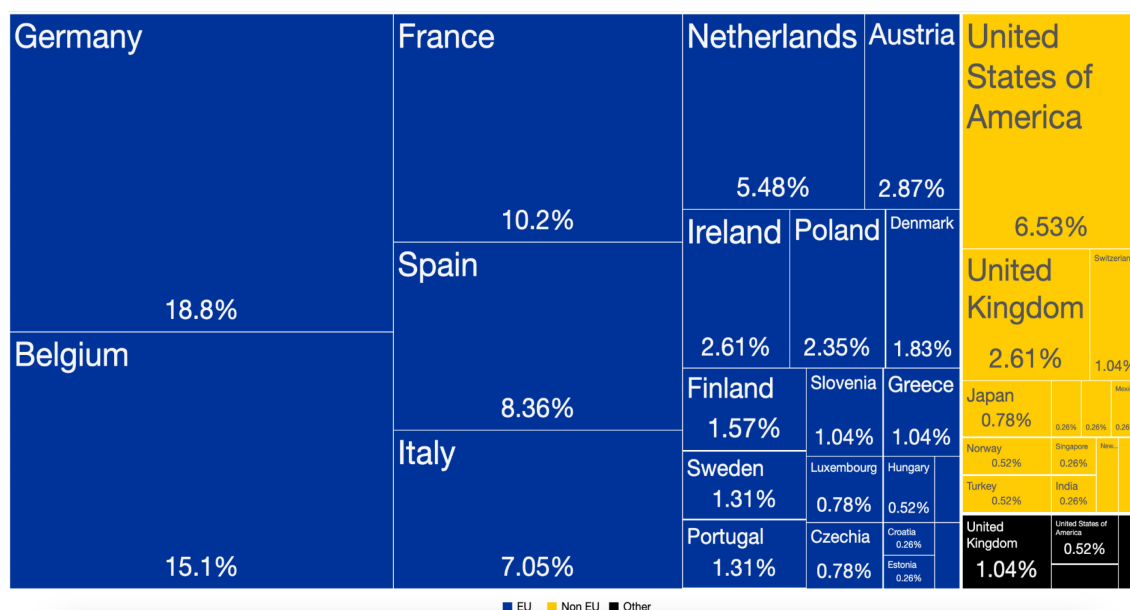


Figure 3: Country of headquarters/residence

Question 0.9: "Do you have an office or other kind of representation in the EU?"

The results regarding EU business/physical presence show that a relative majority (46.7%) of organisations maintain a subsidiary, branch office, or similar establishment within the European Union. About one-third of respondents (32.3%) indicated this question was not applicable to their situation. A smaller portion (13.4%) explicitly stated they have no EU presence, while a small segment (7.61%) reported having some other form of EU presence not covered by the main categories.

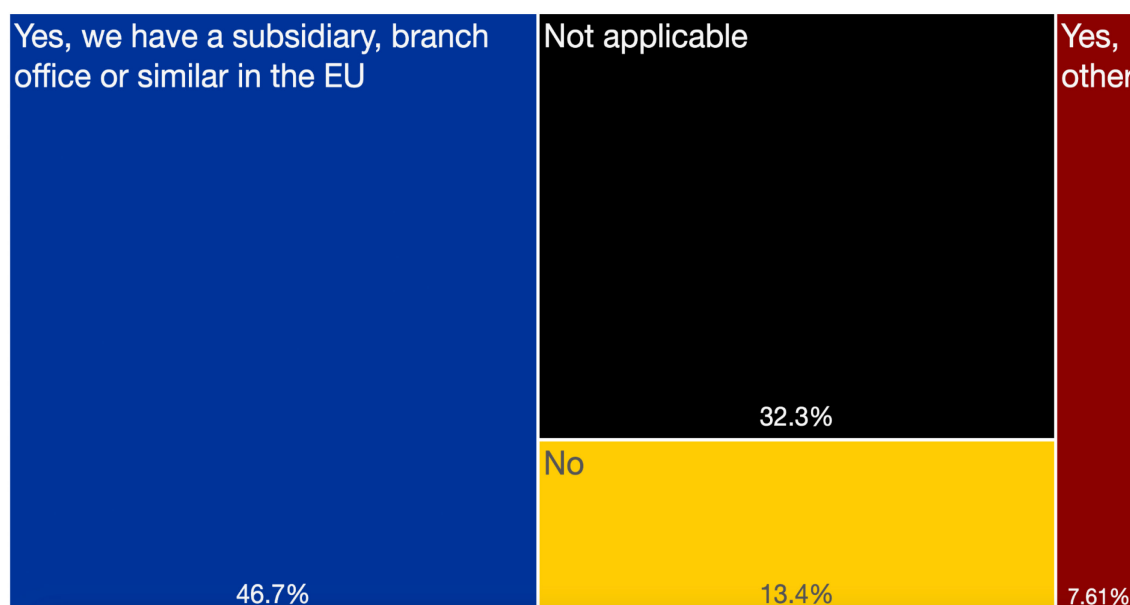


Figure 4: Distribution of the geography of respondents

The follow-up Question 0.10 "If applicable, please specify" yielded limited analysable data, with 75% of responses being NA (not available). The remaining free-text responses did not provide substantive information relevant for this analysis.

Question 0.11: "If you are an organisation, what is the size of your organisation and does it qualify as a small or medium sized enterprise according to the EU recommendation 2003/361, if applicable?"

The responses to the question related to organisation size shows that the largest segment is "Not applicable" (28.2%), followed by large organisations at 24.3%, small enterprises at 20.6%, and medium-sized organisations at 10.7%, with the remaining 16.2% falling into the "Other" category. This distribution suggests that among the applicable respondents, there's a relatively balanced mix of organisation sizes, though SMEs (Small and Medium Enterprises) have a slightly higher representation compared to larger enterprises when combined (31.3% vs 24.3%).

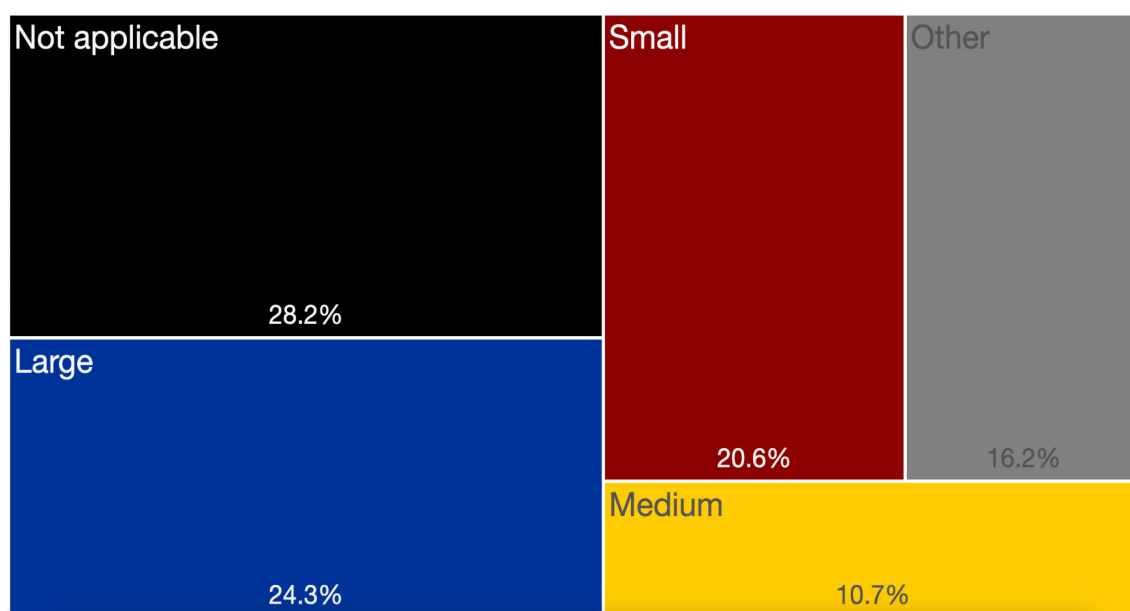


Figure 5: Distribution of organisation size

The follow-up Question 0.12 "If other, please specify" yielded limited analysable data, with 82.7% of responses being NA (not available). The remaining free-text responses did not provide substantive information relevant for this analysis.

Question 0.13: "Which stakeholder category would you consider yourself in?"

The stakeholder distribution shows a diverse range of participants, with the largest group being "Other industry organisation or acting on behalf of such organisations" at 20.9%, followed by "Provider of an AI system" at 15.1%, "Civil Society Organisation" at 14.1%, "Others" at 13.3%, "Deployer of an AI system" at 11.2%, "Academia" at 10.7%, "Public authority" at 8.09%, and "Citizen" at 5.74% - indicating that the survey captured input from across the AI ecosystem, with particularly strong representation from industry stakeholders and AI system providers who collectively make up over 26% of respondents. The stakeholder representation in this survey appears significantly skewed and not representative of the general population. Industry organisations, AI providers, and deployers collectively make up 47.2% of respondents, while ordinary citizens - who represent the vast majority of those affected by AI systems - account

for only 1 every 20 responses. This suggests a heavy over-representation of industry and technical stakeholders, and an under-representation of the public in this consultation exercise.

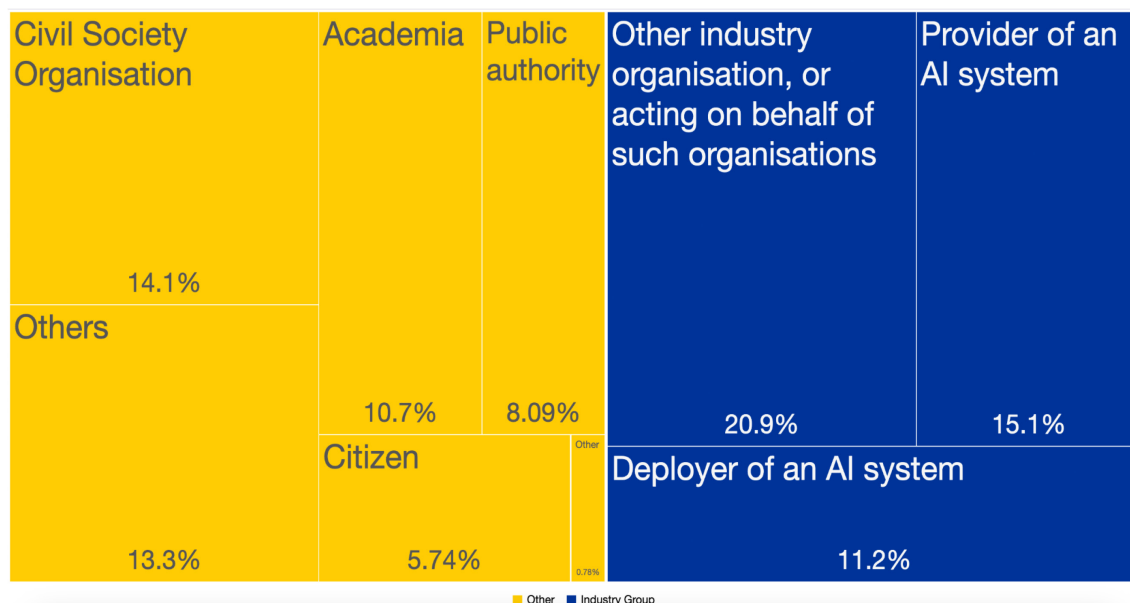


Figure 6: Stakeholder distribution

The follow-up Question 0.14 "If other, please specify" yielded limited analysable data, with 78% of responses being NA (not available). The remaining free-text responses did not provide substantive information relevant for this analysis.

Question 0.15: "In which sector do you operate?"

The sector distribution reveals a strong tech-centric bias, with Information Technology dominating at 15.4%, followed by Banking/Finance (7.93%), Healthcare (7.19%), and public sector (8.18%). This distribution shows an over-representation of tech and digital sectors compared to their actual share of the overall economy, while traditional sectors like Manufacturing (4.96%), Retail (2.48%), and Transport (3.22%) are notably under-represented relative to their economic importance. The high representation of IT and finance sectors aligns with the earlier finding that industry organisations and AI providers were the dominant respondents, suggesting the survey may have primarily reached technology-focused stakeholders rather than capturing a balanced cross-section of the economy's sectors that are or will be impacted by AI systems.

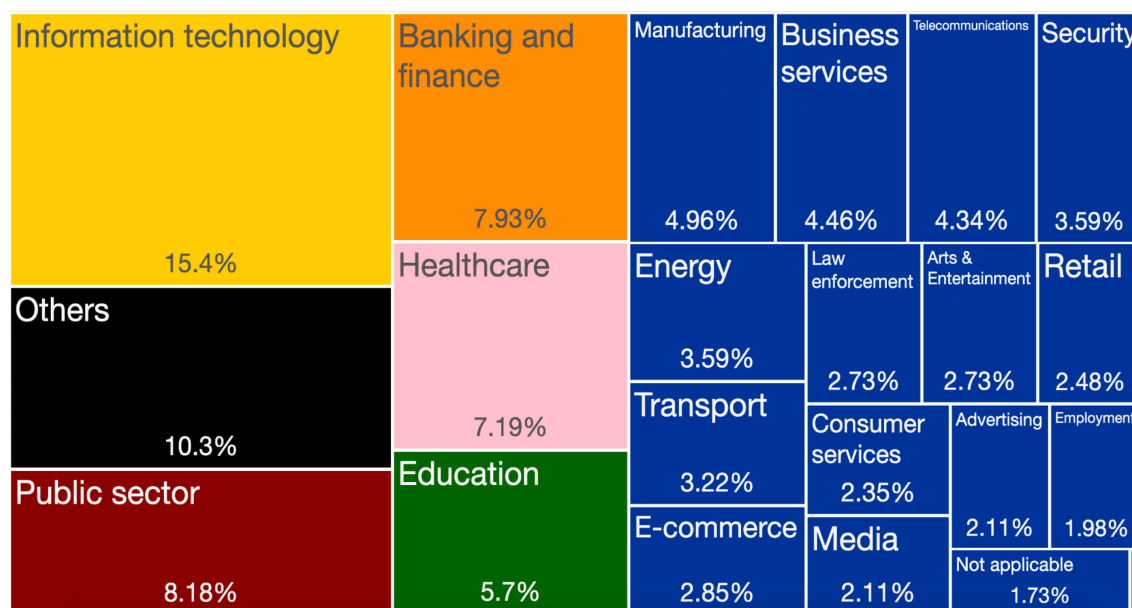


Figure 7: Sectoral distribution

The follow-up Question 0.16 "If other, please specify" yielded limited analysable data, with 78% of responses being NA (not available). The remaining free-text responses did not provide substantive information relevant for this analysis.

Question 0.17: "Please briefly describe the activities of your organisation or yourself:"

The responses confirm the earlier observation of a diverse array of organisations and individuals engaged in various sectors, primarily focusing on technology, finance, healthcare, and public policy. Many organisations, such as IDEMIA Public Security and Microsoft, emphasise their roles in developing and implementing advanced biometric and AI technologies to enhance security and efficiency in public and private sectors. Financial institutions and associations highlight their involvement in shaping AI regulations and compliance frameworks, particularly in response to the EU AI Act. Additionally, several NGOs and advocacy groups focus on human rights, ethical AI use, and the protection of vulnerable populations, indicating a strong commitment to social responsibility.

The responses also indicate a strong emphasis on technological innovation and research, particularly in AI and data analytics.. Academic institutions and research groups are actively engaged in exploring the intersection of AI, law, and ethics, contributing to the discourse on responsible AI governance. This focus on innovation suggests a recognition of AI's transformative potential and the importance of aligning technological advancements with societal values.

Many respondents emphasise the importance of collaboration among stakeholders, including industry, academia, and civil society, to address the challenges posed by AI technologies. Initiatives like the AI Policy Lab at Umeå University and the Digital New Deal think tank aim to foster dialogue and cooperation in shaping AI policies that reflect European values. The involvement of trade associations and NGOs in advocacy efforts highlights a collective

approach to navigating the complexities of AI regulation and ensuring that diverse perspectives are considered in policy-making processes.

Question 0.18: "Is your organisation submitting a collective answer on behalf of other organisations?"

Most respondents (67%) submitted individual responses rather than collective ones, with only 18.36% representing collective organisational responses, suggesting most feedback came from individual entities rather than industry groups or associations.

The follow-up Question 0.19 "If other, please specify" yielded limited analysable data, with 86.6% of responses being NA (not available). The remaining free-text responses did not provide substantive information relevant for this analysis.

Question 0.20: "For natural persons: Contribution publication privacy settings If you act in your personal capacity: All contributions to this consultation may be made publicly available. You can choose whether you would like your details to be made public or to remain anonymous."

The privacy preferences for individual respondents show that among those who answered, there was a fairly even split between those willing to be publicly identified (24%) and those preferring anonymity (21.3%), while the majority (54.2%) indicated the question was not applicable to them, suggesting they were likely responding on behalf of organisations rather than as individuals.

Question 0.21: "For organisations: Contribution publication privacy settings If you represent one or more organisations: All contributions to this consultation may be made publicly available. You can choose whether you would like respondent details to be made public or to remain anonymous."

Looking at the privacy preferences for organisational responses, there's a relatively balanced distribution between those choosing full transparency (46%) who agreed to publish both organisation and respondent details, and those preferring partial anonymity (39.1%) who agreed to publish organisation details but keep respondent names private. Only 15% indicated the question was not applicable, which aligns with the earlier data showing most respondents were representing organisations. This high level of willingness to be identified (85% total willing to have at least organisational details public) suggests a good degree of transparency and accountability in the consultation process from organisational stakeholders.

3.0 SECTION 1: QUESTIONS IN RELATION TO THE DEFINITION OF AN AI SYSTEM

3.1 Question 1: Elements of the definition of an AI system

Question 1.22: *"Please rate the importance of further clarification from 1 to 10, 10 indicating 'most important' for "a machine-based system"*

The results in Figure 8 show a bimodal distribution in responses, with the highest peak of 125 responses at importance level 1 (least important) and a smaller peak of 37 responses at level 10 (most important), suggesting a strong polarisation in views about the need for clarification of *machine-based systems*. While most respondents see little need for further clarification, a significant minority considers it critically important.

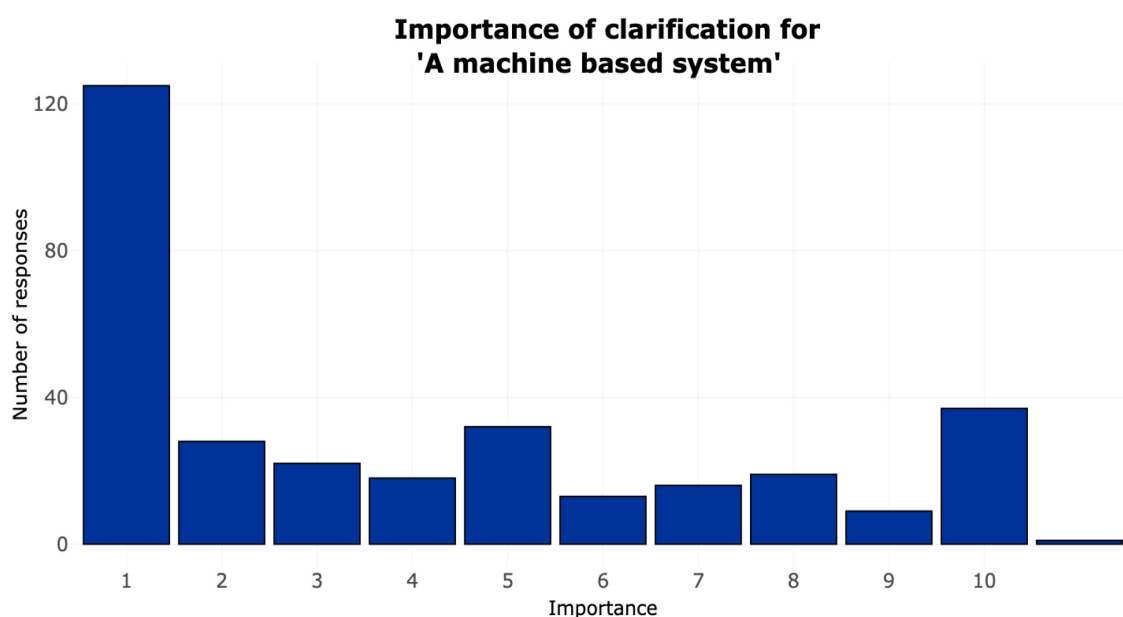


Figure 8: Clarification Needs #1

Question 1.23: *"Please rate the importance of further clarification from 1 to 10, 10 indicating 'most important' for "designed to operate with varying levels of autonomy"*

The results in Figure 9 show a strong consensus that clarification is crucial for systems with varying levels of autonomy - while a small minority see little need for clarification, the overwhelming majority of respondents consider it highly important, with particularly strong representation in the 8-10 range (205 responses, which represents more than 61% of the responses), indicating broad agreement on the need for clear understanding of autonomous system operations. This pattern is notably different from the previous graph, showing that

respondents view clarification as significantly more important when dealing specifically with autonomous systems compared to general machine-based systems.

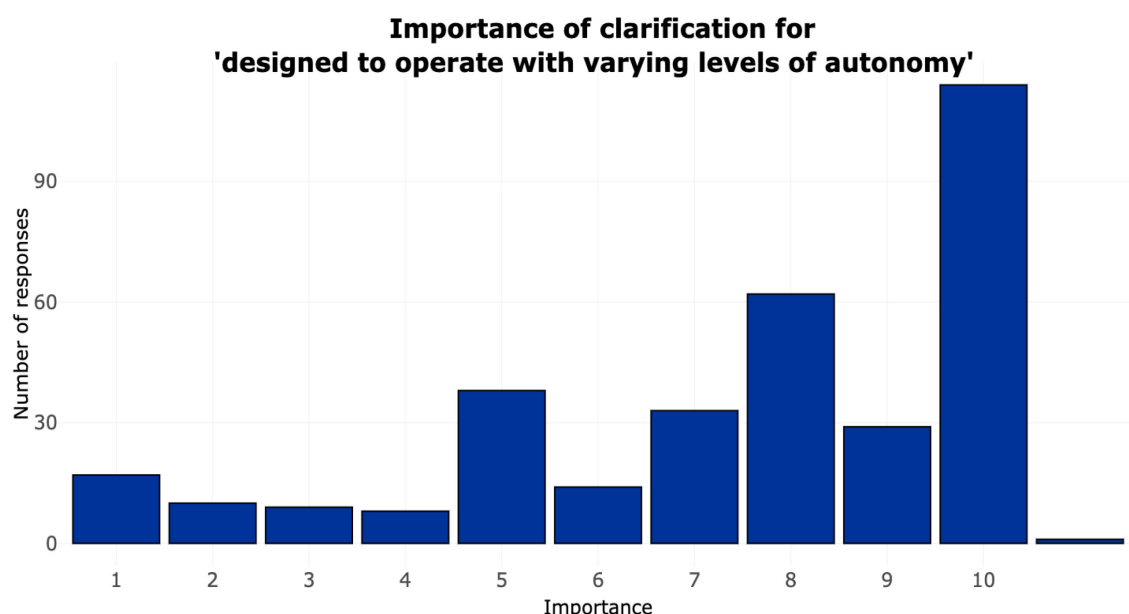


Figure 9: Clarification Needs #2

Question 1.24: "Please rate the importance of further clarification from 1 to 10, 10 indicating 'most important' for "may exhibit adaptiveness after deployment"

The graph (Figure 10) shows an interesting bimodal distribution with two distinct peaks - one of 48 responses at importance level 1 (least important) and a dominant peak of 111 responses at level 10 (most important), suggesting that while there is some polarisation, there's a clear majority favouring the critical importance of clarification for systems that may exhibit adaptiveness after deployment. The middle ranges (2-9) show a gradual upward trend, with most values between 20-35 responses, indicating that respondents generally lean towards higher importance when dealing with adaptive systems, likely due to the unpredictable nature of post-deployment changes. This pattern reveals that respondents are particularly concerned about understanding systems that can change their behaviour after deployment, reflecting potential risks and complexities associated with adaptive capabilities.

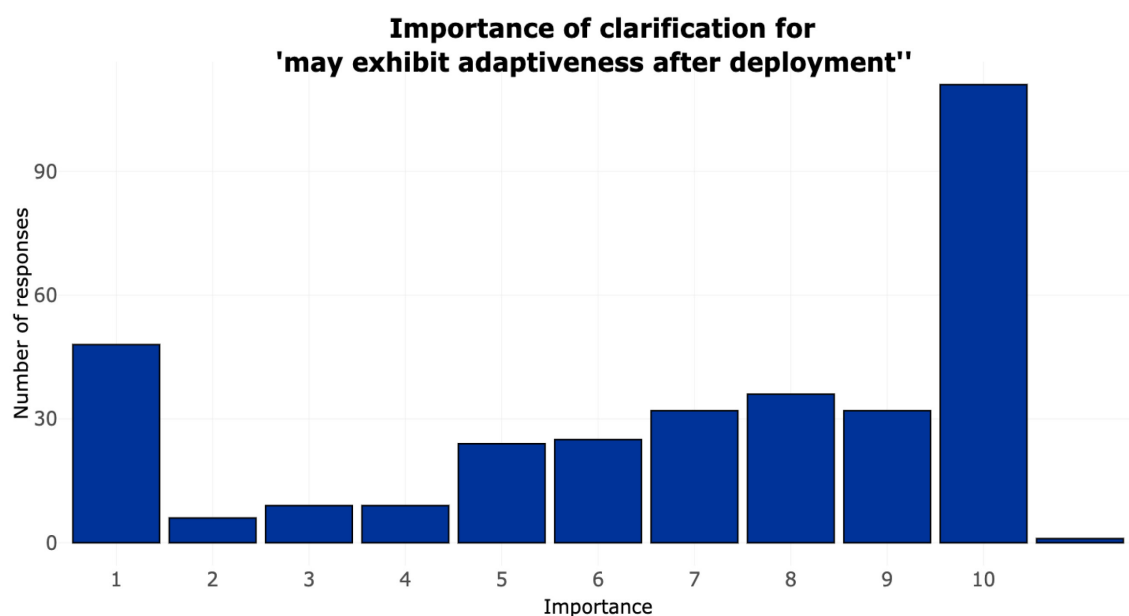


Figure 10: Clarification Needs #3

Question 1.25: "Please rate the importance of further clarification from 1 to 10, 10 indicating 'most important' for "for explicit or implicit objectives"

The graph (Figure 11) shows another bimodal distribution, but with a different pattern - the highest peak of 69 responses occurs at importance level 1 (least important), with a secondary peak of 46 responses at level 10 (most important), and a notable plateau of 30-40 responses across the middle range (5-8). This distribution suggests significant disagreement about the importance of clarifying explicit or implicit objectives - while the largest group considers it minimally important, there's a substantial spread of opinions across the spectrum, with a significant cluster viewing it as critically important. The relatively high number of responses in the middle range indicates more nuanced views about the need for clarification of system objectives compared to other AI aspects.

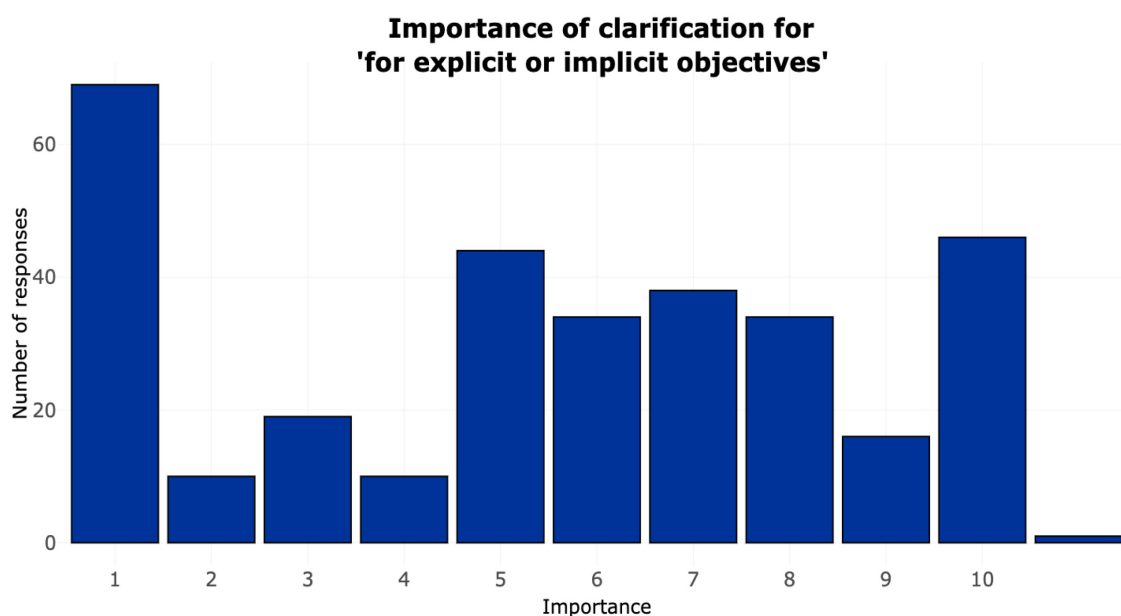


Figure 11: Clarification Needs #4

Question 1.26: *"Please rate the importance of further clarification from 1 to 10, 10 indicating 'most important' for "infers, from the input it receives, how to generate outputs"*

The responses to this question in Figure 12 show a strong skew towards high importance, with the most dramatic peak of 142 responses at importance level 10 (most important), while maintaining a small cluster of 32 responses at level 1 (least important). The distribution shows a clear trend favouring high importance, with a notable secondary peak around level 8 (42 responses), suggesting that respondents overwhelmingly believe clarification is crucial for systems that infer outputs from inputs. The middle ranges (2-7) show consistently low responses (10-20 each), indicating that professionals have strong convictions about this aspect, with most viewing it as a critical area requiring clarification, likely due to the complex nature of inference-based processing and its potential implications.

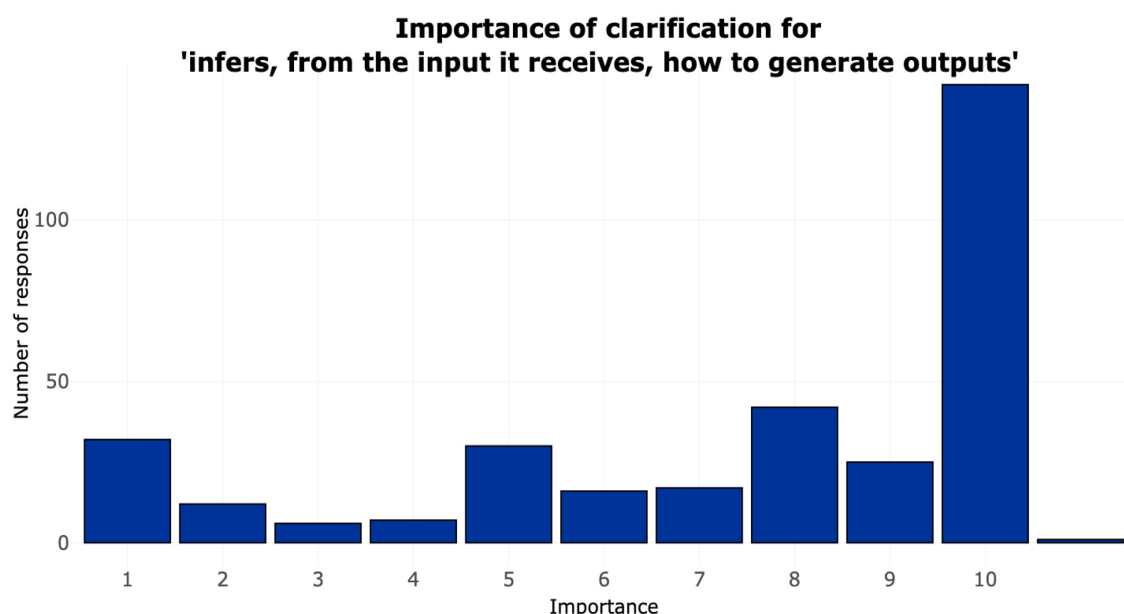


Figure 12: Clarification Needs #5

Question 1.27: *"Please rate the importance of further clarification from 1 to 10, 10 indicating 'most important' for "predictions, content, recommendations, or decisions"*

The largest group considers clarification minimally important (80 responses at importance level 1 - least important). However, the results show a distinctive bimodal distribution with the secondary peak of 43 responses at level 10 (most important). The middle range (5-8) maintains a consistent plateau of around 30-40 responses, suggesting a complex split in opinions about the importance of this element. Clarifying system predictions, content, recommendations, and decisions.

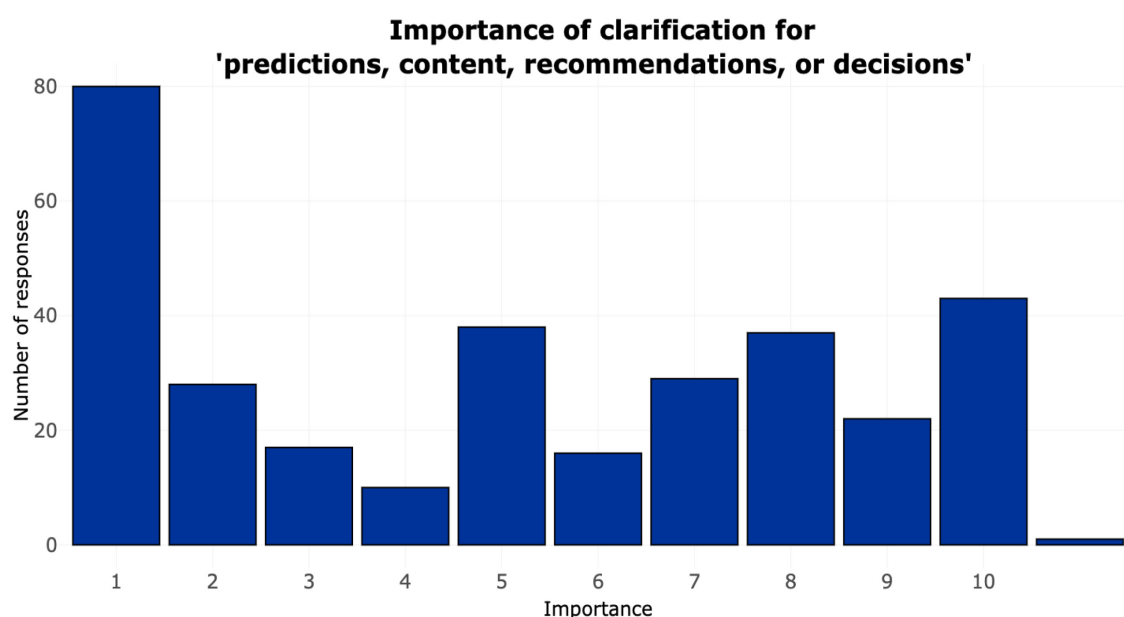


Figure 13: Clarification Needs #6

Question 1.28: "Please rate the importance of further clarification from 1 to 10, 10 indicating 'most important' for "can influence physical or virtual environments"

The results show again a bimodal distribution with the highest peak of 78 responses at importance level 1 (least important) and a relatively even spread of 40-50 responses across levels 8 and 10 (high importance). The middle range (5-8) shows a consistent level of around 35-45 responses, indicating a complex division in views about clarifying systems that can influence physical or virtual environments. While the largest single group considers clarification minimally important, there's a substantial and fairly consistent distribution across higher importance levels, suggesting that many respondents recognise the potential risks and implications of systems that can directly affect real or virtual environments.

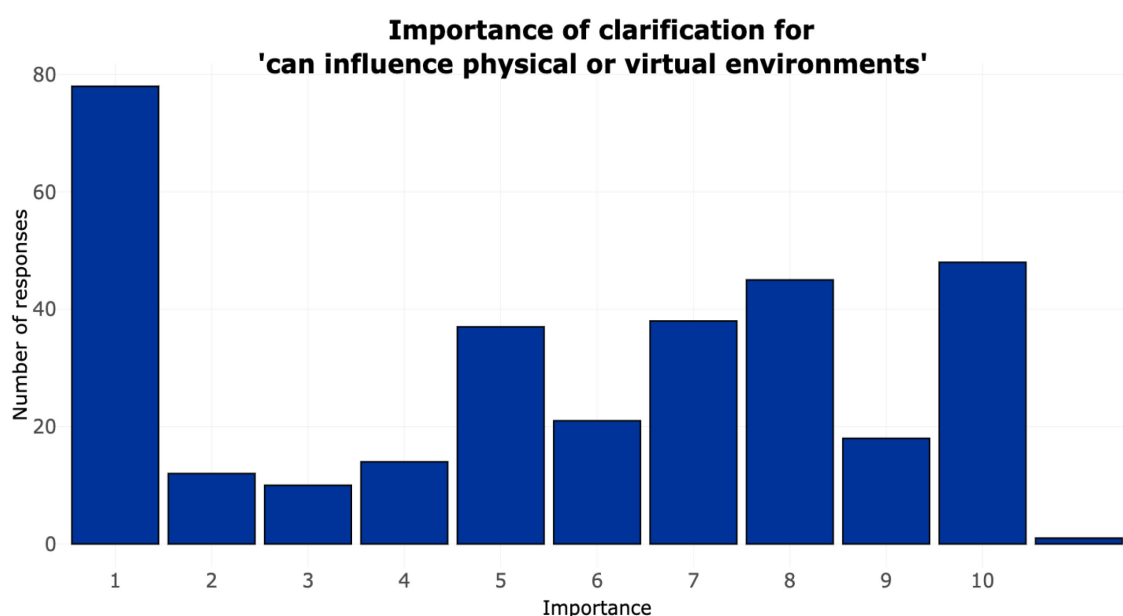


Figure 14: Clarification Needs #7

Question 1.29: "Explain why one or more of these elements require further clarification and what part of this element needs further practical guidance for application in real world applications?"

The responses indicate a strong consensus that the definitions of "adaptiveness," "inference," and "autonomy" need further elaboration. Many stakeholders expressed concerns that the current definitions could inadvertently include traditional software systems, such as rule-based systems and basic statistical models, which do not exhibit the characteristics of AI. For instance, logistic regression and decision trees were frequently cited as examples of systems that, while they may produce outputs based on input data, do not learn, or adapt post-deployment. Stakeholders suggested that clearer guidelines are needed to delineate between traditional software and AI systems, particularly regarding what constitutes "inference" and "adaptiveness."

A common theme across the responses is the distinction between traditional deterministic systems and AI systems. Many respondents emphasised that systems relying on fixed rules or simple statistical methods should not be classified as AI. For example, rule-based decision support systems, basic automation scripts, and traditional statistical analysis tools were frequently mentioned as non-AI systems. The emphasis on transparency and interpretability in traditional systems was also highlighted, contrasting with the opacity often associated with AI systems.

Respondents requested practical examples and clear guidelines to help differentiate between AI and non-AI systems. Many suggested that the European AI Office provide a list of specific software systems or programming approaches that fall outside the AI definition. Stakeholders also called for clarification on the implications of using statistical methods in decision-making processes, particularly in high-risk contexts, to ensure that traditional methods are not inadvertently subjected to AI regulations.

The main challenges raised by respondents include the potential for regulatory loopholes, where developers might reframe AI systems as traditional software to evade scrutiny. Concerns were also expressed about the broadness of the current definition, which could encompass a wide range of software systems, leading to unnecessary regulatory burdens. Additionally, the ambiguity surrounding the terms "adaptiveness" and "inference" was noted as a significant challenge, as it could lead to inconsistent interpretations and applications of the AI Act.

Based on the analysis, it is recommended that the European AI Office refine the definitions of key elements such as "adaptiveness," "inference," and "autonomy" to provide clearer guidance. Establishing a complexity threshold for models, where only those exhibiting significant complexity and learning capabilities are classified as AI, could help delineate traditional software from AI systems. Furthermore, creating a comprehensive list of examples of non-AI systems, along with a methodology for assessing whether a system qualifies as AI, would enhance legal certainty and facilitate compliance for stakeholders. Finally, focusing on the potential societal impacts of systems, rather than solely their technical characteristics, could ensure that genuinely high-risk applications are appropriately regulated.

3.2 Question 2: Simple software systems out of scope of the definition of an AI system

Question 2.30: "Please provide examples of software systems or programming approaches that does not fall under the scope of the AI system definition in Article 3(1) AI Act and explain why, in your opinion, the examples are not covered by one or more of the seven main elements of the definition of an AI system in Article 3(1) AI Act."

A significant number of respondents emphasised that rule-based systems, which operate strictly on predefined human-defined rules without the ability to learn or adapt, should not be classified as AI systems. Examples include basic automation scripts, tax calculation software, and traditional database management systems. These systems lack the essential characteristics of AI, such as inference and adaptiveness, as they produce consistent outputs based on fixed logic and do not exhibit any form of autonomous decision-making.

Many responses highlighted that traditional statistical models, such as linear and logistic regression, do not meet the criteria for AI systems as defined in Article 3(1) of the AI Act. These models are deterministic, operate based on fixed parameters, and do not adapt or learn post-deployment. Respondents argued that while these models can generate predictions, they do so without the complexity and learning capabilities associated with AI, thus should be excluded from the Act's scope.

Respondents pointed out that simple automation tools, such as spreadsheet macros and basic data processing scripts, should also be considered outside the AI definition. These tools execute predefined tasks without any inference or learning capabilities, relying solely on explicit instructions from users. Their deterministic nature means they do not adapt, or change based on new data, further distinguishing them from AI systems.

Several respondents raised concerns about potential loopholes in the AI Act that could allow developers to circumvent regulations by reclassifying AI systems as traditional software. They argued for a clear distinction between AI and non-AI systems based on their operational characteristics, emphasising the need for guidelines that prevent the misclassification of systems that could still pose risks to individuals or society.

There was a strong consensus on the need for clearer definitions and illustrative examples to delineate what constitutes an AI system versus traditional software. Respondents suggested that the Act should explicitly state which systems are excluded, such as basic statistical tools, rule-based systems, and deterministic algorithms, to provide legal certainty and avoid unnecessary regulatory burdens on established methodologies. This clarity would help ensure that the focus remains on genuinely high-risk AI applications that require oversight.

4.0 SECTION 2. QUESTIONS IN RELATION TO THE PROHIBITIONS (ARTICLE 5 AI ACT)

4.1 QUESTIONS IN RELATION TO HARMFUL SUBLIMINAL, MANIPULATIVE OR DECEPTIVE PRACTICE

Question 3.31: Taking into account the provisions of the AI Act, what elements of the prohibition of harmful manipulation and deception do you think require further clarification in the Commission guidelines?

The response distribution clearly indicates that respondents identify multiple aspects of the AI Act's provisions on harmful manipulation and deception that require substantial clarification. The highest concern, with 204 responses, focuses on "deploying subliminal, purposefully manipulative or deceptive techniques," closely followed by concerns about "causing significant harm" (193 responses) and "materially distorting behaviour" (185 responses). The notably lower response rate (77) for clarification needs regarding "placement on the market, putting into service or use" suggests this aspect is relatively well understood. The small number of respondents (28) selecting "none of the above" reinforces that the vast majority of respondents see a pressing need for clearer guidelines, particularly around the technical and ethical boundaries of manipulation and its potential harmful effects.

Question 3.32: Please explain why the elements selected above require further clarification and what needs to be further clarified in the Commission guidelines?

Respondents expressed a strong need for clearer definitions of critical terms such as "subliminal techniques," "manipulative techniques," and "significant harm." Many highlighted that the current language is vague and subjective, leading to potential misinterpretations. For instance, the term "subliminal" lacks a precise definition, making it difficult to determine what constitutes imperceptible influence. Similarly, "materially distorting behaviour" requires clearer thresholds to differentiate between acceptable influence and harmful manipulation. Respondents suggested that concrete examples and measurable standards would enhance understanding and compliance.

There is a consensus that the guidelines should delineate between lawful marketing practices and those that are considered manipulative or deceptive. Many respondents emphasised that common advertising techniques, such as nudging or targeted marketing, should not fall under the prohibition if they do not cause significant harm. The need for clarity on what constitutes "lawful persuasion" versus "prohibited manipulation" was frequently mentioned, with calls for examples to illustrate acceptable practices. This distinction is crucial to avoid stifling legitimate business activities while ensuring consumer protection.

Respondents raised concerns about how "significant harm" and "reasonable likelihood" are assessed. The ambiguity surrounding these terms could lead to inconsistent applications of the law. Many suggested that the guidelines should provide specific criteria for evaluating harm, including physical, psychological, and financial impacts. Additionally, the concept of "reasonable likelihood" needs clarification to ensure that it is not interpreted too broadly, potentially leading to overregulation. Clear benchmarks for what constitutes significant harm would help stakeholders understand their responsibilities and the potential risks associated with AI systems.

The guidelines should clarify the roles and responsibilities of AI providers and deployers, particularly in multi-jurisdictional contexts. Respondents noted that the current language does not adequately address how responsibilities shift between different actors in the AI value chain. For example, it remains unclear who is accountable when an AI system is used in ways that lead to prohibited practices. Clear definitions of "placement on the market," "putting into service," and "use" are necessary to delineate these responsibilities and ensure compliance across various scenarios.

Respondents emphasised the importance of providing practical examples and illustrations to accompany the guidelines. Many suggested that real-world scenarios would help clarify how the prohibitions apply in different contexts, particularly in marketing, healthcare, and social media. Examples of both acceptable and prohibited practices would aid stakeholders in navigating the complexities of the regulations. Additionally, a decision tree or framework for evaluating compliance could enhance understanding and facilitate adherence to the guidelines, ultimately fostering responsible AI development and deployment.

Question 4.33: Do you have or know concrete examples of AI systems that in your opinion fulfil all elements of the prohibition described above?

The largest segment of respondents (38.9%) did not provide any response, suggesting either uncertainty or reluctance to take a position. Among those who did respond, there was a fairly even split between those who claimed to know concrete examples (26.4%) and those who did not (34.7%). This relatively balanced distribution between "yes" and "no" responses, combined with the high no-response rate, could indicate that there is significant ambiguity or complexity surrounding the identification of AI systems that meet the referenced prohibition criteria, or possibly a lack of clear understanding about what constitutes fulfilment of these elements.

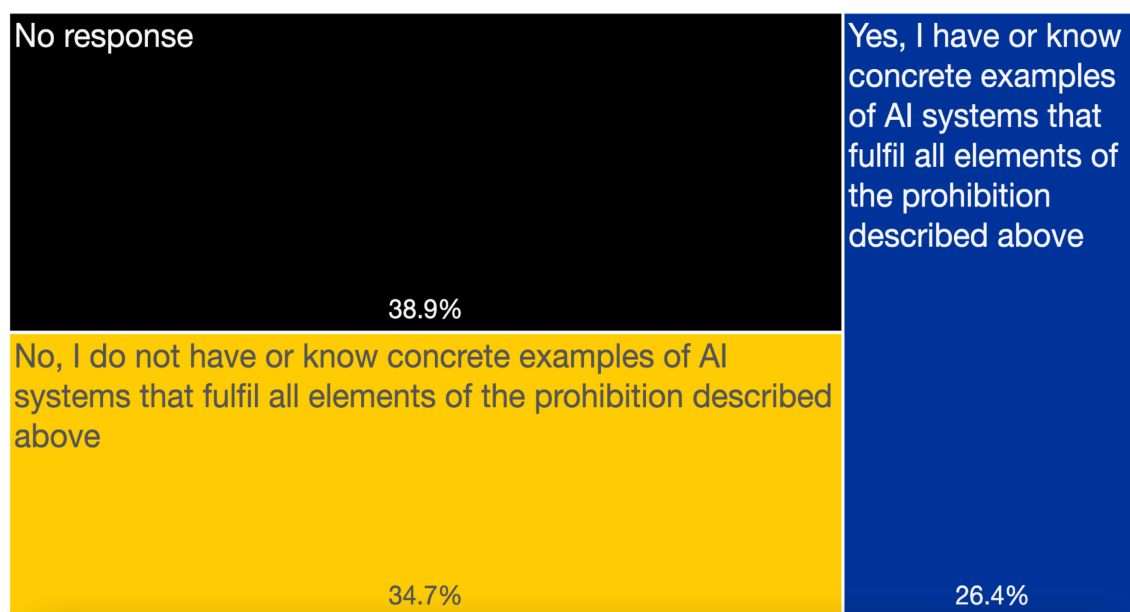


Figure 15: Do you know examples of AI systems

Question 4.34: Please specify the concrete AI system, how it is used in practice and how all the necessary elements described above are fulfilled

The responses highlight a variety of AI systems that are perceived to fulfil the criteria of manipulation and deception as outlined in the AI Act. Key examples include deepfake technology used for political manipulation and scams, conversational bots that influence public opinion, and recommendation algorithms on social media platforms that exploit cognitive biases. These systems are often described as employing subliminal techniques to distort user behaviour, leading to significant harm, particularly in the context of elections and public trust.

Respondents detailed how these AI systems operate through various manipulative techniques. For instance, deepfake technology can create convincing but false representations of individuals, while social media algorithms can amplify emotionally charged content to keep users engaged. The use of subliminal messaging, such as imperceptible visual or auditory cues, is frequently mentioned as a method that undermines users' autonomy and decision-making capabilities, often without their conscious awareness.

The analysis reveals widespread concern about the societal and individual impacts of these AI systems. Many respondents noted that the manipulation of public opinion through social media can lead to polarisation, misinformation, and erosion of trust in democratic processes. Specific examples include the Cambridge Analytica scandal and recent electoral manipulations in various countries, which illustrate how AI can significantly distort public behaviour and decision-making, resulting in psychological and social harm.

Question 5.35: Do you have or know concrete examples of AI systems where you need further clarification regarding certain elements of this prohibition to determine whether the AI system is in the scope of the prohibition or not?

The largest group (41%) provided no response, which could indicate general hesitation or uncertainty about the topic. Among those who did respond, 35% indicated they do not need further clarification about AI systems in relation to the prohibition, while 24% acknowledged they have specific examples where they need more clarity. When compared to the previous question, there's a consistent pattern of high non-response rates, suggesting that the complexity of AI system regulations and prohibitions may be challenging for many respondents to definitively assess. The fact that nearly a quarter of respondents specifically identified cases needing clarification highlights the real-world challenges in interpreting and applying AI prohibition criteria.

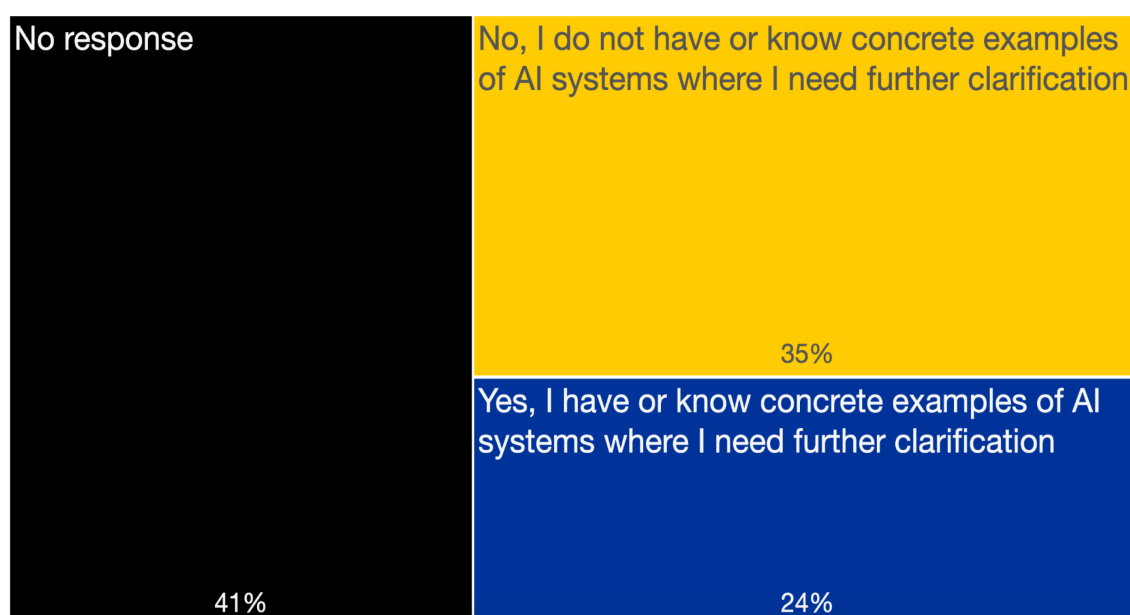


Figure 16: Examples of AI systems needing clarification

Question 5.36: Please specify the concrete AI system, how it is used in practice as well as the specific elements you would need further clarification in this regard

The responses highlight a diverse range of AI systems currently in use, including emotion-based AI, machine learning models for refugee eligibility assessments, and content recommendation algorithms on social media platforms. Specific examples include Affectiva's Emotion AI, which infers user emotions for marketing and customer service, and AI systems used in healthcare for triaging radiology cases. These systems often aim to enhance user experience or operational efficiency but raise ethical concerns regarding manipulation and potential harm.

A significant theme across the responses is the ethical implications of AI systems, particularly regarding manipulation and deception. Respondents' express concerns about systems that exploit cognitive biases, such as social media algorithms that create echo chambers or targeted advertising that nudges users toward impulsive decisions. The ambiguity surrounding what

constitutes "manipulative" or "deceptive" techniques is a recurring issue, with calls for clearer definitions to distinguish between acceptable personalisation and harmful manipulation.

Many respondents seek clarification on what constitutes "significant harm" in the context of AI systems. Examples include the psychological distress caused by biased algorithms in recruitment or the potential for AI-driven content to influence vulnerable populations negatively. The need for clear thresholds and metrics to assess harm is emphasised, particularly in sensitive areas like healthcare, where the consequences of AI misinterpretation can be severe.

4.2 QUESTIONS IN RELATION TO HARMFUL EXPLOITATION OF VULNERABILITIES

Question 6.37: Taking into account the provisions of the AI Act, what elements of the prohibition of harmful exploitation of vulnerabilities do you think require further clarification in the Commission guidelines?

The elements of the AI Act's prohibition on harmful exploitation of vulnerabilities that most urgently require clarification are 'exploiting vulnerabilities due to age, disability or specific socio-economic situation' (176 responses), followed closely by the need to clarify what constitutes "significant harm" (166 responses) and "materially distorting behaviour" (164 responses). Only 37 respondents responded, 'none of the above'.

Question 6.38: Please explain why the elements selected above require further clarification and what needs to be further clarified in the Commission guidelines?

Many respondents expressed the need for clearer definitions of what constitutes "vulnerabilities" in the context of AI systems, particularly regarding age, disability, and socio-economic status. Respondents highlighted the importance of recognising situational vulnerabilities, such as temporary financial instability or emotional distress, and emphasised that vulnerabilities can be context-dependent, requiring a nuanced approach.

The terms "exploiting vulnerabilities" and "materially distorting behaviour" were frequently cited as needing further clarification. Overall, there is a call for specific examples and criteria to distinguish between legitimate uses of AI and exploitative practices. Respondents sought guidance on the thresholds for what constitutes exploitation versus lawful practices, such as targeted assistance or marketing. There is a concern that the current definitions are too broad or subjective, which could lead to misinterpretation and unintended consequences. Clear examples of what constitutes material distortion, particularly in vulnerable populations, were requested to help delineate acceptable from harmful practices.

The concept of "significant harm" was identified as vague and requiring more precise definitions. Respondents called for clarity on what types of harm are considered significant, including psychological, financial, and societal impacts. There is a desire for a framework to assess harm, particularly in cases where cumulative or indirect effects may occur over time. Many emphasised that even minor harms could be significant for vulnerable groups, necessitating a broader interpretation of what constitutes harm.

Several responses highlighted the need for the Commission's guidelines to align with existing EU regulations, such as the GDPR and consumer protection laws. Respondents expressed concerns about potential overlaps and conflicts between the AI Act and other legal frameworks, advocating for a harmonised approach to ensure consistent interpretation and

application across different sectors. This alignment is seen as crucial for preventing regulatory confusion and ensuring that protections for vulnerable groups are robust and effective.

There was a strong call for the inclusion of practical examples and case studies in the guidelines to illustrate how the concepts of vulnerability, exploitation, and significant harm apply in real-world scenarios. Respondents emphasised that clear, actionable guidance would aid stakeholders in understanding their responsibilities and the implications of the regulations. This practical approach is viewed as essential for fostering compliance and encouraging the development of ethical AI systems that respect individual rights and promote social good. In summary, the responses indicate a consensus on the need for clearer definitions, practical examples, and alignment with existing laws to effectively address the complexities of AI systems and their impact on vulnerable populations. The emphasis on context, specificity, and practical guidance reflects a desire for a balanced regulatory framework that protects individuals while allowing for innovation in AI technologies.

Question 7.39: Do you have or know concrete examples of AI systems that in your opinion fulfil all elements of the prohibition described above?

The response data shows that only 21% of respondents could identify concrete examples of AI systems that fulfil all elements of the prohibition, while 35% could not, and a significant 44% were uncertain or did not respond, suggesting there may be considerable ambiguity in how to interpret and apply the prohibition criteria in practice.

Question 7.40: Please specify the concrete AI system, how it is used in practice and how all the necessary elements described above are fulfilled

Many respondents highlighted how AI systems exploit socio-economic vulnerabilities, particularly targeting low-income individuals and marginalised groups. Examples include AI-driven micro-lending platforms that offer high-interest loans to financially distressed individuals, and targeted advertising systems that manipulate users based on their mental health states or emotional vulnerabilities. This exploitation often leads to users making decisions that are not in their best interest, such as accepting predatory loan terms or engaging with harmful content.

The responses frequently noted that AI systems materially distort user behaviour by leveraging psychological profiling and emotional manipulation. For instance, social media algorithms and recommendation systems were cited for promoting addictive behaviours, particularly among children and adolescents. These systems often encourage excessive engagement or spending, leading users to make impulsive decisions that can have long-term negative consequences on their financial and mental well-being.

A recurring theme in the responses was the significant harm caused by AI systems, which can manifest in various forms, including financial instability, psychological distress, and social isolation. Specific cases, such as again the Cambridge Analytica scandal and the use of AI in

political micro-targeting, were mentioned as examples of how these systems can influence public behaviour and decision-making in harmful ways. Additionally, the tragic outcomes associated with AI companionship apps and financial platforms underscored the potential for severe emotional and financial repercussions.

The analysis revealed a wide range of AI applications across various sectors, including finance, healthcare, social media, and gaming. Each application presents unique challenges and risks, necessitating a nuanced understanding of how AI systems interact with different demographics. The responses emphasised the importance of considering the specific contexts in which these technologies operate, as well as the diverse vulnerabilities they may exploit, to develop effective regulatory measures that protect individuals and communities from harm. In summary, the responses collectively underscore the urgent need for comprehensive oversight and ethical considerations in the deployment of AI systems, particularly those that have the potential to exploit vulnerabilities and cause significant harm to individuals and society at large.

Question 8.41: Do you have or know concrete examples of AI systems where you need further clarification regarding certain elements of this prohibition to determine whether the AI system is in the scope of the prohibition or not?

19.3% of respondents indicated they needed further clarification on elements of the AI prohibition, while 36.3% did not require additional clarification. A significant portion (44.4%) did not respond. The high NA rate could indicate unfamiliarity with the specific prohibition in question.

Question 8.42: Please specify the concrete AI system, how it is used in practice as well as the specific elements you would need further clarification in this regard

Many respondents expressed concerns about AI systems that exploit the vulnerabilities of specific groups, such as children, the elderly, and socio-economically disadvantaged individuals. Examples include recommender systems on social media that amplify harmful content, financial chatbots that may mislead users, and educational tools that could exploit cognitive vulnerabilities. Respondents called for clearer definitions of what constitutes exploitation and significant harm, emphasising the need for guidelines that protect these vulnerable populations.

There was a strong demand for clarification on the definitions of "exploitation," "material distortion," and "significant harm" within the context of the AI Act. Respondents highlighted the ambiguity surrounding how these terms apply to various AI systems, particularly in sectors like finance, healthcare, and education. The need for a nuanced understanding of how AI systems interact with vulnerable groups was a recurring theme, with many suggesting that existing frameworks may not adequately address the complexities involved.

Several responses pointed out the potential for AI systems to perpetuate biases and discrimination, particularly in recruitment and credit scoring. Concerns were raised about how

historical biases in training data could lead to unfair treatment of marginalised groups. Respondents urged for guidelines that ensure AI systems do not inadvertently reinforce existing inequalities, calling for a focus on fairness and transparency in AI development and deployment.

Some respondents advocated for exemptions for AI systems designed to comply with existing regulations, particularly in the financial sector. They argued that systems aimed at protecting vulnerable clients should not fall under prohibitions if their intent is to enhance client protection rather than exploit vulnerabilities. This highlights a tension between regulatory oversight and the need for responsible AI use in compliance with legal standards.

4.3 QUESTIONS IN RELATION TO UNACCEPTABLE SOCIAL SCORING PRACTICES

Question 9.43: Taking into account the provisions of the AI Act, what elements of the prohibition of social scoring do you think require further clarification in the Commission guidelines?

In response to the question regarding the elements of the prohibition of social scoring in the AI Act that require further clarification in the Commission guidelines, a significant number of respondents highlighted specific areas of concern. The most prominent issue, with 168 responses, focused on the evaluation or classification of individuals based on their social behaviour or inferred characteristics. Other notable areas included the implications of social scoring leading to detrimental treatment (142 responses), the use of data in unrelated social contexts (146 responses), and the need for clarity on unjustified or disproportionate treatment (134 responses). Additionally, 42 responses addressed the aspects of market placement and service use of AI systems, while 46 respondents indicated that none of the options required further clarification.

Question 9.44: Please explain why the elements selected above require further clarification and what needs to be further clarified in the Commission guidelines?

A significant number of respondents expressed concerns about the vague definitions of terms such as "social behaviour," "personal or personality characteristics," and "detrimental or unfavourable treatment." Many called for clearer, more precise definitions to avoid misinterpretation and ensure that legitimate practices, such as credit scoring and fraud prevention, are not inadvertently prohibited. Respondents emphasised the need for examples to illustrate acceptable versus unacceptable practices, particularly in contexts like employment and financial services.

There is a strong desire for clarification on the scope of the prohibition against social scoring, particularly regarding its application to private entities and various sectors. Respondents highlighted the potential for legitimate business practices, such as marketing segmentation and risk assessment in finance, to be impacted by overly broad interpretations of the guidelines. Many urged the Commission to delineate between harmful social scoring practices and lawful evaluations that serve important societal functions.

Respondents raised concerns about the interpretation of "unrelated social contexts" and how data collected in one context might be used in another. There is a call for clear guidelines on when data repurposing becomes unjustifiable, especially in cases where it could lead to discrimination or unfair treatment. The need for a nuanced understanding of context was emphasised, particularly in relation to the merging of public and private data sources.

The concept of "unjustified or disproportionate treatment" was frequently mentioned as needing further clarification. Respondents sought guidance on how to assess what constitutes unjustified treatment and the criteria for determining proportionality in relation to social behaviour. Many suggested that the guidelines should establish clear thresholds and examples to help organisations navigate these complex assessments.

Finally, there was a recurring theme regarding the potential impact of the guidelines on innovation within the AI sector. Respondents urged the Commission to ensure that the guidelines do not stifle technological advancement or hinder the ability of businesses to operate effectively. They emphasised the importance of balancing the protection of fundamental rights with the need for clear, practical regulations that allow for responsible AI use in various industries. In summary, the responses highlight a collective call for clearer definitions, a well-defined scope of prohibition, contextual considerations, established thresholds for treatment, and a balanced approach that fosters innovation while protecting individual rights.

Question 10.45: Do you have or know concrete examples of AI systems that in your opinion fulfil all elements of the prohibition described above?

A total of 74 respondents affirmed having concrete examples, while 139 indicated they did not have such examples. Additionally, 170 respondents chose not to answer the question, suggesting a significant portion of the audience either lacks knowledge on the topic or prefers not to engage with it. This distribution highlights a divide in awareness or understanding of AI systems in relation to the specified prohibition.

Question 10.46: Please specify the concrete AI system, how it is used in practice and how all the necessary elements described above are fulfilled

A significant number of responses highlighted the use of AI systems in welfare fraud detection, particularly in the Netherlands and Denmark. The Netherlands Childcare Benefit scandal was frequently cited, where algorithms flagged claims based on social behaviours and personal

characteristics, leading to racial profiling and unjust treatment of vulnerable groups, such as single mothers and low-income families. Similarly, the Danish welfare authority's system evaluated residents over time, leading to detrimental outcomes for marginalised populations, including people with disabilities and migrants. These systems often relied on irrelevant data, resulting in invasive monitoring and violations of privacy rights.

Respondents expressed concerns about discriminatory practices embedded in various social scoring systems, particularly in France and the UK. The French National Family Allowance Fund (CNAF) was noted for using algorithms that assigned risk scores based on personal characteristics, leading to disproportionate scrutiny of low-income families. The UK Home Office's visa streaming algorithm was criticised for profiling applicants based on nationality, resulting in biased treatment and unfair rejections. These examples illustrate how social scoring can perpetuate systemic inequalities and infringe on individuals' rights.

Several responses discussed the application of AI in employment and credit scoring, emphasising the potential for unjust treatment based on social media activity and inferred personality traits. For instance, AI systems used by companies like Uber evaluate drivers based on passenger ratings, which can lead to penalties for low scores, impacting their livelihoods. Similarly, credit scoring systems that analyse social media behaviour to determine creditworthiness were criticised for using irrelevant data, leading to unfair financial decisions. These practices raise ethical concerns about the use of personal data in unrelated contexts.

The responses underscored the ethical implications of AI systems that engage in social scoring, with many calling for stricter regulations to prevent discriminatory practices. Concerns were raised about the lack of transparency and accountability in these systems, particularly regarding how data is collected and used. Respondents emphasised the need for clear guidelines to delineate acceptable practices from prohibited ones, particularly in sensitive areas like welfare, employment, and credit. The potential for AI to reinforce existing biases and inequalities was a recurring theme, highlighting the urgency for regulatory frameworks.

The analysis of these responses indicates a growing awareness of the broader implications of AI systems in society. Many respondents pointed to the need for a comprehensive understanding of how AI can impact various sectors, including healthcare, education, and social services. The examples provided illustrate the potential for AI to exacerbate social inequalities and infringe on individual rights. As discussions around the AI Act continue, there is a clear call for more robust protections against unjustified and disproportionate treatment, ensuring that AI systems are designed and implemented with ethical considerations at the forefront.

Question 11.47: Do you have or know concrete examples of AI systems where you need further clarification regarding certain elements of this prohibition to determine whether the AI system is in the scope of the prohibition or not?

A total of 85 respondents indicated that they require further clarification to determine whether specific AI systems fall under the scope of these prohibitions. In contrast, 126 respondents felt that no additional clarification was necessary, while 172 respondents chose not to answer the question. This indicates a notable interest in understanding the boundaries of the prohibitions among a segment of the respondents, highlighting potential areas for further discussion and guidance in the regulation of AI technologies.

Question 11.48: Please specify the concrete AI system, how it is used in practice as well as the specific elements you would need further clarification in this regard

Many respondents expressed apprehension about AI systems that predict behaviours or outcomes, particularly in sensitive areas like crime prediction, credit scoring, and welfare automation. They highlighted that such systems could lead to unfair treatment, discrimination, and violations of the principle of presumption of innocence. The potential for these systems to label individuals based on inferred traits raises ethical concerns about their impact on personal freedoms and rights.

A recurring theme in the responses is the lack of transparency from organisations deploying AI systems. Respondents noted that vague definitions and guidelines regarding what constitutes social scoring hinder the ability to assess whether specific AI applications fall under the prohibition outlined in the AI Act. There is a strong call for clearer guidelines that delineate the responsibilities of authorities, and the burden of proof required to demonstrate compliance.

Respondents requested concrete examples of AI systems that either qualify for exceptions under the AI Act or fall under the prohibition of social scoring. They emphasised the importance of understanding the nuances between lawful evaluations and prohibited practices, particularly in contexts like credit scoring, employment screening, and public service delivery. Clarifications on terms such as "unfavourable treatment" and "unrelated social contexts" are deemed essential for proper implementation.

Many responses highlighted concerns about the disproportionate impact of AI systems on marginalised groups, such as refugees, low-income individuals, and those with limited access to resources. The potential for these systems to exacerbate existing inequalities was a significant concern, with calls for guidelines to ensure fairness and equity in AI applications, particularly in welfare and migration contexts.

The responses covered a wide range of AI applications, from employee monitoring systems to healthcare prioritisation algorithms. While some respondents acknowledged the potential benefits of AI in improving efficiency and decision-making, they stressed the need for ethical considerations and safeguards to prevent misuse. The distinction between beneficial uses of AI and those that could lead to social scoring practices was a critical point of discussion, with a call for a balanced approach that respects individual rights while leveraging technological advancements. Overall, the analysis underscores the urgent need for comprehensive

guidelines and clarifications to navigate the complex landscape of AI applications and their implications for social scoring and individual rights.

4.4 QUESTIONS IN RELATION TO INDIVIDUAL CRIME RISK ASSESSMENT AND PREDICTION

Question 12.49: Taking into account the provisions of the AI Act, what elements of the prohibition of harmful manipulation and deception do you think require further clarification in the Commission guidelines?

The majority, with 119 responses, emphasised the need for clarity on AI systems that profile individuals based on their traits and characteristics. Additionally, 112 responses pointed to the necessity for guidance on AI systems used for risk assessment or predicting criminal behaviour. Meanwhile, 116 respondents noted the importance of excluding AI systems that support human assessments based on objective and verifiable facts related to criminal activities. A smaller group of 40 responses focused on the aspects of market placement and service use, while 56 respondents indicated that none of the mentioned elements required further clarification.

Question 12.50: Please explain why the elements selected above require further clarification and what needs to be further clarified in the Commission guidelines?

Many respondents emphasised the need for clearer definitions of terms such as "risk assessment," "profiling," "objective and verifiable facts," and "supporting human assessment." There is a consensus that the current language is ambiguous, leading to uncertainty about what constitutes prohibited practices versus acceptable uses of AI in law enforcement and other contexts. Respondents called for specific examples to illustrate the boundaries of these definitions, particularly in relation to the use of AI systems in predicting criminal behaviour.

A significant number of responses highlighted the necessity to differentiate between AI applications in law enforcement and those used in private sectors, such as financial institutions. Many argued that systems designed for fraud detection and anti-money laundering should be explicitly exempt from the prohibitions outlined in the guidelines. This distinction is crucial to ensure that legitimate and necessary uses of AI for public safety and financial integrity are not hindered by overly broad regulations.

Respondents expressed concerns regarding the potential for AI systems to perpetuate discrimination, particularly against marginalised communities. There were calls for robust safeguards to prevent the misuse of AI in profiling individuals based on sensitive characteristics such as race, socio-economic status, or migration status. The need for independent oversight

and accountability mechanisms was emphasised to ensure that AI systems do not reinforce existing biases in law enforcement practices.

The role of human judgment in AI-assisted decision-making was a recurring theme. Respondents sought clarity on what constitutes adequate human oversight and how to ensure that AI systems do not unduly influence human decisions. There were suggestions for establishing clear thresholds for human involvement in assessments and for defining the extent to which AI can support rather than replace human judgment in sensitive contexts like criminal justice.

Many responses pointed to the practical challenges of implementing the guidelines, particularly for small and medium-sized enterprises (SMEs) and organisations operating across borders. Respondents called for guidance on compliance responsibilities, especially in scenarios where AI systems are repurposed or integrated into broader tools. The need for a clear methodology to assess compliance with the guidelines was highlighted, along with the importance of providing tangible examples to aid understanding and application of the regulations. In summary, the responses reflect a strong desire for clarity, specificity, and practical guidance in the Commission's guidelines to ensure that AI systems are used responsibly and ethically, while also safeguarding fundamental rights and enabling legitimate uses in various sectors.

Question 13.51: Do you have or know concrete examples of AI systems that in your opinion fulfil all elements of the prohibition described above?

A total of 53 respondents affirmed having concrete examples, while a significantly larger group of 134 respondents indicated they did not have such examples. Additionally, 196 respondents chose not to answer the question, suggesting a notable uncertainty or lack of knowledge about the specific criteria of the prohibition among the majority of participants.

Question 13.52: Please specify the concrete AI system, how it is used in practice and how all the necessary elements described above are fulfilled

The responses highlight various AI systems, such as the Netherlands' "Top 600" program, RisCanvi in Catalonia, and ProKid in the Netherlands, which are designed to predict criminal behaviour or assess recidivism risk. These systems use personal data, including socio-economic factors, historical crime data, and demographic information, to generate risk scores. The applications of these systems often lead to increased surveillance and intervention without concrete evidence of criminal activity, raising ethical and legal concerns.

A significant number of respondents expressed concerns about the inherent biases in these AI systems. Many systems, such as COMPAS and VioGén, have been criticised for perpetuating systemic discrimination, particularly against marginalised communities. The reliance on historical data, which may reflect past biases in policing, can lead to unfair targeting of

individuals based on their socio-economic status, ethnicity, or neighbourhood, undermining the principles of justice and equality.

The responses frequently mention the lack of transparency in the algorithms used by these AI systems. Many systems operate with minimal human oversight, making it difficult to understand how decisions are made or to hold them accountable for errors. This opacity raises serious questions about the reliability of the predictions made and the potential consequences for individuals flagged by these systems.

The use of AI systems for predictive policing raises significant legal and ethical issues, particularly concerning the presumption of innocence and the right to a fair trial. Many respondents pointed out that these systems often lead to punitive actions without formal judicial processes, such as police raids or increased surveillance, based solely on algorithmic predictions rather than verified evidence of wrongdoing.

Question 14.53: Do you have or know concrete examples of AI systems where you need further clarification regarding certain elements of this prohibition to determine whether the AI system is in the scope of the prohibition or not?

A total of 47 respondents indicated that they require additional information to determine whether their AI systems fall within the scope of the prohibition. In contrast, 138 respondents felt that no further clarification was necessary, while 198 respondents chose not to answer the question. This indicates a significant portion of respondents are either confident in their understanding of the prohibition or do not see it as applicable to their AI systems.

Question 14.54: Please specify the concrete AI system, how it is used in practice as well as the specific elements you would need further clarification in this regard

Many respondents expressed concerns that the AI Act could inadvertently hinder fraud prevention efforts. They emphasised that AI systems used by banks and payment service providers (PSPs) for fraud detection should not be penalised for profiling based on behavioural patterns or banking behaviours. Respondents argued that these systems are essential for identifying fraudulent activities and should be explicitly excluded from the prohibitions outlined in the AI Act. The need for clear guidelines that differentiate between legitimate fraud prevention and prohibited profiling was a recurring theme.

Several responses highlighted the complexities surrounding predictive policing systems, such as PredPol and COMPAS, which analyse historical crime data to forecast criminal activity. Respondents sought clarification on whether these systems, which may indirectly profile individuals based on location or demographic data, fall under the AI Act's prohibitions. There was a strong call for guidelines that define what constitutes "objective and verifiable facts" in the context of risk assessments, particularly regarding the use of historical data and its potential biases.

The ethical implications of using AI systems for criminal risk assessments were a significant concern. Respondents pointed out that tools like COMPAS, which evaluate personal characteristics and historical data to predict recidivism, could perpetuate systemic biases and discrimination. They requested clearer definitions of profiling and personality traits to ensure that AI systems do not unfairly target marginalised groups. The need for transparency and accountability in AI decision-making processes was emphasised to prevent unjust outcomes.

Responses also addressed the use of AI in healthcare settings, particularly for predicting aggressive behaviour in patients with mental health issues. Clarifications were sought on whether such predictive tools, which aim to enhance safety without direct links to criminal activity, are permissible under the AI Act. Additionally, concerns were raised about AI systems monitoring animal behaviour, with respondents requesting guidance on how animal welfare laws intersect with the AI Act's prohibitions, particularly regarding neglect and abuse.

Question 15.55: Do you have or know concrete examples of AI systems that fulfil all necessary criteria for the prohibition to apply, but fall under the exception of systems that support the human assessment of the involvement of a person in a criminal activity, based on objective and verifiable facts linked to a criminal activity?

A significant majority of respondents (141) indicated that they do not have examples of such systems. In contrast, 41 respondents affirmed the existence of these AI systems, while 201 chose not to answer the question. This indicates a general uncertainty or lack of awareness about specific AI systems that could fit within the outlined exception.

Question 15.56: Please specify the concrete AI system, how it is used in practice and which exception would apply and why

Respondents identified various AI systems used in law enforcement, including AI-assisted criminal investigation tools, fraud detection systems, and surveillance analysis tools. These systems are employed to analyse data such as surveillance footage, transaction patterns, and historical crime data to support human investigators in identifying potential criminal activities. The emphasis is on using objective and verifiable facts to assist rather than replace human judgment, ensuring that the final decisions remain with law enforcement personnel.

A significant concern raised by respondents is the potential for over-reliance on AI systems, which could lead to automation bias. Many highlighted instances where algorithms, such as those used in welfare fraud investigations, resulted in individuals being flagged as high-risk based on flawed assumptions or uncorroborated data. This raises ethical questions about the presumption of guilt and the need for robust human oversight to prevent unjust outcomes.

Several respondents provided examples of AI systems that could qualify for exceptions under the proposed guidelines. These include systems that analyse verified data, such as transaction records in fraud detection or acoustic data in gunshot detection (e.g., ShotSpotter). The

consensus is that these systems should only assist human decision-making and not autonomously determine guilt or risk, thereby maintaining procedural fairness.

4.5 QUESTIONS IN RELATION TO UNTARGETED SCRAPING OF FACIAL IMAGES

Question 16.57: Taking into account the provisions of the AI Act, what elements of the prohibition of untargeted scraping of facial images do you think require further clarification in the guidelines?

The majority of respondents highlighted the need for further clarification on specific elements. The most significant area of concern was the need for clarity on untargeted scraping itself, with 126 responses, followed closely by the creation or expansion of facial recognition databases with 106 responses. Other aspects, such as the placement on the market and use of AI systems, received 51 responses, while the use of images from the internet or CCTV footage attracted 79 responses. A total of 53 respondents indicated that none of these elements required further clarification.

Question 16.58: Please explain why the elements selected above require further clarification and what needs to be further clarified in the guidelines?

A significant number of respondents emphasised the need for clear definitions of terms such as "untargeted scraping," "facial recognition databases," and "placement on the market." Many expressed confusions over what constitutes "targeted" versus "untargeted" scraping, with calls for specific examples to delineate acceptable practices from prohibited ones. The ambiguity surrounding these definitions could lead to misinterpretation and potential misuse of AI technologies, particularly in sensitive areas like facial recognition.

Respondents highlighted the necessity to clarify the scope of the prohibition on creating or expanding facial recognition databases. Questions arose regarding whether incidental data collection during unrelated operations would trigger the prohibition and whether databases created for legitimate purposes, such as humanitarian efforts, would be exempt. This indicates a concern for balancing regulatory measures with the need for innovation and legitimate use cases in AI.

Many responses pointed to the ethical and legal implications of scraping practices, particularly regarding privacy rights and data protection. Respondents urged that scraping practices should not infringe on individuals' rights, especially concerning the collection of images from public spaces without consent. The potential for misuse of facial recognition technologies, particularly in surveillance contexts, was a recurring theme, with calls for stringent guidelines to prevent abuse.

There were calls for clarification on the responsibilities of different stakeholders involved in the deployment of AI systems. Respondents questioned who would be held accountable if general-purpose scraping tools were repurposed for facial recognition database creation. This concern reflects a desire for clear guidelines that delineate liability across various stages of AI system deployment, ensuring that all parties understand their obligations under the law.

Finally, several respondents expressed concern that overly broad prohibitions could stifle innovation and research in AI. They advocated for guidelines that would allow for the responsible use of AI technologies, particularly in academic and research contexts, while still protecting individual rights. This highlights the need for a balanced approach that fosters technological advancement without compromising ethical standards and privacy protections. In summary, the responses indicate a strong demand for clarity in definitions, scope, legal implications, accountability, and the balance between regulation and innovation in the context of AI and facial recognition technologies

Question 17.59: Do you have or know concrete examples of AI systems that in your opinion fulfil all elements of the prohibition described above?

A total of 60 respondents affirmed having concrete examples, while a significantly larger group of 121 respondents indicated they did not have such examples. Additionally, 202 respondents chose not to answer the question, suggesting a notable level of uncertainty or lack of knowledge about the specific criteria of the prohibition among the participants.

Question 17.60: Please specify the concrete AI system, how it is used in practice and how all the necessary elements described above are fulfilled

The responses predominantly highlight two AI systems, Clearview AI and PimEyes, as concrete examples of technologies that engage in problematic practices related to facial recognition. Clearview AI is noted for its extensive scraping of publicly available images from social media and other online sources to create a vast facial recognition database, which is marketed to law enforcement agencies. Similarly, PimEyes operates as a face search engine that scans the internet for images of individuals, allowing users to upload a photo and find matches. Both systems are criticised for their indiscriminate data collection methods, which violate principles of consent and data minimisation under the GDPR.

Respondents express significant ethical and legal concerns regarding the practices of these AI systems. The mass scraping of biometric data without explicit consent is highlighted as a violation of privacy rights and a breach of GDPR regulations. The responses emphasise that such practices not only infringe on individual rights but also contribute to a culture of mass surveillance, raising alarms about potential misuse of personal data. The need for clear prohibitions against these practices is underscored, with calls for regulatory frameworks to ensure accountability and protect fundamental rights.

In practice, the AI systems discussed are used for various purposes, primarily by law enforcement and private entities for identification and surveillance. Clearview AI, for instance, allows law enforcement to match facial images against its database, which has led to documented cases of wrongful arrests, particularly affecting marginalised communities. The indiscriminate nature of data collection means that individuals are often unaware of their images being used in this manner, leading to significant implications for personal privacy and civil liberties.

Beyond the specific cases of Clearview AI and PimEyes, the responses reflect a broader concern about the implications of AI technologies in society. The potential for algorithmic bias, particularly against marginalised groups, is highlighted, alongside the risks of creating discriminatory databases. The discussion also touches on the need for ethical considerations in the development and deployment of AI systems, including the impact on non-human subjects in research contexts. As AI technologies continue to evolve, the responses call for ongoing scrutiny and adaptation of regulatory frameworks to safeguard individual rights and promote ethical practices in AI development.

Question 18.61: Do you have or know concrete examples of AI systems where you need further clarification regarding certain elements of this prohibition to determine whether the AI system is in the scope of the prohibition or not?

A total of 42 respondents indicated that they have or know of concrete examples of AI systems requiring further clarification on specific elements of a prohibition to ascertain their compliance, while 127 respondents stated they do not have such examples. Additionally, 214 respondents chose not to answer the question, suggesting a significant portion of the audience either lacks familiarity with the topic or prefers not to engage with it. This indicates a divide in understanding and awareness regarding the implications of the prohibition on AI systems among the respondents.

Question 18.62: Please specify the concrete AI system, how it is used in practice as well as the specific elements you would need further clarification in this regard

A significant number of respondents expressed confusion over the definitions of "targeted" and "untargeted" scraping, particularly in the context of facial recognition technologies. Many examples, such as AI systems used in refugee camps or by law enforcement, illustrate the ambiguity surrounding whether data collection methods that focus on specific locations or groups can still be classified as untargeted. Respondents called for clearer guidelines to delineate these terms, especially in mixed-use scenarios where both targeted and untargeted methods may be employed.

The ethical ramifications of AI systems that scrape facial images without explicit consent were a recurring theme. Concerns were raised about the potential misuse of data, particularly in cases involving vulnerable populations, such as refugees or individuals monitored by law

enforcement. Respondents highlighted the need for robust safeguards to prevent abuses, such as the creation of non-consensual content or the violation of privacy rights, emphasising the importance of aligning AI regulations with existing data protection laws like the GDPR.

Various use cases were presented, including AI systems for monitoring livestock, analysing customer demographics in retail, and law enforcement applications. Respondents sought clarification on whether these systems, which may not explicitly aim to create facial recognition databases, still fall under the prohibition of untargeted scraping. The need for guidance on compliance, particularly for systems that may inadvertently contribute to database expansion, was emphasised to ensure that organisations can operate within legal boundaries.

Questions regarding the implications of "placing on the market" and the responsibilities of developers, vendors, and end-users were frequently mentioned. Respondents were concerned about the liability associated with AI systems that use unlawfully collected data, particularly in cases where the systems are sold or offered in the EU. Clear definitions and responsibilities are needed to navigate the complexities of compliance and accountability in the AI landscape.

The responses highlighted a wide range of AI applications, from facial recognition in public safety to academic research using web scraping. **The most common use cases being law enforcement surveillance, public space monitoring refugee/border control systems, social media facial recognition and retail analytics.** This diversity underscores the necessity for regulatory oversight that can adapt to various contexts while protecting individual rights. Respondents advocated for the inclusion of civil society representatives in oversight processes to ensure that the deployment of AI technologies is transparent and accountable, particularly in sensitive areas like law enforcement and public surveillance. In summary, the analysis reveals a pressing need for clearer definitions, ethical considerations, compliance guidelines, and regulatory oversight in the deployment of AI systems, particularly those involving facial recognition and data scraping practices.

4.6 QUESTIONS IN RELATION TO EMOTION RECOGNITION

Question 19.63: Taking into account the provisions of the AI Act, what elements of the prohibition of emotion recognition in the areas of workplace and education do you think require further clarification in the Commission guidelines?

A significant number of respondents highlighted the need for further clarification on specific elements. The majority emphasised the importance of guidelines related to the identification or inference of emotions in natural persons (155 responses) and the application of these provisions within workplace and educational institutions (154 responses). Additionally, 143 respondents noted the necessity of clarifying exceptions for medical and safety reasons, while 43 responses focused on the aspects of market placement and usage of AI systems. A smaller group (45 responses) indicated that no further clarification was needed.

Question 19.64: Please explain why the elements selected above require further clarification and what needs to be further clarified in the Commission guidelines?

A significant number of respondents expressed the need for clearer definitions of terms such as "emotion," "identifying," and "inferring." Many highlighted the ambiguity surrounding what constitutes emotion recognition, particularly in distinguishing between biometric data and non-biometric data, such as text or voice analysis. There is a consensus that the guidelines should explicitly define what types of emotional expressions are included or excluded, particularly regarding observable expressions like smiles and physical states like fatigue.

Respondents raised concerns about the interpretation of "workplace" and "educational institutions," particularly in the context of remote work and hybrid learning environments. Many argued that the current definitions may not adequately cover modern work arrangements, such as gig work or online education, and called for examples to clarify how these contexts are treated under the prohibition. There is a strong desire for the guidelines to specify whether the prohibition applies only to employees or extends to customers and other individuals present in these settings.

The exceptions for medical and safety reasons were frequently cited as vague and potentially prone to misuse. Respondents requested clearer criteria for what qualifies as acceptable use under these exceptions, emphasising the need to prevent broad interpretations that could allow for invasive monitoring practices. Many expressed concern that the exceptions could be exploited to justify continuous emotional monitoring in workplaces and educational settings, which could infringe on individual rights.

There were calls for clarification on the responsibilities of AI system providers and deployers regarding the use of emotion recognition systems. Respondents sought guidance on how liability is distributed when these systems are misused or deployed beyond their intended purpose. This includes questions about the accountability of developers versus organisations that implement these systems in sensitive environments.

Many responses highlighted the tension between fostering innovation in AI technologies and protecting fundamental rights. While some respondents acknowledged the potential benefits of emotion recognition systems in enhancing user experiences and improving mental health support, they stressed the importance of establishing strict guidelines to prevent misuse. The need for a risk-based approach that prioritises ethical considerations and safeguards against discrimination was a recurring theme, with calls for the Commission to provide concrete examples of permissible and prohibited uses of emotion recognition technologies. In summary, the responses indicate a strong demand for clarity in definitions, contextual applications, exceptions, accountability, and the balance between innovation and individual rights in the guidelines surrounding emotion recognition systems.

Question 20.65: Do you have or know concrete examples of AI systems that in your opinion fulfil all elements of the prohibition described above?

A total of 74 respondents affirmed that they have or know of such examples, while 125 respondents indicated they do not. Additionally, 184 respondents chose not to answer the question.

Question 20.66: Please specify the concrete AI system, how it is used in practice and how all the necessary elements described above are fulfilled

The responses highlight various AI systems designed for emotion recognition, primarily used in workplace and educational settings. Notable examples include iBorderCTRL, which assesses emotional states of travellers during immigration interviews, and Rosalyn's StableSight Model, used for proctoring exams. These systems leverage biometric data, such as facial expressions and vocal tones, to infer emotions like stress or engagement. The respondents emphasise that these applications often exist within contexts characterised by significant power imbalances, raising ethical concerns about privacy, consent, and the potential for misuse.

A recurring theme in the responses is the ethical implications of using emotion recognition AI systems, particularly in environments where power dynamics are skewed, such as workplaces and educational institutions. Many respondents' express concerns about the potential for these systems to infringe on individual privacy and autonomy, leading to undue pressure on employees and students to conform to perceived emotional norms. The fear of constant surveillance and the impact on mental health and well-being are significant points of contention, with many arguing that these technologies could exacerbate existing inequalities.

The analysis reveals that many AI systems discussed fulfil the criteria outlined in the prohibition under the AI Act. These systems are commercially available, actively marketed, and used to infer emotions without medical or safety justifications. For instance, systems like HireVue, which analyse video interviews to assess candidates' emotional states, are highlighted as particularly problematic due to their discriminatory potential against individuals with disabilities. The responses collectively underscore the need for stringent regulations to prevent the deployment of such technologies in sensitive contexts.

Several responses point to the discriminatory nature of emotion recognition systems, particularly in hiring processes and educational assessments. The use of AI tools like HireVue has been criticised for disproportionately screening out candidates with disabilities, as these systems often misinterpret non-normative emotional expressions. This raises questions about the scientific validity of emotion recognition technologies and their ability to fairly assess individuals. The potential for bias and the ethical implications of using such systems in decision-making processes are significant concerns raised by respondents.

Question 21.67: Do you have or know concrete examples of AI systems where you need further clarification regarding certain elements of this prohibition to determine whether the AI system is in the scope of the prohibition or not?

A total of 94 respondents indicated that they require additional information to determine whether certain AI systems fall within the scope of the prohibition. In contrast, 112 respondents felt that no further clarification was necessary, while 177 respondents chose not to answer the question.

Question 21.68: Please specify the concrete AI system, how it is used in practice as well as the specific elements you would need further clarification in this regard

A significant number of respondents expressed the need for clearer guidelines regarding the prohibition of emotion recognition systems under the AI Act. Many emphasised that certain applications, particularly those related to health and safety, should not be classified as emotion recognition. For instance, systems that monitor physiological indicators for health emergencies or safety compliance were highlighted as essential tools that should be exempt from the prohibition, as they do not infer emotions but rather assess physical states.

Respondents raised concerns about the use of adaptive chatbots in customer service and employee support contexts. These systems, which adjust their responses based on inferred emotional states, were seen as beneficial for enhancing communication and resolving issues. However, there was a call for clarification on whether such systems would fall under the emotion recognition prohibition, especially when their primary function is to improve user experience without engaging in surveillance or performance evaluation.

The use of AI systems in educational settings to monitor student engagement and emotional states was a recurring theme. Respondents sought guidance on whether these applications, which analyse facial expressions and body language, would be considered emotion recognition. There was a particular interest in understanding how these systems could be justified under the medical or safety exceptions, especially when aimed at improving learning outcomes and mental well-being.

Several responses highlighted the potential for bias and discrimination in AI systems that analyse emotional states, particularly in recruitment and customer service contexts. Examples were provided of systems that misinterpret cultural expressions or rely on flawed datasets, raising ethical concerns. Respondents urged for stricter guidelines to prevent the deployment of such systems that could lead to discriminatory practices.

Question 22.69: Do you have or know concrete examples of AI systems that fulfil all necessary criteria for the prohibition to apply, but fall under the exception of medical and safety reasons?

A total of 64 respondents affirmed the existence of such examples, while 128 indicated they did not know of any. Additionally, 191 respondents chose not to answer the question,

highlighting a significant divide in awareness or knowledge about AI systems that could potentially fall under this exception.

Question 22.70: Please specify the concrete AI system, how it is used in practice and which exception would apply and why

The responses highlight a variety of AI systems employed in health and safety contexts, particularly focusing on emotion recognition and monitoring technologies. These systems are designed to analyse physiological and behavioural data, such as facial expressions, voice patterns, and biometric signals, to assess emotional states and potential risks. Examples include AI tools for monitoring fatigue in drivers, emotional distress in patients, and stress levels in high-risk jobs like air traffic control. The overarching goal of these systems is to enhance safety and well-being by providing timely interventions based on real-time data.

A significant theme in the responses is the distinction between legitimate medical systems and emotion recognition systems. Respondents argue that while medical systems, such as those monitoring vital signs or detecting health emergencies, are grounded in scientific evidence, emotion recognition systems often lack such rigor and can be considered pseudoscientific. This distinction is crucial for regulatory clarity, as it determines which systems may qualify for exceptions under the AI Act. The emphasis is on ensuring that systems used for health and safety purposes are based on objective physiological data rather than subjective emotional assessments.

Several respondents outline specific exceptions under which AI systems may operate without falling under prohibitive regulations. These include systems used for medical purposes, such as monitoring patients' emotional states during therapy, and safety applications like fatigue detection in drivers. The rationale for these exceptions is that they directly contribute to health care and safety, thereby justifying their use despite potential privacy concerns. Respondents stress the importance of maintaining strict oversight and transparency in the deployment of these systems to prevent misuse and protect individual rights.

4.7 QUESTIONS IN RELATION TO BIOMETRIC CATEGORISATION

Question 23.71: Taking into account the provisions of the AI Act, what elements of the prohibition of biometric categorisation to infer certain sensitive characteristics do you think require further clarification in the Commission guidelines?

The majority emphasised the need for clearer definitions related to the individual categorisation of natural persons based on biometric data (109 responses) and the implications of deducing sensitive attributes such as race, political opinions, and sexual orientation (also 109 responses). Additionally, 44 responses pointed to the importance of clarifying the

conditions surrounding the market placement and use of AI systems, while 93 responses noted that the exclusion of lawful labelling or filtering of biometric datasets, particularly in law enforcement, also requires further explanation. A smaller group of 61 respondents indicated that none of these elements needed clarification.

Question 23.72: Please explain why the elements selected above require further clarification and what needs to be further clarified in the Commission guidelines?

Many respondents emphasised the need for clearer definitions of terms such as "biometric categorisation," "placement on the market," and "putting into service." There is a consensus that the guidelines should delineate between individual and group categorisation, as well as clarify what constitutes lawful labelling or filtering of biometric datasets. Respondents expressed concerns about the ambiguity surrounding the application of these terms, particularly in relation to general-purpose AI systems that may inadvertently engage in prohibited practices.

A significant number of responses highlighted the necessity to differentiate between permissible uses of biometric data, especially in law enforcement contexts, and those that could lead to discrimination or violation of rights. Respondents called for explicit examples of acceptable practices, particularly regarding the labelling and filtering of lawfully acquired biometric datasets. The potential for misuse in law enforcement was a recurring concern, with calls for strict limitations and oversight mechanisms.

Respondents raised questions about the definitions of "deduction" and "inference," particularly regarding how these concepts apply to sensitive characteristics such as race, political opinions, and sexual orientation. There is a demand for clarity on whether indirect inferences, which may occur unintentionally, fall under the prohibition. Many emphasised the need for thresholds or criteria to determine when an inference crosses into prohibited territory, as well as guidance on how to assess the risks associated with such inferences.

The responses indicated a strong interest in understanding how the guidelines would apply to specific sectors, particularly healthcare and financial services. Respondents expressed concerns that the prohibition on biometric categorisation could hinder beneficial applications in these fields, such as patient management and compliance with financial regulations. Clarification on how the guidelines interact with existing regulations, such as the GDPR, was also deemed essential to ensure that legitimate uses of biometric data are not unduly restricted.

Many respondents underscored the ethical implications of biometric categorisation and the potential for discrimination. There was a call for the guidelines to reflect a commitment to protecting fundamental rights and preventing harm, particularly considering historical abuses associated with biometric data. Respondents urged the Commission to consider the broader societal impacts of biometric categorisation and to ensure that the guidelines promote

fairness, transparency, and accountability in AI systems. In summary, the responses reflect a strong desire for clarity, specificity, and ethical considerations in the Commission guidelines regarding biometric categorisation and its implications under the AI Act.

Question 24.73: Do you have or know concrete examples of AI systems that in your opinion fulfil all elements of the prohibition described above?

A majority of respondents (130) indicated that they do not know of any such examples, while 56 respondents affirmed that they do have concrete examples. Additionally, a significant number of respondents (197) chose not to answer the question, suggesting a lack of clarity or knowledge on the topic among a substantial portion of the participants. Overall, the data reflects a prevailing uncertainty or scepticism about the identification of AI systems that meet the specified criteria.

Question 24.74: Please specify the concrete AI system, how it is used in practice and how all the necessary elements described above are fulfilled

Respondents highlighted various AI systems, particularly those involved in biometric categorisation, such as facial recognition technologies (e.g., Viso AI's Deepface) and dialect recognition systems used in migration contexts. These systems analyse biometric data to infer sensitive characteristics like race, religion, and political affiliation, often for purposes such as targeted advertising, tenant screening, and migration assessments. The use of these technologies raises significant ethical and legal concerns, particularly regarding their potential to reinforce existing biases and discrimination.

The practical applications of these AI systems span multiple sectors, including housing, migration, and marketing. For instance, AI systems are employed by landlords for tenant screening, by migration authorities to assess asylum claims, and by companies for targeted advertising based on inferred characteristics. The responses indicate that these systems are actively used in decision-making processes that can significantly impact individuals' lives, often without their consent or awareness.

Many respondents emphasised that the identified AI systems fulfil the necessary elements of the prohibition against biometric categorisation as outlined in Article 5(1)(g). This includes their placement in the market, the categorisation of individuals based on biometric data, and the inference of sensitive characteristics. The responses collectively argue that these systems violate fundamental rights, such as privacy and dignity, by using biometric data to make potentially harmful inferences about individuals.

A recurring theme in the responses is the concern that these AI systems can lead to discriminatory outcomes, particularly against marginalised groups. For example, facial recognition technologies have been criticised for their higher error rates when identifying individuals with darker skin tones, leading to wrongful arrests and other negative

consequences. The potential for biased datasets to produce repeated discriminatory deductions was also highlighted, indicating a need for stricter regulations and oversight.

Question 25.75: Do you have or know concrete examples of AI systems where you need further clarification regarding certain elements of this prohibition to determine whether the AI system is in the scope of the prohibition or not?

A total of 40 respondents indicated that they require additional information to determine whether their AI systems fall within the scope of the prohibition. In contrast, 147 respondents felt that no further clarification was necessary, while 196 respondents chose not to answer the question.

Question 25.76: Please specify the concrete AI system, how it is used in practice as well as the specific elements you would need further clarification in this regard

A significant number of respondents highlighted the use of AI systems for biometric data collection and categorisation, particularly in sensitive contexts such as refugee processing, law enforcement, and healthcare. Concerns were raised about the potential for these systems to infringe on fundamental rights, especially when they categorise individuals based on sensitive traits like race, ethnicity, or religion. Respondents called for clearer guidelines on what constitutes prohibited categorisation under Article 5(1)(g) of the AI Act, particularly regarding indirect inferences made by AI systems.

Many responses expressed a need for further clarification on the scope of prohibitions related to biometric categorisation. Respondents sought guidance on whether systems that infer sensitive characteristics incidentally, rather than explicitly, fall under the prohibition. This ambiguity was particularly noted in contexts such as marketing analytics, customer segmentation, and workplace monitoring, where AI systems may inadvertently categorise individuals based on physical attributes or behaviours.

There was a recurring theme regarding the distinction between high-risk AI systems and those that are outright prohibited. Respondents requested detailed guidance on how to classify systems that may not directly infer sensitive traits but could do so as a byproduct of their operation. This distinction is crucial for organisations to navigate compliance with the AI Act effectively and to ensure responsible use of biometric technologies.

Question 26.77: Do you have or know concrete examples of AI systems that fulfil all necessary criteria for the prohibition to apply, but fall under the exception of medical and safety reasons?

A significant majority of respondents indicated a lack of concrete examples. Specifically, 149 respondents answered "no," while only 29 affirmed the existence of such examples, and 205 chose "not applicable." This suggests that while there is some recognition of potential exceptions, the overwhelming consensus is that clear instances of compliant AI systems are not readily identifiable.

Question 26.78: Please specify the concrete AI system, how it is used in practice and which exception would apply and why

The responses predominantly highlight the use of AI systems in law enforcement, particularly for processing and categorising biometric data such as fingerprints, facial recognition, and DNA. These systems are employed to enhance investigative efficiency by labelling and filtering lawfully acquired datasets. Examples include the Next Generation Identification (NGI) system, NEC's NeoFace, and Europol's Prüm II framework, which use AI to assist in identifying suspects and managing biometric information. The emphasis is on the lawful acquisition of data, often through judicial means, ensuring that the systems operate within legal frameworks.

Respondents provided various practical applications of AI systems, including analysing CCTV footage, filtering missing persons databases, and aiding in family reunification efforts in refugee camps. For instance, AI systems are used to identify individuals based on specific characteristics from lawfully obtained footage, thereby supporting criminal investigations. Additionally, systems like Clearview AI were mentioned, albeit with caution regarding their data acquisition methods. The overarching theme is the necessity of these systems in facilitating law enforcement activities while adhering to legal standards.

A significant portion of the responses focused on the exceptions that allow for the use of AI systems in processing biometric data. Many respondents argued that these systems could fall under exceptions for labelling or filtering lawfully acquired datasets, provided they do not infer sensitive characteristics such as race or religion. The justification for these exceptions hinges on the lawful acquisition of data and the systems' alignment with national and EU laws. Respondents stressed the importance of maintaining compliance with data protection regulations and ensuring that the systems are used strictly for investigative purposes.

4.8 QUESTIONS IN RELATION TO REAL-TIME REMOTE BIOMETRIC IDENTIFICATION

Question 27.79: Taking into account the provisions of the AI Act, what elements of the prohibition of real-time remote biometric identification for law enforcement purposes do you think require further clarification in the Commission guidelines?

The most significant focus was on the definition and implications of "real-time" identification, with 99 responses emphasising this aspect. Additionally, 86 respondents sought clarity on what constitutes a "remote biometric identification system," while 72 responses each addressed the terms "for law enforcement purposes" and "in publicly accessible spaces." 58 respondents indicated that none of the specified elements required further clarification.

Question 27.80: Please explain why the elements selected above require further clarification and what needs to be further clarified in the Commission guidelines?

A significant number of respondents expressed the need for clearer definitions of critical terms such as "real-time," "remote biometric identification system," and "publicly accessible spaces." Many highlighted the ambiguity surrounding what constitutes "real-time" processing, with calls for specific thresholds to distinguish between instantaneous and delayed processing. Additionally, the definition of "publicly accessible spaces" requires clarification to address contexts like semi-public areas, private venues, and the implications of access control.

Respondents sought more precise guidelines on what qualifies as "law enforcement purposes." There were concerns about the potential for overreach, particularly regarding the use of biometric identification for general surveillance under the guise of public safety. Specific examples and criteria for permissible uses were requested to prevent misuse and ensure that the guidelines do not inadvertently allow for broad surveillance practices.

The concept of "active involvement" of individuals in biometric identification processes was another area of concern. Respondents called for clarification on what constitutes active involvement and how it relates to the use of biometric systems. Additionally, there were warnings against the misuse of the term "authentication," emphasising that technical verification should not be conflated with authentication, which should be clearly defined in the guidelines.

Many responses highlighted the need for detailed technical specifications regarding the operation of RBI systems, particularly concerning processing delays and the potential for circumvention of the prohibition through minor delays. Respondents urged the Commission to establish minimum time thresholds for what constitutes "significant delay" and to outline necessary safeguards to protect fundamental rights.

Finally, there was a call for the guidelines to align with existing regulations, such as the GDPR, to ensure consistency across legal frameworks. Respondents emphasised the importance of delineating the responsibilities of various stakeholders, including technology providers and law enforcement agencies, to prevent ambiguity in compliance and accountability. Clear guidance on the interaction between the AI Act and other legal obligations was deemed essential for effective implementation and protection of individual rights. Overall, the responses reflect a strong desire for clarity and specificity in the Commission guidelines to ensure that the implementation of RBI systems respects fundamental rights while addressing legitimate security concerns.

Question 28.81: Do you have or know concrete examples of AI systems that in your opinion fulfil all elements of the prohibition described above?

A total of 36 respondents indicated that they have or know of concrete examples of AI systems requiring further clarification on specific elements of a prohibition to assess their applicability,

while 134 respondents stated they do not have such examples. Additionally, 213 respondents chose not to answer the question, suggesting a significant portion of the audience either lacks relevant examples or prefers not to engage with the inquiry.

Question 28.82: Please specify the concrete AI system, how it is used in practice and how all the necessary elements described above are fulfilled

Respondents provided a variety of examples of AI systems in practice, including real-time biometric identification at borders, airports, and public events, as well as monitoring systems in refugee camps and wildlife conservation efforts. These applications highlight the growing reliance on AI for security and identification purposes, but they also raise questions about the ethical implications and potential misuse of such technologies, particularly in sensitive environments.

A significant portion of the responses sought clarification on the legal framework surrounding the use of AI systems, particularly in relation to the AI Act's prohibition on real-time biometric identification. Respondents expressed uncertainty about what constitutes "real-time" processing, the definition of "publicly accessible spaces," and the distinction between law enforcement and private security purposes. This ambiguity complicates compliance for organisations using these technologies.

Question 29.83: Do you have or know concrete examples of AI systems that fulfil all necessary criteria for the prohibition to apply, but which could fall under one or more of the exceptions of Article 5(1)(h)(i) to (iii) AI Act?

A total of 28 respondents indicated "yes," while 139 answered "no," and 216 chose "not applicable." This indicates a significant majority of respondents do not identify specific AI systems that would fit within the exceptions.

Question 29.84: Please specify the concrete AI system, how it is used in practice and which exception would apply and why.

Respondents expressed concerns regarding the ambiguity of the term "specifically targeted individual" in the context of real-time remote biometric identification systems. The lack of clarity on whether an individual must be identified prior to an incident or can be identified post-incident raises significant interpretative challenges. For example, in scenarios like a terrorist attack, the ability to identify suspects based on prior knowledge or through real-time identification of individuals in the vicinity of the incident is debated. This ambiguity necessitates further clarification to ensure consistent application of the law and to protect individual rights.

Several responses highlighted the use of AI systems in humanitarian contexts, particularly in refugee camps for locating missing persons or victims of trafficking. These systems utilise real-time biometric identification to cross-reference individuals with existing databases, thereby

facilitating targeted searches rather than mass surveillance. Respondents emphasised that such applications align with the exceptions outlined in Article 5(1)(h)(i) of the AI Act, as they focus on vulnerable individuals and operate under strict geographic and temporal constraints, ensuring compliance with data protection laws.

A recurring theme in the responses was the potential for AI systems, particularly facial recognition technologies, to normalise mass surveillance practices. While some systems may technically qualify under exceptions for preventing imminent threats, respondents cautioned against the risks of broad surveillance capabilities being misused. Examples included city-wide facial recognition networks and transport hub scanning systems, which could infringe on privacy rights and lead to unwarranted monitoring of the general public, thus necessitating stringent oversight and clear limitations.

Respondents provided various examples of AI systems used by law enforcement, such as real-time biometric identification during public events to monitor for potential threats. These applications were generally seen as compliant with the exceptions in the AI Act, particularly when they are geographically and temporally limited and authorised by judicial oversight. However, there were calls for clearer definitions of the conditions under which these systems can be activated, especially in high-traffic areas where incidental surveillance of uninvolved individuals may occur.

Question 30.85: Do you need further clarification regarding one or more of the exceptions of Article 5(1)(h)(i) to (iii) AI Act or the conditions or safeguards under Article 5(2) to (7) AI Act?

A majority of respondents indicated they do not require additional information, with 122 responses stating "no." However, a significant minority of 54 respondents expressed a need for further clarification, while 207 respondents chose not to answer the question.

Question 30.86: Please specify the concrete condition or safeguard and the issues for you need further clarification; please provide concrete examples

A significant number of responses express apprehension regarding the impact of real-time remote biometric identification systems (RBIS) on fundamental rights. Respondents argue that the use of such systems often leads to mass surveillance, which is contrary to the Charter of Fundamental Rights. They emphasise the need for clear guidelines to ensure that any limitations on rights are necessary and proportionate, particularly considering the Italian Data Protection Authority's findings on mass surveillance.

Many respondents call for precise definitions and criteria related to terms such as "imminent threat," "substantial threat," and "targeted search." There is a consensus that vague language could lead to misuse of the exceptions outlined in the AI Act, allowing for overreach in surveillance practices. Respondents seek clarity on how these terms should be interpreted and applied in various contexts, including public safety and law enforcement.

The responses highlight a need for detailed guidance on the processes for obtaining prior authorisation for the use of RBIS, especially in urgent situations. Respondents question whether expedited processes are permissible and what documentation is required for judicial oversight. There is a call for transparency in how these authorisations are granted and the criteria that must be met to ensure accountability.

Respondents' express confusion regarding the temporal and geographic limitations of RBIS deployment. They seek clarification on how these limitations are defined and enforced, particularly in dynamic situations such as public events or emergencies. Specific examples are requested to illustrate how these limitations would be applied in practice, ensuring that the use of RBIS does not lead to excessive surveillance.

There is a strong demand for comprehensive guidelines on conducting fundamental rights impact assessments (FRIAs) when deploying RBIS. Respondents emphasise the importance of balancing public safety with individual rights and seek standardised methodologies for assessing potential infringements. Additionally, they call for clear safeguards to prevent misuse of data collected during RBIS operations, including retention periods and data minimisation practices. Overall, the responses reflect a deep concern for the implications of AI technologies on civil liberties, emphasising the need for robust regulatory frameworks that prioritise transparency, accountability, and the protection of fundamental rights.

4.9 QUESTION IN RELATION TO INTERPLAY WITH OTHER UNION LEGISLATION

Question 31.87: Do you have or know concrete examples of AI systems where you need further clarification regarding the application of one or more of the prohibitions under the AI Act in relation to other Union legislation?

A total of 102 respondents indicated that they have concrete examples requiring clarification, while 99 respondents stated they do not have such examples. Additionally, 182 respondents chose not to answer the question, highlighting a significant interest in understanding the interplay between the AI Act and existing legislation among a portion of the respondents.

Question 31.88: Please specify the concrete AI system and the prohibition under the AI Act, the relevant provision of a specific Union legislation and where further clarification is needed

A significant number of respondents emphasised the need for clearer guidelines on how the AI Act's prohibitions interact with existing EU legislation, particularly the GDPR, Digital Services Act (DSA, 2022), and consumer protection laws. Many highlighted the complexities arising from overlapping regulations, especially concerning biometric data processing, emotion recognition,

and manipulative practices. Respondents expressed concerns that the lack of clarity could lead to legal ambiguities and hinder compliance efforts.

A recurring theme in the responses was the importance of grounding the AI Act in international human rights law and the EU Charter of Fundamental Rights. Respondents urged that the prohibitions should serve a preventative purpose to protect fundamental rights, particularly against discrimination and bias. Examples were provided, such as the use of scoring systems in visa applications that perpetuate racial profiling, underscoring the need for a broader interpretation of harm prevention in the context of AI systems.

Several responses detailed specific AI systems, such as emotion recognition tools in workplaces and AI-driven credit scoring systems, highlighting the regulatory challenges they face under the AI Act and GDPR. Respondents called for guidance on how these systems can comply with both sets of regulations, particularly regarding consent and the use of sensitive data. The interplay between AI systems in healthcare and existing medical device regulations was also a point of contention, with calls for clearer definitions and risk assessment protocols.

The responses indicated a strong concern regarding the potential for AI systems to manipulate consumer behaviour, particularly through targeted advertising and dark patterns. Respondents sought clarification on how the AI Act's prohibitions on manipulative practices align with existing consumer protection laws, emphasising the need for a coherent regulatory framework that protects consumers without stifling innovation.

Many respondents expressed concerns about the compliance burdens that the AI Act may impose, particularly on small and medium-sized enterprises (SMEs). There were calls for the European Commission to streamline regulations to avoid conflicting obligations and ensure that compliance is feasible for smaller organisations. The need for practical guidance and illustrative examples to clarify the application of the AI Act in various contexts was highlighted as essential for fostering a balanced regulatory environment that promotes innovation while ensuring consumer protection and fundamental rights. Overall, the responses reflect a complex landscape of legal, ethical, and practical considerations surrounding the implementation of the AI Act, with a strong emphasis on the need for clarity, coherence, and alignment with existing laws to protect individual rights and promote responsible AI use.

5.0 References

European Commission, *Commission launches consultation on AI Act prohibitions and AI system definition*, 13 November 2024, <https://digital-strategy.ec.europa.eu/en/news/commission-launches-consultation-ai-act-prohibitions-and-ai-system-definition>

Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation), <http://data.europa.eu/eli/reg/2016/679/oj>

Regulation (EU) 2022/2065 of the European Parliament and of the Council of 19 October 2022 on a Single Market For Digital Services and amending Directive 2000/31/EC (Digital Services Act), <http://data.europa.eu/eli/reg/2022/2065/oj>

Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act), <http://data.europa.eu/eli/reg/2024/1689/oj>

**Europe Direct is a service to help you find answers
to your questions about the European Union.**

Freephone number (*):

00 800 6 7 8 9 10 11

(*) The information given is free, as are most calls (though some operators, phone boxes or hotels may charge you).

HOW TO OBTAIN EU PUBLICATIONS

Free publications:

- one copy:

via EU Bookshop (<http://bookshop.europa.eu>);

- more than one copy or posters/maps:
from the European Union's representations
(http://ec.europa.eu/represent_en.htm);
from the delegations in non-EU countries
(http://eeas.europa.eu/delegations/index_en.htm);
by contacting the Europe Direct service
(http://europa.eu/eurodirect/index_en.htm) or calling 00 800 6 7 8 9 10 11
(freephone number from anywhere in the EU) (*).

(*) The information given is free, as are most calls (though some operators, phone boxes or hotels may charge you).

Priced publications:

- via EU Bookshop (<http://bookshop.europa.eu>).

Priced subscriptions:

- via one of the sales agents of the Publications Office of the European Union
(http://publications.europa.eu/others/agents/index_en.htm).

Place du Congrès 1, B-1000 Brussels

Tel. +32 2 229 39 11

www.ceps.eu



**Publications Office
of the European Union**