

Stakeholder consultation on the Draft Guidelines on the classification of high-risk AI systems under Article 6 of the AI Act

Fields marked with * are mandatory.

Stakeholder consultation on the Draft Guidelines on the classification of high-risk AI systems under Article 6 of the AI Act

Disclaimer: This document is a working document of the AI Office for the purpose of consultation and does not prejudice the final decision that the Commission may take on the final guidelines. The responses to this consultation paper aim to provide relevant stakeholder input to the Commission when finalising the guidelines.

This consultation is targeted to stakeholders of different categories. These categories include, but are not limited to, providers and deployers of (high-risk) AI systems, other industry organisations, as well as academia, other independent experts, civil society organisations, and public authorities and the citizens.

The Artificial Intelligence Act (the 'AI Act'), which entered into force on 1 August 2024, creates a single market and harmonised rules for trustworthy and human-centric Artificial Intelligence (AI) in the EU. It aims to promote innovation and uptake of AI, while ensuring a high level of protection of health, safety and fundamental rights, including democracy and the rule of law.

The AI Act follows a risk-based approach classifying AI systems into different risk categories, one of which is the high-risk AI systems (Chapter III of the AI Act).

The AI Act distinguishes between two categories of AI systems that are considered as 'high-risk', as set out in Article 6(1) and 6(2) AI Act. Article 6(1) AI Act covers AI systems that are embedded as safety components in products or that themselves are products covered by Union legislation in Annex I, which could have an adverse impact on health and safety of persons. Article 6(2) AI Act covers AI systems that in view of their intended purpose are considered to pose a significant risk to health, safety or fundamental rights. The AI Act lists eight areas in which AI systems could pose such significant risk to health, safety or fundamental rights in Annex III and, within each area, lists specific use-cases that are to be classified as high-risk. Article 6(3) AI Act provides for exemptions for AI systems that are intended to be used for one of the use-cases listed in Annex III, but which do not pose significant risk since they fall under one of the exceptions listed in Article 6(3) AI Act.

With the [AI Omnibus](#) the entry into application of the requirements and obligations for high-risk AI systems under Article 6(2) and Annex III will be postponed to 2 December 2027. As regards the AI systems under Article 6(1) and Annex I the rules will apply from 2 August 2028. Other changes introduced to the classification rules under Article 6(1) and the definition of a 'safety component' are considered in the draft guidelines.

The draft Guidelines relate exclusively to the classification of high-risk AI systems. Further guidance from the Commission will follow, as indicated in the list of forthcoming guidelines published under: [Supporting the implementation of the AI Act with clear guidelines | Shaping Europe's digital future](#).

The draft Guidelines are prepared pursuant to Article 6(5) AI Act and focus on the rules for the high-risk classification. It is further required that these draft Guidelines should be accompanied with a comprehensive list of practical examples of use-cases of AI systems that are high-risk and not high-risk.

These draft Guidelines take into account the outcome of a stakeholder consultation organised by the Commission in June-July 2025 and input provided by the Member States in the context of the European Artificial Intelligence Board ('AI Board').

In addition to the written consultation process, the AI Office organised a series of targeted stakeholder workshops in December 2025 in order to gather further practical and sector-specific feedback on the draft Guidelines.

To complement these discussions and ensure the inclusion of a broad range of stakeholders, including representatives from industry, academia, civil society, public authorities and other relevant organisations, the AI Office also circulated a follow-up questionnaire aimed at collecting additional written input on specific aspects of the high-risk classification framework and its practical application.

On the basis of this input, the Commission prepared the draft guidelines which are now open for stakeholder feedback.

This second consultation accompanies the draft Guidelines on the classification of high-risk AI systems under Article 6 AI Act that aim to support consistent interpretation and effective implementation of the AI Act classification rules for high-risk AI systems across the EU.

This consultation seeks to collect targeted feedback on the clarity, completeness and practical usability of the draft Guidelines prepared by the Commission based on broad stakeholder input.

In particular, the Commission seeks feedback on whether the explanations, methodologies and examples provided sufficiently support stakeholders in determining whether an AI system qualifies as high-risk under the AI Act.

The draft Guidelines address the classification of high-risk AI systems in accordance with Article 6 AI Act and are structured as follows:

- **Section I** presents an introduction outlining the purpose, scope and legal context of the draft Guidelines. It recalls the objectives and risk-based structure of the AI Act, situates the concept of high-risk AI systems within that framework, and explains the legal basis and role of the draft Guidelines pursuant to Article 6(5) AI Act.
- **Section II** sets out general principles for the classification of high-risk AI systems, including the requirement that the system qualifies as an AI system and the role of the intended purpose for its classification.
- **Section III** explains the classification of high-risk AI systems under Article 6(1) AI Act and Annex I, including the key elements (such as the notion of a safety component and third-party conformity assessment).
- **Section IV** addresses the classification of high-risk AI systems under Article 6(2) AI Act and Annex III, including:

horizontal issues applicable across Annex III use cases (such as human involvement, the “intended to be used” criterion, and the interaction with national and EU law);

the application of the filter mechanism under Article 6(3) AI Act, including its conditions, exceptions, and safeguards;

detailed guidance on each high-risk use case listed under the eight areas in Annex III (biometrics; critical infrastructure; education and vocational training; employment; access to essential private and public services; law enforcement; migration, asylum and border control management; and administration of justice and democratic processes).

- **Section V** explains the transitional and grandfathering provisions of the AI Act, clarifying the entry into application of the high-risk AI system requirements and the conditions under which AI systems placed on the market or put into service before the relevant application dates remain subject to the new obligations.

All participants are invited to provide feedback on the particular sections of the draft Guidelines they are

interested in. In particular, respondents are encouraged to focus on:

- whether specific sections or subsections require further clarification or refinement to ensure effective implementation and compliance; and
- whether additional or current examples or use-cases should be included or excluded, taking into account the legal provisions of the AI Act.

The feedback collected through this consultation will support the Commission in refining and finalising the draft Guidelines, with the objective of ensuring they are clear, comprehensive, and practically useful for all stakeholders involved in the development, deployment and supervision of AI systems under the AI Act.

As not all sections may be relevant for all stakeholders, respondents may reply only to the section(s) they would like. Respondents are encouraged to provide **explanations and practical cases** as part of their responses to support the practical usefulness of the draft Guidelines. We kindly ask that the respondents to specify the exact section and paragraph of the draft Guidelines to which their comments refer.

The consultation also aims to collect input for the annual review of the list of prohibitions (Article 5 AI Act) and the high-risk use cases in Annex III AI Act to inform the Commission review pursuant to Article 112(1) AI Act.

The targeted consultation is available in **English only** and will be **open for 5 weeks starting on 19 May until 23 July 2026 (22:00 CET)**.

All contributions to this consultation may be made publicly available. Therefore, please do not share any personal or confidential information in your contribution (your written feedback). It is your responsibility to avoid personal data and any reference in your written feedback itself that would reveal your identity.

For information on how the Commission processes your personal data please read our Privacy Statement

[20250122_Public_Consultation_classification_of_high-risk_AI_systems_PrivacyStatement.pdf](#)

Information about the respondent

* First name

Stephanos

* Surname

Hadjistyllis

* Email address

info@actuary.eu

* Do you agree that we may publish your identity together with your contribution in the instance that all contributions are made publicly available?

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 respondent category that you selected for this consultation, your answer regarding residence, and your contribution may be published as received. Should you choose to remain anonymous, your name will not be published. Please do not include any personal data in the contribution itself.

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. Only the following organisation details may be published: The respondent category that you selected for this consultation, the name of the organisation on whose behalf you reply as well as its size, its presence in or outside the EU and your contribution as received. Should you choose to remain anonymous, your name will not be published. Please do not include any personal data in the contribution itself if you want to remain anonymous.

- ☒ Yes
☐ No

* Do you agree that we may contact you in the event of follow-up questions or if we want to learn more about your responses?

- ☒ Yes
☐ No

* Do you represent an organisation (e.g., a public organisation, a company, a think tank or a civil society /consumer organisation) or act in your personal capacity (e.g., independent expert)?

- ☒ Organisation
☐ In a personal capacity

* If you are representing an organisation, please specify the name of the organisation:

Actuarial Association of Europe

* Type of organisation

Association

* Is a representation of the organisation located in the EU?

- ☐ The organisation's headquarter is located in the EU
- ☒ A branch office, or any representation of the organisation is located in the EU
- ☐ None of the representations of the organisation is located in the EU

* Select the EU Member State where the organisation's headquarter, or representation is located

BE - Belgium

* Select the size of the organisation

Othe (e.g. multiple organisations)

* Sector(s) of activity

- | | | |
|--|--|--|
| ☐ Information technology | ☐ Employment | ☐ Transport |
| ☐ Public administration | ☐ Education and training | ☐ Telecommunications |
| ☐ Law enforcement | ☐ Consumer services | ☐ Retail |
| ☐ Justice sector | ☐ Business services | ☐ E-commerce |
| ☐ Legal services sector | ☒ Banking and finances | ☐ Advertising |
| ☐ Cultural and creative sector, including media | ☐ Manufacturing | ☐ Consumer protection |
| ☐ Healthcare | ☐ Energy | ☒ Others |

* Please, specify

30 character(s) maximum

Insurance, Pensions, Risk

* Describe the activities of your organisation or yourself

1300 character(s) maximum

The Actuarial Association of Europe (AAE) was established in 1978 under the name Groupe Consultatif to represent actuarial associations in Europe. Its purpose is to provide advice and opinions to the various organisations of the European Union – the Commission, the Council of Ministers, the European Parliament, EIOPA and their various committees – on actuarial issues in European legislation. The AAE currently has 38 member associations in 37 European countries, representing over 29,000 actuaries. Advice and comments provided by the AAE on behalf of the European actuarial profession are totally independent of industry interests.

* **On which part(s) of the draft Guidelines would like comment?** *Multiple answers are possible. Please note that selecting a particular answer will direct you to a set of questions specifically related to subject specified.*

- ☒ **General principles for classification of high-risk AI systems under Article 6.** (Section II)
- ☒ **High-risk classification according to Article 6(1) – Annex I.** (Section III)
- ☒ **High-risk classification according to Article 6(2) – Annex III.** (Section IV)
- ☒ **Questions in relation to Article 112 paragraph 1 (Evaluation and Review)**

General principles for classification of high-risk AI systems under Article 6 (Section II of the draft Guidelines)

Select the relevant subsection(s):

- ☒ II.1 The system must be an AI system
- ☒ II.2 Intended purpose(s) of the AI system

II.1 The system must be an AI system

Question A. Are there any aspects / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance? Please indicate specific parts of the draft Guidelines (i.e., paragraph(s) number).

1500 character(s) maximum

We welcome the Guidelines confirming that a system must first meet the definition of an AI system under Article 3(1) before any high-risk classification arises (paras 8-9). This gateway is important for our sector. Many established actuarial and statistical methods - for example generalised linear models (GLMs), generalised additive models (GAMs), Cox proportional hazards models, mortality-table methods and standard optimisation techniques - are manually configured, transparent and do not operate autonomously. Consistent with the Commission's Guidelines on the definition of an AI system (C(2025) 5053), such methods generally do not 'transcend basic data processing' and so fall outside the AI system definition. We suggest the Guidelines cross-refer to this explicitly. In our view, classification should turn on whether a system is an AI system, its degree of autonomy and human oversight, and its individual impact - not on the line of business in which it is used. On that logic, the same treatment should apply to transparent, human-supervised methods wherever used, including in life, health, pension and annuity pricing, so that the Annex III use cases (in particular point 5(c)) are not read as capturing traditional actuarial techniques that are not AI systems. This also aligns with EIOPA's letter of 13 April 2026, which the AAE has endorsed, and with ECB Opinion CON/2026/10 — reflecting a convergent European supervisory position.

II.2 Intended purpose(s) of the AI system

Question A. Are there any aspects / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance? Please indicate specific parts of the draft Guidelines (i.e., paragraph(s) number).

1500 character(s) maximum

We welcome the clarification in paras 10-12 that intended purpose drives classification and must be described clearly and coherently across all materials. In insurance, the same model may support several functions (e.g. underwriting, claims, customer service), each carrying different risks. We suggest the Guidelines clarify that intended purpose should reflect actual and reasonably foreseeable use, so that purpose is assessed on the basis of actual and foreseeable use, while also recognising that a tool used only to inform or support a human decision (decision-support) differs from one that determines an outcome autonomously. It would help to clarify how routine, pre-planned maintenance - such as recalibration or refreshing data and reference tables (e.g. mortality tables) without changing the model's logic or purpose - relates to the intended-purpose assessment, and that such updates do not in themselves create or change a high-risk classification. Practical, sector-specific examples would support consistent interpretation across Member States.

High-risk classification according to Article 6(1) – Annex I (Section III of the draft Guidelines)

Select the relevant subsection(s):

- ☒ III.1 The rationale
- ☐ III.2 Main elements for classification as high-risk

III.1 The rationale

Question A. Are there any aspects / sections / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance?

- ☒ Yes
- ☐ No
- ☐ N/A

Relevant paragraph number

Only values of at most 1000 are allowed

Please specify

500 character(s) maximum

We suggest clarifying that AI systems used for pricing, underwriting, reserving or capital assessment in insurance do not, in themselves, perform a product 'safety component' function within the meaning of Annex I, since they do not address physical safety of products. Confirming that such analytical and administrative systems are generally outside the scope of Article 6(1) would prevent misinterpretation, while genuine physical-safety functions remain covered.

Do you have more suggestions?

- ☐ Yes

☒ No

Question B. Do you think there are additional examples or use-cases , relevant for this Section of the draft Guidelines, that you find important that the Guidelines include or exclude from the scope of high-risk classification, taking into account the legal provisions in the AI Act?

- ☐ Yes
☒ No
☐ N/A

High-risk classification according to Article 6(2) – Annex III (Section IV of the draft Guidelines)

Select the relevant subsection(s):

- ☒ IV.1 Rationale and approach
☒ IV.2 Horizontal issues applicable to Annex III use cases
☒ IV.3 High-risk AI systems in the areas of Annex III

IV.1 Rationale and approach

Question A. Are there any aspects / sections / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance?

- ☒ Yes
☐ No
☐ N/A

Relevant paragraph number

Only values of at most 1000 are allowed

Please specify

500 character(s) maximum

We welcome the proportionate, risk-based approach in Section IV. We suggest the Guidelines state clearly that classification under Annex III turns on the intended purpose and the system's individual impact and degree of autonomy, rather than on the statistical complexity of the model or the line of business in which it is used. This helps ensure transparent, human-supervised actuarial tools are treated consistently, and that supervisory attention focuses on genuinely autonomous systems.

Do you have more suggestions?

- ☐ Yes
☐ No

Question B. Do you think there are additional examples or use-cases , relevant for this Section of the draft Guidelines, that you find important that the Guidelines include or exclude from the scope of high-risk classification, taking into account the legal provisions in the AI Act? If so, please explain why you think they should be treated as such, in order to help the Commission refine the draft Guidelines and ensure they are comprehensive.

- ☒ Yes
- ☐ No
- ☐ N/A

Please specify:

2500 character(s) maximum

No additional use cases are proposed here; our specific comments are provided under section 3.5 and the Article 112 questions. As a general point on scope, we suggest the rationale make clear that high-risk classification is intended to capture AI systems whose outputs have a direct and significant impact on individuals' access to, or the terms of, essential services, and which operate with a degree of autonomy. Established actuarial models that are transparent, interpretable and subject to meaningful human oversight - and that inform rather than determine decisions - should not be classified as high-risk by reason of statistical sophistication alone. The same proportionate treatment should apply consistently across lines of business, including life, health, pensions and annuities, so that transparent, human-supervised methods are not high-risk merely because of the product to which they relate. Framing the rationale in these terms would support proportionate and consistent application across the eight areas.

IV.2 Horizontal issues applicable to Annex III use cases

Question A. Are there any aspects / sections / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance?

- ☒ Yes
- ☐ No
- ☐ N/A

Relevant paragraph number

Only values of at most 1000 are allowed

Please specify

500 character(s) maximum

On the use of AI within only part of a process: we suggest clarifying that where an AI system is used for one step (e.g. data preparation, clustering or analysis) but the pricing or risk-assessment decision is produced by a separate non-AI model or by a human, the use of AI upstream does not automatically render the entire process high-risk. Classification should follow the boundaries of the AI system as defined in Article 3(1), assessed case by case.

Do you have more suggestions?

- ☒ Yes
☐ No

Question A. Are there any aspects / sections / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance?

- ☒ Yes
☐ No
☐ N/A

Relevant paragraph number

Only values of at most 1000 are allowed

Please specify

500 character(s) maximum

The Solvency II Actuarial Function provides a concrete, existing model for the meaningful human oversight the AI Act requires — actuaries validate, challenge and adjust model outputs before use under binding professional and regulatory obligations. We propose the Guidelines recognise actuarial sign-off under the Actuarial Function as constituting meaningful human oversight for Article 14 and Article 6(3) purposes, and invite the Commission to leverage existing governance rather than duplicate it

Do you have more suggestions?

- ☐ Yes
☒ No

Question B. Do you think there are additional examples or use-cases , relevant for this Section of the draft Guidelines, that you find important that the Guidelines include or exclude from the scope of high-risk classification, taking into account the legal provisions in the AI Act?

- ☒ Yes
☐ No
☐ N/A

If so:

- a) please explain why do you think they should be treated as such and
b) if your suggested example(s) or use case(s) are to be considered as in or out of scope.

3000 character(s) maximum

We suggest the horizontal guidance address three recurring questions in insurance. First, the interaction between the use of AI in part of a workflow and the classification of the workflow as a whole. A clear statement that classification attaches to the AI system as defined in Article 3(1), and not automatically to every downstream or upstream step, would prevent over-broad capture - for instance, clustering claims data with AI to inform risk groups, where premiums are then set by a transparent, human-supervised method. Second, the role of human oversight. Where a human meaningfully reviews and can override outputs, and the system informs rather than determines the outcome, this should be reflected in the assessment, including under the Article 6(3) filter. Third, the interaction between the Article 6(3) filter and GDPR profiling. The Guidelines clarify that Article 6(3) - the mechanism through which AI systems may avoid high-risk classification - is unavailable where a system performs profiling within the meaning of GDPR Article 4(4) (paras 89, 109-111), which individual insurance pricing models plausibly do. The Guidelines should provide practical guidance on how this interacts with insurance pricing in practice, and what alternative safeguards are available. Sector-specific examples illustrating these boundaries would be valuable, and we would welcome engagement with the actuarial profession and EIOPA in developing them.

IV.3 High-risk AI systems in the areas of Annex III

Please select the sector(s) you wish to comment on. *Multiple answers are possible*

- ☐ 3.1 Biometrics
- ☐ 3.2 Critical infrastructure
- ☐ 3.3 Education and vocational training
- ☐ 3.4 Employment, workers' management and access to self-employment
- ☒ 3.5 Access to and enjoyment of essential private services and essential public services and benefits
- ☐ 3.6 Law enforcement
- ☐ 3.7 Migration, asylum and border control management
- ☐ 3.8 Administration of justice and democratic processes

Question A. Are there any aspects / sections / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance?

- ☒ Yes
- ☐ No
- ☐ N/A

Relevant paragraph number

Only values of at most 1000 are allowed

317

Please specify

500 character(s) maximum

Para 317 extends point 5(c) to personal pension products. We do not object to treating comparable risks consistently, but suggest two clarifications. First, the same proportionate approach must apply: traditional, transparent and human-supervised methods used in pension and annuity pricing (e.g. mortality tables and GLMs) are not high-risk, just as in life and health. Second, the term 'pension' is not defined in the AI Act; the boundary of covered products should be clarified.

Do you have more suggestions?

- ☒ Yes
☐ No

Question A. Are there any aspects / sections / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance?

- ☒ Yes
☐ No
☐ N/A

Relevant paragraph number

Only values of at most 1000 are allowed

321

Please specify

500 character(s) maximum

The Guidelines confirm that claims-management AI falls outside point 5(c), and we recognise that policy reasons may support this approach. We would nonetheless welcome guidance on why the same impact-on-livelihood rationale does not extend to post-contract decisions: a claims denial — particularly in health or life insurance — can affect a person as materially as an underwriting exclusion. Clarification of this distinction would support consistent interpretation.

Do you have more suggestions?

- ☒ Yes
☐ No

Question A. Are there any aspects / sections / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance?

- ☒ Yes
☐ No
☐ N/A

Relevant paragraph number

Only values of at most 1000 are allowed

314

Please specify

500 character(s) maximum

Para 314 draws the boundary by product 'form' (life/health in, non-life out) rather than economic 'substance'. Since point 5(c) uses 'life and health insurance' as an autonomous concept with no reference to Solvency II or national branch classifications, identical AI models applied to identical risks may be high-risk in one legal wrapper and out of scope in another. The Guidelines should acknowledge this regulatory arbitrage risk and apply a substance-based test rather than product label alone.

Do you have more suggestions?

- ☒ Yes
☐ No

Question A. Are there any aspects / sections / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance?

- ☒ Yes
☐ No
☐ N/A

Relevant paragraph number

Only values of at most 1000 are allowed

325

Please specify

500 character(s) maximum

We welcome the clarification (paras 322-328) that systems used solely for prudential capital calculation are not high-risk, and note the key principle that a system's 'main purpose' is irrelevant once any high-risk use is present (para 325). We note that the examples in this section draw primarily on banking analogues; we would welcome equivalent guidance on Solvency II internal-model governance, and suggest EIOPA be consulted to ensure the treatment of insurance internal models is fully address

Do you have more suggestions?

- ☒ Yes
☐ No

Question A. Are there any aspects / sections / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance?

- ☒ Yes

- ☐ No
- ☐ N/A

Relevant paragraph number

Only values of at most 1000 are allowed

310

Please specify

500 character(s) maximum

Consistent with Section II.1, we suggest the Guidelines confirm that transparent, human-supervised actuarial techniques (e.g. GLMs, GAMs, Cox proportional hazards models and mortality-table methods) used for risk assessment and pricing are not high-risk where they do not meet the AI system definition, or inform rather than autonomously determine decisions. This applies across life, health, pensions and annuities, and reflects EIOPA's letter of 13 April 2026, which the AAE endorsed.

Question B. Do you think there are additional examples or use-cases , relevant for this Section of the draft Guidelines, that you find important that the Guidelines include or exclude from the scope of high-risk classification, taking into account the legal provisions in the AI Act?

- ☒ Yes
- ☐ No
- ☐ N/A

If so:

- a) please explain why do you think they should be treated as such and
- b) if your suggested example(s) or use case(s) are to be considered as in or out of scope.

2500 character(s) maximum

On examples and scope, we suggest the Guidelines: (a) distinguish more clearly between fully autonomous pricing or underwriting engines and decision-support tools that assist actuaries or underwriters, with high-risk treatment focused on the former; (b) clarify the treatment of hybrid systems serving both a covered use (risk assessment or pricing) and an excluded use (e.g. fraud detection, claims handling or product design), including how the boundary between distinct systems is drawn under Article 3(1); (c) confirm that internal models and reinsurance-specific tools are only in scope where they are also intended to be used for individual-level risk assessment or pricing; and (d) consider whether AI systems used to manage investment strategies in insurance-based investment products (IBIPs) fall within the rationale for point 5(c), given the potential for material individual impact on policyholders through investment decisions. A worked example showing a transparent, human-supervised pricing process that falls outside scope, alongside the existing examples of systems that fall within it, would aid consistent interpretation.

Question A. Are there any aspects / sections / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance?

- ☐ Yes
- ☒ No

☐ N/A

Do you have more suggestions?

☐ Yes

☒ No

Question B. Do you think there are additional examples or use-cases , relevant for this Section of the draft Guidelines, that you find important that the Guidelines include or exclude from the scope of high-risk classification, taking into account the legal provisions in the AI Act?

☐ Yes

☒ No

☐ N/A

3.5 Access to and enjoyment of essential services and benefits3.4 Employment

Question A. Are there any aspects / sections / paragraphs of this Section of the draft Guidelines that you believe require further clarification, and if so, how would you suggest they be improved to ensure effective implementation and compliance?

☐ Yes

☒ No

☐ N/A

Question B. Do you think there are additional examples or use-cases , relevant for this Section of the draft Guidelines, that you find important that the Guidelines include or exclude from the scope of high-risk classification, taking into account the legal provisions in the AI Act?

☐ Yes

☒ No

☐ N/A

Questions in relation to Article 112 paragraph 1 (Evaluation and Review)

Under Article 112(1) AI Act, the Commission is tasked with annually assessing the need to amend the list of high-risk AI systems set out in Annex III AI Act and the list of prohibited AI practices laid down in Article 5 AI Act. In this context, the Commission aims to gather insights from a wide range of stakeholders and invites you to share your views on whether amendments to the list of prohibited practices and the high-risk use cases in Annex III are needed.

Question 1. Are you aware of concrete use cases of AI systems that in your opinion need to be added to the list of use cases in Annex III - within the existing 8 areas - to ensure that all AI systems having equivalent risk of harm (in accordance with the criteria laid down pursuant to Article 7(2) AI Act) are classified as high-risk under the Annex III AI Act?

☒ Yes

- ☐ No
- ☐ Not applicable

If yes, please specify why you consider that this use case(s) should be classified as high-risk taking into account the criteria in Article 7(2) and provide evidence that justify such a change.

1500 character(s) maximum

Yes. Subject to our overarching point that classification should track significant individual impact and autonomy rather than model complexity, we suggest the Commission consider explicitly capturing certain use cases where individual impact can be significant: (a) fully automated underwriting in life, health and credit-related insurance; (b) fraud-detection or anti-abuse systems where their output has legal or similarly significant effects, such as placement on shared industry fraud databases that can exclude a person from cover; and (c) claims-management systems where automated decisions materially affect access to promised benefits. In each case a clear human-in-the-loop safeguard should apply. These additions would close gaps where automation errors could cause serious individual harm, while leaving transparent, human-supervised actuarial tools outside scope.

Question 2. Do you consider that some of the existing use cases listed in Annex III AI Act need to be amended in order to fulfil the conditions laid down pursuant to Article 7(2) AI Act?

- ☒ Yes
- ☐ No
- ☐ Not applicable

If yes, please specify your arguments and provide evidence that justify such a change.

1500 character(s) maximum

Yes. We suggest amendments that improve clarity and consistency rather than expand scope: (a) amend point 5 (c) to incorporate an explicit carve-out for AI systems relying solely on generalised linear models (including linear or logistic regression) and generalised additive models operated under human supervision, as proposed in EIOPA's letter of 13 April 2026, which the AAE has endorsed; (b) distinguish fully autonomous pricing /underwriting systems from decision-support tools used under human oversight; (c) where pension or long-term savings products are covered (see para 317), confirm that traditional, transparent and human-supervised methods such as mortality-table and GLM-based annuity pricing receive the same proportionate treatment as in life and health, and are not high-risk; (d) clarify the rationale for the life/health versus non-life boundary, recognising that significant individual impact, autonomy and human oversight - rather than product line alone - are the relevant criteria; and (e) confirm that internal models, reinsurance-specific tools and claims-adjustment systems are in scope only where they directly determine individual access, pricing or coverage. These clarifications would support proportionate and consistent classification.

Question 3. Do you consider that some of the existing use cases listed in Annex III AI Act should be removed, since they do not pose anymore equivalent risk of harm (in accordance with the criteria laid down pursuant to Article 7(2) AI Act)?

- ☒ Yes
- ☐ No
- ☐ Not applicable

If yes, please specify your arguments and provide evidence that justify such a change.

1500 character(s) maximum

No. We do not propose removing any of the existing use cases. We suggest instead that the boundaries of the existing use cases be clarified, in particular to confirm that transparent, human-supervised actuarial techniques which do not meet the AI system definition - including traditional mortality-table and GLM-based pricing across life, health, pension and annuity products - and Solvency II internal models used solely for prudential capital purposes, fall outside the high-risk classification. This is a matter of clarification rather than removal.

Question 4. Do you consider that the list of prohibited AI practices in Article 5 should be expanded? In particular, are there concrete examples of AI practices that, in your view, conflict with Union values and are not yet adequately addressed by the AI Act or other Union legislation?

- ☒ Yes
- ☐ No
- ☐ Not applicable

If yes, please specify your arguments and provide evidence that justify such a change.

1500 character(s) maximum

Yes, the Commission could consider whether two practices warrant clearer safeguards, as possible gaps not fully addressed by existing Union law. First, proxy discrimination: using correlates of sensitive attributes (e.g. postcode or certain demographic data) for risk scoring can produce indirect discrimination even where sensitive data are not used directly. The use of such correlates could be required to be demonstrably justified, monitored for bias and proportionate. Second, behavioural price optimisation: setting prices on consumers' perceived willingness to pay rather than on risk can disproportionately affect vulnerable groups and erode fairness. We offer these as areas that may merit clearer ethical guardrails or sector-specific guidance, rather than necessarily new prohibitions.

Question 5. Do you consider that some of the prohibitions listed in Article 5 AI Act are already sufficiently addressed by other Union legislation and should therefore be removed from the list of prohibited practices in Article 5 AI Act?

- ☒ Yes
- ☐ No
- ☐ Not applicable

If yes, please specify your arguments and provide evidence that justify such a change.

1500 character(s) maximum

No. The AAE does not propose removing any existing prohibition under Article 5. Where prohibited practices overlap with other Union legislation — including GDPR, the Insurance Distribution Directive and Solvency II — Article 5 performs an independent function in identifying AI practices considered fundamentally incompatible with Union values. We recommend the Commission provide guidance on the interaction between Article 5 and existing sectoral legislation, with sector-specific examples, to prevent legal uncertainty and duplication rather than reducing the scope of the prohibitions themselves.

Contact

[Contact Form](#)