

Brussels, 20 July 2026

German Chamber of Commerce and Industry

Statement

Draft Commission guidelines on the classification of high-risk AI systems under Article 6 of Regulation (EU) 2024/1689 (AI Act) for stakeholder consultation

Thank you for the opportunity to present a statement on Draft Commission guidelines on the classification of high-risk AI systems under Article 6 of Regulation (EU) 2024/1689 (AI Act) for stakeholder consultation.

A. Outline of main facts

According to the [German Chamber of Commerce and Industry Digitalisation Survey 2026](#), 41 percent of companies already using AI expect a strong positive impact on their productivity, demonstrating that AI is increasingly becoming a key competitive factor for businesses. At the same time, legal uncertainty remains one of the most significant barriers to AI adoption and data-driven innovation.

Against this background, the Draft Guidelines should provide legal certainty, clear classification criteria and practical guidance for businesses. A risk-based and innovation-friendly interpretation of the AI Act, combined with reduced compliance complexity and unnecessary administrative burdens, is essential to enable companies, particularly SMEs, to confidently develop, deploy and scale AI solutions.

The Draft Commission Guidelines are an important tool for ensuring a consistent and proportionate application of the AI Act's high-risk classification rules. The DIHK welcomes the risk-based approach but sees a need for greater legal certainty, particularly regarding intended purpose, safety functions, safety components and the allocation of responsibilities within complex AI value chains.

The DIHK calls for a more practical, proportionate and legally certain application of the high-risk classification framework. In particular, the Guidelines should provide clearer classification criteria, additional practical examples and implementation-oriented support, while reducing unnecessary compliance burdens and enabling the adoption of trustworthy AI by businesses, particularly SMEs:

- **Introduce a more differentiated classification of AI systems based on autonomy:** The Guidelines should distinguish between different types of AI systems, while taking their level of autonomy into account when determining whether they are high-risk.
- **Clear distinction between intended purpose and foreseeable misuse:** The Guidelines should better distinguish between an AI system's intended purpose and foreseeable misuse, and include practical examples, templates, checklists, and sample wording to help companies perform self-assessments with legal certainty.
- **Reduce compliance burdens and provide SME-friendly implementation support:** The Commission should ensure a proportionate, risk-based approach, especially for SMEs, by offering standardized templates and more flexible conformity-assessment procedures to avoid excessive administrative costs.
- **Narrow and clarify the scope of high-risk AI systems:** High-risk classification should focus on genuinely safety-critical and autonomous AI systems. Productivity tools, decision-support systems, predictive maintenance applications, cybersecurity tools, and other assistive technologies that remain under meaningful human oversight should generally not be classified as high-risk.
- **Increase legal certainty:** The Guidelines should provide clearer definitions of concepts such as safety functions, safety components, malfunctioning, material influence, and product/component distinctions, while also clarifying responsibilities across the value chain and adding sector-specific positive and negative examples to help businesses determine whether an AI system falls within the high-risk regime.

B. Detailed explanation (or "Detailed evaluation").

General Observations

The Draft Commission Guidelines on the classification of high-risk AI systems under Article 6 AI Act are intended to clarify when an AI system falls within the AI Act's high-risk regime and is therefore subject to the corresponding regulatory requirements. In this context, the Guidelines provide interpretative guidance on the classification criteria laid down in Article 6 AI Act, including AI systems that qualify as high-risk because they are safety components of regulated products under Annex I or because they are used in high-risk use cases listed in Annex III. The Guidelines are therefore a key instrument for ensuring a consistent, proportionate and legally certain application of the AI Act across the European Union and for supporting businesses in determining their compliance obligations.

General principles for classification of high-risk AI systems under Article 6

a. The system must be an AI system

The Draft Commission Guidelines on the classification of high-risk AI systems under Article 6 of the AI Act seek to differentiate between high-risk and non-high-risk AI systems based on their

areas of application and the potential risks they pose to the safety of individuals. In doing so, however, the Guidelines fail to take into account the specific type of AI system concerned. AI systems are designed to operate with varying levels of autonomy (Article 3(1) AI Act). The level of autonomy may, however, affect the trustworthiness of an AI system. Consequently, this factor should be taken into account when assessing the risk category of the AI system in question. After all, Recital 27 of the AI Act encourages providers of AI systems to develop AI systems as tools that serve people, respect human dignity and personal autonomy, and function in a manner that can be appropriately controlled and overseen by humans, thereby contributing to the development of trustworthy and human-centric AI.

Against this background, the German Chamber of Commerce and Industry proposes drawing a distinction between different types of AI systems based on their respective levels of autonomy (such as truly learning autonomous AI systems, autonomous decision-making AI systems, assistive AI systems, adaptive software, statistical pattern-recognition systems, and classical algorithmic systems) and establishing corresponding definitions for these types of AI systems.

b. Intended purpose(s) of the AI system

The Guidelines state in Section 12 that where the instructions for use, contractual arrangements, terms of service, usage policies, promotional and sales materials, or technical documentation present an AI system as broadly applicable across a wide range of contexts and functions, and do not consistently limit its application, the system's intended purpose will be deemed to encompass high-risk use cases and, therefore, qualify as high-risk. This will be the case in particular where such uses are feasible and reasonably foreseeable given the system's functionalities and capabilities.

These statements make it difficult to distinguish the intended purpose of an AI system from its "reasonably foreseeable misuse", as defined in Art. 3 Nr. 13 AI Act, which, by definition, refers to a use outside the system's intended purpose. Moreover, use cases and applications may evolve and cannot be defined in advance for all hypothetical scenarios. After all, AI systems may exhibit adaptiveness after deployment (Art. 3 Nr. 1 AI Act). The resulting uncertainties also affect the assessment of whether a substantial modification within the meaning of Art. 25 has occurred and may have repercussions for the grandfathering clause in Section 449 of the Guidelines. For these reasons, the Guidelines should provide practical sample wording, positive and negative examples, a simple self-assessment template, and a checklist specifying the minimum information required to define the intended purpose.

High-risk classification according to Article 6(1) – Annex I. (Section III)

a. The rationale

The rationale underlying Article 6(1) AI Act is generally appropriate, as it seeks to limit high-risk classification to AI systems that perform genuine safety-related functions within already

regulated products. However, additional clarification is needed regarding the concept of preventive safety functions.

In particular, it remains unclear when AI systems used for predictive maintenance move from operational optimisation to a safety-relevant function. Many predictive maintenance applications aim primarily to improve efficiency, reliability and resource planning rather than directly mitigating safety risks. The Guidelines should therefore provide clearer criteria distinguishing operational support functions from genuine safety functions.

Furthermore, the Guidelines would benefit from additional examples covering AI modules, software add-ons and cloud-based services integrated into regulated products after market placement. Such examples should clarify when these functionalities constitute independent products, safety components or merely supplementary features.

Finally, greater legal certainty is needed regarding responsibilities across the value chain. Additional guidance should clarify the respective obligations of software providers, product manufacturers, integrators and operators in mixed hardware-software environments.

Examples:

The Guidelines would benefit from additional examples illustrating both in-scope and out-of-scope applications under Article 6(1) AI Act.

The following applications should be explicitly identified as falling outside the scope of high-risk classification where they do not perform safety-related functions:

- internal knowledge-management systems;
- customer-support chatbots;
- automated route-planning tools for field-service personnel;
- AI systems generating draft emails and other text content.

These applications primarily support productivity and business efficiency and do not normally affect the safety of regulated products.

Conversely, the Guidelines should explicitly identify as potentially high-risk:

- AI systems directly controlling safety-critical functions, such as emergency-stop mechanisms or collision-avoidance systems in machinery;
- AI systems that autonomously execute safety-relevant operational actions without meaningful human intervention.

Including such examples would improve legal certainty and help companies distinguish ordinary business software from genuinely safety-relevant AI applications.

b. Main elements for classification as high-risk

The requirements on data governance under Art. 10 AI Act may create disproportionate implementation challenges, particularly for SMEs. Ensuring that training, validation and testing data are fully relevant, sufficiently representative, accurate and complete is often difficult in practice, especially for AI systems developed on the basis of very large, dynamic or continuously evolving datasets. Excessively rigid expectations regarding dataset quality and related documentation may result in significant compliance costs without necessarily leading to proportionate improvements in risk mitigation.

In addition, providers of high-risk AI systems classified under Art. 6 AI Act face extensive technical documentation obligations. Given the iterative nature of AI development, documentation may need to be updated frequently to reflect changes. This can create substantial administrative burdens that are difficult to reconcile with agile software development methodologies and may inadvertently discourage innovation, system improvement and the timely deployment of beneficial updates.

To facilitate effective and proportionate compliance, the Commission should provide standardized, free-to-use templates, checklists and digital compliance tools that support providers in meeting documentation, data governance and conformity assessment requirements. Particular attention should be given to the needs of SMEs, which often have limited compliance resources despite being key drivers of innovation.

c. Component of a product or itself a product

The Guidelines should provide further clarification on the distinction between AI systems that constitute products in their own right and AI systems that merely form part of a product.

Particular uncertainty exists in sectors such as medical technology, where AI functionalities may include signal filtering, physiological pattern recognition, adaptive training control, embedded sensor intelligence and biofeedback-related applications. In practice, companies often struggle to determine whether such functionalities should be regarded as independent products, product components or supporting software features.

The Guidelines should also clarify how this distinction interacts with other applicable regulatory frameworks, including the Medical Devices Regulation, the Cyber Resilience Act and data-protection requirements. Clearer guidance on responsibilities, testing obligations, conformity-assessment requirements and documentation duties would significantly improve legal certainty, particularly for SMEs.

Examples:

The Guidelines should include additional examples involving AI-based signal filtering, physiological pattern-recognition tools, adaptive training-control systems, biofeedback applications and embedded AI in sensor systems.

Many of these applications support monitoring, training or therapeutic processes without autonomously taking safety-critical decisions. Including such examples would help distinguish supportive technologies from genuinely safety-relevant AI components and facilitate a more proportionate and predictable application of Article 6(1) AI Act.

d. Safety component

The Guidelines should further clarify the concepts of “malfunctioning” and “safety component”, particularly regarding AI-enabled cybersecurity solutions.

The current interpretation appears broad enough to encompass false positives and false negatives that are inherent in many cybersecurity environments. A distinction should therefore be made between ordinary performance limitations and genuinely safety-critical malfunctions. Otherwise, a wide range of cybersecurity tools could inadvertently fall within the scope of high-risk classification despite not performing safety-critical functions in the sense of Article 6(1).

The Guidelines would also benefit from additional examples involving AI-enabled retrofit solutions integrated into existing machinery. Where AI systems directly influence steering, braking or other safety-related functions, clearer guidance is needed regarding the allocation of responsibilities among software providers, manufacturers and operators.

Examples

The Guidelines should include examples distinguishing autonomous medical AI systems from assistive, therapeutic or training-related AI applications. Many systems support therapy, rehabilitation or prevention activities without replacing human judgment and therefore have a substantially different risk profile.

The Guidelines should also address AI integration scenarios where standard AI models are incorporated into existing ERP or CRM systems. Additional examples would help clarify responsibility allocation between model providers, software vendors and system integrators.

e. Third-party conformity assessment

The Guidelines should acknowledge the practical challenges arising from the limited availability of harmonised standards. In their absence, providers may frequently need to rely on costly third-party conformity-assessment procedures despite the risk profile of the system not necessarily justifying such an approach.

This may create disproportionate burdens for SMEs, software developers and IT service providers. Additional clarification is therefore needed on how software updates, security patches and iterative improvements can be handled efficiently within the conformity-assessment framework without creating unnecessary delays or compliance costs.

High-risk classification according to Article 6(2) – Annex III. (Section IV)

a. Rationale and approach

Although the classification of AI systems as high-risk under Art. 6(2) AI Act does not render their use unlawful, it may nevertheless discourage businesses from deploying such systems, even where AI could generate substantial operational benefits. This reluctance is often driven by concerns about compliance burdens and the risk of sanctions. In particular, many SMEs fear lacking the resources and expertise needed to comply with obligations associated with a high-risk classification.

The structure of Art. 6 may reinforce this perception. Under Art. 6(2), certain AI systems are initially presumed to be high-risk, while only subsequently assessed under the exemption mechanism in Art. 6(3). As a result, stakeholders are first confronted with a high-risk designation before considering whether an exemption applies.

To mitigate these effects, the German Chamber of Commerce and Industry recommends adjusting the approach reflected in the Guidelines. First, the Guidelines should clarify that companies which transparently and in good faith document the purpose, functionality, and interaction of their AI components may rely on their resulting self-assessment. Sanctions should be reserved for deliberate misrepresentation rather than divergent ex post assessments of honestly described configurations. Second, the Guidelines should place greater emphasis on practical SME-relevant examples, including clear negative examples confirming when a system does not qualify as high-risk.

b. Horizontal issues applicable to Annex III use cases

Where several AI systems form part of a more complex AI system, such that their combined intended purpose or joint outputs materially influence an individual decision, the Guidelines treat the combined configuration as a single AI system. By contrast, AI systems performing strictly procedural or preparatory functions may remain exempt from high-risk classification, provided they are genuinely separable and do not generate, structure, or feed outputs that materially influence the assessment of an individual case (sections 75/76).

However, “material influence” remains a vague concept in the Guidelines. In the absence of clear criteria, businesses will struggle to determine system boundaries and allocate responsibilities among different actors. Yet, in modern software environments, AI agents from

different vendors, open-source components, and rule-based tools increasingly operate together. The issue is compounded by the fact that, once an overall assessment results in a finding of profiling, each AI system within the interconnected setup may be classified as high-risk. While preventing circumvention of the AI Act is a legitimate objective, the current approach risks being interpreted too broadly. Many agentic AI solutions merely automate routine tasks, such as information transfer, document handling, or workflow coordination, while final decisions remain subject to human review. Such systems should not be presumed to be high-risk solely because multiple AI-enabled components interact.

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

(1) Biometrics

The scope of high-risk classification for biometric AI systems appears overly broad and risks creating unnecessary barriers to innovation in this rapidly evolving field. The Guidelines should clarify that biometric applications with no meaningful impact on individual rights or decisions remain eligible for exemption under Article 6(3) AI Act. For example, emotion-recognition systems used in call centres solely for aggregated statistical analysis, service-quality monitoring, or training purposes should not be classified as high-risk, provided that their outputs are not used to make, support, or influence decisions concerning individual employees or customers. Such systems do not create the type of risks that the high-risk regime is intended to address and should therefore fall within the exemptions set out in Article 6(3)(b) and (d) AI Act.

(2) Critical infrastructure

The Guidelines risk classifying a broad range of AI systems used in critical digital infrastructure as high-risk. High-risk status should be limited to AI systems that perform a direct safety function or have a direct, foreseeable and non-speculative impact on risks to life, health or physical property. AI used for network optimisation, performance management, load balancing, energy efficiency, anomaly detection, operational planning, customer-specific adaptation or technical assistance should be expressly excluded unless it directly controls a safety-critical function or protection mechanism.

A key concern is the ambiguity of the “failure or malfunctioning” criterion in the definition of a safety component. As drafted, it could create the presumption that any AI used in network operations becomes high-risk if its failure might indirectly contribute to service disruption. The Commission should clarify that safety-component status only applies where the AI exercises material control over a safety-critical function, the potential harm is direct and reasonably foreseeable, the causal chain does not depend on remote contingencies or exceptional events, and safeguards such as redundancy, fallback mechanisms, human oversight and independent safety layers are taken into account.

The concepts of “directly” and “direct protective function” also remain insufficiently defined and risk overly broad interpretation. They should be limited to AI systems that autonomously execute or control safety-critical operational measures whose failure would immediately compromise the physical integrity of critical infrastructure.

Examples:

AI systems whose primary purpose is to improve efficiency, performance, resilience, or planning should generally not be classified as high-risk. This includes AI used for network optimisation, quality-of-service improvements, load balancing, network planning, resource allocation, customer-experience enhancement, and energy optimisation, provided it does not directly control a safety-critical function or protection mechanism. The same applies to cybersecurity applications, such as systems for detecting, classifying, and preventing cyberattacks or fraud, managing vulnerabilities, conducting penetration testing and attack simulations, or protecting digital infrastructure through AI-enabled security tools.

In the transport sector, non-high-risk examples should include AI used exclusively for traffic-data analysis and optimisation purposes, such as urban planning, historical traffic assessment, and route recommendations that do not directly intervene in traffic control. Predictive-maintenance systems used for wear forecasting, preventive maintenance, and optimisation of maintenance cycles should likewise remain outside the scope of high-risk classification, as their primary purpose is operational efficiency rather than the management of safety-critical functions.

(3) Employment, workers’ management and access to self-employment

Many SMEs use AI tools to support administrative tasks such as shift scheduling, working-time coordination, and workforce planning. Where such systems merely improve efficiency and do not materially influence employment-related decisions affecting individuals, they should remain eligible for exemption under Art. 6(3) AI Act.

Examples of AI applications that may warrant high-risk classification include AI-driven recruitment and candidate-screening tools (e.g. CV screening, candidate ranking, profile matching, and automated pre-selection), platform-based systems that significantly influence work allocation or performance evaluation, and automated online exam proctoring tools using technologies such as webcam monitoring, gaze tracking, or cheating detection.

By contrast, tools that provide only ancillary support and do not materially affect employment opportunities or working conditions should not be considered high-risk. Examples include inclusion and bias checkers for HR texts, language assistants for job advertisements, and employee onboarding chatbots used for FAQs, workplace orientation, and internal knowledge management. These applications primarily support communication and administrative efficiency rather than decision-making with legal or similarly significant effects on individuals.

The Guidelines should also consider a proportionate approach for smaller companies, taking into account both company size and the number of individuals affected by the AI system.

(4) Access to and enjoyment of essential private services and essential public services and benefits

The AI Act explicitly excludes AI systems used for financial fraud detection from the high-risk classification applicable to creditworthiness assessments (Annex III, Point 5(b)). However, in modern banking practice, fraud detection, credit assessment, and anti-money laundering/counter-terrorist financing (AML/CFT) functions are increasingly integrated into shared technological infrastructures and software solutions. This creates legal uncertainty as to whether systems primarily designed for fraud detection could nevertheless be classified as high-risk due to their interaction with credit-related processes. The Guidelines should, therefore, provide a practical and targeted safeguard, ensuring that AI systems whose primary purpose is fraud detection are not indirectly reclassified as high-risk solely because they are embedded in integrated compliance or risk-management frameworks.

(5) Additional Examples

Clearer guidance is needed on the treatment of dual-use technologies and general-purpose communication tools. Companies require practical, predictable and legally certain criteria to distinguish benign use cases from genuinely high-risk applications. Such clarification is essential to provide legal certainty and to prevent overly broad interpretations that could unnecessarily constrain innovation, particularly among European SMEs.

More generally, the Guidelines would benefit from additional sector-specific examples reflecting industrial and SME use cases, including digital twins, AI-enabled retrofitting of existing equipment, visual quality inspection, production planning, predictive maintenance, and AI-assisted machinery. Practical examples would help companies distinguish between applications that primarily enhance efficiency, quality, or operational performance and those that are genuinely safety-critical and therefore warrant classification as high-risk AI systems.

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

The current high-risk framework risks capturing too many low-risk and supportive AI applications. A more risk-based approach is needed, particularly for SMEs. Systems that operate in a limited context, provide decision support only, or remain subject to meaningful human oversight should generally not be classified as high-risk. Consideration should also be given to introducing *de minimis* exemptions for SMEs to avoid disproportionate compliance burdens. After all, the size of the enterprise in question may also impact the risks that the AI system poses for the safety of human beings.

Against this background, the high-risk categories should be narrowed and clarified to distinguish genuinely safety-critical and autonomous systems from low-risk applications. From a practical perspective, high-risk classification should focus on AI systems whose malfunction can directly cause substantial harm or have some other significant impact on health, safety or fundamental rights, such as autonomous driving or autonomous medical diagnosis and treatment, real-time biometric surveillance, safety-critical machine control, or fully automated decisions regarding employment, creditworthiness or access to essential services.

By contrast, systems that merely support human decision-making or whose outputs remain subject to human review should be assessed under a more proportionate framework. This includes therapeutic training tools, biofeedback and neurofeedback applications, adaptive training software, AI-assisted coaching and motivation systems, pattern-recognition tools supporting professionals, data visualization solutions, embedded AI without autonomous decision-making powers, and wellness or prevention applications without diagnostic claims. There is a particular risk that supportive neurocognitive, training and assistance technologies could be regulated in the same way as autonomous diagnostic systems despite fundamentally different risk profiles.

C. Additional information

a. Contact persons with contact details

Jonas Wöll, Referatsleiter Digitaler Binnenmarkt, EU-Verkehrspolitik, Regionale Wirtschaftspolitik, Bereich Digitale Wirtschaft, Infrastruktur, Regionalpolitik, woell.jonas@dihk.de; Tel: +32 2 286 1639

Arian Siefert, Referatsleiter Wirtschaft digital, Bereich Digitale Wirtschaft, Infrastruktur, Regionalpolitik, siefert.arian@dihk.de; Tel.: +49 30 20308 2118

Jennifer Evers, Bereich Recht, Referatsleiterin Alternative Konfliktlösung (SGH), Recht der digitalen Wirtschaft und Legal Tech | Syndikusrechtsanwältin, evers.jennifer@dihk.de; Tel. +49 30 20308 2719

Julian Kulaga, Referatsleiter Geistiges Eigentum, Verbraucherrecht, Bereich Recht, kulaga.julian@dihk.de; Tel.: +49 30 20308 2704

b. Description of DIHK

Who we are:

The 79 Chambers of Commerce and Industry (IHKs) are incorporated under the umbrella of the German Chamber of Commerce and Industry (DIHK). Our joint aim: to obtain the best conditions for successful business.

The DIHK represents the interests of the entire commercial economy in dealings with decision makers, administrations and the public at federal and European level. Several million companies representing trade, industry and services are legal members of an IHK, ranging from kiosk owners to DAX companies. DIHK and IHKs therefore provide a platform for a whole range of corporate concerns. We group them together in an organised procedure on a statutory basis to formulate the general interest of the commercial economy and thus contribute to the formation of opinions in economic policy.

Our statements are based on economic policy positions and position papers adopted by the DIHK taking into account the comments received by the DIHK from IHKs and their member companies prior to the submission of the statement.

The DIHK also coordinates the network of 150 German Chambers of Commerce Abroad, delegations and representative offices of the German economy in 93 countries.

The DIHK is registered with the European Union's Transparency Register under registration number 22400601191-42.