



**EVOLVING AI
SYSTEMS UNDER THE
AI ACT: SUBSTANTIAL
MODIFICATION AND AI
VALUE CHAIN**

ISSUE PAPER

June 2026

Daniel Schnurr



As provided for in CERRE’s bylaws and procedural rules from its “Transparency & Independence Policy”, all CERRE research projects and reports are completed in accordance with the strictest academic independence.

The CERRE Regulatory Lab series, of which this note is part, received the support and/or input from several CERRE members including Apple, Booking.com, and Microsoft. Each draft note is discussed in Chatham House meetings between the CERRE corporate members, national regulatory authorities, EU institutions and civil society organisations. However, they bear no responsibility for the contents of this report. The views expressed in this CERRE report are attributable only to the authors in a personal capacity and not to any institution with which they are associated. In addition, they do not necessarily correspond either to those of CERRE, or of any sponsor or of members of CERRE. AI-based tools were used to assist with text editing and language refinement in this paper. The author retains full responsibility for its content and conclusions.

© Copyright 2026, Centre on Regulation in Europe (CERRE)

info@cerre.eu – www.cerre.eu



Table of Contents

TABLE OF CONTENTS	1
1. INTRODUCTION	2
2. ILLUSTRATIVE USE CASES	4
2.1 SUBSTANTIAL MODIFICATION OF A HIGH-RISK AI SYSTEM USED FOR HR AND EVALUATION OF JOB CANDIDATES..	4
2.2 SUBSTANTIAL MODIFICATION OF A HIGH-RISK AI SYSTEM USED AS A MEDICAL DEVICE TOGETHER WITH OTHER AI SYSTEMS	5
2.3 USE OF A GPAI SYSTEM IN HIGH-RISK DOMAINS AND ADDITION OF AGENTIC CAPABILITIES.....	6
2.4 EXTERNAL TESTING OF GPAI MODELS AND AI SYSTEMS	7
3. MODIFICATIONS OF AI SYSTEMS	9
3.1 SCOPE OF THE AI ACT AND HIGH-RISK CLASSIFICATION	9
3.2 MODIFICATIONS THAT MAY TRIGGER NEW COMPLIANCE OBLIGATIONS	9
3.3 SUBSTANTIAL MODIFICATION OF A HIGH-RISK AI SYSTEM	10
3.3.1 MODIFICATIONS AFFECTING COMPLIANCE WITH HIGH-RISK REQUIREMENTS	11
3.3.2 MODIFICATIONS OF THE INTENDED PURPOSE	13
3.4 RESEARCH, TESTING, AND DEVELOPMENT ACTIVITIES.....	15
4. INSTITUTIONAL FRAMEWORK AND AI VALUE CHAIN IMPLICATIONS	17
4.1 IMPLICATIONS FOR THE AI VALUE CHAIN	17
4.2 INSTITUTIONAL FRAMEWORK AND CROSS-SECTOR APPLICATION	18
5. RECOMMENDATIONS FOR IMPLEMENTATION.....	19
ABOUT CERRE.....	22
ABOUT THE AUTHOR	23



1. Introduction

The European AI Act (AIA) aims to ensure a high level of protection for health, safety, and fundamental rights by promoting the uptake of human-centric and trustworthy AI.¹ To ensure proportionate and effective regulation, it adopts a risk-based approach that tailors requirements and compliance obligations to the level of risk posed by different use cases. In particular, it identifies certain high-risk AI systems that are subject to more stringent compliance requirements. This risk-tiered framework also seeks to avoid overregulation and promote innovation.

Although the AIA entered into force in 2024, the EU is still deep in the implementation process, particularly concerning the regulation of high-risk AI systems. The recently agreed Digital Omnibus on AI reshapes the timeline.² Rather than tying the application of the high-risk rules to the availability of harmonised standards and guidance, as originally proposed by the Commission, the co-legislators settled on fixed application dates, pushing compliance to late 2027 or 2028.³ In parallel, the AI Office is developing a series of guidelines to operationalise the high-risk regime, including on the practical application of the high-risk classification, the rules for responsibilities along the AI value chain, and the provisions related to substantial modification.⁴

To make the AIA work in practice, it is important to recognise that AI systems are subject to frequent change as the underlying technology and commercial practices evolve. Implementation must therefore account for the entire lifecycle of AI systems. Systems may be modified or applied to new contexts, potentially triggering new compliance requirements. Such evolution may cause a system to cross the risk tiers established by the AIA, for example when a system designed for non-high-risk use is repurposed for a high-risk area. Operators must therefore anticipate the conditions under which a modification will trigger new obligations. General-purpose AI (GPAI) and agentic systems present particular challenges in this respect, as their intended purpose and risk classification can be more difficult to discern.

Changes to an AI system or its application domain may be made not only by the provider of the system, but also by other third parties active in the AI value chain. Where such a modification results in a high-risk AI system, the AIA specifies that the modifier may become the new provider of the modified system.⁵ This triggers new compliance requirements for the new provider, but may also have implications for other actors in the value chain, specifically the provider of the original system. Again, the effect may extend across the risk tiers, as operators active outside high-risk areas may be drawn into supporting downstream actors' high-risk compliance. To maintain the original risk-based approach and to ensure effectiveness, proportionality, and innovation, it is essential to clarify the extent of such obligations and the conditions under which they arise.

¹ Regulation 2024/1689 of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act), <http://data.europa.eu/eli/reg/2024/1689/oj>.

² Proposal for a Regulation amending Regulations 2024/1689 and 2018/1139 (Digital Omnibus on AI), COM(2025)836 (19 November 2025); As of June 2026, the Digital Omnibus on AI had been politically agreed but not yet formally adopted or published in the Official Journal.

³ Council of the European Union (2026). Artificial Intelligence: Council and Parliament agree to simplify and streamline rules. <https://www.consilium.europa.eu/en/press/press-releases/2026/05/07/artificial-intelligence-council-and-parliament-agree-to-simplify-and-streamline-rules/>.

⁴ European Commission (2025). Supporting the implementation of the AI Act with clear guidelines. <https://digital-strategy.ec.europa.eu/en/news/supporting-implementation-ai-act-clear-guidelines>.

⁵ Art. 25 AIA.



This paper provides an overview of the main issues related to the modification of AI systems, illustrated by practical use cases, and discusses the implications for obligations under the AIA’s high-risk requirements. It is part of the CERRE Regulatory Labs, which serve as forums to explore key regulatory and policy challenges related to the digital economy and society. The present Lab explored how to make some of the AIA’s key definitions work in practice: When does an update to a high-risk AI system or a GPAI system trigger new compliance obligations? How is a change in the intended purpose of an AI system monitored and operationalised in practice? When does internal use, or research and development, cross the line into “placing on the market” for AI systems and GPAI models? And how should these questions be approached in high-risk sectors such as medical devices, employment, and critical infrastructure? These questions are closely connected to the AIA’s rules on responsibilities along the AI value chain, which this paper also addresses. A more thorough discussion of these issues is provided in the companion CERRE paper on “Information Sharing and Cooperation along the AI Value Chain under the AI Act”⁶. The analysis in this paper is intended to inform the implementation of the AIA’s provisions on substantial modification and cooperation along the AI value chain, and to highlight key issues for the AI Office to address in its forthcoming guidelines.

The remainder of this paper proceeds as follows. Section 2 presents example use cases illustrating the practical challenges of applying the AIA’s modification and high-risk rules. Section 3 addresses the key questions of what constitutes a substantial modification and when modifications of AI systems trigger new compliance obligations, also covering the scope of general research, testing, and development exemptions. Section 4 points to further key implementation issues, with a particular focus on the value chain implications. Finally, Section 5 concludes with overarching recommendations on the implementation of the AIA’s provisions on the modification of AI systems and GPAI models.

⁶ <https://cerre.eu/publications/rules-for-high-risk-ai-systems-under-the-ai-act/>



2. Illustrative Use Cases

This section presents illustrative stylised use cases of changes to AI systems and GPAI models. These examples are meant to ground the further discussion in practical scenarios and may help to illustrate the grey areas that emerge from the legal framework.

2.1 Substantial modification of a high-risk AI system used for HR and evaluation of job candidates

A retailer procures a high-risk AI system from a vendor that automatically screens job applications. The system has undergone a conformity assessment for the intended purpose of ranking candidates against a defined competency profile for entry-level retail roles. After some time, the HR department reconfigures it so that it also evaluates existing employees for internal promotions. To this end, they feed it internal company data, such as performance-review and attendance data, on which the original model was not trained or validated on. Furthermore, they enhance the existing system by adding a feature that can analyse video recordings provided by candidates to generate a recommendation about the candidate's fit for the position.

In this case, the changes to the high-risk AI system shift its intended purpose from external to internal candidates, adding a distinct use case of a high-risk system listed in Annex III AIA. Moreover, both the addition of qualitatively new data sets and functionalities may change the performance and the risks associated with the AI system. When these changes are not anticipated or assessed in the initial conformity assessment, they are likely to qualify as a substantial modification (see Section 3).

Now consider several variants of this scenario that illustrate how the assessment can become more ambiguous when only one element of the change is present, or when the modification is closer to routine maintenance, raising the question whether they represent a sufficient change to the system's risk profile to constitute a substantial modification.

First, consider a variant where the high-risk AI system affects a *broader group of persons*: the company continues to use the system only for external candidates and does not add any functionality or data, but extends its application to all job roles advertised by the company. A second variant involves *routine retraining* on fresh data of the same kind. The HR department periodically retrains the model on the most recent twelve months of applications, with the possible side effect that performance may slightly shift across demographic groups. A third variant concerns *reconfiguration of the system*. Without introducing new data or functionalities, the HR department may adjust the scoring weights for ranking candidates and lower the shortlist cut-off, using the configuration interface the vendor provides. In this case, the system's risk profile may shift even though no code or training data has changed. Finally, in a fourth variant, the *vendor pushes an upgrade of the AI system*. The vendor swaps the underlying GPAI model for a newer version, which significantly affects the predictive accuracy of the system. Whether this qualifies as substantial depends on whether such an update falls outside the anticipated changes in the initial conformity assessment or the continuous learning carve-out provided by the AIA (see Section 3.2). It also raises questions about the interplay between modifications at the GPAI-model



layer (potentially involving systemic risks) and the downstream high-risk AI systems that integrate those models.

These variants illustrate that in practical use cases, many common update practices or changes to a system can affect its risk profile or the scope of its application, raising the question whether these practices and changes reach the threshold of a substantial modification. This first illustrative use case already highlights that the initial conformity assessment plays a major role in this assessment, as it establishes a reference point for assessing whether the changes have been foreseen or planned. Providers of high-risk AI systems therefore have a strategic choice in how broadly to define the scope of the initial conformity assessment, since this determines how the burden of re-assessment is distributed across the AI value chain before any modification occurs.

2.2 Substantial modification of a high-risk AI system used as a medical device together with other AI systems

A large hospital has deployed several high-risk AI systems across its surgical pathway. Each is individually conformity-assessed, each has a defined intended purpose, and each performs acceptably in isolation. They are integrated through the hospital's clinical information system so that outputs from one AI system feed into the next. The pathway includes four systems provided by different vendors:

1. *System A for preoperative risk stratification*: Estimates 30-day mortality and major complication risk for new patients from initial diagnostic results and electronic health data. The system produces a risk score that assists physicians in scheduling and planning the operation.
2. *System B for intraoperative anaesthesia advisory*: Embedded in the anaesthesia workstation in the operating room. It monitors the patient's vital signs in real time and suggests how much anaesthetic drug to give. It takes System A's risk rating into account, so that a patient flagged as higher-risk is treated more cautiously.
3. *System C for early warning of postoperative deterioration*: A ward-level AI system that continuously monitors patients' vital signs and flags patients at risk if measurements deteriorate. Its alert threshold is calibrated dynamically using a patient's preoperative risk class from System A: high-risk patients get tighter thresholds, low-risk patients get looser ones to reduce alarm fatigue.
4. *System D for bed and discharge management*: A hospital-operations AI system that supports allocation of ICU beds, regular ward beds, and discharge decisions, using deterioration probability from System C as one input.

Now suppose that the provider of System A releases a routine update that recalibrates the underlying AI model on a larger and more recent patient population, leading to improved accuracy for elderly patients. The update is covered by a pre-declared change control plan. In this case, the update is likely to fall outside the AIA's definition of a substantial modification when judged on its own. However, the



update may nevertheless materially alter the behaviour of the downstream systems. A high-risk patient score produced by System A is now anchored to a different patient distribution, and Systems B and C may operate outside the input range against which they were validated. Consequently, System C's performance may degrade, even though the system itself has not been touched. This illustrates that where AI systems are chained to support interdependent decisions, there is an inherent risk of error propagation and amplification. Even though no single system may have undergone a substantial modification, the composite system may exhibit changes in the risk profile. This raises questions about how to operationalise the concept of substantial modification where risks emerge from the interaction or orchestration of multiple AI systems, and how responsibility should be allocated among the actors in the value chain.⁷

2.3 Use of a GPAI system in high-risk domains and addition of agentic capabilities

A large language model is placed on the market as a GPAI model and integrated into a chat-style GPAI system. The GPAI model provider complies with the GPAI obligations under the AIA. The GPAI system itself is not placed on the market as a high-risk AI system.

A company licenses the GPAI system and uses it to build a recruiting tool that automatically filters job applicants based on their CVs. The tool accesses the GPAI system through an API and takes a job description and a batch of CVs, prompts the underlying model with an engineered template, and returns a ranked shortlist with justifications for each candidate. In this case, the downstream company becomes the provider of a high-risk AI system under Art. 25(1)(c) AIA, which applies when the intended purpose of an AI system is modified in such a way that it becomes high-risk (here, employment and recruitment under Annex III, point 4(a) AIA). The relevant legal trigger is therefore not a technical change of the GPAI system, but its purpose-modification to deploy it in a high-risk area.

Now consider instead that a company licenses the GPAI system as a general-purpose assistant for all of its employees. The system is also available to staff in the HR department. While HR employees continue to screen and rank job candidates manually, they may use the GPAI system to summarise documents such as CVs and transcripts. Some may also use it to create systematic comparisons between candidates (e.g., regarding skills, experience, and degrees). In these cases, it is less clear whether the use of the GPAI system is sufficiently connected to a high-risk purpose to trigger Art. 25(1)(c) AIA and make the downstream company a new provider of a high-risk AI system.

Consider as a further variant a university where teaching staff use the GPAI system to draft new assignments and exam exercises based on existing course materials, and to generate sample solutions against which students' answers are subsequently graded manually. In this case, it may not be straightforward to determine whether the supporting activities executed by the GPAI system are sufficiently connected to the high-risk use cases listed for education in Annex III. Again, the threshold for a modification of the intended purpose is not straightforward, also because in these cases there does not exist an initial conformity assessment for the GPAI system. In both the HR and the university variants, a low threshold may classify many deployers as providers of high-risk systems, which could

⁷ Art. 25(4) and Art. 72 AIA are particularly relevant for the AIA's framework on responsibilities along the AI value chain and on post-market monitoring in such multi-system settings, as further discussed in Section 4.1. See also the companion CERRE paper, Schnurr (2026), and Meyers, Schnurr and Larouche (2025).



run counter to the original proportionate, risk-based approach of the AIA, given that GPAI systems are by design used for a wide variety of purposes across many application areas.

Finally, consider a variant in which the GPAI system provider adds agentic features, such as external tool use, planning, and reasoning capabilities, so that the system can act significantly more autonomously. When integrated into organisation-wide information systems, such an agentic GPAI system may independently engage in activities that contribute to purposes considered high-risk under the AIA, even where no operator has explicitly deployed it for those purposes. This raises the question of who bears responsibility when an agentic system autonomously takes actions that would, if performed by a purpose-built system, fall within a high-risk area. Further guidance may be needed on this question as Art. 25(1)(c) AIA is not well-suited to systems whose effective purpose emerges from their own autonomous behaviour rather than from a deliberate deployment choice.

These examples highlight that GPAI systems generally lack the anchoring reference point of an initial conformity assessment that defines an intended purpose. The substantial modification framework, which is designed around comparison with an assessed baseline, is therefore not identically applicable to GPAI systems and other non-high-risk AI systems.

2.4 External testing of GPAI models and AI systems

The AIA recognises product-oriented research, testing, and development activity on AI systems and models as important steps in the AI lifecycle, typically conducted before a system or model is released to the wider public.⁸ To support innovation, the AIA's provisions therefore do not apply to such activity prior to a system or model being put into service or placed on the market, subject to an exclusion for testing in real-world conditions (further discussed below).⁹

In practice, providers of GPAI models and AI systems regularly test new versions of their models and systems with real-world users to collect feedback that goes beyond static benchmarks and more limited testing environments.

Many new and updated GPAI models are placed on specialised online platforms, such as the LMSYS Chatbot Arena/LMArena or Yupp Leaderboard, before their commercial release.¹⁰ These platforms are dedicated to community-driven evaluations that allow providers to collect real-world human feedback and to compare the performance of their models against others. They are by now well-established institutions in the AI research and industry community and play a prominent role in evaluating GPAI models.

In addition, providers of both GPAI models and AI systems may test updated versions within systems already placed on the market in order to gather additional real-world user feedback. Typically, this allows providers to run comparative A/B tests against current versions by randomly routing selected

⁸ Recital 25 AIA.

⁹ Art. 2(8) AIA.

¹⁰ LMSYS Chatbot Arena / LMArena : <https://www.lmsys.org/blog/2024-03-01-policy/> and <https://arena.ai/leaderboard/text>; Yupp Leaderboard: <https://yupp.ai/leaderboard>.



users to outputs from the updated model or system. Users may opt in to confirm that they are willing to receive outputs from models still under development and testing.

A further variant is “private previews” of updated models and systems made accessible to selected organisational customers. As with the other testing approaches described above, the goal is to collect feedback on performance, quality, and potential risks before the commercial release. Because these previews target specialised users and organisations, the feedback may capture use-case-specific and evolving requirements in specific domains that are difficult to replicate in other testing environments.

A fundamental question that emerges from these common testing practices is how they relate to the AIA’s general scope and to the exemption for any research, testing, or development activity regarding AI systems or AI models prior to their being placed on the market or put into service (Art. 2(8) AIA). This is particularly noteworthy because Art. 2(8) AIA explicitly excludes “testing in real-world conditions” from the exemption. As test-driven development and short update cycles are common in the AI industry, treating any external testing of a GPAI model or AI system as testing in real-world conditions, or as placing on the market or putting into service, could lead to significant compliance burdens for providers. Conversely, an overly permissive reading of the exemption could undermine the effectiveness of the AIA.

For high-risk AI systems, Art. 60 AIA offers a regulated route for testing in real-world conditions outside the Art. 2(8) exemption, subject to specified conditions. No equivalent pathway seems to exist for GPAI models or for non-high-risk AI systems. The threshold under Art. 2(8) AIA is therefore particularly consequential for the testing of GPAI models, as it determines whether a given activity falls within pre-market research and development or already triggers the full set of AIA obligations.

Furthermore, these threshold questions may also have implications for the rules on substantial modification and high-risk classification. Where the testing practices described above are applied to AI systems used in high-risk areas, the scope of the Art. 2(8) exemption could directly affect the assessment of whether a change qualifies as a substantial modification, or whether the intended purpose of an AI system has been modified such that it qualifies as high-risk.

Specifically, a provider may classify an Annex III system as non-high-risk under Art. 6(3) AIA where it assesses that the system does not pose a significant risk of harm. This includes AI systems intended to perform a preparatory task to an assessment relevant for the use cases listed in Annex III (Art. 6(3)(d) AIA). On the basis of such a self-classification, neither Art. 60 AIA nor Art. 25(1)(b) AIA would apply to the testing practices described above, since both provisions are tied to the high-risk classification. However, if a market surveillance authority disagrees with the provider’s self-classification, Art. 80 AIA establishes a procedure for addressing systems wrongly classified as non-high-risk and for ensuring necessary corrective measures. This places significant weight on the robustness and documentation of the initial Art. 6(3) assessment.



3. Modifications of AI Systems

This section addresses the key questions: When does an update to an AI system trigger new compliance obligations? When is it a substantial modification? It sets out the relevant legal framework, specifically Art. 25(1)(b) and (c) AIA, and discusses the main conceptual and practical issues that arise in their application, pointing to areas where regulatory guidance could support implementation. It further discusses at what point research, testing, or development activities may trigger placing on the market obligations.

3.1 Scope of the AI Act and high-risk classification

The general scope of the AIA comprises AI systems placed on the market or put into service, as well as GPAI models placed on the market in the European Union (Art. 2(1)(a) AIA). “Placing on the market” refers to the first supply of an AI system or GPAI model for distribution or use on the Union market in the course of a commercial activity (in return for payment or free of charge), while “putting into service” means the supply of an AI system for first use directly to the deployer or for own use in the Union for its intended purpose.¹¹ Research, testing, or development activities regarding AI systems or GPAI models may be excluded from this scope (see Section 3.3).

Within this general scope, the AIA establishes a risk-tiered framework, with the most stringent requirements applying to high-risk AI systems. Under Art. 6(1) AIA, an AI system is classified as high-risk if it is intended to be used as a safety component of a product, or is itself a product, that is covered by Union safety legislation listed in Annex I and this component or product is required to undergo a third-party conformity assessment. In addition, an AI system may be classified as high-risk under Art. 6(2) AIA in conjunction with Annex III, which specifies selected intended purposes in eight application areas, including critical infrastructure, education, employment, and law enforcement. High-risk AI systems must comply with the requirements set out in Section 2 of Chapter III AIA, including those concerning risk management, data governance, technical documentation, transparency, human oversight, and accuracy and robustness.

3.2 Modifications that may trigger new compliance obligations

There are two main cases to distinguish in which a modification of an AI system may trigger new compliance obligations under the high-risk regime.

The first is a *substantial modification of a high-risk AI system* already placed on the market or put into service, where the system remains high-risk (Art. 25(1)(b) AIA). The second is a *modification of the intended purpose of a non-high-risk AI system that causes it to become high-risk* (Art. 25(1)(c) AIA). The latter is particularly relevant for GPAI systems, which are by design used across a wide variety of application areas.

¹¹ Art. 3(9), (10) and (11) AIA.



In both cases, the modification may be carried out by the original provider or by a third party. Where a third party makes such a modification, that party becomes the provider of the modified high-risk AI system (Art. 25(1)(b) and (c) AIA).

3.3 Substantial modification of a high-risk AI system

For a high-risk AI system that has already been placed on the market or put into service, a substantial modification triggers new compliance obligations.¹² In particular, the modified AI system is to be considered a new AI system and must undergo a new conformity assessment (Art. 43(4); Recital 128 AIA).

A substantial modification is defined in Art. 3(23) AIA as a change to an AI system, not foreseen or planned in the initial conformity assessment, as a result of which either (i) the compliance of the AI system with the high-risk requirements is affected, or (ii) the intended purpose for which the AI system has been assessed is modified. The two cases of this definition correspond to two distinct triggers, which are discussed in turn below.

The concept of a substantial modification draws on the established notion in EU product safety regulation under the New Legislative Framework.¹³ “The ‘Blue Guide’ on the implementation of EU product rules 2022”, issued as a Commission notice, establishes three criteria for a substantial modification that results in a new product: (i) [a product’s] original performance, purpose or type is modified, without this being foreseen in the initial risk assessment; (ii) the nature of the hazard has changed or the level of risk has increased in relation to the relevant Union harmonisation legislation; and (iii) the product is made available (or put into service if applicable).¹⁴ Reduced to a rule of thumb, a modification is substantial where a post-market change takes the product outside the envelope of the original risk assessment, creates a new hazard, or raises the level of risk, and requires new protective measures.¹⁵

A substantial modification triggers a full new conformity assessment. In this context, the Blue Guide clarifies that it is not necessary to repeat tests or produce new documentation in relation to aspects not affected by the modification, and that the technical documentation must be updated insofar as the modification has an impact. The Blue Guide further recognises that software may be an essential component of a product, or a product in itself.¹⁶ It sets out a parallel three-criteria test for substantial modifications of software, while noting that software updates and repairs may be treated as maintenance operations provided they do not affect compliance with applicable requirements.

¹² In general, the definition of substantial modification in Art. 3(23) AIA is direction-neutral: a substantial modification might also result in a high-risk AI system becoming a non-high-risk AI system, thereby removing the high-risk compliance obligations under the AIA. This paper does not cover such changes explicitly, although similar questions may emerge, in particular with regard to the intended purpose (see the discussion below).

¹³ Recital 128 AIA; On the connection between the AIA and the New Legislative Framework, see also Larouche, P. (2025). Legal Framework for an Effective Implementation of the AI Act. CERRE. https://cerre.eu/wp-content/uploads/2025/02/Legal-Framework-for-an-Effective-Implementation-of-the-AI-Act_FINAL.pdf.

¹⁴ Commission notice – The ‘Blue Guide’ on the implementation of EU product rules 2022, OJ C 247, 29.6.2022, p. 17.

¹⁵ While the Blue Guide represents soft law, the new EU Machinery Regulation, which will apply from January 2027, codifies a similar definition of substantial modification; see Art. 3(16) of Regulation (EU) 2023/1230 of the European Parliament and of the Council of 14 June 2023 on machinery and repealing Directive 2006/42/EC of the European Parliament and of the Council and Council Directive 73/361/EEC.

¹⁶ Blue Guide, p. 18.



3.3.1 Modifications affecting compliance with high-risk requirements

Echoing this framework, the first trigger for a substantial modification under the AIA is a change, unforeseen in the initial conformity assessment, that affects compliance with the high-risk requirements set out in Section 2 of Chapter III AIA. Such changes may, for instance, introduce new functionalities or alter the behaviour or outputs of the system in ways that meet functional or non-functional requirements to a different degree, thereby changing the risks associated with the system. For example, a system may be retrained on new datasets that affect accuracy across different subgroups of affected individuals, or it may be equipped with additional functions that allow it to use external tools and operate more autonomously.

Art. 43(4) AIA qualifies this trigger for AI systems that continue to learn after being placed on the market or put into service. Changes to the system and its performance that the provider has pre-determined in the initial conformity assessment, and documented in the technical documentation, do not constitute a substantial modification. Performance changes resulting from “online learning”, for example, based on user interactions or feedback from the operating environment, may therefore fall outside the substantial modification regime, provided they were anticipated and assessed at that earlier stage. They must, however, stay within the limits set in the risk assessment, as changes beyond that envelope could give rise to new risks.

Taken together, these provisions raise the question of where to draw the line between a routine update and a substantial modification that affects compliance with the high-risk requirements (see the stylised use cases provided in Section 2). To support providers and other actors in this assessment, regulatory guidance could include the following elements:

As a first element, regulatory guidelines may establish conceptual criteria for identifying when compliance with the high-risk requirements is likely to be affected by the modification of a high-risk AI system. These may include:

- New functionalities that give rise to new or heightened risks compared to the original system.
- Changes in system behaviour or outputs such that performance against functional requirements is significantly affected.
- Changes in system behaviour or outputs such that non-functional requirements (such as bias, robustness, or interpretability) are significantly affected.

Where feasible, these conceptual criteria may be supported by technical proxies or quantitative thresholds, for example, in terms of performance thresholds, to make their application predictable. However, such proxies or quantitative thresholds should not interfere with technological neutrality.¹⁷

A second element is a typology of the technical and organisational changes that may change the risk profile of a high-risk AI system and thus affect compliance with high-risk requirements. Such a typology could distinguish between several layers at which changes typically occur:

- Changes to the AI model at the core of the AI system, such as fine-tuning on different data

¹⁷ Meyers, Z., Schnurr, D., & Larouche, P. (2025). The AI Act and Technological Neutrality. CERRE Issue Paper. Available at <https://cerre.eu/publications/the-ai-act-and-technological-neutrality/>.



sources, post-training on new human feedback, model distillation, architectural modifications, the addition of new data modalities, or unlearning interventions.

- Changes at the system level, such as the addition of external tools, integration of additional data sources, agent scaffolding, or revised system prompts that unlock new capabilities.
- Changes to safeguards, such as the reconfiguration of input or output filters, the replacement of guardrails, or changes to safety thresholds.
- Organisational and accountability changes, such as modifications to human oversight arrangements, shifts in responsibility for post-market monitoring, or transfers of control between provider and deployer.

A third element could be indicative positive and negative lists of changes that do, and do not, qualify as substantial modifications. Such lists could draw on practical experience from sector-specific regulatory frameworks, in particular the regulation of medical devices. For example, under the EU Medical Device Regulation, the qualification of software changes has been the subject of detailed guidance from the Medical Device Coordination Group.¹⁸ This guidance presents specific examples for both significant and non-significant changes to software, combined with a decision flowchart. Lessons from such sectors could inform the calibration of the AIA's regime, while taking into account the specific features of AI systems and the cross-sectoral nature of the AIA.

A further instructive analogue comes from the Cyber Resilience Act (CRA), whose definition of substantial modification in Art. 3(30) CRA mirrors the structure of Art. 3(23) AIA: a post-market change that either affects the product's compliance with the essential requirements or modifies the intended purpose for which the product was assessed.¹⁹ Recent draft Commission guidance on the application of the CRA works through this test specifically for software updates, distinguishing substantial modifications from routine ones by reference to whether the update introduces new or increased risks not anticipated in the original risk assessment.²⁰ To this end, it sets out a non-exhaustive list of factors for the case-by-case assessment, including whether an update introduces new threat vectors, enables new attack scenarios, or changes the likelihood or impact of previously identified risks, and illustrates the test with worked examples on both sides of the line. Because the CRA focuses on cybersecurity rather than the AIA's broader high-risk requirements, the value of this guidance for the AIA lies mainly in its method for delineating substantial from routine software changes, rather than in the substantive thresholds themselves. Nonetheless, alignment of the underlying concept of substantial modification across these regimes is desirable, so far as possible, to minimise redundant or divergent compliance assessments for operators that are subject to both regulations.

The proposed elements for regulatory guidance would not eliminate the need for case-by-case judgement, but would provide more predictable reference points for providers, deployers, notified bodies, and authorities. Returning to Use Case A sketched in Section 2, such guidance would help to

¹⁸ Medical Device Coordination Group Document (2023). MDCG 2020-3 Rev.1. Guidance on significant changes regarding the transitional provision under Article 120 of the MDR with regard to devices covered by certificates according to MDD or AIMDD. https://health.ec.europa.eu/document/download/800e8e87-d4eb-4cc5-b5ad-07a9146d7c90_en?filename=mdcg_2020-3_en_1.pdf.

¹⁹ Regulation (EU) 2024/2847 of the European Parliament and of the Council of 23 October 2024 on horizontal cybersecurity requirements for products with digital elements and amending Regulations (EU) No 168/2013 and (EU) 2019/1020 and Directive (EU) 2020/1828 (Cyber Resilience Act), <http://data.europa.eu/eli/reg/2024/2847/oj>

²⁰ Draft Commission guidance on the application of Regulation (EU) 2024/2847 (Cyber Resilience Act), Section 4.3, https://ec.europa.eu/info/law/better-regulation/have-your-say/initiatives/16959-Draft-Commission-guidance-on-the-Cyber-Resilience-Act_en



clarify, for example, whether retraining on new data sources, an adjustment of scoring weights through a vendor-provided configuration interface, or a vendor-pushed swap of the underlying GPAI model amounts to a substantial modification, or whether such changes can be accommodated within a sufficiently broad initial conformity assessment.

3.3.2 Modifications of the intended purpose

The second trigger for a substantial modification is a change that modifies the intended purpose for which the high-risk AI system has been assessed.

As described above, a parallel rule applies to non-high-risk AI systems, including GPAI systems: where a third party modifies the intended purpose of such a system in a way that causes it to become a high-risk AI system, that third party becomes the provider of the resulting high-risk AI system and must fulfil the corresponding obligations, including carrying out a conformity assessment (Art. 25(1)(c) AIA).

In the context of systems classified as high-risk under Art. 6(2), the AIA ties the high-risk classification to the list of intended purposes set out in Annex III. In consequence, a modification under Art. 25(1)(b) and (c) AIA will be relevant where the modified intended purpose falls within one of these listed purposes, or the modified system falls under the scope of Art. 6(1) AIA.

A central interpretative question is to what extent and in what role the AI system is connected to that intended purpose (see Use Case C in Section 2). Where the AI system constitutes the final product or system as a whole, the assessment is more straightforward. The assessment becomes more complex where the AI system is one component of a broader system or product, since its contribution to the overall intended purpose may vary considerably depending on system design and use case. At one extreme, an AI system may form part of a broader (non-AI) system used for a purpose listed in Annex III without materially affecting the process or outcome of the overall system in relation to that purpose. More commonly, the role of the AI system as a sub-component, and its contribution to the overall output, will fall somewhere on a spectrum.

Several settings and constellations are of particular practical relevance and may thus be considered by future regulatory guidance:

Significance threshold for the change of an intended purpose: For high-risk AI systems falling under Annex III, the classification structure suggests that a shift of the intended purpose from one listed use case to another can be assumed to be substantial. The assessment is more ambiguous, however, where the intended purpose is expanded within a single listed use case (see also Use Case A in Section 2), and additional guidance on the relevant threshold of significance would be desirable.

Integration of GPAI systems into systems and products in high-risk areas: GPAI systems will frequently be used in high-risk areas, and may be integrated as a sub-component into a broader high-risk AI system (see Use Case C in Section 2). A key question in this context is under what conditions such use or integration brings the GPAI system itself within the high-risk regime.

Orchestration and interaction of multiple AI systems in high-risk areas: Where multiple AI systems interact and rely on each other's outputs, or where a larger AI system consists of multiple components that are themselves AI systems, the high-risk nature may emerge only from the interaction and orchestration of these components (see Use Case B in Section 2). In such situations, it is not obvious



how to localise the substantial modification, nor how to allocate responsibility among the providers of the individual components and the operator orchestrating them.

Features that confer autonomy and general-purpose capabilities: A further difficulty arises where a system could, in principle, be applied to high-risk domains, but this is not its intended purpose, and where the system may autonomously enter such a domain as part of a broader objective. As illustrated in Use Case C, agentic GPAI systems may independently engage in activities that contribute to high-risk purposes, without any deliberate deployment decision by an operator to that effect. Further guidance may consider whether and how the framework of Art. 25(1)(c) AIA could accommodate such emergent purposes, and how responsibility should be allocated between the provider of the agentic system and the operator within whose environment it acts.

Safety components in products covered by Union safety legislation: A modification will also trigger high-risk requirements where the intended purpose of the AI system shifts to use as a safety component of a product covered by Union safety legislation. Under Art. 3(14) AIA, a safety component is a component of a product or of an AI system which fulfils a safety function for that product or AI system, or the failure or malfunctioning of which endangers the health and safety of persons or property.²¹ The term is also used in sectoral safety legislation, such as the Machinery Regulation, but with a narrower scope in some respects. For example, the Machinery Regulation additionally requires that a safety component is “not necessary in order for that product to function or for which normal components may be substituted in order for that product to function” and refers only to the safety of persons rather than also to property.²²

The recently published draft guidelines on the classification of high-risk AI systems address this divergence directly.²³ They explicitly state that the Art. 3(14) definition is “an autonomous definition of the AIA that has its own meaning, independent of definitions of ‘safety component’ included in other Union harmonisation legislation,”²⁴ and that it applies independently of, and to the exclusion of, those sectoral definitions, so that high-risk classification under Art. 6(1) is assessed solely against the AIA definition. The draft guidelines also provide indicative examples of safety functions, comparable to the indicative list of safety components in Annex II of the Machinery Regulation,²⁵ to clarify which AI systems may qualify as safety components, which could support implementation.

For the classification question itself, the autonomous-definition approach resolves the risk of parallel assessments: an operator does not apply both the AIA and the sectoral definition to determine high-risk status, but the AIA definition alone. A coherence problem, nevertheless, remains downstream. An operator subject to both the AIA and sectoral legislation will encounter the same term, “safety component,” carrying different meanings for the purpose of its substantive obligations under each regime. Regulatory guidance should therefore aim to give businesses subject to several regimes practical direction on how to comply across them without maintaining parallel compliance processes. From a longer-term perspective, the legislator should work towards aligning the meaning of identical

²¹ Note that the Digital Omnibus on AI (provisional agreement of 7 May 2026) amends the definition of “safety component”. As the consolidated text was not yet available at the time of writing, this paper refers to Art. 3(14) AIA as enacted and the relevant passages should be checked against the final text.

²² Art. 3(3) of the Regulation (EU) 2023/1230 of the European Parliament and of the Council of 14 June 2023 on machinery and repealing Directive 2006/42/EC of the European Parliament and of the Council and Council Directive 73/361/EEC.

²³ European Commission (2026). Draft Commission guidelines on the classification of high-risk AI systems under Article 6 of Regulation (EU) 2024/1689 (AI Act), Annex I, <https://ec.europa.eu/newsroom/dae/redirection/document/128560>

²⁴ Ibid., para. (33).

²⁵ Notably, Annex II already explicitly lists “safety components with fully or partially self-evolving behaviour using machine learning approaches ensuring safety functions,” reflecting the convergence of machinery safety practice with AI-specific concerns.



terms across legislation, in order to facilitate implementation and realise the intended benefits of the AIA as a cross-cutting horizontal regulation.

3.4 Research, testing, and development activities

A general exemption from the scope of the AIA applies to AI systems and AI models that are specifically developed and put into service for the sole purpose of scientific research and development (Art. 2(6) AIA). In addition, the regulation exempts any research, testing, or development activity regarding AI systems or AI models prior to their being placed on the market or put into service (Art. 2(8) AIA). However, this exemption explicitly does not extend to testing in real-world conditions.

Art. 3(57) AIA defines testing in real-world conditions as the temporary testing of an AI system for its intended purpose in real-world conditions outside a laboratory or otherwise simulated environment, with a view to gathering reliable and robust data and to assessing and verifying the conformity of the AI system. It further stipulates that real-world testing does not qualify as placing the AI system on the market or putting it into service, provided that the conditions set out in Art. 57 on AI regulatory sandboxes or Art. 60 on testing of high-risk AI systems are fulfilled.

Thus, for high-risk AI systems, Art. 60 AIA establishes a regulated route for testing in real-world conditions outside the scope of the Art. 2(8) exemption, subject to a real-world testing plan, registration, oversight by market surveillance authorities, informed consent of subjects, and other safeguards. No equivalent pathway exists for GPAI models, which makes the threshold under Art. 2(8) AIA particularly consequential in those cases.

The delineation of these exemptions is important for GPAI model providers, AI system providers, and downstream actors alike, as it determines the point at which the full set of AIA obligations begins to apply. As illustrated in Use Case D in Section 2, industry practices such as the release of new GPAI model versions on community evaluation platforms, A/B testing of updated models within systems already on the market, and private previews to selected organisational customers are central to the evolution of AI systems and GPAI models, and some of these practices may sit between pre-market development and post-market deployment.

Furthermore, the Art. 2(8) exemption may interact with the rules on substantial modification. Where a provider tests an updated version of an AI system already on the market, for example through A/B testing within an existing service, the activity may be characterised either as research and development on the new version or as a deployment of the modified system. The latter characterisation could trigger the substantial modification regime where the system in question is high-risk, with corresponding obligations to perform a new conformity assessment. The assessment is further influenced by the extent to which the relevant updates fall within the carve-out under Art. 43(4) AIA for AI systems that continue to learn after being placed on the market, and therefore do not constitute a substantial modification.

A related boundary question concerns significant modifications of GPAI models (such as through fine-tuning) that are exclusively used internally. A modifier whose modification leads to a sufficiently



significant change in the model may become the provider of the resulting GPAI model.²⁶ However, Recital 97 AIA indicates that the obligations for GPAI models do not apply where an own model is used for purely internal processes, provided that those processes are not essential for providing a product or service to third parties, the rights of natural persons are not affected, and the model is not a GPAI model with systemic risk. Internal use of a self-modified GPAI model therefore should, in principle, not trigger model-provider obligations, even where the modification would otherwise be significant enough to do so. Yet the principle is expressed conditionally, and each condition is open-textured. It is, for instance, not always clear when an internal process is “essential” for providing a service to third parties: a model used in a back-office function may come to underpin a customer-facing service without being supplied to anyone, which would remove the exemption. The systemic-risk exception is particularly consequential, as it withdraws the general principle where a modification has pushed a model into the systemic-risk tier. This may lead to considerable uncertainty about organisations’ ability to modify GPAI models for internal purposes without inadvertently facing GPAI model-provider obligations.

Further regulatory guidance on these issues would be particularly helpful, as they concern widespread industry practices that are highly relevant both for the effectiveness of risk mitigation and for innovation.

²⁶ Guidelines on the scope of the obligations for providers of general-purpose AI models under Regulation (EU) 2024/1689 (AI Act), 18 July 2025. <https://ec.europa.eu/newsroom/dae/redirection/document/118340>, Sec. 3.2 establishes that downstream modifiers may become the providers of GPAI models when the modification leads to a significant change in the model’s generality, capabilities, or systemic risk.



4. Institutional Framework and AI Value Chain Implications

4.1 Implications for the AI value chain

The modification of AI systems by third parties has important implications for the AI value chain and the responsibilities of the actors within it. This section highlights three main issues that arise in this context. A more detailed treatment of cooperation and information sharing obligations along the AI value chain is provided in a companion CERRE paper.²⁷

Scope of cooperation duties between original and new providers: Where a third party modifies an AI system in a way that triggers Art. 25(1)(b) or (c) AIA and thereby becomes the new provider of a high-risk AI system, Art. 25(2) AIA requires the original provider to closely cooperate with the new provider. In particular, the original provider must make available the necessary information and provide the reasonably expected technical access and other assistance required for the fulfilment of obligations under the AIA, in particular as regards the conformity assessment of high-risk AI systems. However, the scope of such information sharing and assistance, as well as the procedural arrangements for its implementation, remain to be determined. While this leaves room for case-by-case flexibility, regulatory guidance would be desirable to reduce coordination costs and clarify compliance expectations. Such guidance could indicate, for example, what categories of documentation are typically expected, what forms of technical access are reasonable, and how confidentiality and trade-secret protection interact with these duties.

Notification and monitoring along the value chain: The AIA is largely silent on the procedural arrangements through which value chain actors exchange information and notify each other about modifications to AI systems. This raises several questions. What responsibilities do downstream providers or deployers have to inform upstream providers about modifications, and vice versa? While Art. 49 AIA requires providers to register their high-risk AI systems in a centralised EU database (see Art. 71), the AIA does not appear to impose a direct duty on a third-party modifier to notify the original provider that a modification has occurred. Conversely, do AI system providers need to monitor the downstream use of their systems in order to detect potential modifications? Especially for providers of GPAI systems made available through APIs to a large customer base, it may be difficult, if not impossible, to monitor all downstream uses of their systems. Even with a relatively small customer base, a GPAI system provider may have little visibility into how customers integrate and use the system (see also Use Case C in Section 2).

Modification restrictions and residual cooperation duties: Art. 25(2) AIA stipulates that the cooperation duties of the original provider do not arise where the original provider has clearly specified that its AI system is not to be changed into a high-risk AI system. Several questions arise as to the operation of this modification-restriction waiver. First, the form and placement of such modification restrictions should be clarified and harmonised. Second, where additional measures, such as technical safeguards, are considered necessary to give effect to such modification restrictions, further guidance on the expected level of these measures is needed. Third, the relationship between

²⁷ Schnurr (2026). Information Sharing and Cooperation along the AI Value Chain under the AI Act. CERRE.



Art. 25(2) and Art. 25(4) AIA warrants further clarification. Art. 25(4) AIA requires providers of AI systems, tools, services, components, or processes that are used in or integrated into a high-risk AI system to specify, by written agreement with the high-risk provider, the information, capabilities, technical access, and other assistance necessary for compliance. Crucially, the AIA does not include a corresponding waiver for these general input supplier duties. As a result, it is unclear whether an original provider that has effectively restricted high-risk modification under Art. 25(2) AIA may nonetheless face cooperation duties under Art. 25(4) AIA where its system is integrated as a component of a high-risk AI system.

4.2 Institutional framework and cross-sector application

Beyond the substantive issues highlighted in the preceding sections, the implementation of the AIA rules on the modification and evolution of AI systems raises additional questions regarding the accompanying institutional framework.²⁸ A central question concerns the respective roles of notified bodies, market surveillance authorities, and the AI Office in clarifying and enforcing the boundaries set by the modification regime. Notified bodies will play a key role in conformity assessment and re-assessment following substantial modifications, while market surveillance authorities will be responsible for enforcement at the national level. The interaction between these actors, particularly where modifications cut across Member States and involve AI systems with Union-wide reach, will need to be clarified to avoid fragmentation and inconsistent outcomes. Thus, the AI Office should play a central role in issuing cross-cutting guidance to support consistent application across Member States and by market surveillance authorities. This may be further supported by institutions intended to promote harmonised application of the AIA, specifically the European AI Board (Art. 65–66 AIA).

This institutional question is closely connected to how the modification rules will be applied and enforced across sectors. As discussed in Section 3, sector-specific regulation directly interacts with the AIA rules, such as in the case of the Medical Device Regulation, and in some cases uses somewhat different definitions for the same key terms, such as the notion of a safety component under the Machinery Regulation. Recital 84 AIA clarifies that the AIA provisions apply without prejudice to more specific provisions established in sector-specific Union safety legislation. Regulatory guidance and implementation will therefore need to account for sector-specific characteristics and pre-existing regulatory frameworks. At the same time, for the AIA to realise its benefits as a horizontal regulatory framework, implementation should be aligned as much as possible across sectors. This is of particular relevance for GPAI systems, which by their very nature may be deployed across a wide range of sectors and use cases.

²⁸ For a broader discussion of the institutional framework of the AIA and its implementation, see Larouche (2025).



5. Recommendations for Implementation

Taken together, the issues discussed in this paper highlight that the AIA's provisions on the modification of AI systems and the responsibilities along the value chain leave substantial room for interpretation in their practical operation. Forthcoming guidance from the AI Office on the practical application of the high-risk classification, the responsibilities of value chain actors, the practical application of the substantial modification regime, and the coordination with sector-specific frameworks will play an important role in addressing these uncertainties and in ensuring that the AIA operates in a way that is both effective and proportionate.

On this basis, the paper offers the following overarching recommendations on how regulatory guidance can support the implementation of the AIA's provisions on substantial modification:

Recommendation 1: Operationalise the threshold for modifications that affect compliance with the high-risk requirements. Guidance from the AI Office should give providers and other actors workable reference points for identifying when a change affects compliance with the high-risk requirements. As set out in Section 3.2, this could combine conceptual criteria (new functionalities, and changes to functional or non-functional performance), a typology of the technical and organisational layers at which changes occur, and indicative positive and negative lists of qualifying and non-qualifying changes. In developing these, the AI Office can draw on the Blue Guide's established approach to substantial modification under the New Legislative Framework, as well as the method of sector-specific analogues, such as the Medical Device Coordination Group's guidance on software changes under the Medical Device Regulation and the draft Commission guidance under the Cyber Resilience Act.

In designing this guidance, the AI Office should ground it in real-world use cases that reflect the complexity and dynamic nature of value creation in AI value chains. The primary intended audience for operational guidance of this kind should be engineering and product teams as much as legal departments, and the guidance should be drafted to be usable by those who actually implement and modify AI systems. This implies a trade-off between guidance that is short and crisp and guidance that is concrete and workable through detailed examples. The two should be balanced and combined, rather than the one pursued at the expense of the other. At the same time, because the AIA functions as a horizontal framework spanning many sectors and use cases, the guidance should retain sufficient flexibility to allow compliance solutions to be tailored to specific contexts and to accommodate rapid technical progress and the associated market developments, rather than fixing rigid criteria that may quickly fall out of step with the technology and market practice.

Recommendation 2: Clarify the significance threshold for changes to the intended purpose. Guidance should address when a modification of the intended purpose is substantial, in particular the more ambiguous cases in which the purpose is expanded within a single Annex III use case rather than shifted from one listed use case to another. It should also address how to assess an AI system's role in, and contribution to, the intended purpose of a broader product or system of which it forms part.

In operationalising these thresholds, the AI Office should be guided by proportionality. Frequent updating is intrinsic to how AI systems are built and maintained, and agile, iterative development is common practice across the industry. Guidance should therefore avoid an interpretation of



substantial modification so expansive that it discourages providers and deployers from improving their systems for fear of triggering a fresh round of compliance obligations. Routine lifecycle changes should as a rule remain outside the substantial modification regime, consistent with the continuous learning carve-out in Art. 43(4) AIA and the post-market monitoring duties that apply in any case. Where there is concern that a permissive approach to individual updates could hollow out the substantial modification framework, the AI Office could consider a backstop mechanism: a periodic reassessment of whether a series of smaller updates, each non-substantial on its own, has cumulatively amounted to a substantial modification. Such a mechanism would preserve the framework's integrity against incremental drift while leaving everyday development unimpeded, although this should be weighed against the additional compliance burden such a mechanism would create.

Recommendation 3: Address GPAI and agentic systems that lack an assessed baseline. Because GPAI systems are not anchored to an initial conformity assessment that defines an intended purpose, the substantial modification framework does not apply to them in the same way. At the same time, the use of GPAI systems is growing rapidly and diffusing across sectors, including areas the AIA treats as high-risk, which makes the question of when such systems fall within the high-risk regime increasingly consequential in practice. Guidance should therefore clarify the conditions under which the use or integration of a GPAI system brings it within the high-risk regime, how to localise the modification and allocate responsibility where high-risk character emerges from the orchestration of multiple systems, and how Art. 25(1)(c) AIA should be applied to agentic systems whose effective purpose emerges from their own behaviour rather than from a deliberate deployment decision.

Such guidance should also address the case in which a non-high-risk AI system gradually comes to be used to support high-risk purposes, for example, through the everyday actions of a deployer's staff, rather than through any deliberate decision to deploy it in a high-risk area (see the HR-assistant and university variants in Use Case C, Section 2). It should also clarify when and how GPAI systems can be used in high-risk areas without themselves being considered high-risk AI systems under the AIA; this will depend on the significance of the AI system's contribution to the intended purpose of the broader system or process.

Recommendation 4: Delineate the boundary between pre-market testing, internal use, and placing on the market. Guidance should clarify the scope of the Art. 2(8) exemption and its exclusion of testing in real-world conditions, with particular attention to GPAI models, for which no equivalent to the Art. 60 pathway exists. It should address how common practices such as community-platform evaluations, A/B testing within deployed systems, and private previews relate to that boundary. It should also clarify how the boundary interacts with the substantial modification regime and with self-classification under Art. 6(3) AIA. For GPAI models used exclusively internally, guidance should further address the uncertainty about when modification triggers model-provider obligations and how the internal-use principle in Recital 97 AIA, including its conditions and systemic-risk exception, applies.

In doing so, guidance should recognise that testing AI systems and GPAI models with real-world users is essential and common market practice, and a key means by which providers identify and mitigate risks and improve quality before wider release. The AIA's provisions should not be interpreted so as to impair the quality and safety of high-risk AI systems through unduly strict limitations on such testing. At the same time, guidance should ensure that testing does not itself expose individuals to harm, in particular where it is conducted with real users outside controlled environments.



Recommendation 5: Clarify cooperation and notification duties along the value chain. Guidance should specify the scope of the original provider’s cooperation duties under Art. 25(2) AIA, specifically the categories of documentation typically expected, the forms of reasonable technical access, and the interaction with confidentiality and trade-secret protection. It should further clarify the notification and monitoring responsibilities of upstream and downstream actors and how such arrangements might work in practice. This applies both to the relationship between the providers of the original and the modified AI system in the scenarios set out in Art. 25(1) AIA, and to the more general relationship between GPAI model providers, GPAI system providers, high-risk AI system providers, and deployers. Particular attention should be given to the relationship between the input-supplier duties in Art. 25(4) AIA and the modification-restriction waiver in Art. 25(2) AIA, since Art. 25(4) contains no corresponding waiver, leaving it unclear whether a provider that has restricted high-risk modification could nonetheless face cooperation duties where its system is integrated into a high-risk AI system.

Recommendation 6: Align shared concepts across sectoral regimes and ensure institutional consistency. Where the AIA and sectoral legislation use the same terms with different meanings, for instance, “safety component”, guidance should help operators subject to several regimes comply without maintaining parallel compliance processes, and the legislator should, over the longer term, work towards aligning the meaning of identical terms across legislation. The AI Office, supported by the European AI Board, should issue cross-cutting guidance to ensure consistent application across Member States, market surveillance authorities, and notified bodies, and to prevent fragmentation where modifications have Union-wide reach. Decentralised enforcement at national level can bring responsiveness, but centralised coordination is essential to prevent divergent or overlapping action across authorities and Member States.

The development of AIA guidance should also be coordinated with the broader revision of the EU product-safety framework now underway. The New Legislative Framework, on which the AIA’s substantial modification concept largely builds, is itself under review as part of the forthcoming European Product Act, expected in 2026. The AI Office should track this revision and ensure its guidance remains capable of adjusting to any resulting changes, so that the two frameworks evolve consistently.



About CERRE

Providing high quality studies and dissemination activities, the Centre on Regulation in Europe (CERRE) is a not-for-profit think tank. It promotes robust and consistent regulation in Europe's network, digital industry, and service sectors. CERRE's members are regulatory authorities and companies operating in these sectors, as well as universities.

CERRE's added value is based on:

- its original, multidisciplinary and cross-sector approach covering a variety of markets, e.g., energy, mobility, sustainability, tech, media, telecom, etc.;
- the widely acknowledged academic credentials and policy experience of its research team and associated staff members;
- its scientific independence and impartiality; and,
- the direct relevance and timeliness of its contributions to the policy and regulatory development process impacting network industry players and the markets for their goods and services.

CERRE's activities include contributions to the development of norms, standards, and policy recommendations related to the regulation of service providers, to the specification of market rules and to improvements in the management of infrastructure in a changing political, economic, technological, and social environment. CERRE's work also aims to clarify the respective roles of market operators, governments, and regulatory authorities, as well as contribute to the enhancement of those organisations' expertise in addressing regulatory issues of relevance to their activities.



About the Author



Daniel Schnurr is a CERRE Research Fellow and a Professor of Information Systems at the University of Regensburg, where he holds the Chair of Machine Learning and Uncertainty Quantification. Previously, he led the Data Policies research group at the University of Passau. He received his Ph.D. in Information Systems from the Karlsruhe Institute of Technology, where he also completed his B.Sc. and M.Sc. in Information Engineering and Management. Daniel Schnurr has published in leading journals in Information Systems and Economics on competition and data sharing in digital markets, regulation of data-driven market power, and competition and cooperation in telecommunications markets. His current research focuses on the role of artificial intelligence in competition, privacy and data sharing in digital markets as well as regulation of AI, cloud computing and the data economy.

cerre



Avenue Louise 475 (box 10)
1050 Brussels, Belgium
+32 2 230 83 60
info@cerre.eu
www.cerre.eu

 Centre on Regulation in Europe (CERRE)

 CERRE Think Tank

 CERRE Think Tank

