

The EU AI Act

A Practitioner's Guide for Korean Companies

Kim Kyung-jin, Attorney at Law

Prologue

August 2, 2026 - A Call from Brussels

August 2, 2026, was a Sunday. That afternoon, I received a phone call from the CEO of an artificial intelligence company based in Pangyo. His voice was different from usual. He said he had received an email from a European partner indicating that their chatbot seemed to be violating European law starting today. He thought that could not possibly be the case - after all, they had no office in Europe and had merely attached a customer support chatbot to their service.

I asked him whether the chatbot informs European users that it is an artificial intelligence when interacting with them. A brief silence ensued. He asked back why on earth they would need to inform users of such a thing. That counter-question contained the entire reason why this book came to be written.

That day marked the full enforcement of the transparency obligations under Article 50 of the EU Artificial Intelligence Act. From that day on, any AI system interacting with human beings was required to inform users that it was an AI. The company in Pangyo had no idea it fell within the scope of the law. They believed that having no legal entity in Europe meant European law had nothing to do with them. Yet their chatbot was conversing with European users, and in that moment, the company's registered address did not matter. What mattered was what that AI system was doing in Europe.

I have received calls like this multiple times. For one company, a promotional video created using generative AI became an issue; for another, an evaluation system supplied to a European bank was caught in the regulatory net. They all shared something in common: none of them knew they were subject to regulation. And by the time they realized it, those regulations were already in effect.

Materials covering the EU Artificial Intelligence Act flood the market - law firm newsletters, consulting firm white papers, and guidebooks summarizing legal provisions. Most of them set out what the articles stipulate and list the dates of application. Even if he had read them, the CEO in Pangyo would not have realized that his company was affected. Those materials were either written assuming European enterprises or stopped short at summarizing the legal provisions.

This book is written from a different vantage point. It assumes, right from the start, a company headquartered in Seoul. Rather than simply setting out what the provisions dictate, it addresses how enforcement authorities will view those provisions. And where judgment is required, it does not shy away from making a call. I will write sentences such as: *This is how I read this provision*, *Views diverge on this point*, or *This part remains uncertain*. Even if the neutrality expected of a reference document is somewhat softened, I chose to leave the distinct imprint of a book written by a human.

I hope this book serves two purposes. First, as a decision-making tool. I hope executives of companies considering expansion into Europe can spend 30 minutes with it and answer whether they are subject to this law and, if so, under which classification. Second, as a practical manual. I hope compliance officers can print out the checklist in the appendix and keep it on their desks for daily use. That is why the opening section of each chapter is designed to be read as narrative prose, while the closing section serves as reference material.

There is one thing I should point out in advance. The subject of this book remains a moving target. Even while writing this text, enforcement dates shifted and new guidelines were issued. On June 29, 2026, amendments redesigning several key dates received final approval. For this reason, the dates and figures in this book have been managed separately, and I advise readers to double-check those details before embarking on practical preparations. The ground underfoot is still shifting slightly.

I had a long phone conversation with that CEO in Pangyo later that evening. We walked through what had been triggered, what needed to be done first, and what still allowed some time. Near the end of the call, he said, "If only I had known sooner." This book is written for that "sooner." My hope is that it will be a book opened and read before August 2, rather than a phone call received on it.

Contents

Prologue

Part I Why Korean Companies Must Read Europe's Law

Chapter 1 The Reality of Extraterritorial Application

- 1.1 Three Paths Through Which a Seoul-Based Company Becomes Subject to EU Law
- 1.2 Distinguishing Providers, Deployers, Importers, and Distributors
- 1.3 When Output Is Used in the EU: Article 2(1)(c)
- 1.4 EU Authorised Representative: Articles 22 and 54
- 1.5 Self-Diagnostic Checklist

Chapter 2 Framework of the Law

- 2.1 The Design Philosophy Called the Risk-Based Approach
- 2.2 The Four Tiers: Prohibited, High-Risk, Transparency, and Minimal Risk
- 2.3 Relationship with GDPR, the Product Liability Directive, and the Data Act
- 2.4 Sanction Levels: Fines Calculated Based on Global Turnover

Chapter 3 The 2026 Timeline, Redrawn Maps

- 3.1 Entry into Force on August 1, 2024, and the Design of Phased Application
- 3.2 February 2, 2025: Prohibited Practices and AI Literacy Obligations
- 3.3 August 2, 2025: Obligations for Providers of General-Purpose AI Models
- 3.4 August 2, 2026: Full Applicability of Article 50 Transparency Obligations and Enforcement Powers over General-Purpose AI
- 3.5 The Digital Omnibus Amendment: What Changed with the Council's Final Approval on June 29, 2026
- 3.6 Postponement of Annex III Standalone High-Risk AI to December 2, 2027
- 3.7 Postponement of Annex I Embedded High-Risk AI to August 2, 2028
- 3.8 The Line Dividing What Was Postponed and What Was Not

Part II What You Must Comply With Now: Transparency

Chapter 4 Article 50 Splits into Four Directions

- 4.1 Chatbots and Conversational Agents: How to Disclose That It Is Not a Human
- 4.2 Machine-Readable Marking of Synthetic Content: How Far Has Watermarking Come?
- 4.3 Deepfake Disclosure Obligations and Deployer Responsibility
- 4.4 Notification Obligations for Emotion Recognition and Biometric Categorisation Systems
- 4.5 To What Extent Can Limited Exceptions for Art, Satire, and Safety Be Invoked?

Chapter 5	Watermarking Technology and Law
	5.1 December 2, 2026 Grace Period for Existing Systems, and Zero Grace Period for Newly Launched Systems
	5.2 The C2PA Standard and Content Provenance Authentication
	5.3 Robustness Requirements: Can It Withstand Re-encoding and Screenshots?
	5.4 Ban on Generating NCII and CSAM: Article 5 Amendments and December 2, 2026 Entry into Force
Chapter 6	Transparency Code of Practice
	6.1 Practical Benefits of Signing: Shift in the Burden of Proof
	6.2 Compliance and Its Relationship to the Presumption of Conformity
	6.3 Alternative Pathways for Non-Signatory Companies
Part III For Companies Handling General-Purpose AI Models	
Chapter 7	Are We a Provider of General-Purpose AI?
	7.1 The Criterion of Significant Generality
	7.2 The Starting Point of the Lifecycle: Large-Scale Pre-training
	7.3 Fine-tuning and Downstream Modifications: When Does One Become a New Provider?
	7.4 Discussion of the One-Third Threshold and Practical Judgment
Chapter 8	Basic Obligations of General-Purpose AI Providers
	8.1 Technical Documentation: What Annex XI Requires
	8.2 Information Provided to Downstream Providers: Annex XII
	8.3 EU Copyright Compliance Policy: Reservation of Rights under Article 4(3) of the Directive
	8.4 Public Summary of Training Content: The Boundary Between Templates and Trade Secrets
Chapter 9	Systemic Risk Models
	9.1 The 10^{25} FLOP Threshold and Compute Estimation
	9.2 Hardware-Based Estimation and Architecture-Based Estimation
	9.3 Model Evaluation and Adversarial Testing
	9.4 Serious Incident Reporting
	9.5 Securing Cybersecurity and the July 2026 Cybersecurity AI Action Plan
Chapter 10	For Companies Seeking to Use the Open-Source Exception
	10.1 The Four Freedoms of Free and Open-Source Licenses
	10.2 Typical Non-Compliance Cases
	10.3 Practical Application of the Non-Monetization Requirement
	10.4 Two Obligations That Remain Despite the Exception

Part IV High-Risk AI: Preparing by 2027

Chapter 11 Is Our Product High-Risk?

11.1 Annex III: Eight Areas Covering Standalone Systems

11.2 Annex I: AI Embedded Within Existing Products

11.3 The Filter under Article 6(3) and the Trap of the Burden of Proof

11.4 Applying the Framework to Six Key Korean Industries

Chapter 12 Obligation System for High-Risk System Providers

12.1 Risk Management System: The Cyclical Structure Required by Article 9

12.2 Data Governance and Training Data Quality: Article 10 and the Issue of Bias

12.3 Technical Documentation, Automatic Logging, and Information to Be Provided to Deployers

12.4 Human Oversight: The Control Vision of Article 14

12.5 Accuracy, Robustness, and Cybersecurity: The Technical Baseline Required by Article 15

12.6 Conformity Assessment, CE Marking, and EU Database Registration

Chapter 13 Obligations of Deployers

13.1 Compliance with Instructions for Use, Assignment of Human Oversight, and Six-Month Log Retention

13.2 Deployment of AI for Workers and Prior Notification

13.3 Fundamental Rights Impact Assessment: Who Must Perform It?

Chapter 14 How to Use the Delay Period

14.1 Background: Delays in Establishing Harmonized Standards

14.2 Alternatives to Missing Standards: ISO/IEC 42001 and NIST AI RMF

14.3 Data Governance Void Created by Delays and Voluntary Responses

Part V Prohibited AI Practices

Chapter 15 The Line That Must Not Be Crossed

15.1 Subliminal Manipulation and Exploitation of Vulnerabilities

15.2 Social Scoring

15.3 Real-Time Remote Biometric Identification: Prohibition in Principle and Exceptions

15.4 Emotion Recognition: Outright Banned in Workplaces and Educational Institutions

15.5 Allowed Domestically, Prohibited in the EU

Part VI Intersection with Korean Law

Chapter 16 Dual Regulation of the Framework Act on AI and the EU AI Act

	16.1 Two Layers of the Amended Framework Act on AI Taking Effect in 2026
	16.2 High-Impact AI and High-Risk AI: Similar Concepts, Different Mesh Sizes
	16.3 Deepfakes: Visible to the Human Eye or Readable by Machines?
	16.4 Domestic Representatives and EU Authorised Representatives: Thresholds vs. No Thresholds
	16.5 AI Product and Service Verification System and Procurement Preferences
Chapter 17	Intersection with Data Protection
	17.1 Simultaneous Application of the GDPR and the Personal Information Protection Act
	17.2 Cross-Border Data Transfers and Adequacy Decisions
	17.3 Legal Basis for Processing Personal Data in Training Datasets
Part VII Execution: Building a Compliance Program	
Chapter 18	90-Day Implementation Plan
	18.1 Building an AI Inventory: Department by Department, Down to Invisible AI
	18.2 System-by-System Classification and Role Determination
	18.3 Gap Analysis
	18.4 Setting Priorities: By Dates and Fine Severity
Chapter 19	Documentation Systems
	19.1 What, in What Format, and for How Long Should Records Be Retained?
	19.2 Audit Response: How to Prove Consistency Between Documentation and Actual Operation
	19.3 Procedures for Responding to Information Requests from EU Authorities
Chapter 20	Contracts and the Supply Chain
	20.1 Provisions to Include in AI Supplier Contracts
	20.2 Information to Receive from Upstream Model Providers
	20.3 Information to Provide to Downstream Customers
	20.4 Liability Allocation and the Limits of Indemnification Clauses
Chapter 21	Organization and People
	21.1 AI Literacy Obligations: Article 4 and Article 26 Are Distinct Provisions
	21.2 Creating Roles Outside the Org Chart
	21.3 Legal Focuses on Law, Development on Technology, Business on the Market, and the Table Where All Three Decide Together
Chapter 22	Regulatory Sandboxes and Prior Consultation
	22.1 Member State Sandbox Requirements
	22.2 Measures for SMEs
	22.3 Communication Channels with the AI Office

Epilogue	Should Regulation Be Seen as a Cost or Used as a Barrier to Entry?
Appendix	
Appendix 1	Integrated Implementation Schedule Timeline
Appendix 2	System Classification Flowchart
Appendix 3	Self-Assessment Checklist
Appendix 4	Sample Disclosure Statements for Article 50 (Korean / English)
Appendix 5	Technical Documentation Table of Contents Template
Appendix 6	Sample AI Provider Contractual Clauses
Appendix 7	Terminology Comparative Table (Korean · English · Legal Provisions)
Appendix 8	EU Authorities and Domestic Support Agencies
Appendix 9	Summary of Key Articles

Part I

Why Korean Companies Must Read Europe's Law

Chapter 1. The Reality of Extraterritorial Application

Let us start with the story of a medical device company in Busan. This company produces AI diagnostic assistance software that reads chest X-ray images and highlights suspected nodule areas. Development, training, and server operations all take place domestically within South Korea. The company has no branch offices, no employees, and has signed no contracts in the EU. Yet, the moment this software is supplied to the department of radiology at a university hospital in Germany under a license agreement and begins to be used in actual diagnostic assistance, this company becomes subject to the European Union Artificial Intelligence Act (Regulation (EU) 2024/1689, hereinafter the AI Act). This holds true even though the server is located in Busan and only the outputs are transmitted to Germany. It holds true even though there is not a single line of EU address listed in its corporate registry.

Many legal counsel at Korean companies stumble at this very point. This is because it takes time to digest the fact that the statement "We have never expanded into the EU" and the statement "We are subject to the EU Artificial Intelligence Act" can both be true simultaneously. I view this point as the first threshold that runs through this entire book. It is not the company that crosses borders, but its output. And the AI Act draws its jurisdiction precisely based on the destination of that output.

1.1 Three Paths Through Which a Seoul-Based Company Becomes Subject to EU Law

Article 2(1) of the AI Act is the provision that determines to whom this law applies. There are three paths through which Korean companies fall within its scope. In order of the provisions, they are points (a), (b), and (c).

The path under point (a) applies to providers who place an AI system or a general-purpose AI model (hereinafter GPAI model) on the market or put it into service in the EU. The provision deliberately inserts the phrase "regardless of whether they are established or located within the Union." If a Korean company develops an AI system or a GPAI model at its Seoul headquarters and sells it directly or provides it via API to EU clients, it incurs provider obligations even without a single physical branch in the EU. In practice, people frequently overlook the fact that placing on the market can occur in the form of a software library, an API, a direct download, or even a physical copy. This means that placing on the market can be established through a single SaaS contract.

The path under point (b) applies to deployers of AI systems established or located within the EU. Seoul headquarters are rarely caught directly under this path. However, if a Korean company establishes a subsidiary or branch in the EU and that EU entity uses an AI system in the course of its business, that EU entity becomes a deployer. It is a structure where the local entity, rather than the headquarters, is regulated - though this is hardly grounds for headquarters to rest easy. This is precisely where holding-company-level compliance policies become necessary.

The remaining path under point (c) is the true subject of this chapter. Where the output produced by an AI system is used in the EU, both its provider and deployer fall under the law even if they are located outside the EU. The case of the Busan medical device company falls precisely under this path. I read this provision as follows: EU lawmakers determined that defining jurisdiction by physical borders no longer works in the AI era. In an industry where services flow through the cloud and cross borders with a single API call, using the company's location as the benchmark renders regulations an empty net. Therefore, they shifted the benchmark to the destination of the output - that is, the place where the results are actually used.

To summarize the three paths in one sentence: if any one of the three points - the seller, the user, or the place where the output arrives - is inside the EU, the law attaches. Most

Korean companies get caught under the first and third paths.

1.2 Distinguishing Providers, Deployers, Importers, and Distributors

Article 3 of the AI Act divides roles within the supply chain into four categories. Depending on the transaction structure, which of these four a company falls under will vary even for the same entity, and the weight of obligations will change completely.

A provider is defined in Article 3(2) (AI systems) and Article 3(3) (GPAI models). It is an entity that develops, or has developed, an AI system or GPAI model and places it on the EU market under its own name or trademark. Among the four roles, its obligations are the heaviest. A provider of a high-risk AI system must continuously operate a risk management system, verify bias in training datasets, and embed technical documentation and automatic logging capabilities from the design stage. A provider of a GPAI model bears four obligations starting August 2, 2025: drawing up and maintaining technical documentation, providing information to downstream integrators, establishing a policy to comply with EU copyright law, and publishing a summary of training content. These four obligations remained untouched even in the Digital Omnibus amendment. This means they remain alive and in force.

If a GPAI model has an accumulated computation threshold exceeding 10^{25} FLOPs and is presumed to present a systemic risk, an additional layer of obligations attaches: model evaluation aligned with state-of-the-art standards, adversarial testing (red-teaming), evaluation and mitigation of systemic risks, reporting serious incidents to the AI Office, and ensuring an adequate level of cybersecurity. Upon reaching this threshold, the provider must notify the AI Office without delay and at the latest within two weeks. If a domestic generative AI startup is planning to independently develop a large model and supply it overseas, I believe calculating this threshold should be placed in the main body of the business plan rather than in a footnote.

A deployer is defined in Article 3(4). It refers to any natural or legal person, public authority, agency, or other body that uses an AI system under its authority for professional activities. It is also referred to as the buyer or user of an AI system. A deployer of a high-risk AI system must operate the system in accordance with the instructions for use provided by the provider, manage the relevance of input data, and maintain a human oversight framework where a human retains the final judgment. As mentioned earlier, it is rare for Seoul headquarters itself to be a direct deployer; the common structure is that the local EU entity assumes the role of deployer.

An importer is defined in Article 3(6) as the entity that places an AI system produced outside the EU onto the EU market for the first time. A distributor is defined in Article 3(7) as any person in the supply chain, other than the provider or importer, who makes an AI system available on the EU market. Both act as intermediaries in the supply chain. If an EU partner purchases an AI system developed by a Korean company and introduces it to the EU market for the first time, the Korean company is the provider and the partner is the importer. Although responsibility is divided, division does not make it light. Importers and distributors bear the obligation to verify, prior to placing a high-risk AI system on the market, whether the CE marking is affixed, whether the EU declaration of conformity and instructions for use are accompanied, and whether the provider and importer have fulfilled their respective duties. The practical habit of merely writing "The Korean company is the provider, and the EU partner is the importer" in the contract and calling it a day is dangerous. If the partner fails to actually fulfill its importer duties, that non-compliance will ultimately return as an incident that shakes the entire supply relationship on the Korean company's side.

I would add one more provision. Where an AI system undergoes a substantial modification or its intended purpose changes such that it becomes a high-risk system, the entity making that modification - rather than the original provider - assumes obligations as the new provider. If an EU partner modifies a general-purpose image analysis engine made by a Korean company to use it for medical diagnosis, this means provider status can shift to that partner. The outcome here depends on how granularly modification and sublicensing provisions are drafted in the contract.

1.3 When Output Is Used in the EU: Article 2(1)(c)

I have already discussed why this provision is particularly critical for Korean companies. Here, let us delve a bit deeper.

Article 2(1)(c) explicitly states that the law applies to "providers and deployers of AI systems that have their output used in the Union, regardless of whether they are established or located within the Union." Neither physical location nor the provider's nationality serves as the criterion. The benchmark is the final destination of the output or where its impact reaches. I view this design as the product of the borderless nature of AI technology combining with legislative intent to safeguard the health, safety, and fundamental rights of EU citizens.

In practice, the paths through which this provision attaches are far more diverse than expected. Direct sales are not the only issue. Providing APIs, supplying software libraries, and embedding models within mobile apps all serve as conduits for outputs flowing into the EU. If an EU resident uses an AI translation service created by a Korean company via web or app and that translation output is used within the EU, this service alone triggers applicability. The same goes for the diagnostic assistance software of the Busan medical device company. Even though the engine runs in a Korean data center and only the outputs cross over, the moment those outputs actually influence the judgment of German medical personnel, they come under the scope of legal application.

To be candid, the exact boundaries of what the phrase "output used in the Union" covers have not yet fully solidified. In structures where results are delivered through multiple stages, such as AI-as-a-Service, the scope of liability for entities at each stage will likely be evaluated on a case-by-case basis. I do not believe we need to view this uncertainty purely as bad news. At this stage, where precedent and enforcement records remain thin, taking a conservative interpretation and preparing in advance will position companies favorably once cases accumulate.

Let us discuss enforcement a bit further. Compliance with the AI Act is supervised by the AI Office, which holds exclusive jurisdiction over providers of GPAI models. The AI Office bases its compliance assessments on whether harmonised standards have been observed and whether a Code of Practice for AI-generated content transparency has been signed and actually followed. While signing a Code of Practice is a voluntary measure rather than a legal obligation, I view it as closer to an insurance policy than an optional extra. This is

because it has the potential to serve as a mitigating factor when determining fine amounts later on.

Let us also clarify dates and monetary amounts precisely. The AI Act entered into force on August 1, 2024, and full enforcement of regulations on prohibited practices (Article 5) - such as human behavior manipulation and social credit scoring - began on February 2, 2025. Obligations for GPAI model providers (Articles 51 to 55) started applying from August 2, 2025, and this part remains untouched in the amendments discussed later. These are past dates, leaving no room for dispute.

The issue lies with what follows: the full enforcement date for high-risk AI system obligations (Annex III, covering eight areas including recruitment/HR management, essential public services, biometrics, and administration of justice), which was originally scheduled for August 2, 2026. Materials prepared during the initial compilation of NotebookLM primary sources listed this date as fixed, but a fresh check via search reveals that the situation has shifted. The European Parliament adopted the Digital Omnibus on AI amendment on June 16, 2026, the Council granted final approval on June 29, and final signature of the text took place on July 8. Under this amendment, the standalone enforcement date for Annex III high-risk systems was pushed back to December 2, 2027, while the enforcement date for AI systems embedded in products already regulated under EU product safety legislation (Annex I) was deferred to August 2, 2028. As of late July 2026 when this text is being written, this amendment awaits only publication in the Official Journal, entering into force three days post-publication; thus, it is expected to take effect before the end of this month. However, the date August 2, 2026 has not vanished entirely. Most transparency obligations under Article 50 and the AI Office's power to impose fines on GPAI models remain alive and active from that date. While high-risk regulations have been delayed, transparency obligations and GPAI enforcement powers will switch on as scheduled. This is where I believe judgments diverge. One must not look at news headlines alone and conclude that "high-risk regulations have been postponed in their entirety, so we can skip 2026." What is deferred is merely one section under Annex III; the remaining obligations march forward unchanged.

Fine amounts are structured across three tiers. Violations of prohibited practices (Article 5) carry fines of up to 35 million EUR or 7% of total worldwide annual turnover, whichever is higher. Most other violations - including non-compliance with obligations by providers, deployers, importers, distributors, and GPAI model providers - carry fines of up to 15 million EUR or 3% of total worldwide annual turnover, whichever is higher. Supplying

incorrect, incomplete, or misleading information to supervisory authorities carries fines of up to 7.5 million EUR or 1% of total worldwide annual turnover, whichever is higher. One frequently encounters materials listing this last amount as 750,000 EUR, but based on the exact original text of Article 99, it is 7.5 million EUR. Though it is a single digit's difference, this is a figure reported to the board of directors or executive management, so it must be pinpoint accurate. Small and medium-sized enterprises (SMEs) and startups receive preferential treatment across all three tiers, with the lower of the fixed amount or percentage of turnover applying.

1.4 EU Authorised Representative: Articles 22 and 54

To place high-risk AI systems or GPAI models on the EU market, providers established outside the EU must designate an Authorised Representative within the EU by written mandate. High-risk AI systems draw their legal basis from Article 22, and GPAI models from Article 54. Both provisions share the requirement of designation by written mandate, and both serve as mechanisms to forcibly create a legal point of contact for providers with no physical presence.

The duties of an authorized representative go beyond storing documents. They must verify whether the provider's technical documentation satisfies AI Act requirements and retain those records for ten years. They must respond when EU market surveillance authorities or the AI Office request information, and they bear cooperation duties including risk mitigation measures. If necessary, they also handle EU registration procedures on the provider's behalf. Here lies a point that practitioners often miss: if non-compliance with the AI Act by the provider is identified, the authorized representative has an obligation to terminate the mandate on its own initiative and immediately notify the relevant authorities. The authorized representative is not merely the provider's spokesperson; it is also a window for monitoring the provider. Passing over the weight of this clause lightly when drafting a representative agreement will lead to trouble later.

There are exceptions. Where a GPAI model is released under an open-source license, the obligation to designate a representative itself is exempted, provided the model is not classified as presenting systemic risk. However, obligations regarding copyright compliance policy (Article 53(1)(c)) and publishing a summary of training content (Article 53(1)(d)) are not exempted. Releasing software as open source does not mean one is completely free from copyright issues.

To summarize the preparation timeline once more: GPAI model provider obligations have been applying since August 2, 2025. For models placed on the market prior to that date, a transitional provision - a so-called grandfathering clause - requires compliance measures to be completed by August 2, 2027, and this applies throughout the model's entire lifecycle. If a Seoul-headquartered company plans to expand into the EU, I believe these dates must be placed in contract negotiation schedules rather than in business roadmap footnotes. I have seen many instances where rushing to find a representative late in the process led to a disadvantageous position in contract terms. Securing time is

synonymous with securing bargaining power.

One more thing: the relationship with domestic regulation must also be examined. South Korea brought into force the Framework Act on Artificial Intelligence (AI Framework Act) on January 22, 2026, and starting July 21 of the same year, amended enforcement decrees specifying delegated matters came into application, giving rise to new mechanisms such as the AI Product and Service Verification System. Administrative fines under the domestic AI Framework Act are capped at around 30 million KRW - a completely different weight class compared to the EU AI Act, which levies up to 7% of global turnover. Still, domestic regulations should not be taken lightly. A structure has already been established where a single AI system must simultaneously align with dual regulations: the AI Framework Act domestically, and the EU AI Act for outputs reaching the EU. It should also be noted that the risk classification criteria and obligation items of the two laws do not overlap completely. One cannot simply copy and paste a domestic compliance checklist over to the EU side.

1.5 Self-Diagnostic Checklist

To enable immediate verification in practical work, the contents read so far are summarized into five questions. If you answer "Yes" to even one, please return to the corresponding section to verify your obligations.

- 1) Does our company develop an AI system or GPAI model under our own name or trademark and sell it directly to EU clients, or supply it in the form of APIs or licenses? If so, you are a provider under Article 3(2) and (3). Article 2(1)(a) applies, and you must review the obligation to designate an EU representative under Article 22 or Article 54.
- 2) Has your company established a subsidiary or branch in the EU, and does that entity use an AI system for business purposes? If so, that EU entity is a deployer under Article 3(4). If it is a high-risk system, human oversight frameworks and input data management obligations must be established at that entity's level.
- 3) Is an EU partner purchasing our company's AI system and placing it on the EU market for the first time? If so, the partner is an importer under Article 3(6), and our company is the provider. Please recheck whether your contract includes clauses verifying that the importer side fulfills CE marking, declaration of conformity, and instructions for use obligations.
- 4) Even if our company has no physical presence in the EU, do the results or services generated by our system reach users inside the EU via APIs, libraries, or mobile apps and get actively used? If so, you must review provider obligations in their entirety pursuant to Article 2(1)(c). The case of the Busan medical device company is precisely this situation.
- 5) Does our system touch upon the domains of recruitment/HR management, essential public services, biometrics, or administration of justice under Annex III, or is it embedded in products already regulated by EU product safety legislation (medical devices, machinery, toys, aircraft, etc.)? For the former, the enforcement date was deferred to December 2, 2027; for the latter, August 2, 2028. However, it must not be overlooked that transparency obligations under Article 50 and the AI Office's power to impose fines on GPAI models, both taking effect on August 2, 2026, apply as scheduled regardless of this postponement.

If the legal counsel of the Busan medical device company sits down before these five questions again, they will likely raise their hand for questions 1 and 4 simultaneously. A

company that believed it operated strictly domestically turns out to be shouldering two sets of obligations at the same time on the pages of EU law provisions. When this fact is realized - whether after a hospital contract is already signed or before putting ink to paper - that difference determines the weight of the stories that will follow in the chapters ahead.

Chapter 2. Framework of the Law

It was last month when a CEO of a healthcare startup based in Pangyo came to see me. He said he wanted to deploy a service across a few European hospitals that reads patients' facial expressions and vocal tones in hospital emergency rooms to assess their risk levels. He brought only one question: Can we just use this in Europe? Before opening the statutory provisions, I asked him back: Is that service closer to surveillance or to assistance? The CEO paused for a moment. That silence is the starting point of this chapter. The European Union Artificial Intelligence Act (Regulation (EU) 2024/1689, hereinafter the AI Act) established its entire structure upon that very single question: To what degree of risk does this system subject human beings?

2.1 The Design Philosophy Called the Risk-Based Approach

The AI Act entered into force on August 1, 2024. It contains 113 articles alone, and together with its annexes, it runs as long as a thick book. Yet, the framework underpinning this law is surprisingly straightforward: the Risk-Based Approach, where higher risks entail greater obligations, and no risk entails no obligations. To put it in American legal terms, rather than establishing a single principle and stacking exceptions on top of it, it is a method of first measuring the system against a ruler of risk level and then allocating regulatory intensity according to those graduation marks.

The meaning of this design in practice is clear. Even between two AI products made by the same company, one can be sold without any reporting obligations, while the other can only be sold after undergoing a conformity assessment and obtaining CE marking. The criteria that divide them are not the sophistication of the technology, but the magnitude of the impact that technology exerts on human health, safety, and fundamental rights. This is why the emotion recognition service of the Pangyo startup's CEO receives entirely different legal treatment depending on whether it functions as a medical diagnostic aid or a worker surveillance tool.

The AI Act broadly categorizes systems into four tiers: prohibited practices that must never be used, high-risk systems subject to strict ex-ante obligations, transparency obligation targets requiring specific disclosures, and minimal-risk systems that can essentially be used freely. Rather than drawing these four tiers as a pyramid, I prefer to illustrate them as a corridor with four thresholds. It is easier to understand if you view it as a structure where a system walks down this corridor and is checked for whether it passes through the first threshold and second threshold in sequence.

2.2 The Four Tiers: Prohibited, High-Risk, Transparency, and Minimal Risk

The first threshold is a door that permits no passage at all. The prohibited AI practices under Article 5 of the AI Act fall here. Representative examples include systems that manipulate behavior without a person realizing it to cause harm; systems that impair judgment by targeting vulnerabilities such as age or disability; social scoring where governments or companies score individuals' social behavior to inflict detriments; and real-time remote biometric identification systems targeting unspecified individuals in publicly accessible spaces. This provision entered into full force on February 2, 2025, six months after the law took effect. The emotion recognition service that the Pangyo CEO intends to build must cross this threshold first. If it is intended for attendance management in workplaces or reading student emotions in educational settings to reflect them in evaluations, there is a high likelihood of triggering this provision unless it falls under specific exceptions (such as safety or medical purposes). When incorporating emotion recognition features, I always advise narrowing the purpose to clinical assistance and maintaining written documentation as evidence.

Crossing the second threshold leads to High-Risk AI systems. There are two entry routes here. One consists of products subject to existing EU product safety legislation (Annex I) pursuant to Article 6(1) of the AI Act - such as safety components integrated into AI-based medical devices or automotive parts. The other comprises the eight use areas listed in Article 6(2) and Annex III of the AI Act, such as recruitment and HR management, access to essential public services, biometrics, law enforcement, migration, asylum, and border control, and administration of justice. There is also an exception clause (Article 6(3)) allowing a system listed in Annex III to be excluded from high-risk classification if it does not actually pose a significant risk to human health, safety, or fundamental rights. The Commission guidelines on how to interpret this exception were scheduled to be released by February 2, 2026, but at the time of writing, I have not yet confirmed a final finalized text. In practice, it is safer to act conservatively until these guidelines are issued - that is, assuming it is high-risk if its name touches Annex III.

Providers classified as high-risk shoulder heavy burdens pursuant to Article 16. These include: a risk management system to identify and mitigate risks throughout the entire lifecycle of the system; technical documentation detailing the development process and evaluation results; record-keeping including automatically generated logs; transparency to provide sufficient information to deployers; human oversight mechanisms allowing humans to ultimately override results; accuracy, robustness, and cybersecurity; a quality

management system; conformity assessment and CE marking prior to market placement; and post-market monitoring that continues even after launch. Deployers cannot idle either. Pursuant to Article 26, they must operate the system according to the instructions provided by the provider, manage input data, and actually maintain human oversight mechanisms.

The enforcement timing of this second threshold is the point that has fluctuated most significantly at the time of writing this book. Originally, obligations regarding Annex III high-risk systems were set to apply fully from August 2, 2026, 24 months after the law took effect. However, on June 29, 2026, the Council of the European Union reached a final agreement on the so-called Digital Omnibus package, which was signed on July 8, 2026, and currently awaits publication in the Official Journal. According to this package, obligations for Annex III high-risk systems are deferred by 16 months to December 2, 2027, while AI systems acting as safety components under existing product safety legislation in Annex I are delayed to August 2, 2028. The surface reason is that member state notifying authorities and harmonised standards are not yet ready. The NotebookLM materials I consulted were still organized under the assumption of full application on August 2, 2026, marking a discrepancy between documentation and reality. For someone designing a product today like the Pangyo CEO, I advise interpreting this delay not as a signal for relief, but as a signal that the preparation window has merely expanded. Regulators granted more time not because they lowered their requirements, but because they themselves were unprepared.

The third threshold is Transparency obligations. This obligation under Article 50 applies even to non-high-risk systems. If you build an AI system that interacts directly with natural persons like a chatbot, you must inform them of the fact that they are interacting with AI, unless it is obvious from the circumstances. Deepfakes and AI-generated or manipulated images, audio, or text must be marked so that their artificial origin is disclosed. Deployers using emotion recognition systems or systems classifying people based on biometric categorization must also inform the affected individuals. This obligation applies broadly regardless of high-risk status, whether it is a chatbot, an image generator, a voice synthesizer, or a writing assistant. An interesting point is that even under the Digital Omnibus package, the core framework of Article 50 will still take effect as scheduled on August 2, 2026. Only the deadline for technical marking means of AI-generated content, such as watermarking, was postponed by three months to December 2, 2026. On May 8, 2026, the Commission issued draft guidelines regarding who must fulfill chatbot disclosure and deepfake labeling obligations and how, alongside the

preparation of the Code of Practice on Transparency of AI-Generated Content. If the Pangyo CEO's service displays emotion scores on hospital screens, you must assume that Article 50 disclosure obligations will almost certainly apply, irrespective of any high-risk determination.

The fourth tier is Minimal Risk systems. The vast majority of AI systems that are neither prohibited, high-risk, nor subject to transparency obligations - such as spam filters or recommendation algorithms - fall here. There are virtually no mandatory legal obligations, and adherence to voluntary codes of conduct is merely encouraged. However, I do not advise letting your guard down just because it is minimal risk. EU corporate clients increasingly inquire about voluntary compliance during due diligence, and there is always a lingering possibility that the service scope might later expand into Annex III domains.

2.3 Relationship with GDPR, the Product Liability Directive, and the Data Act

The AI Act is not a standalone statute. One must first examine its relationship with the General Data Protection Regulation (GDPR). Article 2(7) of the AI Act explicitly mandates that existing EU law, including the GDPR, applies unimpaired regarding the processing of personal data. This means the AI Act does not displace the GDPR but rather layers on top of it. In practice, the point of overlap is impact assessments. If a high-risk AI system processes personal data, both a Data Protection Impact Assessment (DPIA) under Article 35 of the GDPR and a Fundamental Rights Impact Assessment under Article 27(4) of the AI Act must be conducted. However, an existing DPIA can complement the Fundamental Rights Impact Assessment to reduce duplication. Rather than creating two separate assessments, I find that companies placing personal data risk items and fundamental rights risk items side by side within a single document convey a far more favorable impression during due diligence.

Attention should also be paid to the relationship with the Product Liability Directive and existing product safety legislation. If AI is embedded into medical devices, machinery, toys, or aircraft components, both existing safety certifications applicable to the product and high-risk obligations under the AI Act apply simultaneously. Article 74(4) of the AI Act specifies that where existing harmonised legislation listed in Section A of Annex I already provides procedures ensuring an equivalent level of protection, those sector-specific procedures shall be followed instead of certain AI Act procedures (Articles 79 through 83). This coordinates regulations to prevent double layering in fields with already stringent rules, such as medical device certification frameworks. Nevertheless, having already obtained a CE mark does not exempt a company from the risk management system or technical documentation obligations under the AI Act. Existing certifications evaluate the overall safety of the product, whereas the AI Act separately assesses the risk of the AI component embedded within it.

The Data Act became applicable across most of its provisions starting September 12, 2025. The core of this law is expanding user and third-party access rights to data generated by connected products. When combined with the AI Act's obligation to publish a summary of training data (Article 53(1)(d)), it functions to exert pressure toward greater transparency regarding data sources. The Digital Services Act (DSA) and Cyber Resilience Act are neighboring statutes that must also be reviewed alongside the AI Act. If an online platform utilizes AI, it must attend to both the DSA and the AI Act; if it is a digital product

containing AI, it must identify overlapping cybersecurity requirements across both laws.

When explaining this relationship to Korean companies, I always use the same analogy: If the GDPR is a law governing the raw ingredient - personal data - the AI Act is a law governing the dish prepared with that ingredient, namely the safety and performance of the system itself. Thoroughly washing your ingredients does not guarantee that the dish is safe, nor does preparing a fine dish legitimize the origin of your ingredients. Companies that review the two laws in isolation end up creating the same document twice.

2.4 Sanction Levels: Fines Calculated Based on Global Turnover

Article 99 of the AI Act requires Member States to lay down effective, proportionate, and dissuasive penalties. The resulting fine structure consists of three tiers. The heaviest tier applies to violations of Article 5 prohibited practices and breaches of data-related requirements. The ceiling is 35 million EUR or 7% of total worldwide annual turnover for the preceding financial year, whichever is higher. If the Pangyo CEO's emotion recognition service is sold in the EU while infringing upon Article 5 prohibited domains, up to 7% of the entire company's global turnover could be levied. I always emphasize twice in consultations that this is calculated on global turnover, not EU turnover.

The middle tier applies to breaches of other obligations under the AI Act - most articles excluding Article 5. The cap is 15 million EUR or 3% of total worldwide annual turnover, whichever is higher. Violations of obligations by general-purpose AI (GPAI) model providers also fall into this tier, with Article 101 of the AI Act laying down detailed fines for GPAI model providers separately. Basic obligations for GPAI providers have applied since August 2, 2025; if cumulative compute used for training exceeds 10^{25} FLOPs, a systemic risk is presumed, triggering additional obligations under Article 55 such as model evaluation, adversarial testing, ensuring cybersecurity, and reporting serious incidents. Exceeding this threshold requires prompt notification to the AI Office within 2 weeks at the latest, and failure to do so itself becomes subject to middle-tier fines.

The lowest tier covers providing inaccurate, incomplete, or misleading information to notified bodies or competent authorities. The ceiling is 7.5 million EUR or 1% of total worldwide annual turnover, whichever is higher. Small and Medium-sized Enterprises (SMEs) receive an exception across all three tiers: the lower amount applies instead of the higher amount. When EU institutions or bodies commit a violation, a different standard applies. The European Data Protection Supervisor (EDPS) imposes the fine, capped significantly lower at up to 1.5 million EUR for Article 5 violations and up to 750,000 EUR for other violations.

The point at which the AI Office actually obtains enforcement authority to impose these fines is August 2, 2026. However, as noted in 2.2 above, now that Annex III high-risk obligations themselves have been deferred to December 2, 2027, it is correct to view the actual activation point for high-risk related fines as naturally delayed accordingly. Conversely, GPAI obligations and Article 50 transparency obligations proceed on their original schedule, meaning fine risks in these two areas are either already active or will

take effect immediately from August 2026. Because dates are scattered across multiple tracks, I have summarized them in a table below.

Factors examined by authorities when determining fine amounts are also enumerated in Article 99(7). These include: the nature, gravity, and duration of the infringement; the number of affected individuals and degree of damage; intentional or negligent character; size, turnover, and market share of the infringing company; financial benefits gained or losses avoided; degree of cooperation with authorities; level of technical and organizational measures implemented; voluntary disclosure; and whether fines were previously imposed for the same violation in other Member States. Here, there is a practically useful card to play. Article 101(1) specifies that commitments following a Code of Practice evaluated as adequate by the AI Office and AI Board can serve as a mitigating factor when calculating fines. I explain this not as a discount coupon for fines, but as a statement declaring upfront to authorities that the company acted in good faith. Failing to follow a Code of Practice does not immediately put you at a disadvantage, but having a record of compliance gives you an extra card at the negotiation table.

Extraterritorial applicability must also be addressed. Pursuant to Article 2(1)(a) and (c) of the AI Act, even providers located outside the EU are subject to this law if the output produced by their system is used within the EU. Even if the Pangyo CEO's company operates offices solely in Seoul, the moment its service integrates into an EU hospital, it falls into the same sanction framework as EU-based competitors. And that fine is calculated based on the company's global turnover, not revenue earned in the EU. Even for a company where EU sales account for less than 1% of total revenue, 7% of its entire global revenue could be at stake if the headquarters generates substantial turnover. EU extra-territorial providers must designate an Authorised Representative in writing within the EU as a contact point for authorities; I have repeatedly observed that companies delaying this appointment suffer delayed responses when an incident actually occurs.

Reference Materials

Enforcement Schedule Table

August 1, 2024: Entry into force of the AI Act. February 2, 2025: Full application of Article 5 prohibited AI practices. August 2, 2025: Application of basic obligations for GPAI model providers (Article 53). August 2, 2026: Application of Article 50 transparency obligations, activation of AI Office fine imposition powers, original scheduled date for full application of Annex III high-risk obligations (deferred to December 2, 2027 under Digital Omnibus).

December 2, 2026: Expiration of grace period for implementing technical marking means under Article 50, such as watermarking. December 2, 2027: Target date for full application of Annex III high-risk obligations reflecting Digital Omnibus adjustments. August 2, 2028: Target date for obligations applicable to AI components subject to Annex I product safety legislation.

The 2026 and 2027 dates in this table are provisional estimates reflecting the status where the Digital Omnibus package was signed on July 8, 2026, but publication in the Official Journal has not been confirmed as of the writing of this book. Readers are advised to recheck Official Journal publication dates and final article numbers immediately following publication.

Checklist Items

Have we verified whether our system triggers the Article 5 prohibited list? If it falls under one of the 8 areas in Annex III, have we documented a review of whether it meets the Article 6(3) exception criteria? If chatbot or generative features are present, have we established compliance measures for Article 50 disclosure obligations (disclosure text, watermarks)? If training a GPAI model, have we calculated whether cumulative compute approaches 10^{25} FLOPs? Have we integrated the DPIA under GDPR and the Fundamental Rights Impact Assessment under the AI Act into a single documentation system? If operating as a non-EU provider, have we designated an Authorised Representative in writing? Are we tracking participation in Codes of Practice and preserving relevant records?

As the Pangyo CEO left after concluding our consultation that day, he asked: "So by when do we need to be ready?" I couldn't pinpoint a single date for him - because the law itself was in the middle of rewriting those dates multiple times.

Chapter 3. The 2026 Timeline, Redrawn Maps

A compliance officer at a Pangyo startup recently reopened the compliance schedule created last year. Half of the cells were highlighted in red. That meant the dates were wrong. Around this time last year, the company's plan was structured around cramming all preparation into a single target date: August 2, 2026. Believing that high-risk AI system regulations would become fully applicable on that day, both the development team and the legal team raced forward with their eyes fixed solely on that date. However, a significant portion of those target dates has been pushed back. To be exact, among the various obligations clustered around August 2, some remained intact while others were deferred by more than sixteen months.

The objective of this chapter is to place that realigned map directly into your hands. Missing these dates compromises the remaining twenty-two chapters of this book. Allocating resources to obligations not applicable until 2027 while neglecting urgent tasks required by August 2 means getting the order completely backward. Conversely, assuming an obligation has been postponed and letting your guard down, only to be caught by one that remains intact, leads to an even greater disaster. In writing this chapter, I cross-checked each date through two paths: official documents issued by the European Commission and the Council of the European Union, and the latest news reports available as of today, the time of writing. Where the two diverge, I will explicitly note the discrepancy as we proceed.

Let us first establish the big picture. The European Union Artificial Intelligence Act (Regulation (EU) 2024/1689) is not a law that takes effect entirely on a single day. Rather, it is structured so that obligations become applicable sequentially with staggered timeframes after entry into force. Furthermore, an amendment known as the Digital Omnibus, initiated in November 2025, readjusted some of these staggered timelines. This amendment completed its legislative procedure with final approval from the Council on June 29, 2026, and as of July 24, 2026, the time of writing, it is awaiting publication in the Official Journal. The timeline presented below is the map of this present moment - overlaying this amendment onto the original statutory text.

3.1 Entry into Force on August 1, 2024, and the Design of Phased Application

The AI Act entered into force on August 1, 2024. Entry into force and applicability are distinct concepts. Entry into force is when the law comes into legal existence, whereas applicability is when individual obligations actually begin to take effect. On the day this law entered into force, not all obligations became applicable immediately. Article 113 specifically set out different dates of application for different provisions.

Why was it structured this way? Because the risk level of the subjects being regulated varies. Practices posing a high risk of harm to humans must be banned swiftly, whereas preparing technical documentation and undergoing conformity assessments require time for companies to build or adapt their systems. Therefore, the timeline was staggered to activate prohibitions early and delay heavier preparatory obligations. Understanding this underlying design principle clarifies why the subsequent dates are arranged in that particular order: addressing urgent risks first, followed later by obligations requiring lengthy preparation.

3.2 February 2, 2025: Prohibited Practices and AI Literacy Obligations

The first to become applicable were the prohibitions. On February 2, 2025, regulations on the prohibited AI practices enumerated in Article 5 became fully applicable. From this date onward, systems employing subliminal manipulation or exploiting vulnerabilities, social scoring by public or private entities, predictive policing based solely on profiling, the untargeted scraping of facial images from the internet or CCTV footage to create facial recognition databases, emotion recognition in workplaces and educational institutions, biometric categorization inferring race, political opinions, trade union membership, religion, or sexual orientation, and real-time remote biometric identification for law enforcement purposes (with limited exceptions) could no longer be placed on the market, put into service, or used in the EU. What is banned and why will be examined article by article in Chapter 15.

On the same date, the AI literacy obligation under Article 4 also became applicable. This duty requires providers and deployers to ensure that their staff and persons dealing with AI systems on their behalf possess sufficient knowledge and skills to handle AI systems responsibly and assess associated risks. Many companies miss the fact that this obligation is already applicable, as it is a provision that can easily pass unnoticed. We will revisit what must actually be put in place in Chapter 21.

Penalties for violating prohibitions are the highest under the entire AI Act: up to 35 million EUR or 7 percent of total worldwide annual turnover, whichever is higher. For a company headquartered in Seoul, cross-checking any AI system used within or placed on the EU market against this list is the very first task that cannot be delayed in this book.

3.3 August 2, 2025: Obligations for Providers of General-Purpose AI Models

On August 2, 2025, exactly one year after entry into force, obligations for providers of general-purpose AI models (GPAI) became applicable. Defined under Article 3(63), a general-purpose AI model refers to a model trained with a large amount of data using self-supervision at scale, displaying significant generality and capable of competently performing a wide range of distinct tasks. Language models are a typical example, though models generating images, video, or audio can also fall under this definition.

The core obligations imposed by Article 53 are as follows: draw up technical documentation required by Annex XI and keep it up to date throughout the model's lifecycle; provide information specified in Annex XII to downstream providers who intend to integrate the model into their own systems; put in place a policy to respect reservations of rights pursuant to Article 4(3) of the European Copyright Directive (2019/790); and draw up and make publicly available a summary of the content used for training. In addition, Article 55 imposes further requirements on providers of models with systemic risk, including model evaluation, red-teaming, risk assessment and mitigation, reporting serious incidents, and ensuring adequate cybersecurity. The threshold for determining systemic risk is a cumulative training compute of 10^{25} FLOPs. If this threshold is met or expected to be met, the AI Office must be notified within two weeks at the latest. Currently, only about twelve to fifteen models worldwide are known to cross this threshold. The details of these obligations will be unpacked one by one in Chapters 7 through 10 of Part 3.

On the same date, governance provisions also began to apply, effectively activating the oversight structure of the AI Office and national competent authorities. However, there is a catch. The date an obligation becomes applicable and the date fines can be imposed for non-compliance are separate. The AI Office established a one-year period of "application without enforcement" from August 2, 2025, to August 2, 2026. For one year, the obligations are fully alive, yet no fines are levied. This period must not be interpreted as an absence of compliance duties. Violations accumulate as facts, and liability will begin to be enforced starting from the date discussed in the next section.

3.4 August 2, 2026: Full Applicability of Article 50 Transparency Obligations and Enforcement Powers over General-Purpose AI

Now we reach the date right around the corner. On August 2, 2026, two distinct events occur simultaneously.

The first is the full applicability of Article 50 transparency obligations. AI systems intended to interact directly with natural persons must clearly inform individuals that they are interacting with AI. Systems generating synthetic audio, image, video, or text content must mark outputs in a machine-readable format to render their artificial origin detectable. Deployers using emotion recognition or biometric categorization systems must inform affected individuals, and deepfakes as well as AI-generated text must be disclosed as manipulated or generated. Notice must be provided at the latest at the time of first interaction or exposure, in a clear, distinguishable, and accessible manner. Discussions targeted at these obligations continue in Chapters 4 through 6 of Part 2.

The second is the activation of the AI Office's enforcement powers regarding general-purpose AI models. While obligations for general-purpose model providers became applicable on August 2, 2025, the authority to actually impose fines for non-compliance takes effect one year later, on this date. The period of "application without enforcement" mentioned earlier ends here. Fines are capped at up to 3 percent of total worldwide annual turnover or 15 million EUR, whichever is higher. For companies dealing with general-purpose models, I interpret this as the day leniency expires.

One caveat to keep in mind: among Article 50 labeling obligations, there is a single exception applicable exclusively to generative systems. Generative AI systems already placed on the market before August 2, 2026, are granted a four-month grace period, allowing them until December 2, 2026, to comply with labeling obligations. Newly released systems do not receive this grace period; they are subject to enforcement immediately starting August 2. We will reinforce this distinction in Section 3.8.

3.5 The Digital Omnibus Amendment: What Changed with the Council's Final Approval on June 29, 2026

Let us now examine the amendment that forced us to redraw this map. On November 19, 2025, the Commission introduced a legislative package titled the Digital Omnibus on AI. The intention was to eliminate delays and uncertainties in the implementation of the AI Act and to refine the rules clearly. Its core proposal was to postpone the date of application for high-risk AI obligations from August 2, 2026.

The progression toward codifying this proposal into law unfolded as follows: a provisional trilogue agreement was reached among the Council, the European Parliament, and the Commission on May 7, 2026; the European Parliament plenary approved it on June 16; and the Council granted final approval on June 29. The Council of the European Union stated that this approval was a measure to streamline regulatory procedures and ensure legal clarity. According to the latest reports verified as of today, July 24, 2026, the final act was signed on July 8, 2026, and is currently awaiting publication in the Official Journal. Observations suggest that publication must occur by July 30 at the latest to avoid a gap with provisions scheduled for application on August 2. This aligns with calculations based on the standard rule that entry into force occurs on the third day following publication. However, the exact publication date remains unconfirmed as of this writing. I advise readers to verify whether publication in the Official Journal has taken place when reading this book.

What this amendment postponed was the obligations surrounding high-risk AI systems. Yet, it did not merely delay timelines. It introduced a brand-new prohibition: a ban on AI systems that generate or manipulate realistic depictions of intimate body parts or sexual acts of real persons without free, specific, and informed consent. This provision targets so-called "nudifier" apps while also addressing risks associated with child sexual abuse material (CSAM). Incorporated into Article 5, this prohibition enters into force on December 2, 2026. While a safe harbor mechanism exists for providers who implement technical safeguards reliably preventing such output, the scope applies broadly - not only to providers who designed systems for that purpose, but also to providers of systems where such output is reasonably foreseeable. Penalties for violations match those of prohibited practices: up to 35 million EUR or 7 percent of global annual turnover. Failing to recognize that an amendment offering delays simultaneously embedded a heavy new prohibition would lead to the mistake of viewing this amendment solely as deregulation.

3.6 Postponement of Annex III Standalone High-Risk AI to December 2, 2027

High-risk AI systems fall broadly into two main tracks. One comprises standalone systems listed in Annex III. This refers to self-contained AI systems used in eight designated areas, including recruitment and HR management, essential public services, biometrics, migration and border control, and the administration of justice. Article 6(2) serves as the legal basis for this classification.

Obligations under Article 16 (for providers) and Article 26 (for deployers) regarding these standalone high-risk systems were originally scheduled to become fully applicable on August 2, 2026, twenty-four months after entry into force. The Digital Omnibus amendment deferred this date to December 2, 2027. This represents a delay of sixteen months, moving from August to December of the following year. While some media reports cite seventeen months, I personally counted the interval between August 2026 and December 2027; sixteen months is correct.

Behind this postponement lies unpreparedness on the implementation side. The specific technical specifications companies must comply with - namely, harmonized standards being developed by CEN-CENELEC JTC 21 - have not yet been fully issued, and the establishment of competent supervising authorities across Member States has also faced delays. Enforcing obligations without established standards leaves companies exposed to sanctions without knowing what to comply with. It is reasonable to interpret the postponement as a measure to mitigate that risk. However, it must be remembered that the authority to impose fines itself remains active starting August 2, 2026. Penalties for violating high-risk obligations reach up to 7.5 million EUR or 1 percent of global annual turnover, whichever is higher, though this upper cap will actually be levied when violations accumulate after December 2, 2027. How to utilize this gap as a preparation period will be addressed separately in Chapter 14.

3.7 Postponement of Annex I Embedded High-Risk AI to August 2, 2028

The second track concerns high-risk AI embedded in existing products falling under Annex I. This applies when AI serves as a safety component within products already regulated by European product safety legislation - such as medical devices, machinery, toys, and aircraft - where the product itself requires a third-party conformity assessment. Article 6(1) serves as the basis for this classification.

The completion deadline for regulatory procedures regarding embedded high-risk AI was originally set for August 2, 2027, thirty-six months after entry into force. The Digital Omnibus deferred this by an additional twelve months to August 2, 2028. This comes eight months later than the standalone track. The difference between the two dates stems not from political consideration, but from practical circumstances tied to existing product certification cycles. Product categories that already operate on multi-year certification cycles, such as medical devices or automotive components, require more time to integrate AI Act requirements into their pre-existing procedures. Article 8(2) allows for this integration by permitting testing and reporting required by the AI Act to be built atop documentation and procedures required under existing safety laws, aiming to minimize duplication.

Medical device companies in Busan or domestic automotive component manufacturers fall into this category. Depending on whether a product falls under Annex I or Annex III, the compliance deadline diverges by eight months, and the regulatory pathway for obtaining CE marking differs altogether. We will analyze this distinction case by case in Chapter 11.

3.8 The Line Dividing What Was Postponed and What Was Not

This is where real-world mistakes happen. Upon hearing that several dates were pushed back by the Digital Omnibus, many mistakenly assume the entire AI Act has been delayed. There is, however, a clear line between what was postponed and what was not.

What was postponed is the obligations for high-risk AI systems: Annex III standalone systems moved to December 2, 2027, and Annex I embedded systems moved to August 2, 2028.

What was not postponed falls into two main groups. First, the vast majority of Article 50 transparency obligations remain set for August 2, 2026 - including chatbot notifications, synthetic content labeling, and deepfake disclosures. The sole exception is a four-month grace period extending to December 2, 2026, applicable exclusively to generative systems already on the market before August 2. Second, the newly introduced prohibition on generating NCII (non-consensual intimate imagery) and CSAM is not delayed; it is a brand-new obligation activated by this amendment, taking effect on December 2, 2026.

Therefore, for Korean companies exporting AI systems or models to Europe, the urgent target requiring immediate action by August 2 is not high-risk regulation, but Article 50. High-risk obligations have bought time; transparency obligations have none. This is where strategic judgment diverges. I use this line as the primary benchmark for prioritizing preparation. I have personally observed several companies that misread this line, pouring resources into high-risk compliance while missing imminent transparency labeling requirements due right in front of them on August 2.

Implementation Schedule at a Glance

August 1, 2024. Entry into force of the AI Act.

February 2, 2025. Start of application for Article 5 prohibited practices. Start of application for Article 4 AI literacy obligations.

August 2, 2025. Start of application for general-purpose AI model provider obligations (Articles 53 and 55). Start of application for governance provisions. Beginning of the "application without enforcement" period (through August 2, 2026).

June 29, 2026. Final approval by the Council of the Digital Omnibus amendment. Awaiting publication in the Official Journal (as of writing deadline).

August 2, 2026. Full application of Article 50 transparency obligations. Activation of the AI Office's authority to impose fines for general-purpose AI model violations.

December 2, 2026. Entry into force of the prohibition on NCII and CSAM generation (Article 5 amendment). Expiration of the grace period for labeling obligations of existing generative systems placed on the market before August 2.

February 2, 2027. Deadline for equipping interoperable solutions for watermark detection (Code of Practice for marking synthetic content).

August 2, 2027. Compliance deadline for existing general-purpose AI models placed on the market before August 2, 2025.

December 2, 2027. Start of application for Annex III standalone high-risk AI obligations (Articles 16 and 26).

August 2, 2028. Completion deadline for regulatory procedures regarding Annex I embedded high-risk AI.

This table is the map of this present moment. The subject matter of this book remains in motion. Even after finishing this manuscript, detailed dates may be adjusted again as official publication dates are finalized, harmonized standards are released, and new guidelines are announced. I continue to update these dates separately, and I encourage readers to verify this table once more before launching actual compliance efforts. The map has been drawn, but the ground is still shifting underfoot.

Part II

What You Must Comply With Now: Transparency

Chapter 4. Article 50 Splits into Four Directions

A game company located in Pangyo attached a chatbot to its customer service desk a few years ago. Initially, it was a system that provided predefined answers when a few buttons were pressed, but last year they upgraded it to an interactive language model. Now, no matter what users ask, it responds like a human. Its sentences are natural, and it even takes jokes. When a user in Europe asks about a payment error at dawn, it is not easy to distinguish whether the entity beyond the screen is a human or a machine.

This company had assumed that their chatbot had nothing to do with regulation. After all, it is not a high-risk AI system. It does not make hiring decisions, nor does it evaluate creditworthiness. It merely handles inquiries from game users. However, starting August 2, 2026, this chatbot becomes subject to Article 50 of the European Union Artificial Intelligence Act (Regulation (EU) 2024/1689, hereinafter referred to as the AI Act). Since today is July 24, 2026, barely ten days remain.

What Article 50 targets is not dangerous artificial intelligence, but artificial intelligence capable of deceiving humans. This distinction is the key to interpreting Article 50. A high-risk classification applies because the system significantly impacts human rights or safety. The transparency obligations under Article 50 are different. They apply because the system can lead people to mistake a machine for a human, or a fake for the real thing. Medical device diagnostic software is high-risk, but it does not deceive people. A game customer support chatbot is not high-risk, but it has the potential to deceive people. Article 50 captures the latter.

Therefore, this chapter represents the first obligation that many Korean companies will actually encounter. A significant portion of the high-risk regulations was deferred under the so-called Digital Omnibus adjustment agreed upon last May. The enforcement of high-risk classifications under Annex III has been postponed to 2027, with some pushed back to August 2028. Unless one is a provider of general-purpose AI models, there is still breathing room on that front. However, Article 50 was not deferred even in this Omnibus adjustment. It applies as scheduled on August 2, 2026. There is one exception. For the machine-readable marking obligation under Article 50(2) - that is, the watermark-related duty - systems already placed on the market prior to that date are granted a four-month grace period until December 2, 2026. Everything else applies on August 2. If a company operates a chatbot or generates images, videos, and text using AI to export to Europe, Article 50 is the very first wall it will hit.

Article 50 itself diverges into four directions from within. It divides based on who bears the obligation and what must be disclosed. A single company may fall under just one of the

four, or under two or three simultaneously. If a company does not know where it stands, it will end up expending effort in the wrong direction.

4.1 Chatbots and Conversational Agents: How to Disclose That It Is Not a Human

Article 50(1) imposes obligations on providers - those who develop and place on the market AI systems intended to interact directly with natural persons. Chatbots, virtual assistants, and automated response calls fall into this category. The text of the provision is brief: inform the natural person that they are interacting with an AI system.

While the wording is brief, practical implementation is not. The issue lies in the exception clause. Article 50(1) specifies that disclosure is not required if it is obvious from the point of view of a reasonably well-informed, observant, and circumspect natural person that they are interacting with AI. The word "obvious" is the trap.

Companies can easily fall into the mindset of thinking: "Our chatbot displays a robot icon on the screen and is named something-bot, so anyone can tell it's AI; therefore, it qualifies for the exception." I do not recommend such a judgment. From an enforcement perspective, the party claiming the exception bears the burden of proving that obviousness, and failing to prove it constitutes a violation of the disclosure obligation. The AI Office under the European Commission views "not obvious" as the default baseline. It is no easy task to prove that a single icon enabled every user in Europe to recognize that it was AI.

The safe path is to eliminate room for dispute. If a clear sentence reading "This consultation is conducted by an AI" is placed visibly on the initial screen where the conversation begins, the debate over obviousness disappears altogether. The cost of this disclosure is a single sentence. The cost of claiming an exception and losing is a fine under Article 101. Article 101 of the AI Act stipulates that violations of obligations such as Article 50 can incur fines of up to 15 million euros, or 3 percent of total worldwide annual turnover for the preceding financial year, whichever is higher. I see the scales tipping overwhelmingly to one side.

The timing of disclosure is also explicitly anchored in the text. Article 50(5) stipulates that disclosure must be made at the latest at the time of the first interaction. Disclosing "In fact, I was AI" after a lengthy conversation does not satisfy this requirement. The manner is also prescribed: it must be clear and distinguishable, and it must comply with accessibility requirements for users with disabilities. Hiding the text in tiny gray font at the very bottom of the screen carries a high risk of falling short of the "clear and

distinguishable" standard. When the AI Office enforces this requirement, it is highly likely it will not deem a visual cue alone to be sufficient. For a voice bot, it may well demand an audio cue as well.

There is an issue that Korean companies frequently overlook: when they do not build the chatbot themselves, but instead integrate someone else's. If a company calls an application programming interface (API) from a foreign language model provider and layers it into its own service, who is the provider under Article 50(1)? The entity that develops the system and places it on the market is the provider. If a company takes someone else's model and releases a chatbot service under its own name, that company is the provider of that chatbot system. The fact that the model was borrowed does not exempt the company from obligation.

At this junction, there is one more circumstance that Korean companies easily miss. Domestically, the Framework Act on Artificial Intelligence Development and Trust Building (hereinafter referred to as the Domestic Framework Act on AI), in effect since January 22 of this year, already contains a similar disclosure obligation. Businesses launching services using generative AI or high-impact AI must inform users of that fact in advance and must mark the output as AI-generated. In terms of structure alone, it is a mirror image of Article 50(1) and (2). Consequently, during domestic consultations, I frequently encounter questions like, "We have already completed compliance for domestic law, so wouldn't a similar approach work for Europe?" Although the structures are similar, the gravity is entirely different. The upper limit of administrative fines for violating the Domestic Framework Act on AI is 30 million won, and the government is even operating a one-year grace period to allow the system to settle. In contrast, the upper limit of fines for violating Article 50 is 15 million euros - well over 20 billion Korean won - or 3 percent of global turnover, with no such thing as a grace period. This is why the relief felt from satisfying domestic law must not be directly transposed to the European market. Double regulation does not merely mean obligations arise twice; it means the exact same matter must be measured again against two distinct yardsticks.

4.2 Machine-Readable Marking of Synthetic Content: How Far Has Watermarking Come?

Article 50(2) is also a provider obligation, but the target is different. It applies to providers of AI systems that generate synthetic audio, image, video, or text content. This also includes general-purpose AI models. The substance of the obligation is to mark the outputs in a machine-readable format so that the fact of artificial generation or manipulation can be detected.

Here, two types of markings must be distinguished. One is a human-visible marking, such as the statement "This video was created with AI." The other is a machine-readable marking embedded within the file, invisible to the human eye but verifiable through validation tools. What Article 50(2) demands is the latter: machine-readable marking. Human-visible disclosure reemerges later under Article 50(4) concerning deepfake disclosure. Because they operate at different layers, companies providing generative services often must implement both.

The legal text describes the required attributes of machine-readable marking in four words: effective, interoperable, robust, and reliable. What these four words mean in practice will be examined in detail alongside watermarking technology in Chapter 5. Here, I simply emphasize that this marking is not decorative. A mark that is erased by a single screenshot or re-encoding fails to meet the standard of robustness.

Developments on this front have moved rapidly over the past few months. On June 10, the European Commission finalized and published the Code of Practice on Transparency of AI-Generated Content. Subsequently, on July 8, the Commission issued an adequacy evaluation opinion stating that this Code sufficiently supports compliance with the obligations under Article 50(2), (4), and (5), and on the following day, July 9, the AI Board approved that assessment. Companies intending to sign this Code must complete the commitment form and submit it via email to the AI Office, with the deadline set for July 27 at 6:00 PM (Central European Time). As of today, July 24, when this manuscript is being written, only three days remain for companies deliberating signature. The list of signatories is scheduled to be disclosed prior to enforcement on August 2.

The Code of Practice is divided into two tiers. Section 1 deals with the provenance marking obligations of providers offering generative AI systems, while Section 2 addresses the disclosure obligations of deployers utilizing that content. It recommends

layering techniques such as watermarking, metadata, and fingerprinting. This is because no single technique is complete on its own. A watermark embedded in an image disappears with a single screenshot, and metadata is lost the moment the file format is changed. Thus, the Code presents a layered approach combining all three as the practical standard.

Signing is voluntary. However, I believe refraining from signing carries greater risk. For signatories, the AI Office has signaled that enforcement focus will center on monitoring compliance with the Code. Non-signatory companies must prove on their own, without the safety net of the Code and relying solely on the abstract text of the provision, that their technical measures are sufficient. Ultimately, the entity judging that adequacy is the AI Office.

It is also worth noting what specific elements the Code of Practice calls to layer. One is the Content Credentials mechanism developed by the Coalition for Content Provenance and Authenticity (C2PA). It embeds cryptographically signed provenance history within the file, recording when, by whom, and with what tool it was created, as well as subsequent modifications. The other is invisible watermarking. This method embeds traces into pixel or acoustic signals so that they persist to some degree even after compression or resizing. The problem is that both possess individual weaknesses. Content Credentials data is wiped out entirely if the file format is changed or captured. An invisible watermark may persist despite format changes, but by itself cannot reveal who created it, when, or with what tool. Therefore, the Code recommends combining the two, augmented by fingerprinting and logging to enable comparison against originals, forming a defense-in-depth structure. Even if one layer is breached, the remaining layers preserve detectability. Relying on a single technology and assuming marking obligations are fulfilled falls short, in my judgment, of the robustness demanded by this provision.

Article 50(2) also contains exceptions: systems performing purely assistive functions for standard editing, and systems lawfully authorized for crime detection and prevention. A grammar correction tool is an example of the former. If the text-refining function does not substantially alter the content, there is no obligation to mark the output as AI-generated. This exception must be read narrowly. Even if labeled an editing aid, if a tool actually generates non-existent paragraphs, it can hardly hide behind the shield of an assistive function.

The area where Korean companies frequently become caught is image and video generation. Marketing agencies generating product imagery or promotional videos via

generative AI to distribute in the European market, as well as startups selling such capabilities as a service, fall here. Even if a company did not train the model from scratch but built its service using another party's generative model, it is not exempt from marking obligations as long as it places the output of that service on the market.

4.3 Deepfake Disclosure Obligations and Deployer Responsibility

Here, the subject of the obligation shifts. Article 50(4) sets forth obligations for deployers. A deployer is an entity that actually operates an AI system under its authority. It is the party utilizing the system, not the party building and selling it. This distinction forms the thick dividing line in interpreting Article 50.

What paragraph 4 targets is deepfakes. The text defines a deepfake as AI-generated or manipulated image, audio, or video content that resembles actual persons, objects, places, or other entities or events and falsely appears to be authentic or truthful. The deployer must disclose that the content has been artificially generated or manipulated. This time, the disclosure must be human-perceptible. Whereas paragraph 2 concerned machine-readable marks inside files, paragraph 4 concerns notifying viewers that the content is not real.

Paragraph 4 also applies not only to deepfakes but also to AI-generated text published with the purpose of informing the public on matters of public interest. If the text has undergone substantive human review and editorial control prior to publication, this disclosure obligation is exempted. An article drafted by AI but fact-checked and responsibly published by an editor qualifies for this exception. If an editor's name is listed merely as a placeholder while the machine publishes the text unvetted, this exception breaks down.

Let us translate this into a real-world scenario. An advertising agency in Seoul creates a campaign video targeting the European market. Using AI, it synthesizes the face of a real, famous actor, making it appear as though the actor is saying things they never said. This agency is not the provider of the generative tool used to create the video; it used a third party's tool. However, it is the deployer who generated the deepfake with that tool and released it to the world. The disclosure obligation under paragraph 4 falls squarely upon this agency.

The status of deployer is of exceptional importance to Korean companies for a specific reason. Many Korean firms are users of AI rather than developers. They take foreign generative models to produce content and distribute it to European users. In that context, the box where their company name is inscribed under Article 50 is not the provider column, but the deployer column. Dismissing provider obligations as someone else's business often leads to missing the deployer obligations that actually apply to them.

Section 2 of the aforementioned Code of Practice addresses these deployer obligations. It contains guidelines on how to label deepfakes and AI-generated public-interest text, specifying the text, icon size, positioning, and display duration. The annex includes three draft options for standardized EU icons that may be used optionally. Adopting these icons rather than designing custom labels from scratch reduces potential disputes down the line over whether the disclosure method was appropriate.

Exceptions to paragraph 4 are covered separately in Section 5. This is because exceptions for art and satire fall here, and the extent to which these exceptions can be invoked remains a central point of practical dispute.

4.4 Notification Obligations for Emotion Recognition and Biometric Categorisation Systems

Article 50(3) is again a deployer obligation. Its targets are emotion recognition systems and biometric categorisation systems. A deployer operating such systems must inform natural persons exposed to them. The provision further specifies that if personal data is processed, the General Data Protection Regulation (GDPR) and relevant EU data protection laws must also be complied with concurrently.

An emotion recognition system is designed to infer emotional states from human facial expressions, vocal tones, or gestures. A biometric categorisation system categorizes persons into specific groups based on biometric data such as faces or fingerprints. Tools in call centers that analyze customer vocal pitch to determine anger, or in-store cameras estimating visitors' age groups or gender to dynamically alter advertisements, can fall under this provision.

Two regulations must not be confused here. Article 50(3) sets forth a notification obligation - it requires informing individuals before use. However, using emotion recognition in the workplace or educational institutions is a prohibited practice altogether under Article 5. It is not a matter of notifying and using; it cannot be used at all. This prohibition is addressed in a later chapter, not in Article 50. This is where employee emotion analysis tools used lawfully in Korea cross the line into outright prohibition in the EU. For now, it suffices to distinguish that the notification duty of Article 50(3) and the prohibition under Article 5 reside in separate provisions.

Paragraph 3 contains an exception for lawful crime detection purposes. However, because this exception assumes law enforcement authorities, there is virtually no room for ordinary Korean businesses expanding into Europe to invoke it. Moreover, this exception carries an additional condition: safeguards must be implemented to protect third parties' rights and freedoms, and relevant EU legislation must be observed. Invoking an exception does not grant blanket immunity from data protection duties.

This is precisely why I advise conducting both a Data Protection Impact Assessment (DPIA) and a Fundamental Rights Impact Assessment (FRIA). Fulfilling notification obligations does not automatically mean the system satisfies data processing principles. The two assessments address different questions: a FRIA asks to what extent the system impacts human fundamental rights, while a DPIA examines how personal data is handled

during that process. For systems that process human body and mind as data, such as emotion recognition and biometric categorisation, it is practically superior to conduct both assessments together within a single procedure rather than handling them separately.

Article 50(3) has no separate grace period. One cannot expect a four-month leeway like the watermarking obligation in paragraph 2. It applies directly starting August 2. If a service operating in Europe incorporates emotion recognition or biometric categorisation elements, this provision must be addressed first.

Quite a few companies attempt to transfer customer service analytics tools developed in Korea directly to their European subsidiaries' call centers. A prime example is voice emotion analysis used for agent training. A feature treated domestically as a mere quality control tool becomes subject to the deployer notification requirement under Article 50(3) the moment it runs against European agents and customers. If applied to agents, the risk of violating the Article 5 prohibition on workplace emotion recognition must be evaluated; if applied to customers, the notification duty under paragraph 3 must be satisfied. That the applicable provision varies depending on who the feature targets is a point I find myself explaining repeatedly during practical consultations.

4.5 To What Extent Can Limited Exceptions for Art, Satire, and Safety Be Invoked?

The first thing companies look for when reading Article 50 is exceptions. They wonder: "Since we are art or satire, shouldn't we be exempt?" I view such expectations with caution. Exceptions exist, but the scope drawn by the statutory text is narrower than companies hope.

Let us examine the deepfake exception in paragraph 4 first. When content forms part of an artistic, creative, satirical, or fictional work, the disclosure obligation does not vanish. Rather, the manner of disclosure is relaxed: it suffices to disclose in an appropriate manner that does not hamper the enjoyment of the work. Listing in a film's end credits that AI synthesis was used in a scene is one such method. It means there is no need to continuously display a large warning banner in the center of the screen that ruins viewing. The true nature of the exception is revealed here: it is not an exception granting exemption from disclosure, but an exception allowing disclosure in a less intrusive manner.

This distinction is significant in practice. It is an error for a company to interpret the law as "We are art, so we need no marking at all." Disclosure remains mandatory, provided it is done without impeding enjoyment. Disputes inevitably arise along the boundary between what constitutes art and what constitutes commercial advertising. Having artistic elements in a brand promotional video does not guarantee that it falls under an artistic work within the meaning of paragraph 4. Where commercial advertising intent is clear, I consider it safer to make formal disclosures rather than relying on the art exception. Since the aforementioned Code of Practice details label size, positioning, and exposure duration, adhering to the Code is preferable to arbitrary judgment.

Exceptions for criminal detection, prevention, investigation, and prosecution are also scattered throughout Article 50. All four duties - paragraph 1 chatbot notification, paragraph 2 machine-readable marking, paragraph 3 emotion recognition/biometric categorisation notification, and paragraph 4 deepfake disclosure - contain law-enforcement-purpose exceptions authorized by law. These exceptions serve the clear public interest of public safety. However, paragraph 3 explicitly stipulates that data protection laws must still be strictly observed even in such cases. Ordinary enterprises should assume there is virtually no room to invoke these exceptions for commercial operations.

The burden of proof runs through all exceptions. The party seeking to rely on an exception must establish that its system satisfies the requirements for that exception. While acknowledging that legally binding final interpretation rests with the Court of Justice of the European Union, the AI Office has maintained that it will assess matters on a case-by-case basis. In these early days before enforcement precedents accumulate, companies must remember that when boundaries are blurred, the risk is borne by the party asserting the exception.

Article 50 at a Glance: Summary Table

Categorizing the obligors and requirements in this manner makes it easy to locate which box your company falls into.

Article 50(1). Obligor: Provider. Must notify users that they are interacting with AI (chatbots, etc.). Exception: obvious AI interaction; burden of proof rests on the provider. Application date: August 2, 2026, no grace period.

Article 50(2). Obligor: Provider. Must mark synthetic audio, image, video, and text output in a machine-readable format. Exceptions: standard editing assistance, lawful crime prevention. Application date: August 2, 2026; existing systems granted grace period until December 2, 2026.

Article 50(3). Obligor: Deployer. Must notify natural persons exposed to emotion recognition or biometric categorisation systems, complying concurrently with data protection rules such as GDPR. Workplace/education emotion recognition is separately prohibited under Article 5. Application date: August 2, 2026, no grace period.

Article 50(4). Obligor: Deployer. Must disclose deepfakes in a human-perceptible manner. AI-generated text of public interest must also be disclosed, unless subject to substantive human editorial responsibility. Art and satire must disclose in a manner that does not hamper enjoyment. Application date: August 2, 2026.

Article 50(5). Timing of notification: at the latest at the time of first interaction or exposure. Manner: clear and distinguishable, adhering to accessibility requirements.

Article 50(7). The AI Office encourages the drawing up of Codes of Practice. The Code of Practice on Transparency of AI-Generated Content was finalized on June 10, 2026, approved by the AI Board on July 9, with the signature deadline set for July 27 at 6:00 PM (CET).

Five Self-Diagnostic Questions

Does our company operate a chatbot or voice bot interacting with European users? If so, paragraph 1 applies. Verify whether a disclosure notice is present on the initial screen.

Does our company generate images, videos, or text using AI and export them to Europe? If selling the creation tool, paragraph 2 machine-readable marking applies; if distributing content created with such tools, paragraph 4 deepfake disclosure applies. Both may apply simultaneously.

Does our company employ a system in Europe that categorizes something based on human emotions or biometric data? If so, paragraph 3 notification applies, and if used in workplace or educational contexts, check for potential Article 5 prohibitions as well.

Among the content our company produces, have we omitted disclosure on grounds of art, satire, or creative works? If so, re-examine whether that determination has been documented and whether commercial advertising intent is mixed in.

Has our company not yet decided whether to sign the Code of Practice? The deadline is July 27 at 6:00 PM (CET). If deciding not to sign, alternative technical measures must be prepared separately to satisfy the requirements of paragraphs 2 and 4.

Let me close with a single date. Article 50 applies starting August 2, 2026. For the watermarking obligations of generative systems already placed on the market and operating prior to this date, a four-month grace period until December 2, 2026, is granted. However, for systems launched anew after that date, there is no grace period, and the notification and disclosure obligations under paragraphs 1, 3, and 4 carry no grace period from the outset. That game company in Pangyo passed an agenda item in last week's meeting to add a single line of text to the chatbot's initial screen. The development team said it would take half a day. What actually took time was not the screen design, but the meeting to convince executive management why they fell under the scope of this provision. If you are currently exporting anything to Europe, please count the days remaining on your calendar until August 2.

Chapter 5. Watermarking Technology and Law

This actually happened at an image-generation startup in Pangyo. During an investor demo, the team lead typed in a prompt and generated a portrait photograph. A small AI label icon appeared in the corner of the screen, and opening the file properties revealed signed information showing when and with which model the photo was created. Up to this point, everything went according to plan. The problem came next. An investor took a screenshot of that screen and uploaded it to his Slack. When opening the captured file, the visible icon remained, but the signed information embedded inside the file had completely vanished. A mundane action like taking a screenshot wiped out the carefully attached provenance information. This scene encapsulates the core issue addressed throughout Chapter 5. While the law mandates labeling AI-generated content, how well that label survives what content routinely undergoes on the internet is an entirely different question.

5.1 December 2, 2026 Grace Period for Existing Systems, and Zero Grace Period for Newly Launched Systems

The European Union Artificial Intelligence Act (Regulation (EU) 2024/1689, hereinafter the AI Act) entered into force on August 1, 2024, and Article 50, which lays down transparency obligations, becomes fully applicable 24 months later, on August 2, 2026. Article 50 is divided into four paragraphs, each targeting a different scope. Paragraph 1 requires AI systems intended to interact directly with natural persons, such as chatbots, to inform users that they are interacting with an AI system. Paragraph 2 requires providers of AI systems generating synthetic audio, image, video, or text content to mark outputs in a machine-readable format and ensure they are detectable. Paragraph 3 requires deployers of emotion recognition or biometric categorization systems to inform natural persons exposed to them. Paragraph 4 requires deployers of systems generating or manipulating deepfakes - such as altered images, audio, or video - to disclose that the content has been artificially generated or manipulated.

The change directly affecting practical operations here is that a separate transitional measure was attached exclusively to Paragraph 2 through the Digital Omnibus amendment agreed upon by EU co-legislators. A political agreement on this amendment was reached between the European Parliament and the Council on May 7, 2026, and signing was completed on July 8, 2026. As of writing, publication in the Official Journal is pending. Pending publication does not mean it can be ignored. Treating an issue fully agreed upon and signed by co-legislators as groundless is a far riskier judgment. The transitional measure works as follows. Generative AI systems placed on the market or put into service in the EU before August 2, 2026, have until December 2, 2026, to comply with the marking and detection obligations required by Article 50(2). That grants a four-month breathing window. Outputs released before this grace period, meaning results including deepfakes produced before August 2, do not need to be marked retroactively.

Newly launched systems, by contrast, get no such cushion. AI systems placed on the market or put into service in the EU for the first time after August 2, 2026, including generative AI, must meet all requirements of Article 50 immediately upon launch. The grace period under Paragraph 2 is a narrow exception applicable solely to existing systems already on the market before August 2. Paragraphs 1, 3, and 4 apply from August 2 without exception, whether existing or new. The principle stated by the EU AI Office aligns with this direction: every AI system covered by Article 50 must comply if placed on the EU market or put into service as of August 2, regardless of when it was first launched. I

read this provision as a penalty structure for new entrants. It gives incumbents four additional months while demanding a fully finished implementation from day one for newcomers trying to break into the market.

For companies headquartered in Seoul, this distinction means the preparation burden varies drastically depending on the set launch date. If you already operate a generative AI service in the EU, you have until December 2 to finalize your watermarking and metadata systems. But if you plan to launch a new service in the EU for the first time after August 2, you must complete the labeling features starting from the design phase before launch. Preparing both on the same timeline guarantees that the new service will fall behind. Fines for non-compliance under Article 99 reach up to 3 percent of total worldwide annual turnover for the preceding financial year or 15 million euros, whichever is higher. Starting August 2, national market surveillance authorities and the AI Office will actively enforce these penal powers. Rather than relaxing over a granted extension, it is safer to treat those four months as time allocated for actual technical implementation.

5.2 The C2PA Standard and Content Provenance Authentication

Article 50(2) of the AI Act itself does not mention the name C2PA anywhere in its text. The statute merely requires providers to mark outputs in a machine-readable format and ensure they are detectable, without dictating which specific technology must implement the marking. Filling this gap is the Code of Practice on Transparency of AI-Generated Content published by the AI Office on June 10, 2026. This code divides marking techniques into three methods: watermarking, metadata, and fingerprinting, translating the statutory demands for effectiveness, interoperability, robustness, and reliability into concrete actionable measures. Detailed Measure 1.1.1 requires digitally signed metadata, while Measure 1.1.2 calls for invisible watermarking. Virtually only one technology satisfies both conditions simultaneously: Content Credentials created by the Coalition for Content Provenance and Authenticity (C2PA).

What Content Credentials actually does is straightforward. The moment content is created, a data bundle recording which tool was used, when, and what edits occurred is cryptographically signed and attached to the file. This signed bundle is called a manifest. Every time the content passes through another tool, a new manifest attaches while pointing back to the previous manifest, building a chained history from creation to distribution within a single file. If even one tool that fails to verify signatures gets introduced into this chain, the provenance breaks right there, leaving subsequent stages incapable of proving origin.

Absence from legal text and being established as a de facto standard in practice are two different matters. The Code of Practice is a voluntary agreement, lacking direct legally binding force. Yet signing the code and adhering to its provisions offers strong grounds to demonstrate compliance with Article 50(2). Viewing this alongside the harmonised standards system under Article 40 clarifies the framework further. High-risk AI systems or general-purpose AI models complying with harmonised standards enjoy a presumption of conformity with relevant requirements. Though C2PA has not been formally designated as a harmonised standard with that official status yet, the fact that the AI Office effectively singled it out as the technical path in its Code of Practice practically dictates what evidence to present to regulators. The AI Office released a standard icon set on the same day. While the designated AI label attached to content created or modified by AI forms the human-visible layer, C2PA Content Credentials operate underneath as the machine-verifiable layer.

Korean companies must tackle two practical takeaways here. First, while alternative technologies remain permissible, companies bear the burden of proving that their solution matches C2PA in detectability and open accessibility. The fact that the code does not mandate a specific technology does not imply that any arbitrary method will pass muster. Second, adopting C2PA is not an issue confined to image or video generation pipelines. Because Content Credentials chain signatures across generation, editing, and distribution stages, a break in the signature chain at any point in an internal system drops provenance metadata entirely from the final output. Handing this issue off solely to the engineering team carries real risk. When legal teams must demonstrate compliance to regulatory authorities, they need the technical capability to pinpoint precisely where any signature chain broke.

5.3 Robustness Requirements: Can It Withstand Re-encoding and Screenshots?

Article 50(2) stipulates that marking technologies must be effective, interoperable, robust, and reliable. Among these four adjectives, robustness proves exceptionally tricky in practice. It demands that marks attached to content survive compression, resolution changes, re-uploading to other platforms, or direct screen captures. The incident at the Pangyo startup illustrates this exact challenge. The icon visible on screen was not drawn directly into the image pixels; it was rendered by the viewer reading embedded metadata. The moment a screenshot was taken, the original file's metadata was severed, leaving behind only the captured pixels on screen. Relying strictly on metadata causes labels to vanish in such scenarios.

Invisible watermarking, which embeds signals directly into the pixels of content, presents a different story. Technologies like Google's SynthID retain a substantial portion of the signal within the image itself even during routine re-encoding, light editing, or full metadata loss caused by screenshots. That resilience stems from embedding the signal inside the content rather than in metadata. The real challenge comes from adversaries intentionally scrubbing marks. Regeneration attacks - adding noise to distort signals before re-rendering via generative AI - and watermark stealing attacks that learn and strip watermarks drop detection reliability significantly without visible loss in image quality. Text watermarking faces similar limits; simply rewriting sentences dilutes the signal. Watermarking currently serves as a mechanism to raise the cost of rendering AI-generated content untraceable, rather than an unbreachable lock that prevents removal entirely.

Screenshots occupy a peculiar position in this robustness debate. A screenshot neither re-encodes the file nor degrades visual quality through editing. Because it simply recaptures displayed pixels, all metadata, signatures, and manifests attached to the original file drop off completely. Methods attaching information alongside files, like Content Credentials, become completely helpless at that instant. Conversely, watermarking that imprints signals directly into pixel arrangements stands a chance of survival because the signal transfers right along with the captured pixels. Rather than viewing the two technologies as competitors, we should see them as complementary mechanisms plugging each other's gaps. Metadata signatures preserve detailed history but disappear with a single screenshot; pixel watermarking survives screenshots but cannot reveal who created what or when. Embedding both ensures that if one layer

breaches, the other leaves at least baseline evidence behind.

Article 50(2) dictates that assessing robustness requires considering the specific characteristics and limitations of content types, implementation costs, and the state of the art at any given time. Yet quantitative benchmarks defining how much re-encoding or screenshotting a mark must endure to qualify as robust remain unestablished. Standards will likely solidify as AI Office guidelines, harmonised standards, and member state enforcement precedents accumulate. The only certainty today is that the state-of-the-art benchmark continually raises its bar over time. A watermarking technology approved today carries no guarantee of being recognized as robust a year later. Approaching labeling as a one-and-done feature embedded into content generation engines risks getting labeled outdated during the next audit cycle. I advise budgeting for this requirement not as a one-time setup, but as an ongoing R&D item requiring continuous updates.

5.4 Ban on Generating NCII and CSAM: Article 5 Amendments and December 2, 2026 Entry into Force

Article 5 acquired this scope only recently. Prior to these amendments, Article 5 functioned merely as a provision banning manipulative, exploitative, and social scoring practices, without explicitly citing non-consensual intimate imagery (NCII) or child sexual abuse material (CSAM). Filling that gap was the Digital Omnibus on AI package politically agreed upon by the European Parliament and the Council on May 7, 2026. That agreement incorporated so-called nudifier applications - AI systems stripping individuals' clothing or manipulating content into sexual images, videos, or audio without consent - along with AI systems producing child sexual abuse material into Article 5's prohibited practices. This package completed formal signing on July 8, 2026, leaving only publication in the Official Journal pending.

The amended Article 5 bans placing such systems on the EU market or putting them into service altogether. The scope extends beyond systems intentionally designed for such purposes; systems where generating such outputs represents a reasonably foreseeable outcome fall under regulation even if providers lacked that intent. A safe harbour provision exists, however. Systems equipped with effective technical safeguards that reliably prevent such outputs escape the prohibition. Companies operating image-generation services must therefore demonstrate the reliability of filters preventing such content creation to steer clear of this clause. Because court precedents and guidelines have not yet fixed the required reliability threshold for the safe harbour, I advise preserving test logs along with false-positive and false-negative metrics from the development stage rather than cobbling together retrospective reports. Regulators weighing safe harbour eligibility will ultimately demand these exact logs. The enforcement date coincides with December 2, 2026, the end of the Article 50(2) grace period discussed in Section 5.1. While the dates overlap, the provisions address distinct matters. Article 50(2) marks the end of a grace period requiring existing generative AI to complete labeling obligations, whereas the amended Article 5 marks the beginning of a fresh rule imposing an absolute ban on specific outputs. One date marks an expiration of grace; the other marks the commencement of a ban, coinciding purely by date.

The gravity of legal violations also differs fundamentally. Article 50 non-compliance carries penalties under Article 99 of up to 3 percent of global annual turnover or 15 million euros. Violations of Article 5 prohibited practices, by contrast, cap out at the highest tier in the AI Act's penalty hierarchy: up to 7 percent of global annual turnover or 35 million

euros, whichever is higher. Criminal laws across EU member states compound this risk. Generating NCII or CSAM already constituted criminal offenses in multiple member states prior to the AI Act; amending Article 5 explicitly cements this longstanding ban against AI systems as new tools. Korean firms building image or video generation services for EU users cannot afford to push content filtering down their priority list. Executive leadership needs to understand directly that a single output breaching content filters can jeopardize 7 percent of the company's total turnover.

Reference Materials

The timeline of dates and obligations summarizes as follows: August 1, 2024: AI Act enters into force. February 2, 2025: Article 5 prohibited practices become fully applicable. August 2, 2026: Article 50 full application date. From this date, newly launched systems must comply with Article 50(1) through (4) without exception. Existing systems must also comply with Paragraphs 1, 3, and 4 from this date. December 2, 2026: Two events converge. First, the Article 50(2) grace period expires for existing generative AI systems launched before August 2, 2026. Second, the ban on nudifier and CSAM generation systems under amended Article 5 takes effect.

Decision criteria narrow down to three questions: When was or will our system be placed on the EU market? Whether it falls before or after August 2 determines grace period eligibility. Does our marking technology rely strictly on metadata or embed signals inside the content itself? That dictates how well it endures screenshots and re-encoding. Does our generation system feature filters screening out sexual imagery or child-related content, and can we document their reliability? That decides whether we qualify for the Article 5 safe harbour.

Penalty tiers are restated here as well: General violations under Article 99 incur up to 3 percent of turnover or 15 million euros. Violations of Article 5 prohibited practices incur up to 7 percent of turnover or 35 million euros. SMEs receive the lower cap in both tiers.

In that photo captured by the investor and posted to Slack, only a small visible icon remains today. The signature underneath vanished, and no one can prove which model or prompt produced that file anymore. The word robustness demanded by law ultimately signifies the capacity to endure this exact moment.

Chapter 6. Transparency Code of Practice

At 6:00 p.m. on July 22, 2026, one day before the Central European Time deadline, the legal team of an image and video generation service startup based in Seongsu-dong held a meeting. They needed to decide whether to sign the Transparency Code of Practice finalized in Brussels right away or wait and see. Since the company's video generation service was also used by European users, there was no way from the outset to escape the transparency obligations set out in Article 50 of the European Union Artificial Intelligence Act (Regulation (EU) 2024/1689, hereinafter the AI Act). However, how to demonstrate compliance with that obligation and where to source the basis for that explanation remained a choice for the company to make.

The Code of Practice covered in this chapter shares its name with the general-purpose AI model (GPAI model) Code of Practice examined in Chapter 5, but it has different origins. The GPAI Code of Practice, stemming from Article 56 of the AI Act, targets a small number of large model providers. Conversely, the Transparency Code of Practice envisaged under Article 50(7) targets a far broader group of companies, ranging from startups creating deepfake videos to customer service centers integrating chatbots and advertising agencies generating synthetic images. The European Commission finalized the ultimate version of this Transparency Code of Practice at the Closing Plenary on June 10, 2026. The multilateral consultation process, which lasted nearly nine months, brought together over 187 participants from industry, academia, civil society, copyright holder organizations, and Member State governments, with six independent experts commissioned by the AI Office leading the drafting process.

One point must be noted in advance: the transparency obligations set forth in Article 50 are not entirely covered by this single Code of Practice. Article 50(1) requires providers of AI systems intended to interact directly with natural persons, such as chatbots, to design them so that users are informed that they are interacting with AI, while Article 50(3) requires deployers of emotion recognition systems or biometric categorization systems to inform exposed individuals of that fact. Among these, Article 50(7) targets only the detection and labeling obligations - specifically Article 50(2) and Article 50(4) - as the scope of the Code of Practice. Consequently, even if a company operating a chatbot or utilizing emotion recognition functions signs this Code of Practice, the obligations under those two provisions will not be automatically resolved. This means that a separate notification design remains necessary for each.

6.1 Practical Benefits of Signing: Shift in the Burden of Proof

The scope covered by the Transparency Code of Practice encompasses two obligations under the AI Act: the marking and detection obligation under Article 50(2) and the labeling obligation under Article 50(4). Article 50(2) imposes an obligation on providers of AI systems generating synthetic audio, image, video, or text content to ensure that outputs are marked in a machine-readable format and detectable as artificially generated or manipulated. Article 50(4) requires deployers of systems that generate or manipulate deepfakes, or systems that generate text published with the purpose of informing the public on matters of public interest, to inform users that the output has been artificially generated or manipulated.

Both obligations apply as scheduled from August 2, 2026, regardless of whether a company signs the Code of Practice. What the Code of Practice changes is not the existence of the obligation itself, but how compliance with that obligation is demonstrated. A signatory company does not need to explain from scratch why it believes its measures are sufficient. Instead, it can simply go through the items set out in the Code of Practice - such as which watermarking standard it uses and what metadata framework it attaches - as a checklist to show compliance. From the perspective of Member State market surveillance authorities, the subject of review also becomes clear. Judging that a company has faithfully followed a standardized framework created through EU multilateral consultation is far easier than scrutinizing each company's proprietary methodology from square one every single time. I view this as the true practical benefit of signing the Code of Practice. It is not the creation of a new statutory provision, but the shift of the center of gravity in proof from the company to the Code - this is precisely the shift in the burden of proof referenced in the title of this section.

In addition, there are practical gains. Market surveillance authorities are likely to focus their supervisory efforts on verifying whether companies that have declared compliance are adhering to the Code. This means spending fewer individual investigative resources on signatory companies compared to non-signatory ones. However, one point must be made clear: unlike the GPAI Code of Practice, there is no explicit statutory provision stating that compliance with the Transparency Code of Practice serves as a mitigating factor to reduce fines in the event of an AI Act violation. The fine mitigation provision under Article 101 of the AI Act applies exclusively to providers of GPAI models, whereas sanctions for violations of Article 50 follow the general fine scheme under Article 99 - namely, up to 3 percent of total worldwide annual turnover or 15 million euros, whichever

is higher. It is accurate to view the weight given to Code compliance in calculating these fines as being left to the discretion of Member State supervisory authorities.

The signing procedure itself also differs in nature from that of the GPAI Code of Practice. The GPAI Code of Practice takes effect only after undergoing a formal assessment and approval procedure by the European Commission. The Transparency Code of Practice, by contrast, operates by publishing the outcome of multilateral consultations and receiving signatures, making the weight of the approval procedure relatively light. The deadline of July 22, 2026, served as the window to be included on the initial list of signatories, and it appears highly likely that the window will remain open to continuously accept companies wishing to sign later. Given the precedent set by the GPAI Code of Practice, where additional signatories were added on multiple occasions after its initial release, I anticipate a similar trajectory this time. Nevertheless, as this expectation could prove wrong, any company seeking to sign at a later date should check AI Office announcements regularly.

6.2 Compliance and Its Relationship to the Presumption of Conformity

There is one distinction that must be clearly drawn here: signing and complying with the Transparency Code of Practice does not give rise to a presumption of conformity under the AI Act. The presumption of conformity is a benefit attached solely to harmonised standards defined under Article 40 of the AI Act. When an entity complies with a harmonised standard drafted by European standardisation organisations (CEN, CENELEC) and published in the Official Journal of the European Union, the requirements covered by that standard are presumed to comply with the law. Although it is a rebuttable presumption, it has the powerful effect of shifting the burden of proof to the regulatory authorities. The Transparency Code of Practice does not reach this level.

This difference becomes stark when compared with the GPAI Code of Practice. For the GPAI Code of Practice, statutory provisions themselves - namely Article 53(4) and Article 55(2) of the AI Act - explicitly designate compliance with the Code of Practice as a means of demonstrating fulfillment of obligations. The statutory text effectively singles out the Code of Practice as an acceptable route. In contrast, Article 50(7) merely states that the AI Office shall encourage and facilitate the drawing up of Codes of Practice, without embedding any statutory link that equates compliance with the Code to proof of fulfilling obligations. Thus, if one were to establish a hierarchy, it would look like this: at the top is the full presumption of conformity attached to harmonised standards; one step below lies the evidentiary pathway of the GPAI Code of Practice explicitly recognized by statutory provisions; and further below is the Transparency Code of Practice, which carries weight only through administrative practice without explicit statutory backing. I believe this hierarchy must not be confused. A company that misinterprets signing the Transparency Code of Practice as a legal shield identical to the GPAI Code of Practice may find itself caught off guard the moment a Member State supervisory authority declares that signing alone is insufficient.

More precisely, the weight that the Transparency Code of Practice currently carries is closer to compelling circumstantial evidence reflecting the state of the art. Because it is the product of consensus among more than 187 stakeholders and six independent experts, it serves as a basis to argue that such measures represent the industry standard. However, it is not in itself the key to triggering the legal presumption effect established by law. To add one point, harmonised standards regarding the obligations set out in Article 50 have not yet been published as of the writing of this manuscript. Thus, for now, the Transparency Code of Practice serves as virtually the sole practical benchmark, and once

harmonised standards are issued, that hierarchy may shift again.

Enforcement channels also differ. The GPAI Code of Practice is supervised directly by the AI Office, which exercises fine-imposing authority starting August 2, 2026. The obligations under Article 50 related to the Transparency Code of Practice are enforced by Member State market surveillance authorities. This means the counterpart Korean companies face is not the AI Office in Brussels, but the supervisory authorities of individual Member States where their services are distributed; hence, practical response channels must be prepared in the language and procedural framework of each respective country.

Even within a single Member State, enforcement channels are not consolidated into a single entity. If a deepfake video touches upon politics or elections, broadcasting and media regulatory bodies may examine it jointly; if synthetic content involves personal data, Data Protection Authorities (DPAs) may take a separate interest. A company assuming it only needs to deal with a single market surveillance authority might find itself receiving inquiries from multiple agencies simultaneously. Furthermore, the interpretative guidelines clarifying the scope and applicability of Article 50 remain at the draft stage at the time of writing. This presents a reversed order from the norm, where the practical tool - the Code of Practice - was released first, followed by interpretative guidance. Once the guidelines are published, clearer answers are expected regarding gray areas not covered by the Code of Practice, such as the extent to which exceptions for artistic or satirical expressions will be recognized.

6.3 Alternative Pathways for Non-Signatory Companies

Choosing not to sign is, of course, a valid option. The Code of Practice is voluntary. However, the marking, detection, and labeling obligations set out in Article 50 remain fully intact regardless of signature status. Non-signatory companies must independently explain and document why their watermarking methods or deployer notification wording meet the standard required by law - that is, measures that are effective, interoperable, robust, and reliable. In a context where the Code of Practice has already established itself as the de facto industry standard, conducting such documentation without referencing the language of the Code is arguably more difficult. I expect that even companies choosing this alternative pathway will often keep the Code of Practice close at hand to benchmark their own measures. While not a legal obligation, it serves in practice as a compass.

What cards, then, can non-signatory companies actually hold in practice? I point to three. First, independently adopting metadata or watermarking technologies that preserve content provenance, and documenting why that choice is effective and robust in the form of a technical report. Second, having the validity of their measures evaluated by external auditors or technical experts to secure an independent opinion. Third, codifying marking procedures in internal policies or terms of service to demonstrate that the measures were not accidental but designed into the system from the outset. All three options require considerably more effort than signing the Code of Practice and ticking off a checklist. However, for a company whose business structure does not align well with the standard items of the Code of Practice - such as one building specialized industrial image generation tools - this pathway may represent a more honest choice.

Timelines must also be managed carefully. Article 50 obligations apply generally from August 2, 2026; however, under the provisional agreement reached in May 2026 on the AI Omnibus amendment, generative AI systems already placed on the market prior to that date are granted a grace period until December 2, 2026, limited to the machine-readable marking requirement. The interoperability requirement will become fully applicable later, on February 2, 2027. Companies needing to retroactively adapt products already in service thus have a slight buffer during this transitional period to refine their measures.

Another point to note is the Article 50 interpretative guidelines being prepared separately from the Code of Practice. Ordinarily, interpretative guidelines are issued first, followed by practical tools such as the Code of Practice; here, the order was reversed. The Code of

Practice was finalized first in June 2026, while the guidelines explaining the scope and applicability of statutory provisions remain in draft form. For the time being, non-signatory companies lack reliable official interpretative material to lean on in place of the Code. Companies deciding against signing have no choice but to closely monitor upcoming guidelines from the AI Office and early enforcement cases by Member State supervisory authorities, continuously checking that their measures do not fall behind.

For Korean companies, an additional layer of complexity arises. Domestic AI framework legislation in Korea already contains marking obligations for generative AI outputs, creating a dual compliance burden to satisfy both EU and domestic regulations simultaneously. If the marking formats or wording requirements of the two laws do not fully overlap, companies must first evaluate whether a single watermarking and notification framework can fulfill both regimes at once. For companies headquartered outside the EU, it is advisable to establish an in-EU contact point - namely an authorized representative - to handle inquiries or investigations related to Article 50.

At the meeting that day, the legal team of the Seongsu-dong startup leaned toward signing. The proportion of European users exposed to their video generation service was non-negligible, and dedicating internal personnel to wrestle individually with supervisory authorities was unfeasible. Yet the legal director's closing remark was telling: signing is not the end, but rather the beginning of an internal review to verify whether the company actually fulfills all the items demanded by the Code. When the development team could integrate the watermarking module, and how to embed the deepfake notification text into customer support scripts - these were answers that each team ultimately had to supply, regardless of whether the company signed.

References

Legal Basis: Article 50(2) of the AI Act (marking and detection), Article 50(4) (labeling of deepfakes and public interest text), Article 50(7) (legal basis for Codes of Practice)
Finalization Date: June 10, 2026 Closing Plenary Initial Signatory Deadline: July 22, 2026, 18:00 CET; subsequent participation subject to AI Office guidance
General Applicability: August 2, 2026
Grace Period for Existing Systems Marking Requirement: December 2, 2026 (reflecting the AI Omnibus provisional agreement)
Applicability of Interoperability Requirement: February 2, 2027
Basis for Fines: Article 99 (General, 3% of global turnover or 15 million euros, whichever is higher), distinct from Article 101 fine mitigation provisions exclusive to GPAI
Enforcing Body: Member State market surveillance authorities (contrasted with direct enforcement by the AI Office for the GPAI Code of

Practice) Legal Status Comparison: Harmonised standards (Article 40, full presumption of conformity) at the top; GPAI Code of Practice (Article 53(4) & Article 55(2), statutorily recognized evidentiary path) next; Transparency Code of Practice (Article 50(7), functioning through administrative practice without explicit statutory backing) below

Part III

For Companies Handling General-Purpose AI Models

Chapter 7. Are We a Provider of General-Purpose AI?

Consider an artificial intelligence startup in Daejeon. Its development team downloaded a language model released overseas. The company re-trained that model using chemical process data accumulated over several years, and the resulting output reached a level where it can explain anomalies in chemical equipment in full sentences. The company began providing this service for a fee to two chemical companies in Germany. The CEO thinks: The model was created by someone else, and we merely customized it to suit our work; therefore, from a regulatory standpoint, we are closer to being a user.

Whether this judgment is correct is the entire focus of this chapter. If this company is classified as a provider of a general-purpose AI model (GPAI model), the entire bundle of obligations under Article 53 of the AI Act applies. It must maintain technical documentation, establish a copyright policy, and publish a summary of the content used for training. The distance between the perception of "having merely customized it" and the legal status of "provider" - that distance is the risk this company is missing.

Let us first unpack the term "GPAI model." Article 3(63) of the AI Act defines a GPAI model as an AI model that displays significant generality and is capable of competently performing a wide range of distinct tasks regardless of the way the model is placed on the market and that can be integrated into a variety of downstream systems or applications. Whether it has been trained with a large amount of data using self-supervision at scale does not matter. Language models are a common example. A single model translates, summarizes, and writes code. This broad utility is what significant generality means. The party that creates this and places it on the market is the provider.

7.1 The Criterion of Significant Generality

The term "significant generality" is not listed exhaustively in the articles themselves. Recitals 97 and 99 of the AI Act cite large generative AI models as typical examples of GPAI, pointing to their ability to flexibly generate various forms of content such as text, audio, images, and video, and to easily handle a broad range of tasks. Models that only recognize faces or models that only evaluate credit scores are excluded here because they are bound to a single narrow task. Conversely, large language models that generate plausible answers no matter what is asked are considered to possess significant generality.

The Commission itself acknowledges that it is impossible to lock down this determination with an exact list of capabilities. Instead, it uses computational power as an indicative criterion, measuring the amount of floating-point operations, or FLOPs, used to train the model. According to practical commentaries circulating in the market, models whose training computation exceeds 10^{23} FLOPs are treated as having crossed the indicative threshold to be considered a GPAI model. In addition, derived models reduced in size by pruning or quantizing a model that has crossed this threshold cannot escape being treated as GPAI models if the original model exceeded that threshold. This means making a model smaller does not allow one to evade regulation. This figure of 10^{23} occupies a different place from the threshold for systemic risk - namely, 10^{25} FLOPs - which we will examine later in Chapter 9. One threshold distinguishes whether a model is a GPAI model or not, while the other imposes heightened obligations on those GPAI models that present systemic risks. The two must not be confused.

The guidelines containing this indicative criterion were first released by the Commission after approving their content on July 18, 2025, and were formally adopted on November 19, 2025, as Communication Document C(2025) 7719. They are not legally binding. The provisions of the AI Act and subsequent interpretations by the Court of Justice of the European Union take precedence over these guidelines. Nevertheless, in practice, because the Commission and the AI Office will actually assess cases against this criterion, it is not a document that can be ignored simply because it lacks binding force.

Where Korean companies get caught up in practice with this standard is in the middle ground. When building a model that is neither exceptionally large nor very narrow, whether it has crossed the threshold of significant generality becomes ambiguous. In such ambiguous territory, I believe it is safer to assume that the threshold has been crossed and prepare accordingly. If a company concludes that it has not crossed the

threshold and it is later determined to have done so, the company will have missed the entirety of the Article 53 obligations. Conversely, if it assumes it has crossed the threshold and prepares, but later turns out not to have, it has merely prepared somewhat in advance. The weight of risk tilts decisively toward one side. Additionally, it is worth keeping in mind that signing the Code of Practice for providers of GPAI models makes the process of demonstrating compliance to the AI Office considerably easier.

The Commission makes a point of explicitly stating within the guidelines that the determination of significant generality must always undergo a case-by-case assessment, and that it will review and revise these guidelines in light of technological, societal, and market trends, as well as rulings by the Court of Justice of the European Union. This means there is no guarantee that the line drawn today will remain in the same place three years from now. Therefore, I advise against viewing this threshold as a one-and-done measurement; instead, establish it as an ongoing procedure to be reassessed whenever a model is newly trained or substantially modified.

7.2 The Starting Point of the Lifecycle: Large-Scale Pre-training

To determine when provider obligations attach, one must first establish where the lifecycle of a model begins. In its guidelines, the Commission takes a broad view of the model and its lifecycle. The various development stages of a model - for instance, main training following initial small experimental training runs, followed by fine-tuning for specialization or refinement - are considered not as separate models, but as activities occurring within a single lifecycle. The point at which that lifecycle begins is identified by the Commission as large-scale pre-training. This is the moment when the foundational training performed on large amounts of data to build the model's general capabilities commences.

There is a reason why this starting point is critical: certain obligations must be in place from this very moment. The publicly available summary of training content required under Article 53(1)(d) - a document that must comply with a template established by the AI Office - cannot be reconstructed retroactively unless records of what data went in are kept from the moment training begins. The same applies to the copyright policy. To comply with the text and data mining exception under EU Copyright Directive 2019/790, a policy must already be established to filter out works with reserved rights at the exact moment data is scraped. The obligation to draw up and keep up to date technical documentation under Article 53(1)(a) and (b) must also be maintained across the entire lifecycle; failing to build documentation habits from the early stages of training will only accumulate debt over time.

Regarding systemic risk, there is a requirement that goes one step further. The Commission considers that even before large-scale pre-training begins - that is, at the planning stage - a provider can reasonably predict whether its model will reach the systemic risk threshold of 10^{25} FLOPs. Thus, the notification under Article 52(1) of the AI Act - which must be made without delay and at the latest within two weeks - must occur at the point when reaching the threshold is anticipated, rather than after training is completed. Notifying after training is finished and observing the result is too late.

Therefore, if a Korean company is planning large-scale pre-training, it must establish two things before breaking ground on training: a record-keeping system to log what data was ingested, when, and in what volume, and a copyright policy defining what data to exclude. Detailed strategies for both will be examined in Chapter 8. The AI Act itself entered into force on August 1, 2024, and the obligations for GPAI model providers have

applied since August 2, 2025; thus, a company preparing for large-scale pre-training at this very moment is already within the clock of legal obligations.

7.3 Fine-tuning and Downstream Modifications: When Does One Become a New Provider?

Here lies the question raised by the Daejeon startup that opened this chapter: when a company takes someone else's model and fine-tunes it - that is, re-trains and modifies it with its own data - does that company become a new provider of that modified model?

The Commission's answer is clear: not every modification of a model turns a downstream modifier into a new provider. This aligns with the principles of the EU Blue Guide. The Blue Guide states that a product subjected to important changes or substantial modifications intended to alter its original performance, purpose, or type may be considered a new product. The same logic applies to GPAI models. A downstream modifier becomes a new provider of a modified GPAI model only when the modification brings about a significant change in the model's generality, capabilities, or systemic risk.

A point must be made precise here: the criterion determining whether one becomes a new provider is whether a significant change occurred across these three dimensions - generality, capabilities, and systemic risk. Whether the output of that modification serves an independent intended purpose or was created for a distinct use case separate from the original model is not a criterion for this determination. In practice, these two concepts are often conflated. It is easy to assume that because a fine-tuned model acquired a narrow, specific purpose such as diagnosing chemical equipment, it has become a distinct product and therefore makes the modifier a new provider. However, the Commission's criterion evaluates not what the output is used for, but how drastically the modification altered the original model. Even if fine-tuned for a narrow purpose, if there was no significant change in generality, capabilities, or risk profile, the modifier is not a new provider. Conversely, even if a modification aimed at a broad purpose, if it substantially altered the model's generality or capabilities in practice, the modifier can become a new provider. The balance weighs the magnitude of the change, not the breadth or narrowness of the application.

How should practitioners gauge whether a change is significant? I suggest asking three questions: Has the modified model acquired the ability to perform new tasks the original model could not do, or does it perform the same tasks as the original with only improved accuracy? Have its benchmark scores jumped substantially into a different performance tier compared to the original, or is it an improvement within the margin of error? Have safeguards been bypassed or new risky behaviors emerged, or has it inherited the safety

profile of the original intact? If the answer is yes to all three questions, it leans toward a significant change; if no, the opposite. The narrow application of chemical equipment diagnosis itself does not factor into any of these three questions.

At this point, Korean companies often harbor a misconception: the idea that because they did not build the original model, they bear no responsibility. The law does not look for the person who created the model; it looks for the person who places it on the market. If a company takes someone else's model, substantially modifies it, and places it on the market for European customers under its own name, a significant portion of responsibility for that placement shifts to the modifying company. However, becoming a new provider does not mean inheriting all the obligations of the original provider wholesale. The obligations under Article 53(1) of the AI Act - technical documentation, copyright policy, and training data summary - are limited to that specific modification or fine-tuning. The documentation and information required relate only to the modification process itself and the data used in that process. There is no obligation to re-document the training process of the original model itself. Nevertheless, there is a necessary flow of information that must be obtained from the original provider, the mechanics of which are addressed in Chapters 8 and 20.

If the original model was a GPAI model with systemic risk, and the downstream modification satisfies the criteria for becoming a new provider, the modified model is also presumed to be a GPAI model with systemic risk. In this case, the downstream modifier must notify the AI Office without delay and at the latest within two weeks, and assumes new, heightened obligations such as model evaluation, adversarial testing, risk mitigation, reporting serious incidents, and ensuring cybersecurity. Furthermore, if a downstream modifier located outside the EU becomes a new provider in this manner, it must appoint an authorized representative within the EU pursuant to Article 54 of the AI Act. The moment the Daejeon startup supplies its model to German companies, this obligation becomes directly incumbent upon it.

7.4 Discussion of the One-Third Threshold and Practical Judgment

Regarding fine-tuning, one specific figure frequently comes up in practice: the value of one-third. As an indicative criterion for assessing significant change, the Commission cites instances where the training compute used for modification exceeds one-third of the compute used to train the original model. The underlying rationale behind this threshold is straightforward: if modification consumes compute of that scale, it is considered highly probable that the modification brought about a meaningful change in the model's characteristics and performance.

This one-third benchmark must not be misunderstood. It is not an automatic toggle for whether one becomes a provider. The Commission itself admits that this criterion is forward-looking and that its approach may evolve alongside technological and market developments. Even if the compute spent on fine-tuning does not exceed one-third, a modifier can still become a provider if the modification significantly alters generality or capabilities; conversely, exceeding one-third does not automatically make the model one with systemic risk. It is merely an indicator, not a legally binding threshold.

The real sticking point in practice lies elsewhere: cases where the downstream modifier has no knowledge of the original model's training compute. If the original provider does not share that information, there is no basis upon which to calculate the one-third benchmark itself. In such situations, the AI Office may require the downstream modifier to demonstrate compliance through other appropriate means. Lacking information does not exempt one from obligations; rather, it can shift the trajectory toward a heavier burden of proof. Therefore, any Korean company utilizing another party's model is advised to request training compute information from the original provider as early as the contracting stage. Claiming ignorance after the fact will not serve as a defense.

To speak candidly, the precise relationship between this one-third threshold and a significant change has not yet been solidified through administrative practice. What constitutes a significant change and how the one-third value will be applied in individual cases will become clearer as enforcement precedents accumulate. Drawing a definitive line at this very moment is difficult. Therefore, I recommend that any company building a business on fine-tuning maintain self-generated records of how much its modifications changed the original, what proportion of compute was used relative to the original, and what data was utilized during that process. Should a dispute later arise over whether the company is a provider, those records will serve as the evidentiary basis for determination.

The AI Office's power to impose fines takes full effect on August 2, 2026. It is worth keeping in mind that the present moment as you read this chapter is not far from that date.

Questions to Determine Which Side Our Company Falls On

Is the model we created widely used across multiple tasks, or does it perform only a single narrow task? If it is widely used and its training compute is around 10^{23} FLOPs or higher, the threshold of significant generality must be re-examined.

Are we planning large-scale pre-training? If so, before breaking ground on training, a data logging system and a copyright policy must be established. If it appears likely to reach the systemic risk threshold, notification must also be prepared in advance.

Are we fine-tuning someone else's model and placing it on the market in Europe? We must examine whether the modification significantly altered the model's generality, capabilities, or systemic risk. The mere fact that the resulting product is used for a narrow purpose is not a criterion for judgment.

Does the fine-tuning compute exceed one-third of the original? If so, the metrics regarding systemic risk must also be examined. If the original compute is unknown, a clause requiring the provision of that information should be added to the contract immediately.

Key relevant dates collected separately: Entry into force of the AI Act: August 1, 2024. Application of GPAI model provider obligations: August 2, 2025. Initial publication of guidelines: July 18, 2025; formal adoption: November 19, 2025. Grace period for compliance for models placed on the market prior to August 2, 2025: until August 2, 2027. Full activation of the AI Office's fine-enforcement powers: August 2, 2026.

Returning to the Daejeon startup: the company claimed it had merely customized someone else's model. Yet it has not measured how much that customization actually altered the model's generality and capabilities, nor what fraction of the original compute was used. Without measuring, it has already delivered the model to two European corporate clients. If it were up to me, I would tell this company what it needs to do right now: measure the numbers first before reaching a conclusion; only after those numbers emerge can one state whether the company is a user or a provider.

Chapter 8. Basic Obligations of General-Purpose AI Providers

I received a phone call from the head of the legal team at a startup in Pangyo. He said a contract to supply their proprietary language model to a German manufacturer was in its final stages, but the counterpart's legal team had inserted an annex into the agreement. Its title was "Confirmation of Compliance with Article 53 of the AI Act." Facing an unfamiliar article number, he asked: "Our model isn't even that large - do we really have to comply with this?"

The answer is generally yes. The European Union Artificial Intelligence Act (Regulation (EU) 2024/1689, hereinafter the "AI Act") entered into force on August 1, 2024, and the obligations for providers of general-purpose AI models (hereinafter "GPAI models") began to apply on August 2, 2025. Article 53 applies without exception to all providers that develop GPAI models and place them on the EU market, whether a startup in Pangyo or a giant tech firm in Silicon Valley. Unlike Article 55, which applies only to large models with systemic risk, Article 53 attaches the moment a model carries the label of a GPAI model.

Article 53(1) lays down four obligations: point (a) requires drawing up and keeping up-to-date technical documentation in accordance with Annex XI; point (b) requires providing information in accordance with Annex XII to downstream providers who intend to integrate the model into their systems; point (c) requires establishing a policy to comply with EU copyright law; and point (d) requires drawing up and making publicly available a summary of the content used for training, in accordance with a template provided by the AI Office. The first two are documents handed over to authorities or contractual counterparts, while the latter two are a policy and summary released to the public. The confirmation requested from the Pangyo legal team lead was essentially a document demanded by the counterpart under point (b) to fulfill their own obligations.

8.1 Technical Documentation: What Annex XI Requires

Annex XI establishes the framework for documents that must be prepared before placing a model on the market. Starting from general information such as the name, version, release date, distribution channels, and licensing conditions, it requires detailing the architecture, parameter size, modalities (text, image, audio, etc.), the datasets used for training and their extent, how data was collected, cleaned, and filtered, whether evaluation was conducted internally or by external entities, as well as known limitations and energy consumption during training. While the name "technical documentation" sounds modest, in practice it resembles a comprehensive growth record starting from the model's inception.

The starting point of this record is often a source of confusion in practice. The AI Office views the lifecycle of a GPAI model as commencing at the onset of a large pre-training run - that is, the foundational training conducted with vast amounts of data to establish the model's general capabilities. Subsequent fine-tuning or post-enhancement efforts are not treated as creating a new model, but as subsequent stages within the same model's lifecycle. Training computation is also counted cumulatively, factoring in not only pre-training but also the compute expended on synthetic data generation and fine-tuning. Merely fine-tuning an overseas open-source model does not automatically exempt a provider from these obligations.

Technical documentation must be kept up-to-date throughout the lifecycle, be immediately accessible for submission upon request by authorities, and be retained for ten years after placing the model on the market. Under Article 78 of the AI Act, recipients of the documentation are subject to confidentiality obligations covering intellectual property rights and trade secrets, and must maintain cybersecurity measures; however, one should not rely excessively on this provision. The moment information enters the hands of authorities, an additional leak vector is created, and the legal team must draw an internal line regarding how detailed the disclosures should be.

8.2 Information Provided to Downstream Providers: Annex XII

Article 53(1)(b) looks in a slightly different direction. While the technical documentation under Annex XI is directed at regulatory authorities, the information under Annex XII is directed at downstream system providers who intend to take the model and incorporate it into their own products. The German manufacturer partnering with the Pangyo startup stands precisely in this position. Planning to integrate the Pangyo model into its industrial AI system, the company must receive information on the model's capabilities and limitations, predictability, reliability, reproducibility, and safety in order to fulfill its own obligations under the AI Act. Without understanding the integrated model, it is impossible for downstream providers to satisfy their obligations; thus, the provision of information by GPAI providers serves as a cornerstone supporting the entire value chain.

Although the contents overlap with Annex XI, the focus differs. Information directly relevant to technical integration - such as how to integrate the model into a system, required hardware and software prerequisites, and input/output formats - takes center stage. This information is likewise subject to confidentiality obligations under Article 78. If downstream providers determine that a GPAI provider has supplied inadequate information or failed to provide it altogether, they reserve the right to lodge a complaint pursuant to Article 89(2). The confirmation received by the Pangyo legal team lead was essentially intended by the counterpart to preemptively mitigate this liability, and signing a blank form is a fundamentally different matter from signing it with actual, complete information prepared.

For Korean companies, this obligation intersects with extraterritorial application. If a provider headquartered outside the EU places a GPAI model on the EU market or if its output is used within the EU, it falls under the scope of the law pursuant to Article 2(1)(c) of the AI Act, and must designate an Authorised Representative in writing within the EU pursuant to Article 54 or Article 22. Because the representative often serves as the contact channel for providing information, it is advisable to prepare the information package alongside the designation of the representative.

8.3 EU Copyright Compliance Policy: Reservation of Rights under Article 4(3) of the Directive

Article 53(1)(c) requires GPAI providers to establish a policy to comply with EU copyright law and related legislation. At its core is Article 4 of the Directive on Copyright in the Digital Single Market (Directive (EU) 2019/790, hereinafter the "DSM Directive"). Article 4(1) provides an exception for text and data mining (TDM), meaning that scraping copyrighted works in bulk to train models does not, in principle, constitute copyright infringement. However, Article 4(3) attaches a condition: if rights holders have expressly reserved their rights in an appropriate manner, the TDM exception does not apply to those works.

This is where practical execution diverges. For content made publicly available online, a reservation of rights may be indicated using machine-readable means, such as robots.txt files or metadata, and GPAI providers must employ state-of-the-art technologies to detect and honor such designations. Bypassing paywalls restricting access, breaching technological protection measures under Article 6(3) of the DSM Directive, or crawling websites that do not make content available to the public are strictly prohibited. Unless a provider can demonstrate via logs that its crawlers actually read and filtered out such indications, the compliance policy remains a mere paper declaration.

While this policy must be applied throughout the entire lifecycle of each relevant model, a single policy may be created and applied across all models. Where an entity fine-tunes or modifies an existing GPAI model to become a new provider, the copyright policy obligation is confined strictly to the data used for that modification and fine-tuning, and does not require tracing back responsibility to the original model's training data.

This is where the Code of Practice comes into play. Drawn up by the AI Office pursuant to Article 56 of the AI Act, the final version of the GPAI Code of Practice was released on July 10, 2025, comprising three chapters: transparency, copyright, and safety and security. Twenty-six organizations - including Amazon, Anthropic, Google, IBM, Microsoft, and OpenAI - signed all chapters; xAI signed only the safety and security chapter; while Meta was the sole entity to refuse signing altogether, citing legal uncertainty. Signing provides a streamlined pathway to demonstrate compliance and reduces regulatory burden from the AI Office. However, as the AI Office itself has explicitly noted that adherence to the Code of Practice does not affect the application and enforcement of copyright and related laws, it does not serve as a shield against national copyright legislation in Member States.

8.4 Public Summary of Training Content: The Boundary Between Templates and Trade Secrets

Article 53(1)(d) requires providers to draw up and make publicly available a sufficiently detailed summary of the content used for training the GPAI model. On July 24, 2025, the European Commission released an Explanatory Notice along with a standardized template. The template is divided into three sections: general information identifying the model and provider, the types and sources of data used in training, and the data processing procedures. Under data sources, categories are itemized into large-scale public datasets, data gathered via web crawling, third-party or licensed data, and user-generated or synthetic data.

The objective of this obligation lies in balancing two competing interests: copyright holders and regulatory authorities have a right to know what fed the model's development, while providers possess the right to protect their trade secrets. Although Article 78 of the AI Act imposes an obligation to protect intellectual property rights, confidential business information, and trade secrets, this does not mean the public summary itself may be filled with omissions under the guise of secrecy. Specific details intended to be kept as trade secrets should not be placed in the public summary, but rather submitted in separate documentation provided exclusively to the AI Office or competent authorities of Member States. Since the law does not draw a precise boundary between what to disclose and what to withhold, I advise each company to establish internal guidelines to define this threshold in advance.

Exceptions are recognized where training data information is no longer available or where retrieving such information would impose a disproportionate burden. For models placed on the market prior to August 2, 2025, providers are not required to retrain models or undo past work; however, the absence of information or the existence of a disproportionate burden must be explicitly disclosed and justified within the copyright policy and the summary of training content. Remaining silent about ignorance and formally disclosing lack of knowledge occupy fundamentally different legal ground.

Enforcement and Timeline

The AI Office's authority to impose substantial fines on GPAI providers takes effect on August 2, 2026. Until then, it focuses on exercising its power to request information (Article 91) to assess providers' compliance plans and encourage cooperation. Once fine

enforcement begins, violations of GPAI obligations, including Article 53, may incur administrative fines under Article 101 of the AI Act of up to 3 percent of total worldwide annual turnover for the preceding financial year or 15 million euros, whichever is higher.

Providers of GPAI models already placed on the market before August 2, 2025, are granted a 24-month grace period until August 2, 2027. On November 19, 2025, in officially adopting guidelines interpreting the scope of GPAI obligations, the AI Office clarified that derivative models subject only to minor modifications will not be treated as introducing a separate provider each time; provider obligations shift only to entities that make substantial modifications. Executing a few lines of fine-tuning does not instantly transform an entity into a GPAI provider, though where that boundary lies must be confirmed on a case-by-case evaluation.

Reference Materials

The timeline is as follows: entry into force on August 1, 2024; publication of the final Code of Practice on July 10, 2025; release of the training content summary template on July 24, 2025; start of application of GPAI obligations on August 2, 2025; adoption of guidelines interpreting the scope of obligations on November 19, 2025; activation of fine enforcement powers on August 2, 2026; and expiration of the grace period for existing models on August 2, 2027.

The checklist is as follows: verify whether your model qualifies as a GPAI model and determine when its large pre-training run commenced. Prepare a draft of the technical documentation meeting Annex XI standards and establish a ten-year retention system. Separately prepare the Annex XII information package to be handed to downstream providers alongside contractual documents. Audit your web crawling pipelines to confirm whether they actually read robots.txt standards and rights reservation metadata. Draft the public summary of training content in accordance with the template and document internal criteria distinguishing disclosed from non-disclosed information. Confirm whether an Authorised Representative has been designated within the EU, and ensure the legal team understands that signing the Code of Practice does not substitute for compliance with copyright law.

The Pangyo legal team lead ultimately signed the confirmation. However, he took three extra days before signing. Those three days were spent sitting down with the development team, reviewing line by line whether they possessed the information required under Annex XII. Without those three days, that signature would have been

nothing more than ink on a piece of paper.

Chapter 9. Systemic Risk Models

9.1 The 10^{25} FLOP Threshold and Compute Estimation

I once received a call from an AI startup in Pangyo. They were in the middle of training their next-generation language model, and after calculating their GPU cluster usage time, they realized that near the end of training, they would likely cross a line the company had not anticipated. That line was precisely 10^{25} FLOPs - 1,000 trillion times 10 quadrillion floating-point operations. The number might be hard to grasp intuitively. If you consider that the training compute for GPT-4-class models is estimated to be near or above this level, the scale might become a bit clearer.

Article 51(2) of the AI Act (Regulation (EU) 2024/1689) stipulates that a general-purpose AI model - General-Purpose AI model in English, abbreviated as GPAI model - is presumed to have high-impact capabilities if the cumulative compute used for its training exceeds 10^{25} FLOPs. Systemic risk, defined in Article 3(65), refers to a risk stemming from the capabilities of state-of-the-art GPAI models that has a significant impact on the EU market and can propagate across the entire value chain. A model exceeding this threshold is presumed to be a model with systemic risk, and its provider bears heightened obligations under Articles 52 and 55.

Here was a point that the CEO of the Pangyo startup had misunderstood. When the threshold is reached or anticipated to be reached, notification must be given without delay and at the latest within two weeks, and they believed the recipient was the AI Office. I read this provision differently. The recipient of the notification set forth in Article 52 is the Commission, not the AI Office. Although the AI Office is a newly created organization under the Commission serving as a practical working channel, the text of the provision itself explicitly specifies the recipient as the Commission. Because a considerable portion of practical reference materials uses these two interchangeably, confusion has spread, but knowing where to send written submissions is just as important in practice as the amount of the fine.

The notification must state the compute that triggered reaching the threshold using two significant figures, accompanied by an explanation of the methodology used to estimate that compute. Although a path remains open to challenge the classification by providing evidence that meeting this threshold does not actually generate systemic risk, the Commission has made clear that the burden of proof rests on the provider. It is by no means an easy procedure.

Let me highlight one more date. The obligations of GPAI model providers have already been applicable since August 2, 2025. And on August 2, 2026 - nine days from the time of writing this - the Commission's power to impose fines will take effect. In case of a violation, fines can be up to 3 percent of total worldwide annual turnover for the preceding financial year or 15 million euros, whichever is higher. No company has received a fine notice yet. Starting next week, things will be different.

The number 10^{25} itself is not a fixed value either. Article 51(3) provides that the Commission is empowered to amend the threshold through delegated acts in accordance with technological progress. This means that a company that appears safe under today's standards could see its status change in the next revision.

For a company headquartered in Seoul, the moment the output of its model is used within the EU (Article 2(1)(c) of the AI Act), all of this discussion is no longer someone else's affair. A compute estimation framework should have been established prior to August 2, 2025, and if one has not been built yet, it is not too late to start now. However, you must hurry.

9.2 Hardware-Based Estimation and Architecture-Based Estimation

Let us return to the Pangyo startup case mentioned earlier. The company's engineering team lead was at a loss, not knowing how to estimate the compute. The AI Office guidelines offer two approaches.

One is the hardware-based approach. This method extracts compute through the formula $C \text{ (FLOPs)} = N \times L \times H \times U$ by multiplying four values: the number of GPUs or other hardware units used (N), total usage time (L , in seconds), the theoretical peak performance of the hardware unit (H , FLOPs per second), and the average GPU utilization rate (U). If a mixture of different types of GPUs was used, H must be set using a weighted average. This was precisely where the Pangyo team lead ran into difficulty. They used A100s in the early stage of training and H100s in the later stage, and because the theoretical performance difference between the two devices was several-fold, determining how to calculate the weighted average was the key issue.

Let us look at this with numbers. Assuming 8,000 H100 GPUs were operated for 30 days at an average utilization rate of 40 percent, N is 8,000, L is approximately 2.59 million seconds (converting 30 days to seconds), H is the theoretical peak performance of a single H100 GPU of roughly 1,000 TFLOPs per second, and U is 0.4. Multiplying these four values yields approximately 8.3×10^{24} FLOPs. Since this is one order of magnitude short of the threshold of 10^{25} , one might feel relieved, but adjusting the assumed utilization rate to just 50 percent causes the calculated result to exceed 1.0×10^{25} . This means a single utilization rate cell in an Excel sheet determines whether a notification obligation is triggered.

The other is the architecture-based approach. This method multiplies the number of forward and backward passes performed by the neural network during training by the amount of computation in a single full pass. In practice, the approximation formula $6 \times P \times D$ is widely used. P is the number of parameters, and D is the number of training tokens. Although this approximation formula is simple to calculate, applying certain optimization techniques may cause it to deviate from actual compute.

From the perspective of enforcement authorities, it matters not which of these two methods is used, but the crucial point is that which method was used must be disclosed in the notification. Concealing or glossing over the methodology itself becomes a clue of non-compliance. The AI Office approved the contents of these guidelines on July 18, 2025,

and formally adopted them on November 19 of the same year (C(2025) 7719). Although non-binding interpretative guidance, it is a de facto roadmap providing advance notice of how the Commission will actually interpret the provisions. Because the estimation methods in the guidelines themselves are explicitly noted to be illustrative, there is as yet no judicial precedent on which side the authorities will favor in specific situations.

Both methods share one premise: that all compute used to enhance model capabilities - including synthetic data generation and fine-tuning in addition to pre-training - is included in cumulative compute (Recital 111). In practice, there are companies that miss instances where compute during the fine-tuning stage pushes the total over the threshold. I have actually heard of cases where companies felt reassured after measuring only pre-training compute, only to cross the threshold once fine-tuning volume was added.

In practice, I offer one piece of advice: rather than selecting and using only one of the two methods, cultivate the habit of calculating with both methods from the early stages of model development and comparing them against each other. If the two numbers diverge significantly, that discrepancy itself becomes something to explain in the notification. When building a compute logging system at a Seoul headquarters, recording logs across the entire development lifecycle serves as evidence whether you later file an objection or submit a notification.

9.3 Model Evaluation and Adversarial Testing

What happens after exceeding the compute threshold is even weightier. Article 55(1)(a) of the AI Act requires providers of GPAI models with systemic risk to implement measures including model evaluation, specifically adversarial testing - commonly referred to as red-teaming. Think of it as a process of hiring ethical hackers to actually attack your model. It is a demand to continuously examine throughout the entire lifecycle of the model whether it produces harmful output or has potential for misuse.

This obligation has already been applicable since August 2, 2025. However, providers of models already placed on the EU market prior to that date are granted a 24-month grace period until August 2, 2027. You must not miss the difference in application timing between newly introduced models and existing models.

From the perspective of enforcement authorities, what makes this clause tricky is that standards for what constitutes sufficient evaluation have not yet matured. The AI Office itself admits that the external model evaluator ecosystem is still in a developing stage and that their methodologies may initially be inconsistent. Nevertheless, determinations are made based on the state of the art. I read this clause as follows: the lack of standards does not constitute a ground for exemption. Rather, precisely because there are no standards, documenting your self-established criteria and their grounds will put you in a far more advantageous position later.

Whether to build internal red-teaming capability or delegate to external specialized organizations is also a frequent question in practice. While outsourcing may seem convenient for now, given that the AI Office has publicly mentioned the immaturity of the external evaluator ecosystem, your credibility could be questioned later depending on which institution you choose. I would advise keeping core red-teaming capabilities internal and supplementing with external experts only for specific domains (such as biosecurity or cyberattack simulation).

If information is insufficient or if a warning is received from the scientific panel, the AI Office also has the authority to directly conduct evaluations. A poor internal red-team report means the authorities will inspect it themselves, so the quality of the report effectively determines the intensity of supervision.

The sequence that a Seoul headquarters company must follow is as follows. First, confirm before August 2, 2025 (at this point, since that date has already passed, immediately)

whether your model is a GPAI model meeting the substantial generality criteria and whether its compute exceeds the threshold. Next, embed an evaluation and adversarial testing framework covering the entire model lifecycle starting from the design stage. This approach, known as compliance by design, costs significantly less than retrofitting later.

9.4 Serious Incident Reporting

If red-teaming in the previous section is ex-ante preparation, this section is about the story after an incident actually occurs. Article 55(1)(c) requires providers of GPAI models with systemic risk to track, document, and report information regarding serious incidents, but this time the recipient is different: the AI Office and, where necessary, the competent authorities of the relevant Member States. Do not confuse this with notification to the Commission covered in 9.1. Notification of reaching the compute threshold goes to the Commission, while serious incident reporting goes to the AI Office. Even under the same regulation, the destination for written submissions diverges.

Serious incidents as defined by the AI Office fall into two categories. One is a serious cybersecurity breach affecting the model or its physical infrastructure, including leaks of model parameters or cyberattacks. The other is when a malfunction or accident of a GPAI model directly or indirectly leads to a serious incident concerning an AI system as defined in Article 3(49)(a) through (d). However, the specific content of points (a) through (d) of paragraph 49 could not be confirmed from the materials I secured. It is better to write clearly that what is uncertain is uncertain. This is a space to be filled as additional guidance comes from the Commission in future revised editions.

In November 2025, the Commission officially issued the serious incident reporting template for GPAI models with systemic risk. Available in Word and PDF formats, this template has established itself as a document demonstrating compliance with Commitment 9 of the GPAI Code of Practice, and submission is carried out via the EU SEND platform. A requirement to keep safety and risk management activity records for at least 10 years comes attached as well. The release of the template also means that the excuse of being unprepared will no longer hold up easily.

The AI Office has made clear that it expects proactive reporting - contacting them before an incident actually breaks out. This may be part of commitments under the GPAI Code of Practice or an alternative mode of proving compliance, and it exists separately from the obligation to report serious incidents under Article 55(1)(c). My judgment diverges on this point. Separating legal obligations from practical recommendations, the obligation is to report after an incident occurs, but considering relations with authorities, reaching out first at the stage where signs of an incident appear is advantageous in lowering the intensity of subsequent investigations. Regulatory relations are closer to building trust than to assigning penalty points.

A Korean enterprise headquartered outside the EU must designate an Authorised Representative established within the EU in writing (Article 54 of the AI Act) prior to placing a GPAI model on the EU market. This representative must act as a channel to the AI Office and be capable of representing the company in reporting serious incidents. Because which competent authority in a relevant Member State applies is often framed around the Member State where the representative is established, where you set up your representative effectively determines which country's authority you will sit across from. I recommend checking right now whether incident reporting duties are explicitly stated in your representative agreement.

9.5 Securing Cybersecurity and the July 2026 Cybersecurity AI Action Plan

Article 55(1)(d) requires ensuring an adequate level of cybersecurity protection for the model and its physical infrastructure throughout its lifecycle. Among the serious incidents discussed in the previous section, severe cybersecurity breaches connect directly to this obligation. A point frequently overlooked in practice is that the target of protection includes not only the model itself but also the servers and networks supporting that model.

At the time of writing this manuscript, what must not be missed has already occurred. On July 7, 2026, the Commission released the Cybersecurity AI Action Plan (Communication COM(2026) 577 final). This announcement came around the time I was recording the Pangyo startup case in 9.1, making the timing exquisitely coincidental. This plan is not a new law. It does not place new obligations on companies. It is accurately understood as a coordination framework adding funding and market opportunities on top of the AI Act, NIS2, and the Cyber Resilience Act already in effect.

The core of the Action Plan boils down to three points. First is developing EU-level capacity for third parties to evaluate the capabilities and risks of AI models before they hit the EU market, a measure supporting the AI Office's regulatory function. Second is an open-source software resilience campaign, where ENISA, the Commission, Member States, and the open-source community jointly support the maintenance and security of critical open-source projects. Third is funding. Around 200 million euros have already been injected through Horizon Europe and the Digital Europe Programme until the end of the current Multiannual Financial Framework, with a target set to achieve 100 million euros in investment in cybersecurity and AI startups through the European Innovation Council Fund by the end of 2026.

Reading this Action Plan from the perspective of enforcement authorities makes its meaning explicit. Starting August 2, 2026, the Commission begins exercising supervisory powers over GPAI models, explicitly targeting models that generate systemic cybersecurity risks. The date when fine imposition authority begins and the target targeted by the Action Plan align precisely. This appears to be an intended order rather than a coincidence.

For a Seoul headquarters company, there are two points to note in this Action Plan. One is that pathways are open to participate in calls for building third-party evaluation capacity,

and the other is that models relying heavily on open-source components should monitor trends in the open-source resilience campaign. The final version of the GPAI Code of Practice was established in June 2026, and signing this Code carries the benefit of a presumption of conformity. However, certain specific technical standards and harmonised standards under the Code remain yet to be fully published in the Official Journal of the EU.

In 10 days, the Commission's fine clock starts ticking. I advise placing this Action Plan factsheet on the desk of your compliance officer in Seoul. On the morning of August 2, you may find yourself unfolding that factsheet once again.

Reference Materials

Dates. August 1, 2024: AI Act entry into force. August 2, 2025: Application of GPAI model provider obligations begins. July 18, 2025: Draft GPAI guidelines approved; formally adopted November 19 of the same year (C(2025) 7719). November 2025: Serious incident reporting template published. June 2026: Final version of GPAI Code of Practice confirmed. July 7, 2026: Cybersecurity AI Action Plan (COM(2026) 577 final) announced. August 2, 2026: Commission fine imposition authority takes effect; starting point of grace period for existing models. August 2, 2027: Grace period ends for models launched prior to August 2, 2025.

Criteria for Determining Notification Recipients. Threshold compute notification (Article 52): Commission. Serious incident reporting (Article 55(1)(c)): AI Office and, if necessary, competent national authorities of relevant Member States. Explicitly stipulate in internal regulations that destinations for these two documents are different.

Checklist Items. Confirm whether your model is a GPAI model and whether it exceeds the 10^{25} FLOP threshold. Dual-calculate compute using both hardware-based and architecture-based methods. Establish a procedure for notifying the Commission within 2 weeks upon expecting to reach the threshold. Determine internal red-teaming capability and scope of external evaluation agencies. Build a system for defining serious incidents and submitting reporting templates (EU SEND), and retain related records for 10 years. Inspect cybersecurity systems for model and physical infrastructure. Check whether EU Authorised Representative (Article 54) contract includes provisions for delegated notification and reporting. Conduct a separate review for potential dual-regulation conflicts with the domestic AI Framework Act.

Chapter 10. For Companies Seeking to Use the Open-Source Exception

A development team lead from a startup in Pangyo once came to consult with me. He was trying to upload an in-house trained Korean language model to Hugging Face, but said his company's legal team had been wrestling with a single line of license terms for weeks. They wanted to release it for free while keeping the company name attached, and they wanted to restrict commercial use out of a sense of unfairness if a large corporation were to take it and make money off it. However, his expression hardened when he was told that, due to that single condition, it would not be recognized as open source under the European Union Artificial Intelligence Act (Regulation (EU) 2024/1689, hereinafter the "AI Act"). The reason behind that expression is precisely what this chapter addresses.

The AI Act entered into force on August 1, 2024, and the obligations imposed on providers of general-purpose AI models (hereinafter "GPAI models") are concentrated in Chapter V. These obligations began applying on August 2, 2025, while the fine-enforcement authority of the AI Office will take effect on August 2, 2026. Therefore, at the time you are reading this book, the obligations are already active, with only the blade of fines remaining in its sheath for now. GPAI models already placed on the market before August 2, 2025, receive a grace period for compliance until August 2, 2027, but models newly released after that date are subject to immediate application.

Above this heavy pile of obligations, the AI Act has left a single breathing room. It provides an exemption from certain obligations for models released under a free and open-source license. The underlying provisions are Article 53(2) and Article 54(6). The European Commission deemed that such models contribute to research and innovation and offer growth opportunities for the EU economy, judging that if transparency is high with parameters, architecture, and usage information disclosed, relieving some of the burden is justified. However, conditions are attached. And if a company misses even one of those conditions, no matter how genuinely it believes it is distributing open source, it will be treated under the AI Act exactly the same as a paid, proprietary model.

10.1 The Four Freedoms of Free and Open-Source Licenses

The four rights required by the AI Act are access, use, modification, and distribution. If any one of these four is missing, the license is not recognized as free and open-source under the AI Act.

Access means that any stakeholder must be able to freely obtain the model without payment or other restrictions. As long as there is no unfair discrimination based on factors like nationality, reasonable safety and security measures, such as user authentication, may be implemented. This means steps like account registration or email verification are permissible.

Usage is the right ensuring that the original provider cannot use intellectual property rights as a weapon to prevent use or charge royalties. However, restrictions requiring the attribution of derivative works and their distribution under terms identical or similar to the original are permitted. In other words, so-called copyleft conditions are acceptable, but conditions that restrict the intended purpose of use itself are not.

Modification is the right of anyone to alter the model without payment or restrictions. Whether fine-tuning or pruning weights, it is entirely up to the user.

Distribution is the right of those who have accessed, used, and modified the model to redistribute the results under limited conditions. Conditions such as retaining attribution or copyright notices can be attached, and it cannot even prevent a recipient from releasing derivative works under a closed, proprietary license. As I interpret it, keeping the door open for redistribution is the core point itself; one must not attempt to control the licensing form of the redistributed outputs.

These four freedoms closely resemble the definitions refined over decades by open-source communities like the Linux Foundation. The difference is that the AI Act has codified them into explicit legal provisions. If the licensing text fails on even one of these four conditions, no matter how publicly available the source code is on GitHub, it is legally not open source.

A point of confusion needs to be addressed here. In the developer community, "open source" and "open weight" are often used interchangeably. It is common practice to call a model open source as long as the weight files are downloadable, but the AI Act takes a far less loose view. Parameters, model architecture, and usage information must all be

publicly accessible, and on top of that, the aforementioned four freedoms must be guaranteed through explicit license text. The common practice of simply releasing weights while leaving usage conditions ambiguous is insufficient to qualify for the exemption discussed in this chapter.

10.2 Typical Non-Compliance Cases

The clauses that frequently cause issues in practice are predictable. A prime example is the restriction on commercial use sought by the Pangyo development team lead. Clauses stating that a model may only be used for non-commercial or non-profit research purposes are viewed as severely restricting the right of usage. This is because the AI Act considers that a license must remain open for use regardless of the purpose.

The second is a prohibition on distribution. A condition that allows using the model but prohibits redistribution violates the right of distribution itself. The third is a user scale threshold. This involves requiring an additional license if monthly active users (MAU) exceed a certain number, a structure famously controversial in Meta's Llama license family. The fourth is demanding a separate commercial license for specific use cases. The common dual structure - where research use is free, but commercial distribution requires a separate contract - falls prey to this.

Though these four cases carry different names, they ultimately stem from a single root: the psychological desire of providers who are reluctant to relinquish full control, yet feel uncomfortable calling their model closed. I believe it is more accurate to call such licenses "source-available" licenses. Code is visible, but it is not open source - and the AI Act is uncompromising at this point. Its stance is that being visible and being free are two entirely different matters.

These four typical restrictions ultimately impinge upon either usage or distribution among the four freedoms discussed in 10.1. Therefore, when reviewing a license, the fastest method is to test each clause against these four freedoms. In practice, I find that this test alone filters out half the issues.

There is one point to clarify. The obligations exempted by the open-source exception are Article 53(1)(a) and (b) - namely, drawing up technical documentation and providing information to downstream system providers - as well as the obligation to appoint an EU authorized representative under Article 54. In effect, if a non-EU provider satisfies the genuine open-source conditions and presents no systemic risk, it does not need to appoint an authorized representative. Where other parts of this book explain that Korean companies must appoint an EU representative without exception, that was under the assumption of closed models or models with systemic risk; I clarify here that models covered by a pure open-source exception may not be subject to the authorized

representative obligation itself. While that is true textually in the provisions, in practice there is always room for dispute over whether a license truly meets the requirements, so I advise that appointing a representative in advance remains the safer path.

10.3 Practical Application of the Non-Monetization Requirement

Even if a license upholds all four freedoms in its text, there are cases where the exception is lost due to the non-monetization requirement. Recital 103 of the AI Act states that financial remuneration must not be demanded anywhere across access, usage, modification, or distribution of the model. The European Commission interprets this concept quite broadly.

A prominent form that the Commission regards as monetization is dual licensing - releasing a model under a free open-source license while demanding a separate paid license for commercial use. Restrictions to non-commercial research, prohibitions on redistribution, user scale thresholds, and requirements for separate commercial licenses for specific use cases all fall into the same category. It is no coincidence that these overlap with the typical non-compliance cases discussed in 10.2, as those restrictions are ultimately mechanisms to extract revenue.

Here, practitioners often harbor a common misconception: believing that if the model itself is distributed free of charge, it does not constitute monetization. However, the Commission notes that revenue models seemingly separate from the model - such as consulting, support services, or cloud hosting fees - may also be scrutinized on a case-by-case basis. The Commission's phrasing is that assessments are always made on a case-by-case basis, and whenever I read this line, I take it as a signal that the grey zone is wide. It is premature to conclude that common open-source business models - where the model itself is free but API access is charged - lie safely in the green zone.

To offer a yardstick: whether money is charged for the model itself - that is, the weights and code - serves as the primary criterion, and charging for ancillary services that assist in using the model comfortably is generally safe. However, if that ancillary service serves as virtually the sole gateway for accessing the model - for instance, if weights are theoretically public but practically usable only through a paid API operated by the company - this is tantamount to gating the right of access behind money. On this boundary line, I advise erring on the side of conservative judgment as model size grows.

There is a crucial point that must be underscored here. Even if a license respects all four freedoms and involves no monetization, the open-source exception does not apply at all to GPAI models presenting systemic risk. Article 53(2) explicitly provides for this. The benchmark for presuming systemic risk is when the cumulative computation used for

training exceeds 10^{25} floating-point operations (10^{25} FLOPs). Once this threshold is crossed, fully releasing the source code and reaping zero commercial gain will not help; the provider must comply with all heightened obligations under Articles 52 and 55 without exception. Companies whose training computation approaches this threshold must consult a separate chapter regardless of open-source status - a topic already covered in Chapter 9.

10.4 Two Obligations That Remain Despite the Exception

Even if a provider fully qualifies for the exception by upholding the four freedoms, avoiding monetization, and presenting no systemic risk, two obligations adhere to the very end: the obligation to establish a copyright policy and the obligation to draw up and disclose a summary of training content.

Article 53(1)(c) requires establishing a policy to comply with EU copyright law and related regulations. This includes procedures to identify and respect the reservation of rights under Article 4(3) of the Directive on Copyright and Related Rights in the Digital Single Market (Directive (EU) 2019/790, DSM Directive). Put simply, if creators have marked their works in a machine-readable format indicating they should not be used for AI training, state-of-the-art technology must be deployed to recognize and filter out those markings. Technical protection measures such as paywalls must not be bypassed, and this policy must apply throughout the entire lifecycle of each model created by the company. It is not necessary to write a new policy for every model; a single policy may span multiple models.

Article 53(1)(d) mandates drawing up and publicly releasing a sufficiently detailed summary of content used for training according to a template defined by the AI Office. On July 24, 2025, the AI Office issued an Explanatory Notice along with a standardized template, which sets the minimum baseline for what must be included in the public summary. While stated to balance trade secret protection against the right to know of copyright holders and regulatory authorities, where to draw the line on disclosure will inevitably vary across companies. At this juncture, I advise establishing internal guidelines in advance regarding non-disclosure scope. While relying on Article 78 obligations of professional secrecy, narrowing the scope too aggressively may lead the AI Office to question the fidelity of the summary itself.

The reason these two obligations were excluded from the exception is clear: releasing a model under a free open-source license does not automatically render transparent what the model was fed for training or how copyright was managed. Even when source code is open, the provenance of training data frequently remains opaque. The Commission retained these two obligations regardless of open-source status precisely to plug that gap.

Stated differently, the exception in this chapter is a partial exemption, not a full exemption. While drawing up technical documentation or appointing an authorized representative may be excused, the two documents - the copyright policy and the training content summary - must be prepared identically whether open source or closed. I often see practitioners at Korean companies mistakenly assume that going open source frees them from paperwork altogether, but that is only half true.

Perspectives of Enforcement Authorities and Fines

The AI Office will supervise compliance by GPAI model providers starting August 2, 2025, and exercise authority to impose fines for violations beginning August 2, 2026. The Commission first released draft guidelines on the scope of GPAI model obligations on July 18, 2025, and the official Communication C(2025) 7719 final containing the same content was adopted on November 19, 2025. There was a four-month gap between the draft and the formal text, during which practice moved based on the draft. I view these guidelines as filling in the margins of the legal provisions, though ultimate legal certainty will only take shape as case law builds up at the Court of Justice of the European Union (CJEU).

The AI Office has signaled that it will interpret the concept of monetization broadly. It intends to scrutinize not only dual licensing and demands for additional licenses following usage restrictions, but also indirect revenue-generating activities. Given the intent to prevent abuse of the open-source exception, simply cleaning up license text while keeping actual revenue structures unchanged is unlikely to succeed. If a violation is established under Article 101, fines of up to 3% of total worldwide annual turnover in the preceding financial year or 15 million euros, whichever is higher, may be imposed.

Frankly speaking, many of the criteria discussed in this chapter have not yet solidified into established case law. Guidelines merely indicate direction and do not adjudicate individual license clauses line by line. Compute estimation methodologies also allow a 30 percent margin of error whether hardware- or architecture-based, meaning companies near the 10^{25} FLOP threshold may face disputes over the calculation methodology itself. Until CJEU jurisprudence accumulates, I consider it safer to interpret conservatively and act accordingly.

What Korean Companies Must Do Now

If a company headquartered in Seoul develops a GPAI model and places it on the EU market, or if its output is used within the EU, the contents of this chapter apply directly

pursuant to Article 2(1)(c). Here, in no particular order, is what practitioners must check now:

First, assess whether your model constitutes a GPAI model and whether its training computation approaches the 10^{25} FLOP threshold. If it crosses or is likely to cross that threshold, no amount of attractive open-source packaging will grant an exception.

Next, evaluate the original license text line by line against the four freedoms of access, usage, modification, and distribution. Check whether restrictions on non-commercial use, distribution prohibitions, MAU thresholds, or commercial licensing requirements for specific uses are lurking within.

Even if the license text is pristine, re-examine whether the actual revenue structure leaves room to be interpreted as monetization. Surrounding revenue models such as support services or hosting fees must also be audited.

Ensure the organization clearly understands that even under the premise of fully qualifying for the open-source exception, the two tasks of establishing a copyright policy and drafting a training content summary remain intact. Neither is solely a job for the legal team; both must be designed collaboratively starting from the data collection stage.

Remember that the requirement to appoint an authorized representative hinges on whether your model genuinely meets the open-source exception conditions. However, if there is potential dispute over meeting those conditions, appointing a representative in advance is, in my view, the path that avoids future complications.

Reference Materials

Criteria for Judging the Four Freedoms Access: Prohibition of payment demands; prohibition of unfair discrimination such as nationality; reasonable authentication procedures allowed. Usage: Prohibition of royalty claims using IP rights; copyleft conditioned on attribution/distribution allowed. Modification: Alteration permitted without payment or restrictions. Distribution: Attribution/copyright notice retention conditions allowed; violating if prohibiting redistribution into closed derivative works.

Typical Clauses Invalidating the Exception The model may only be used for non-commercial, non-profit research purposes. The model may be used, but redistribution is prohibited. A separate license is required if monthly active users exceed a certain number. Specific commercial use cases require entering into a separate contract.

Obligations Not Exempted Article 53(1)(c): Policy establishment to comply with EU copyright law, identification of reservation of rights under Article 4(3) of the DSM Directive. Article 53(1)(d): Drafting a public summary of training content, compliance with the AI Office standardized template.

Timeline August 1, 2024: Entry into force of the AI Act July 18, 2025: Release of draft GPAI guidelines July 24, 2025: Release of standardized template for training content summary August 2, 2025: Application of GPAI obligations begins November 19, 2025: Adoption of official guidelines Communication C(2025) 7719 final August 2, 2026: AI Office fine-enforcement authority takes effect August 2, 2027: Expiration of compliance grace period for models released prior to August 2, 2025

In the end, this is what I told the startup development team lead from Pangyo: If you want to include a non-commercial condition, that is the company's choice, but the moment you do, you must abandon the exception discussed in this chapter. Instead, you must choose one of two paths: either take the closed-source route by diligently preparing technical documentation and appointing an authorized representative, or take the open-source route by completely opening up the conditions while taking on the two tasks of a copyright policy and a training content summary. That day, he fiddled with his business card and remained silent for a long while.

Part IV

High-Risk AI: Preparing by 2027

Chapter 11. Is Our Product High-Risk?

Let's start with the story of a battery management system company located in the Namdong Industrial Complex in Incheon. For convenience, let's call this company Sinheung. Sinheung manufactures battery management systems (BMS) for electric vehicles and supplies them as components to European automakers. The core of their product is safety logic that immediately cuts the circuit if a cell temperature spikes or a voltage deviates. On the afternoon of the same day, a manager at a HR software startup in Songdo also called their legal team. Let me call this company KoreaTalent. KoreaTalent developed an AI recruitment tool that reads applicants' resumes and ranks interview candidates, selling it to a few mid-sized companies in Germany and the Netherlands. Both companies came to me with the exact same question: Is our product high-risk?

While the question seems simple, the answers take completely different paths for the two companies. Sinheung's BMS is a safety component integrated into an existing product - an automobile - whereas KoreaTalent's recruitment tool is a standalone service in its own right. The AI Act treats these two under different articles, different timelines, and different methods of proof. One falls under Annex I, the other under Annex III. Starting from this fork in the road, this chapter places Annex III (split into eight areas) alongside Annex I (embedded within existing products), addresses the burden of proof under the filter clause caught in between, and then applies these rules one by one to six key Korean industries.

11.1 Annex III: Eight Areas Covering Standalone Systems

Article 6(2) of the AI Act stipulates that AI systems listed in Annex III are considered high-risk. The dividing criterion here is not the hardware hosting the system, but where the system's intended purpose is directed. Standalone software offered as an independent service - like KoreaTalent's recruitment tool - is the main target of this annex, which is divided into eight areas.

The first is biometrics. Remote biometric identification systems and biometric categorization systems that infer sensitive attributes fall under this area. However, applications carrying higher risks - such as untargeted scraping of facial images, emotion recognition in workplaces and educational institutions, or biometric categorization inferring race, political opinions, or sexual orientation - move out of Annex III entirely and into the list of prohibited practices under Article 5. The biometric applications remaining in Annex III represent a less severe, yet still heavily regulated category.

The second area comprises safety components used in the safety management and operation of critical digital infrastructure, road traffic, and the supply of water, gas, heating, and electricity. This area becomes important when examining the battery industry, which we will revisit later.

The third is education and vocational training. This covers admissions evaluations, learning outcome assessments, cheating detection during examinations, and determining educational placement.

The fourth area is recruitment, employment, workforce management, and access to self-employment. This includes analyzing job advertisements, screening and ranking applicants, decisions regarding promotion and termination, and task allocation and performance evaluation. KoreaTalent's product falls squarely into this category.

The fifth area is access to essential private and public services. This category is further divided into four branches: (a) evaluating eligibility for public assistance benefits by public authorities, (b) creditworthiness evaluation and credit scoring, (c) risk assessment and pricing in life and health insurance, and (d) prioritizing emergency call dispatches. The financial industry is directly impacted here.

The sixth area is law enforcement. This includes individual risk assessments, lie detection, evaluating the reliability of evidence, and profiling for crime prediction.

The seventh is migration, asylum, and border control management. This encompasses lie detection, risk assessments, examining applications, and individual identification.

The eighth area is the administration of justice and democratic processes. This covers systems that assist in legal research and interpreting facts, or those intended to influence election outcomes and voting behavior.

Among these eight areas, those that Korean companies actually encounter most frequently are recruitment, credit scoring and insurance, education, essential services, and biometrics. Because private companies rarely act as direct providers in law enforcement, migration and border control, or the administration of justice, this chapter will place less emphasis on those areas.

One point must be made clear from the start: if a system performs profiling of natural persons, it is always considered high-risk regardless of which Annex III area it belongs to. Profiling refers to the automated processing of personal attributes - such as work performance, economic situation, health, preferences, reliability, behavior, location, or movements - to evaluate or predict them. KoreaTalent's applicant-ranking tool processes individual histories and responses to predict hiring suitability, fitting this definition precisely. This means the filter clause discussed later is a door that was never open to KoreaTalent in the first place.

As noted in Chapter 3, the full application date of Article 16 (provider obligations) and Article 26 (deployer obligations) for Annex III standalone systems was originally August 2, 2026, but was postponed to December 2, 2027, under the Digital Omnibus amendment. This means KoreaTalent gained extra time, not that the obligations disappeared. I read this grace period not as an excuse to delay preparation, but as leeway to prepare properly.

11.2 Annex I: AI Embedded Within Existing Products

Now let's move on to Sinheung's story. Article 6(1) sets out two conditions entirely different from Annex III. First, the AI system must be used as a safety component of a product - or be the product itself - covered by existing EU harmonization legislation listed in Annex I. Second, the product must be required to undergo a third-party conformity assessment prior to being placed on the market or put into service under that harmonization legislation. Both conditions must be met to classify the system as high-risk.

The harmonization legislation listed in Annex I encompasses the Medical Devices Regulation (2017/745), the In Vitro Diagnostic Medical Devices Regulation (2017/746), the Machinery Regulation, the Toy Safety Directive, the Motor Vehicles Type-Approval Regulation (2018/858), the Agricultural and Forestry Vehicles Regulation (167/2013), the Two- or Three-Wheel Vehicles Regulation (168/2013), the Marine Equipment Directive (2014/90/EU), and aviation-related legislation. Sinheung's BMS is a safety component of an automobile covered by the Motor Vehicles Type-Approval Regulation, and that automobile is already a product subject to third-party conformity assessment in the form of type-approval. Since both conditions are satisfied, it is high-risk.

A noteworthy aspect regarding Annex I is determining who qualifies as the provider. When an AI system is integrated into an Annex I product, the manufacturer of that product is considered the provider under the AI Act. If Sinheung supplies the BMS to be placed on the market alongside the vehicle under the automaker's name or trademark, that automaker bears the Article 16 obligations. If Sinheung merely supplies the component without affixing its own trademark to the final product, Sinheung itself might not be the provider under the AI Act. However, in supply contracts with automakers, component suppliers almost always assume information-sharing obligations to support AI Act compliance, effectively taking on quasi-provider responsibilities in practice. I advise reviewing every single clause in your supply contracts regarding this point.

Article 8(2) permits providers of Annex I products to integrate the testing, reporting, and documentation required by the AI Act into the existing documentation and procedures of sectoral safety legislation. For example, risk management and training data bias verification required by the AI Act can be added directly into the technical documentation under the Medical Devices Regulation. The intent is to minimize duplication - a welcome provision in practice for companies like Sinheung that already possess established

automotive component certification frameworks.

The completion deadline for regulatory procedures governing Annex I embedded high-risk AI was originally August 2, 2027, but was pushed back by one year to August 2, 2028, under the Digital Omnibus amendment. This comes eight months later than Annex III. I view this not as a political concession, but as a practical alignment with existing product certification cycles. Sectoral products like automobiles and medical devices operate on multi-year certification cycles, making it inherently more time-consuming to weave AI Act requirements into their pre-existing procedures.

11.3 The Filter under Article 6(3) and the Trap of the Burden of Proof

Being listed in Annex III does not automatically make an AI system high-risk. Article 6(3) provides that an AI system referred to in Annex III shall not be considered high-risk if it does not pose a significant risk of harm to the health, safety, or fundamental rights of natural persons. This is the so-called "filter clause." There is one critical point to underscore here: this filter is available exclusively to Annex III. No such exception exists for Annex I embedded systems. The moment Sinheung's BMS met the condition of being an automotive safety component, it became high-risk immediately, with zero leeway for a separate exception review. The presence or absence of this filter represents the second major dividing line between Annex I and Annex III.

The conditions for actually applying the filter fall into four scenarios: when the system performs a narrow procedural task; when it improves the result of a previously completed human activity; when it merely detects decision-making patterns or deviations from such patterns without replacing or influencing a human judgment already made without proper human review; or when it performs only a preparatory task to an assessment relevant for the purposes listed in Annex III. Satisfying any one of these four allows a provider to claim the filter exception.

However, these four conditions collapse under a single proviso: the filter does not apply at all if the system performs profiling of natural persons. Consider KoreaTalent's hiring ranking tool again. Processing individual candidate data to predict fit is profiling in itself, and at that moment, the door to the filter clause locks shut. No matter how elaborately one crafts an argument claiming it is a narrow procedural or preparatory task, it is to no avail. In legal consultations, I often see companies trying to convince themselves that their system is merely an "auxiliary recommendation tool." Yet if it automatically evaluates individual attributes to rank candidates, that argument will not hold.

For providers seeking to invoke the filter, the entire burden of proof shifts to them. Article 6(4) requires a provider who concludes that an Annex III system is not high-risk to document the basis for that assessment prior to placing it on the market or putting it into service. This documentation must be submitted immediately upon request by national competent authorities and registered in the EU database pursuant to Article 49(2). Any company that chooses to quietly claim the filter without registering is not complying - it is simply waiting to be caught. Article 80 grants market surveillance authorities the power to re-evaluate systems classified as non-high-risk by providers. If authorities deem a system

to actually be high-risk, the provider will receive an order to comply with AI Act requirements without delay. The filter is not a quiet back door for self-exemption; it functions as a front door contingent on constant readiness to prove the validity of that assessment.

The Article 6 guidelines containing concrete application examples for this filter were originally due by February 2, 2026. That deadline has passed. The European Commission released a draft only on May 19, 2026, and stakeholder feedback closed on June 23. As of July 24, 2026 - the time of writing this manuscript - the final version has yet to be published. The direction indicated by the draft remains consistent with what has been described so far: the filter must be interpreted narrowly, and it does not apply to profiling. However, since the final guidelines will include more detailed examples, I consider it safer to treat current internal assessment documents as provisional and revise them as soon as the final version is released.

11.4 Applying the Framework to Six Key Korean Industries

Let's overlay the framework built so far onto six representative Korean industries. This analysis reflects my practical assessment based on cases handled to date and publicly available information; the actual classification of an individual product can only be finalized by closely inspecting its specific intended purpose.

Semiconductors are Low Risk. A semiconductor chip itself is not an AI system, so it falls outside the direct scope of the AI Act. However, the story changes for AI-based safety systems deployed on semiconductor fabrication lines to prevent robot malfunctions or control toxic gas leaks. If such a control system serves as a safety component of equipment covered by the Machinery Regulation, it could fall under Annex I. I view most of the semiconductor industry as standing outside this narrow exception.

Automotive is High Risk. As Sinheung's case illustrates, safety-related functions in Advanced Driver Assistance Systems (ADAS), emergency braking, driver monitoring, and battery management systems mostly constitute safety components under Annex I. Because this domain directly touches human safety, avoiding a high-risk classification is exceedingly difficult. Although the compliance deadline has expanded to August 2, 2028, automotive type-approval certification cycles themselves run on multi-year timelines. I assess that integrating AI Act requirements into those existing processes must begin now to meet the deadline.

Consumer Electronics are Low Risk. Features like content recommendations on smart TVs, grocery management on smart refrigerators, or autonomous navigation in robot vacuums without collision avoidance are typical auxiliary functions targeted by the filter clause. They serve convenience without materially influencing decision-making outcomes. Exceptions do exist, however. Toys involving child safety, healthcare appliances incorporating medical device functions, or smart door locks using biometrics for access control may fall under Annex I or Annex III. I recommend that electronics companies avoid lumping all products together as low risk and instead screen each item to check whether safety-related add-on functions are present.

Batteries are Medium Risk. This industry serves as a prime example of split classification between Annex I and Annex III. A battery management system in an electric vehicle is an automotive safety component under Annex I. On the other hand, overheat prevention and fire suppression systems for large-scale Energy Storage Systems (ESS) can be interpreted

as safety components involved in operating power grids, potentially falling into the second area of Annex III (critical infrastructure). Because automotive applications and utility-scale ESS follow different annexes within the same battery industry, classifications must be conducted separately by product line.

Medical Devices are High Risk. AI-based computer-aided diagnostic software for medical imaging, surgical robots, and patient monitoring and treatment planning systems are products subject to the Medical Devices Regulation (2017/745) or the In Vitro Diagnostic Medical Devices Regulation (2017/746). Software classified as Class IIa or higher under these regulations is already subject to third-party conformity assessment, almost automatically satisfying both conditions of Annex I under the AI Act. If a Korean medical device company holds CE marking, confirming its risk class virtually determines its high-risk status under the AI Act.

Financial Services are High Risk. Creditworthiness evaluation and credit scoring fall under Annex III 5(b), while risk assessment and pricing in life and health insurance fall under 5(c). AI systems that exert a material impact on personal property and opportunities - such as credit underwriting, insurance underwriting, and robo-advisory chatbots - mostly belong to this domain. The path to escaping via the filter clause is virtually blocked as well. Credit scoring and insurance premium calculation inherently involve automated processing to evaluate an individual's credit, assets, or risk profile, making it nearly impossible to escape the definition of profiling. If a Korean financial institution provides credit scoring or insurance underwriting AI targeting EU customers, I advise preparing under the assumption of high risk from the outset.

Placing all six industries along a single spectrum, automotive, medical devices, and financial services are High, batteries are Medium, and semiconductors and consumer electronics are Low. Note, however, that a Low risk determination applies to current product configurations and carries no guarantee of persisting in future iterations. Adding biometric entry features to a home appliance or integrating automated decision-making tied directly to human safety into semiconductor equipment can alter the risk tier instantly. I suggest treating this overview as a snapshot of today's product offerings.

Allow me to add a note regarding the position of Korean companies. If a company in any of these six industries places products on the EU market or if their output is used within the EU, it becomes subject to extraterritorial application under Article 2(1)(c) of the AI Act. Providers headquartered outside the EU must designate an authorized representative in the EU in writing. If classified as high-risk, this also intersects with requirements under

South Korea's pending Framework Act on Artificial Intelligence. Because the documentation and risk management frameworks required by both laws overlap significantly, I recommend designing EU compliance documentation and domestic compliance under a single unified architecture from the start. Creating two separate sets of documentation is a far greater waste of resources.

Reference Material

Four-Step Classification Flowchart: Step 1, verify whether your product is a safety component of - or the product itself under - existing EU product safety legislation (medical devices, machinery, motor vehicles, toys, etc.). If so, check whether that product is subject to third-party conformity assessment. If both are true, it falls under Annex I: it is high-risk with no filter available. Step 2, if it does not fall under Annex I, check whether it belongs to one of the eight areas in Annex III (biometrics, critical infrastructure, education, employment, essential services, law enforcement, migration and border control, administration of justice). Step 3, if it does, check whether it performs profiling of natural persons. If it performs profiling, it is immediately high-risk and the filter does not apply. Step 4, if it does not perform profiling, evaluate whether it satisfies one of the four conditions: performing a narrow procedural task, improving a completed human activity, acting as an auxiliary pattern detection tool, or performing a preparatory task. If satisfied, document the rationale for that assessment and register it in the EU database.

Summary of Six Key Korean Industries: Semiconductors: Low Risk; exception for safety control equipment. Automotive: High Risk; Annex I; August 2, 2028. Consumer Electronics: Low Risk; exception for toys and biometric access control. Batteries: Medium Risk; EV BMS falls under Annex I, large ESS subject to Annex III review. Medical Devices: High Risk; MDR/IVDR Class IIa or higher. Financial Services: High Risk; Annex III 5(b)/5(c); filter blocked by profiling.

Timeline Table: February 2, 2026: Original deadline for Article 6 guidelines (passed). May 19, 2026: European Commission published draft guidelines on high-risk classification. June 23, 2026: Stakeholder feedback period closed. August 2, 2026: Article 50 transparency obligations become fully applicable and AI Office fine enforcement powers take effect (high-risk obligations not yet applicable). December 2, 2027: Application of Annex III standalone high-risk AI obligations (Articles 16 & 26) begins. August 2, 2028: Deadline for completing regulatory procedures for Annex I embedded high-risk AI.

That afternoon, the BMS manager from Sinheung and the hiring tool manager from KoreaTalent left the meeting room with entirely different conclusions. Sinheung bypassed any debate over the filter and immediately began preparing for high-risk compliance, whereas KoreaTalent took two full weeks just to accept the reality that their product was performing profiling. The difference between the two companies was not a difference in technology, but a difference in how honestly they examined what their product was actually doing to people.

Chapter 12. Obligation System for High-Risk System Providers

Let us take a recruitment solution startup in Pangyo as an example. This company created a service that analyzes resumes and interview videos to rank candidates for hiring managers, and is currently about to integrate this service for clients in Germany and the Netherlands. AI used in recruitment screening is a high-risk AI system falling under Annex III, point 4 of the European Union Artificial Intelligence Act (Regulation (EU) 2024/1689, hereinafter the "AI Act"). From this moment on, this company is no longer just selling a single service; it becomes a provider bearing the entirety of the six obligations required by Chapter III of the AI Act. What this chapter deals with are precisely those six obligations: the risk management system, data governance, technical documentation along with automatic logging and provision of information to deployers, human oversight, accuracy, robustness, and cybersecurity, as well as conformity assessment, CE marking, and EU database registration.

First, let us clarify the timeline. The NotebookLM notes used for researching this book list the full application dates for high-risk obligations as August 2, 2026 for Annex III, and August 2, 2027 for Annex I. Since these are the original timelines calculated as 24 months and 36 months after the AI Act entered into force on August 1, 2024, the calculations themselves are not incorrect. However, this schedule was pushed back by the Digital Omnibus on AI, which received final Council approval on June 29, 2026. As the drafting of harmonised standards lagged behind and the designation of Notified Bodies by Member States failed to keep pace, the obligation application date for Annex III standalone high-risk systems was deferred to December 2, 2027, and for Annex I embedded product systems to August 2, 2028. Therefore, the application dates referred to in this chapter are not the old dates in the notes, but these new dates. The reason for pointing out where the source material and search results diverge without glossing over them is that a company like the Pangyo startup might rush based solely on the old dates and end up making superficial preparations, or conversely, might slack off entirely upon hearing that the dates have been postponed. An interesting point is that Articles 99 and 101 themselves - the provisions serving as the legal basis for fines - will take effect on the original schedule of August 2, 2026. Violations of provisions already in force, such as prohibited practices under Article 5 or transparency obligations under Article 50, will become subject to immediate penalties from that date. However, as Articles 9 to 15, as well as Articles 43, 48, and 49 covered in this chapter, are not yet fully applicable in terms of high-risk obligations themselves, I read this as the actual center of gravity for enforcement shifting to after December 2027 and August 2028.

12.1 Risk Management System: The Cyclical Structure Required by Article 9

The very first provision encountered by the development team lead at the Pangyo startup is Article 9. This article is not a mere formality completed by creating a single document; it requires a continuous procedure operating throughout the entire lifecycle of the high-risk system. One must identify and estimate risks in advance, reduce or eliminate them through means such as design modifications or user guidance, examine separately how the system impacts minors under the age of 18 or vulnerable groups, and undergo testing to verify compliance and performance. The core point is that these four steps do not end with a single execution, but must be repeated every time the system is modified. For a recruitment AI, this means that this cycle must be executed anew whenever the candidate ranking algorithm is updated.

Article 9(10) allows providers who already operate internal risk management procedures under other EU legislation to add or combine the elements of Article 9(1) through (9) into those existing procedures. For products falling under Annex I, such as medical devices or machinery, Article 8(2) allows the AI Act requirements to be embedded within the documentation and procedures of existing harmonisation legislation. Companies building standalone Annex III systems like the Pangyo startup do not have such an existing framework and must design a risk management system under Article 9 from scratch; in such cases, I believe it is better to design it as an integral part of the quality management system (Article 17). There is no reason to maintain two separate management procedures to reflect risk assessment results into system modification processes in real time.

When evaluating this risk management system, the AI Office, the enforcement authority, uses the intended purpose of the system as the benchmark. It is also worth noting that applying harmonised standards grants the benefit of the Presumption of Conformity. However, as noted earlier, since harmonised standards themselves are not yet fully established, at this point it is safer to establish and document in-house procedures by applying international standards (such as ISO/IEC 23894 AI risk management guidelines) rather than relying on standards. Uncertain elements should be documented as uncertain. The specific implementation methods and detailed risk assessment criteria for risk management systems leave significant room to vary depending on the type of system, and this margin will be filled by the Article 6 guidelines, which were scheduled for publication by February 2, 2026, and by future harmonised standards.

12.2 Data Governance and Training Data Quality: Article 10 and the Issue of Bias

This is where the real fuse of recruitment AI lies. If a model is trained on past hiring data, the gender or university biases already embedded within that data are inherited as-is. Article 10 stipulates that training, validation, and testing datasets used for high-risk systems must be relevant, representative, and to the best extent possible, free of errors and complete. Article 10(2) goes a step further, explicitly requiring evaluations of design choices for data collection, the collection process, preparation processes including annotation, labeling, cleaning, and augmentation, governance frameworks, roles and responsibilities of personnel, evaluations for bias detection and mitigation, and even the potential inappropriateness of datasets under specific contexts. For the Pangyo startup, this means work must begin by examining the statistical properties of the past five years of hiring data to identify which job categories or age groups are underrepresented.

Here lies the point of intersection with data privacy issues. Article 10(5) exceptionally allows providers to process special categories of personal data, such as racial origin or health data, strictly to the extent necessary for bias detection and correction. However, this is possible only under the premise that all safeguards required by the GDPR (Regulation (EU) 2016/679) and related directives are fully observed. I interpret this provision as follows: collecting gender or racial data indiscriminately under the guise of eliminating bias becomes a GDPR violation in itself; thus, the exception opened by Article 10 of the AI Act is a cracked door, not an invitation to walk in freely. In practice, it is safer to pass through this narrow door by using anonymization or third-party proxy variables.

Another easily missed aspect is that these requirements apply even to data not directly used for model training, and that the datasets must reflect the specific geographical, contextual, behavioral, or functional settings under which the AI system is intended to be used. If a recruitment AI trained on data collected in South Korea is deployed directly in the German market, it may face criticism that the structure of the German labor market or resume conventions were not reflected in the data. The AI Office also uses the intended purpose as the benchmark when judging data relevance, representativeness, completeness, and the presence of bias, referencing whether harmonised standards or common specifications were applied. Specific technical criteria for quantitatively measuring relevance, representativeness, and completeness have not yet been fully settled via harmonised standards, and interpretations of this aspect may vary depending on the type of system.

12.3 Technical Documentation, Automatic Logging, and Information to Be Provided to Deployers

Article 11 requires technical documentation to be drawn up before market placement and kept up to date thereafter. The items to be included in the documentation are listed in Annex IV, covering a broad scope from the system's intended purpose and design logic, architecture, data requirements, and training methodology, to validation and testing procedures and results, risk management system operation logs, and modification histories. Simplified forms are permitted for small and medium-sized enterprises (SMEs) and startups, so a company of the Pangyo startup's scale would do well to leverage this relaxed framework.

Article 12 mandates automatic logging capabilities. High-risk systems must be capable of automatically recording events while operating, and these logs are used to identify system risks, facilitate the tracking of substantial modifications, and support monitoring throughout the entire lifecycle. For a recruitment AI, this means being able to reconstruct later which resume was input, what score resulted, and which model version was involved in that determination. Without these logs, neither the provider nor the deployer can explain what happened when an incident occurs.

Article 13 requires this log data and the system itself to be handed over to deployers with sufficient information so they can actually operate it. The instructions for use must include measures to be taken for human oversight, the computational resources and hardware required to run the system, the expected duration of the system's lifecycle, maintenance methods, and how to collect, store, and interpret Article 12 logs. Weaving these three provisions into a single sentence yields the following: Article 11 mandates that providers place themselves in a state where they can explain their system; Article 12 mandates that evidence be retained to substantiate that explanation; and Article 13 mandates that both the evidence and explanation be handed over to the deployer. This section of the chapter is shorter than others because the assigned NLM source material contained only an introductory portion and lacked detailed narrative on the statutory text. The specific contents of the three provisions were filled directly into this section based on the published original text of the AI Act and Annex IV.

12.4 Human Oversight: The Control Vision of Article 14

Article 14 stipulates that high-risk systems must be designed and developed, including with appropriate human-machine interface tools, such that they can be effectively overseen by natural persons. The objective is clear: even after all other requirements have been fulfilled, a human must be able to catch remaining risks at the final stage. The article divides the capabilities that the system must enable for deployers into five categories: the capacity to understand the system's capabilities and limitations and monitor for anomalies; the capacity to remain aware of automation bias, where humans uncritically rely on AI output; the capacity to correctly interpret outputs; the capacity to decide not to use the system or to disregard and override outputs in specific situations; and the capacity to interrupt the operation of the system through procedures such as a stop button.

Applying this to recruitment AI yields the following: hiring managers must not blindly trust the rankings assigned by the AI, but must be able to question why a result was produced if the ranking seems abnormal, and override those rankings if necessary. Here lies a mistake frequently committed by the Pangyo startup: displaying only ranks and scores on the interface while concealing the underlying rationale. In this case, human oversight remains in name only, stripped of substance. Article 26, which governs deployer obligations, also mandates that deployers operate the system in accordance with the provider's instructions, manage the relevance of input data, and maintain a human oversight framework where a human retains final judgment. Thus, Article 14 is not an obligation borne by the provider alone, but one that operates in tandem with the deployer.

Phrases such as "appropriate and proportionate" capabilities appear throughout the statutory text, but exactly what level this entails inevitably varies depending on the type of system and its operational environment. The intensity of human oversight required for an emergency room diagnostic support AI cannot possibly be the same as that for a recruitment ranking AI. The AI Office also uses the intended purpose as the benchmark when evaluating appropriateness, referencing harmonised standards or common specifications.

12.5 Accuracy, Robustness, and Cybersecurity: The Technical Baseline Required by Article 15

Article 15 requires high-risk systems to achieve an appropriate level of accuracy, robustness, and cybersecurity, and to maintain these levels consistently throughout their lifecycle. Regarding accuracy, the level of system accuracy and related metrics must be explicitly stated in the instructions for use. For a recruitment AI, this means documenting the metrics used to measure ranking prediction accuracy and stating the corresponding numerical figures. Robustness refers to resilience against errors, inconsistencies, or faults without failing. Even if input values are somewhat anomalous, the system must not produce erratic outputs.

The cybersecurity requirement presents practical complexity. Technical measures against threats such as data and model poisoning, adversarial examples, and confidentiality breaches must be put in place in proportion to the system's risks and context. If someone maliciously tampers with resume data to automatically disqualify certain candidates, that constitutes a data poisoning attack. Failing to implement measures to prevent such threats violates Article 15. It is also worth noting that systems certified or declared compliant under Regulation (EU) 2019/881 (the Cybersecurity Act) are presumed to satisfy the cybersecurity requirements of Article 15. The quantitative benchmarks for measuring "appropriate levels" of accuracy, robustness, and cybersecurity have not yet been fully codified into harmonised standards either. Although the European Commission stated it would develop benchmarks and measurement methodologies through stakeholders and metrology institutes, at present the only recourse for the Pangyo startup is to set internal metrics by applying international guidelines such as ISO/IEC 24029 or the NIST AI RMF.

12.6 Conformity Assessment, CE Marking, and EU Database Registration

While Articles 9 through 15 covered so far deal with internal design requirements, Articles 43, 48, and 49 represent procedures for demonstrating those results externally. Article 43 presents a fork in the road: if harmonised standards are satisfied, providers may use the internal control method (Annex VI) based on self-assessment; if harmonised standards do not exist, or if the system is a biometrics system under Annex III, point 1, providers must undergo third-party assessment (Annex VII) audited by a designated Notified Body. Recruitment AI falls under Annex III, point 4, so in principle it follows internal control. However, if even a minor biometric element is incorporated - for instance, if facial expression or voice analysis features are added to interview videos - the system immediately shifts to Annex III, point 1, and if harmonised standards are not fully applied or common specifications are not used, Notified Body assessment becomes mandatory. This is a point I constantly emphasize to companies, as development teams often overlook how the weight of conformity assessment changes completely the moment a biometric element is subtly introduced.

When AI is embedded into an Annex I product such as a medical device or machinery, it follows the conformity assessment procedures of existing product safety legislation, with Chapter III requirements of the AI Act integrated into those procedures. Article 8(2) provides the legal basis permitting this integration. If a system or its intended purpose is substantially modified, it must undergo a new assessment; however, changes that are predetermined and documented within continuous learning systems are exempt. Here again, the boundary of what constitutes a 'substantial modification' remains ambiguous, as the statutory text alone does not cleanly delineate where the re-assessment obligation is triggered between replacing a recruitment AI's ranking algorithm entirely versus adjusting a few parameters.

Once conformity assessment is passed, CE marking (Article 48) is affixed. It must be affixed visibly, legibly, and indelibly; if physical affixation is difficult, it may be placed on packaging or accompanying documents, or for digitally delivered systems, substituted on the interface or machine-readable code. If a Notified Body was involved in the assessment, its identification number must be displayed alongside the CE marking. An EU Declaration of Conformity pursuant to Article 47 must also be drafted and kept for 10 years; for products subject to multiple EU laws, it can be consolidated into a single declaration.

Finally, there is EU database registration (Article 49). The provider, or an Authorised Representative if the company is located outside the EU, must register themselves and the system prior to market placement. A frequently overlooked fact is that registration applies not only to high-risk systems listed in Annex III, but also to systems determined by self-assessment not to be high-risk pursuant to the Article 6(3) filter clause. In the latter case, the documentation supporting the self-assessment rationale must be submitted. This means that even if the Pangyo startup later concludes that its recruitment AI is 'not high-risk', it must still register and retain the grounds for that determination. Certain high-risk systems used in law enforcement, migration, asylum, or border control are registered in a non-public section of the database, while certain critical infrastructure systems are registered separately at the national level.

For companies headquartered outside the EU - that is, Korean companies like the Pangyo startup - an additional obligation is superimposed above all these procedures: written designation of an Authorised Representative established within the EU prior to market placement. This representative bears the obligation to keep copies of technical documentation and the declaration of conformity for 10 years and produce them upon request by authorities. It is also worth noting that the drafting of harmonised standards by CEN/CENELEC missed its original target deadline of August 2025 and remains in progress. This gap forms the background behind the delays introduced by the Digital Omnibus, and until standards are finalized, specific benchmarks for conformity assessment will inevitably remain fluid.

Reference Materials

Application Dates Summary. AI Act entry into force: August 1, 2024. Penalty provisions (Articles 99 & 101) take effect: August 2, 2026. Full application of obligations for Annex III standalone high-risk systems: December 2, 2027 (postponed from August 2, 2026 via Digital Omnibus). Full application of obligations for Annex I embedded high-risk systems: August 2, 2028 (postponed from August 2, 2027 via Digital Omnibus). Deadline for Notified Body designation applications under the AI Act: around February 2, 2028.

Fine Brackets. Violations of high-risk obligations (Chapter III, including Article 16) for Annex III, point 4 recruitment AI are subject under Article 99(2) to up to 3% of total worldwide annual turnover or 15 million euros, whichever is higher. Violations of prohibited practices under Article 5 incur up to 7% or 35 million euros, whichever is higher, carrying significantly heavier penalties than the obligations in this chapter.

Self-Audit Checklist. Is the risk management system integrated with the quality management system? Are records of statistical properties and bias evaluations of training data retained? If special categories of personal data were handled, are GDPR-compliant safeguards documented separately? Does technical documentation cover all Annex IV items? Is the automatic logging capability at a level capable of actually reconstructing events? Do deployer instructions include methods for human oversight and log interpretation? Are rationale for rankings or determinations exposed on the interface so deployers can genuinely disregard or override them? Are accuracy metrics specified numerically in the instructions for use? Are defensive measures against data poisoning or adversarial attacks in place? Has it been re-verified whether any biometric element was introduced? Has an EU Authorised Representative been designated? Has EU database registration been taken care of, including cases determined not to be high-risk?

How many of these twelve items can be answered with a confident "yes," and how many more "yeses" can be added before December 2027 and August 2028 arrive, now remains the responsibility of the development team lead at the Pangyo startup.

Chapter 13. Obligations of Deployers

A healthcare startup in Bundang decided to implement an AI recruitment screening system for its German subsidiary. From the perspective of a company founder, this scenario seems like no big deal. The software is already finished, and the company just needs to buy and use it. However, the law applies even to those who simply purchase and use software. The AI Act refers to an entity in this position as a "deployer" and sets forth their obligations under Articles 26 and 27. After reading this chapter, it is best to set aside the notion that not having developed the system excuses one from liability.

13.1 Compliance with Instructions for Use, Assignment of Human Oversight, and Six-Month Log Retention

The starting point for a deployer's obligations is to strictly follow the instructions for use provided by the provider. The provider specifies in these instructions for use under what conditions, with what data, and how the results should be interpreted for that recruitment screening AI. Deployers must not operate the system outside of this scope. If the AI system was certified solely for evaluating resumes, but the company also uses it for current employee promotion reviews, the company goes beyond the instructions at that very moment, and liability shifts to the company should any issues arise.

Next is the assignment of human oversight. As seen in Chapter 12, the provider must design the system to enable human intervention. The deployer must assign a person to actually put that design into operation. You cannot assign just anyone. That individual must understand on what grounds the system screens applicants, and possess both the authority and competence to override the results if something seems amiss. Entrusting a junior HR team member with simply watching the screen and clicking an approval button is not oversight - it is a mere pretense of oversight. When enforcement authorities conduct an investigation following an incident, this is the very first point they question: what could the supervisor actually see, and what authority did they genuinely possess?

The duties of monitoring and record-keeping follow. While the system is running, the deployer must monitor for anomalies, and upon discovering a serious malfunction or incident, suspend use and inform both the provider and the competent authorities. In addition, the deployer must retain the logs automatically generated by the system. The period is at least six months. This six-month figure is non-negotiable. Should litigation arise later over recruitment outcomes, the company can only defend itself if it can demonstrate on what basis and what determination the AI made at that time. A company that has erased its logs effectively forfeits any material for its defense.

13.2 Deployment of AI for Workers and Prior Notification

This is a point Korean companies frequently overlook. When deploying a high-risk AI system targeted at workers, the company must inform the workers of that fact in advance. The order is to notify prior to deployment, not to announce after implementation.

This differs fundamentally from domestic practice. In Korea, while companies sometimes hold briefing sessions for employees when adopting a new HR evaluation tool or attendance management system, the law does not mandate it. Europe differs on this point. If it is a high-risk AI system affecting workers' status - such as recruitment, assignment, performance evaluation, or dismissal - prior notification is a legal obligation. If the aforementioned Bundang healthcare startup uses a recruitment screening AI at its German subsidiary, situations arise where it must notify not only the applicants processed by that AI but also existing employee representatives or works councils in advance. I view this notification obligation as something the legal team, rather than HR, should manage. HR is prone to viewing notification merely as public relations or routine guidance, but this is a procedure stipulated by law, and omitting it constitutes a violation in itself.

13.3 Fundamental Rights Impact Assessment: Who Must Perform It?

Among the deployer's duties, the heaviest is the Fundamental Rights Impact Assessment. It is prescribed in Article 27 and commonly abbreviated as FRIA (Fundamental Rights Impact Assessment). It is a procedure to evaluate in advance how a high-risk AI system may impact human fundamental rights before putting it to use.

This obligation does not apply to all deployers. Precisely delineating its scope is the key to this section. Deployers that are public authorities, as well as private deployers providing public services, must perform this assessment when deploying high-risk systems under Annex III. A purely private company using high-risk AI for internal business purposes does not automatically incur this obligation.

However, there are two exceptions, and these exceptions are even more crucial in practice. AI used for credit scoring under point 5(b) of Annex III, and AI used for risk assessment and pricing in relation to life and health insurance under point 5(c), are subject to the Fundamental Rights Impact Assessment even for purely private deployers. If a Fintech company headquartered in Seoul deploys a credit scoring AI for loan screening in Europe as a deployer, it cannot avoid this assessment even though it is neither a public authority nor a provider of public services. The same applies to insurers. If a company concludes that it can skip FRIA simply because it is a purely private enterprise unrelated to public interest, that conclusion is wrong in those two domains: credit evaluation and insurance.

The required contents of the assessment are specified: the purpose, timeframe, and manner of system use; the category of natural persons affected; the specific risks of harm associated with its deployment; how human oversight will be implemented; and the mitigation measures to be taken if risks materialize. These five elements must be documented prior to initial use and updated whenever the manner or context of use undergoes significant change. Upon completing the assessment, the results must be notified to the market surveillance authority. Deployers are expected to use a questionnaire template developed by the AI Office; however, as of my review, the final official version of this template has not been published, and practitioners have been circulating self-drafted templates in the interim.

The relationship with Data Protection Impact Assessments (DPIAs under the GDPR) is another area warranting clarification. Article 27(4) specifies that the Fundamental Rights

Impact Assessment complements the DPIA. It complements it; it does not replace it. A company that has already conducted a DPIA may utilize those materials for its FRIA. Since how applicant personal data is collected and processed will already be addressed in the DPIA, there is no reason to rewrite that portion. However, while a DPIA focuses on personal data protection, a FRIA evaluates fundamental rights extending beyond data protection, such as non-discrimination and the right to a fair trial. Where the two assessments overlap, data can be shared; where they do not, the gaps must be filled separately. Any company assuming FRIA is complete simply because it conducted a DPIA has misread this provision.

At this juncture, a timing issue must be highlighted. The NotebookLM material referenced in this manuscript states that the obligations under Articles 26 and 27 for Annex III standalone high-risk systems take full effect on August 2, 2026. However, on May 7, 2026, the European Parliament, the Council, and the Commission reached a provisional agreement on an amendment known as the Digital Omnibus, which includes a proposal to defer the application date of obligations for Annex III standalone high-risk systems by 16 months, to December 2, 2027. As of the writing of this manuscript, this agreement remains provisional, having not yet undergone formal adoption by the European Parliament and the Council or publication in the Official Journal. In other words, there is a possibility it will be formally finalized before August 2; if so, the original August 2 application date will lapse, shifting to December 2027. My advice to Korean enterprises is this: there is no reason to slow down preparations originally scheduled for August 2. If August 2 passes without formal adoption of the amendment, the original enforcement date remains in effect. Nonetheless, be aware that the possibility of postponement genuinely exists, and monitor the Official Journal of the European Union through early August to verify whether this amendment is published.

Key takeaways to confirm from this chapter:

Compliance with instructions for use. Do not deviate from the conditions and intended purpose specified by the provider.

Human oversight. Assign an actual individual who possesses the competence and authority to understand the system and intervene.

Monitoring and logs. Respond to anomalous signs and retain logs for at least six months.

Worker notification. Inform workers in advance prior to deploying high-risk AI systems targeting them. Ensure the legal team, not HR, manages this.

Fundamental rights impact assessment. Public authorities and private deployers providing public services are subject to mandatory assessment when using Annex III systems. Even purely private enterprises are strictly subject to this requirement for credit evaluation (point 5(b)) and life/health insurance risk assessment and pricing (point 5(c)). Share data with GDPR DPIAs, but do not treat a DPIA as a substitute.

Verification of enforcement date. Obligations under Articles 26 and 27 for Annex III standalone high-risk systems were originally scheduled for August 2, 2026, but may be deferred to December 2, 2027, under the Digital Omnibus agreement. Verify formal adoption directly in the Official Journal.

Let us return to that healthcare company in Bundang. Before its German subsidiary implements a recruitment screening AI, the question the company must first ask is not how sophisticated this AI is. It must first ask: as a deployer, where do we stand under the law?

Chapter 14. How to Use the Delay Period

A compliance officer at a Pangyo-based startup sent me an email last week. The question was: given that the CEN and CENELEC harmonized standards have not yet been released while the obligations for high-risk AI systems will still apply starting August 2, 2026, on what basis are companies supposed to demonstrate conformity? There is no single correct answer to this question. However, depending on how you arrange the tools available in your hands today, your outlook on the morning of August 2 will look very different.

14.1 Background: Delays in Establishing Harmonized Standards

The European Union Artificial Intelligence Act (Regulation (EU) 2024/1689) entered into force on August 1, 2024, imposing obligations set forth in Chapter III on providers of high-risk AI systems. These obligations will become fully applicable on August 2, 2026, the same day the AI Office's power to impose fines takes effect. In the event of a violation, administrative fines can be up to 15 million EUR or 3 percent of total worldwide annual turnover for the preceding financial year, whichever is higher. In effect, both obligations and penalties are triggered simultaneously on a single date.

The mechanism that technically fulfills these obligations is harmonized standards. On May 22, 2023, under Commission Decision C(2023) 3215, the European Commission issued standardization request M/593 to the European Committee for Standardization (CEN) and the European Committee for Electrotechnical Standardization (CENELEC). Developing and verifying systems in accordance with these standards confers the benefit of a presumption of conformity (Article 40), under which the system is presumed to comply with the relevant requirements of the law. Since providers do not need to cross-check every legal clause individually, this is the surest shortcut for them.

However, CEN and CENELEC missed the requested deadline of August 2025. As of July 2026, work on the draft standards remains ongoing. Internally, the standards organizations resolved to take several emergency measures to make up for the delay. If a draft receives a favorable vote during the Enquiry stage, it can be published directly without a separate formal vote, and small drafting teams were formed to tackle the six most delayed drafts. The target is to release the first set of standards by the fourth quarter of 2026. Among them, prEN 18286, which corresponds to the Quality Management System (QMS) under Article 17, has reached the Enquiry stage.

Here is my assessment: even if draft standards are released within 2026, they will not immediately yield the benefit of a presumption of conformity. A separate review procedure remains for the European Commission to publish references in the Official Journal, and this review takes time. Therefore, it is safest to assume that virtually no Korean company will be able to claim a presumption of conformity based on harmonized standards on the morning of August 2. The announcement that a standard has been drafted and the actual point at which legal benefits accrue under that standard are two distinct moments in time.

The Commission is well aware of this gap. The Act allows the Commission to adopt Common Specifications (Article 41) via implementing acts if harmonized standards are not developed in a timely or satisfactory manner. Codes of Practice (Article 56) are also explicitly designated as interim tools to be used until standards are approved. The AI Office maintains the position that a track record of adhering to a Code of Practice will be considered a mitigating factor when calculating fines, and the Commission is even considering granting an additional grace period for the application of high-risk provisions through the Digital Omnibus proposal. They have effectively acknowledged the delay and prepared multiple overlapping alternatives to fill the void.

What this context implies for Korean companies is clear: you must abandon the expectation that you will receive a completed answer key in the form of standards on August 2. Instead, you must begin preparing now for what to present to the evaluators in the absence of an answer key.

14.2 Alternatives to Missing Standards: ISO/IEC 42001 and NIST AI RMF

The first alternative to fill the void left by missing standards is a Code of Practice. A Code of Practice under Article 56 is recognized as a simple means to demonstrate compliance with Article 53(1) (obligations for providers of all general-purpose AI models) and Article 55(1) (obligations for providers of general-purpose AI models with systemic risk). Conversely, failing to adhere to a Code of Practice allows the AI Office to request more detailed information; ignoring the Code of Practice actually increases paperwork.

The second alternative is international standards. ISO/IEC 42001 is a standard certifying an Artificial Intelligence Management System (AIMS) at the organizational level, and certifications have been growing rapidly in 2026. However, there is a point to note here: obtaining an ISO 42001 certificate does not automatically confer the benefit of a presumption of conformity under the AI Act. Presumption of conformity is a privilege granted exclusively to harmonized standards whose references have been published in the Official Journal of the EU, whereas ISO 42001 is an international standard outside the EU harmonization process. A provider holding only ISO 42001 still bears the burden of proof in the event of a dispute. I have seen many practitioners misunderstand this point, mistaking an ISO certificate for a legal shield. It is more accurate to view it not as a shield, but as the raw material for forging one.

Within the same 42000 series ecosystem, prEN 18286 warrants close attention. Targeting the Article 17 Quality Management System, this standard is a candidate harmonized standard focused at the individual system level and is already at the Enquiry stage. While ISO 42001 is a tool for certifying organization-wide governance, prEN 18286, once finalized, will serve as a narrower and more direct tool in its place. Rather than treating both standards as equivalent, I believe a safer two-step strategy is to build the framework using ISO 42001 now, and then layer individual system documentation on top of it once prEN 18286 is published in the Official Journal.

In third place is the NIST AI RMF. To be clear, however, no clause in the AI Act mentions the NIST AI RMF. This framework is a voluntary guidance document released by the U.S. National Institute of Standards and Technology in January 2023, and Version 1.0 remains the current edition as of 2026. It features neither certification procedures nor penalties. Built around four functions - Govern, Map, Measure, and Manage - NIST itself published a crosswalk mapping these subcategories against ISO 42001 clauses, revealing significant overlap. Govern aligns with Clauses 5 and 6 of ISO 42001, Map with impact assessments,

Measure with monitoring, and Manage with operational controls.

Thus, a practical combination in real-world operations works like this: use NIST's four functions as the structural skeleton for categorizing internal control items, and treat the AI Act clauses as the list of obligations that this skeleton must fulfill. If a Korean company has already structured its internal controls around the NIST RMF for U.S. operations, there is no need to discard that framework and start from scratch. A single mapping matrix placing the AI Act clauses, ISO 42001 clauses, and NIST RMF subcategories side-by-side for each control item will suffice. Running this single matrix requires significantly less headcount than managing three separate frameworks and pulling out different documents for every audit.

To summarize, my assessment for this section is as follows: currently, nothing can replace the legal status of harmonized standards. However, by layering ISO 42001, the NIST AI RMF, and Codes of Practice together, you can achieve a sufficient volume of documentation to present before supervisory authorities even without harmonized standards. While volume itself does not equal legal defense, it serves as circumstantial evidence of diligent preparation.

14.3 Data Governance Void Created by Delays and Voluntary Responses

Of the three sections, I view this section as the most urgent. The reason is simple: data is a raw material that cannot be retroactively fixed later.

Article 10 of the AI Act requires that datasets used for training, validation, and testing of high-risk AI systems possess relevance, representativeness, freedom from errors, and completeness. It mandates consideration of data collection and preparation, bias detection and mitigation, as well as the geographical, contextual, and behavioral environments in which the system is intended to be used. This obligation also becomes fully applicable on August 2, 2026, yet without harmonized standards, there is no recognized metric establishing what numerical values demonstrate relevance or representativeness. The rubric indicating what constitutes sufficiency remains blank.

The problem here is not merely that it takes time. The data collected today determines the future performance and limitations of the system. If you postpone data governance and wait for standards to emerge, the datasets accumulated during that period may already harden into a state lacking relevance or representativeness. Legal interpretations can be caught up on later, but a model trained on biased data cannot be rolled back after standards are published. Therefore, I interpret this void not as a matter of legal uncertainty, but as one of technical irreversibility. There is no reason to wait.

There are two areas where voluntary action can fill the gap. One is the Code of Practice for General-Purpose AI Models. It includes the obligation to draw up a publicly available summary of training content, which directly connects to data governance. Complying with the Code of Practice is recognized as a simple means of demonstrating compliance with Articles 53 and 55, and serves as a mitigating factor in calculating fines after August 2, 2026. The other is an internal documentation system covering Risk Management under Article 9, Technical Documentation under Article 11, Record-Keeping under Article 12, Human Oversight under Article 14, and Accuracy, Robustness, and Cybersecurity under Article 15. The absence of standards is no justification for delaying documentation. On the contrary, in the absence of standards, you must meticulously record the rationale behind your internal criteria and the verification logs based on those criteria so that you have something to present when supervisory authorities inquire about your basis later.

The provision allowing the AI Office to request more detailed information from providers who do not follow a Code of Practice is also worth examining carefully. It translates to

this: companies that voluntarily establish internal standards only need to demonstrate those standards, whereas companies that establish no standards at all must produce documentation from scratch according to whatever the AI Office demands. In the end, taking the initiative now spares you trouble later.

The Commission is scheduled to issue guidelines under Article 6 by February 2, 2026, and is also considering Common Specifications and grace period proposals under the Digital Omnibus. While these developments must be continuously monitored, watching trends and establishing a data governance framework today must not be reversed in priority. This was precisely my response to that compliance officer in Pangyo: do not wait for standards; create internal procedures now to self-verify the relevance and representativeness of the data currently accumulating. Standards are merely tools to validate those procedures later; they do not build the procedures for you.

Key References

Timeline August 1, 2024: AI Act enters into force. August 2, 2025: Obligations for general-purpose AI (GPAI) models become applicable. February 2, 2026: Guidelines under Article 6 scheduled to be issued. August 2, 2026: Obligations for high-risk AI systems (Annex III) become fully applicable; AI Office's power to impose fines takes effect. August 2, 2027: High-risk AI systems embedded in products (Annex I) become fully applicable; compliance grace period ends for GPAI models placed on the market before August 2, 2025.

Assessment Criteria A presumption of conformity cannot be claimed until harmonized standards are referenced in the Official Journal. ISO/IEC 42001 certification provides evidence of organizational governance, but does not inherently confer legal presumption benefits. Compliance with a Code of Practice serves as a simplified means of proof and a mitigating factor for fines, while non-compliance leads to requests for additional information. Data governance is not something to postpone while waiting for published standards; it must be applied immediately to data currently accumulating.

Checklist Items Re-verified whether your system falls under Annex III high-risk categories or GPAI categories. Reviewed and voluntarily implemented provisions regarding training data summaries from the GPAI Code of Practice. Organized the four functions of the NIST AI RMF into a control mapping matrix independently of building an ISO/IEC 42001 AIMS. Established internal documentation and verification records for each of Articles 9, 10, 11, 12, 14, and 15. Designated an authorized representative within the EU in writing.

Part V

Prohibited AI Practices

Chapter 15. The Line That Must Not Be Crossed

In a conference room of an HR tech startup in Pangyo, a proposal was tabled to add a facial expression analysis feature to their video interview tool for recruitment. This feature reads an applicant's eye movements and voice tremors to score their level of nervousness and sincerity. In Korea, similar tools are already being used in the recruitment processes of several large conglomerates, and there is no specific statutory provision directly prohibiting them. The problem was that this company happened to be about to use the same tool for hiring at its German branch. After the corporate lawyer checked the statutory provisions, the meeting took a different turn. A tool that was entirely lawful in Korea had, the moment it crossed a single border, turned into a target for fines on a criminal-penalty scale.

Article 5 of the European Union Artificial Intelligence Act (Regulation (EU) 2024/1689, hereinafter the AI Act) is the provision that creates such a scenario. Entered into force on August 1, 2024, this law takes a risk-based approach as its backbone and divides AI systems into four levels; Article 5 occupies the very top tier: prohibited practices that cannot be permitted under any circumstances. While high-risk or limited-risk systems can be placed on the market as long as obligations are fulfilled, if a system falls under Article 5, no amount of documentation will help. It cannot be developed, sold, or used. For this reason, practitioners refer to this provision as the AI Act's first hard compliance gate. Its date of application was also the earliest. It has already been fully applicable since February 2, 2025, six months after entry into force, and in case of a violation, a fine of up to 35 million EUR or 7 percent of total worldwide annual turnover for the preceding financial year, whichever is higher, is imposed pursuant to Article 99(3). This is the heaviest penalty tier in the entire AI Act.

15.1 Subliminal Manipulation and Exploitation of Vulnerabilities

Article 5(1)(a) prohibits the deployment of subliminal techniques beyond a person's consciousness or the intentional deployment of manipulative or deceptive techniques to materially distort a person's behavior and significantly impair their decision-making ability, thereby causing significant harm. Point (b) prohibits exploiting vulnerabilities arising from age, disability, or socio-economic situation to produce the same result. What the two provisions jointly require ultimately comes down to two conditions: material distortion of behavior, and significant harm. A violation is established only when both are confirmed together. Although source materials clarify that minor inconveniences or trivial financial losses do not fall under this, I anticipate that this line will actually be drawn fairly low in practice. Given the repeated statements defining significant harm as a level capable of severely impacting health, safety, or fundamental rights, it seems unlikely that enforcement authorities will set this threshold loosely or generously.

There is an easily overlooked aspect here: the provision focuses on the effect of the system, not the intention of the developer. Even if the Pangyo startup had zero intention to manipulate when creating its recruitment tool, if that tool actually obscures an applicant's judgment and disadvantages them, the provision applies all the same. This is the most frequently misunderstood point in practice. Everyone argues that they had no ill intent, but the question asked by the AI Office is different. They look not only at the intended purpose but also at reasonably foreseeable misuse. If a single nudge embedded in a commerce app's UI induces unwanted payments from children or elderly users, it is difficult to escape a finding of violation even if the designer did not desire such an outcome.

A case experienced by a medical device company in Busan is similar. While building a chatbot that recommends dietary supplements to elderly patients, they structured the conversation flow with repeated exposure and phrases of urgency. While such a design might pass as a marketing technique domestically, in the EU it carries the potential of violating point (b) in that it targets the vulnerability of old age to distort purchasing behavior.

15.2 Social Scoring

Article 5(1)(c) prohibits systems that evaluate or classify natural persons - whether by public or private entities - based on their social behavior or personality characteristics, leading to detrimental treatment in contexts unrelated to the original context or detrimental treatment that is disproportionate to the gravity of the behavior. It is easy to picture a scenario where the state surveillance ranks citizens, but the provision applies equally to private companies. If a credit rating agency scrapes social media posts or online behavioral data to construct a comprehensive trustworthiness score and uses that score for loan screening or insurance enrollment, it falls squarely under this prohibition.

The source materials include a proviso reading "where the legal conditions are met," but precisely what those conditions are remains unconfirmed. I view this as an area where we must await subsequent guidelines from the European Commission. Nevertheless, the practical assessment criteria are clear: Are the evaluated behavior and the context in which the score is used distinct from each other? Is the resulting treatment disproportionate to the gravity of the original behavior? If the answer to these two questions is yes, it must be treated as a red flag. This means that many alternative credit scoring models operated by the domestic fintech industry need to be re-examined against these criteria.

What enforcement authorities focus on when viewing this provision is also the effect actually produced by the system, rather than the developer's explanation. Even if the designer of a credit scoring model states that it was intended solely for loan evaluation, if that score flows into settings disconnected from the original context - such as employment or lease screening - I view that a finding of violation still holds. Since the authority to impose fines has already been open since February 2, 2025, the AI Office will likely proceed with sanctions without hesitation the moment a violation is confirmed.

15.3 Real-Time Remote Biometric Identification: Prohibition in Principle and Exceptions

Article 5(1)(h) prohibits, in principle, the use of real-time remote biometric identification systems in publicly accessible spaces for the purpose of law enforcement. A representative example is a system that scans faces in real time via CCTV in places accessible to anyone - such as public squares, subway stations, or stadiums - to match identities.

The exceptions are not broad. One of the following three conditions must be met, and on top of that, the threshold of being "strictly necessary" must be passed: targeted search for victims of abduction, human trafficking, sexual exploitation, or missing persons; preventing a specific, substantial, and imminent threat to life or physical safety, or a terrorist attack; and searching for suspects of crimes listed in Annex II punishable in the Member State concerned by a custodial sentence of at least four years. The threshold of being "strictly necessary" will likely be interpreted narrowly. I read this to mean that if the objective can be achieved through other less intrusive means, the exception will not be granted.

If a Korean security firm or smart city solution company intends to supply real-time facial recognition CCTV to an EU municipality, this provision is the first hurdle. While facial recognition for non-law enforcement purposes - such as theft prevention in commercial facilities or store marketing - is not directly subject to Article 5(1)(h) itself, separate reviews are needed to determine whether it was designed as post-identification rather than real-time, and whether it falls under the definition of publicly accessible spaces. Since precisely which crimes are listed in Annex II cannot be verified from source materials alone, cross-checking the original text prior to contracting is an absolute necessity.

From the perspective of enforcement authorities, I read real-time biometric identification in public spaces as carrying among the heaviest weight among provisions touching upon privacy and freedom of movement. Therefore, while carefully balancing public safety against fundamental rights, they are expected to apply strict standards to any use that diverges even slightly from exception requirements. For domestic companies contracting with local governments, it is safer to check how the ordering authority has documented the exception conditions. This is to avoid a scenario where the supplier shares liability simply because the ordering authority's assessment was incorrect.

15.4 Emotion Recognition: Outright Banned in Workplaces and Educational Institutions

Article 5(1)(f) prohibits the placing on the market, putting into service, or use of AI systems to infer emotions of natural persons in the workplace and educational institutions. Here, there is a point that domestic practitioners frequently confuse. There are two provisions in the AI Act addressing emotion recognition, and their natures are completely different. Article 50 is a transparency obligation provision requiring that users be informed when an emotion recognition system is used. If a call center uses a tool to analyze customer emotions during calls, simply providing notification suffices under Article 50. Conversely, in workplaces and educational institutions where Article 5(1)(f) applies, it does not end with notification. The act of using it in the first place is prohibited. The very same emotion recognition technology transforms from a matter of disclosure to a prohibited practice subject to criminal-level penalties based solely on location.

Exceptions are narrowly carved out for medical or safety reasons. Systems monitoring pilot or air traffic controller fatigue are cited as examples of such exceptions. Conversely, I view purposes such as measuring employee work engagement or scoring student class concentration as not falling under the exception. Even if wrapped in words like efficiency or productivity, what the provision examines is the actual context of use. It is safer to interpret the interview expression analysis tool examined by the Pangyo startup as falling within the workplace category of this provision, given that recruitment is the gateway to entering the workplace.

Enforcement authorities are unlikely to accept the exception of medical or safety purposes broadly. The moment expressions like efficiency improvement or productivity enhancement appear in contracts or product manuals, they may instead be read as evidence that the system is not eligible for the exception. Korean companies selling services in HR management, edtech, or customer service should start by checking whether emotion inference features can be toggled on and off as an option, and whether the system is structured to prevent that option from being enabled at all for workplace and educational institution clients.

15.5 Allowed Domestically, Prohibited in the EU

When conducting consultations, I am frequently asked to list items that are absent in domestic law but prohibited in the EU. I would highlight three main points first.

First is emotion analysis in the workplace. Domestically, there is no law directly prohibiting this. While indirect regulation is possible through the Labor Standards Act or the Personal Information Protection Act, there is no statutory provision prohibiting emotion recognition itself. In the EU, it is quite the opposite. As seen in 15.4, Article 5(1)(f) directly prohibits it.

Second is private comprehensive scoring. Domestic fintech and e-commerce industries already operate various alternative credit scoring models, and the practice of reflecting online behavioral data into credit scores has generally been allowed within regulatory sandboxes. In the EU, because the prohibition on social scoring applies without distinguishing between public and private entities, introducing the same model intact significantly increases the risk of violating Article 5(1)(c).

Third is the indiscriminate collection of facial images. The method where domestic companies build facial recognition databases by scraping images in bulk from the internet or CCTV footage was not blocked as a business itself, as long as consent requirements under the Personal Information Protection Act were somehow satisfied. Article 5(1)(e) of the AI Act classifies the act of creating or expanding facial recognition databases through such untargeted scraping itself as a prohibited practice.

Beyond these three, Article 5 also prohibits, under point (d), systems that assess the risk of a natural person committing a criminal offense based solely on profiling or personality traits, and under point (g), categorization systems that infer race, political opinions, trade union membership, religious beliefs, or sexual orientation from biometric data. While tools that assist human judgment based on objective crime-related facts remain exempt under point (d), models that calculate risk levels solely through profiling fall outside the exception. I view cases of domestic scoring firms crossing this line as not uncommon.

A point that must be addressed at this juncture is extraterritorial application and the requirement for an authorized representative. Article 2(1)(c) of the AI Act stipulates that even if a company is headquartered outside the EU, it falls within the scope of the Act if the output produced by the system is used in the EU. Even if a Seoul headquarters runs the model exclusively on Korean servers, the moment its output is delivered to a German

branch or French client, Article 5 applies immediately. Non-EU companies must designate in writing an Authorised Representative established in the EU prior to placing a system on the market, and domestically they are also separately subject to the Framework Act on AI, placing them in a position of having to manage a single service under two regulatory frameworks simultaneously. Since the domestic Framework Act on AI does not yet feature an absolute prohibition list like Article 5, the burden of bridging the gap between the two laws ultimately falls upon internal corporate compliance organizations.

To this, a recent update must be added. The Digital Omnibus on AI, provisionally agreed upon by the European Parliament and the Council on May 7, 2026, and finally approved by the Council on June 29, inserted two new points into Article 5: the generation of non-consensual intimate images (NCII), including so-called "nudify" tools that realistically synthesize or manipulate body parts of specific individuals, and the generation of child sexual abuse material (CSAM). AI systems that synthesize body parts without the subject's free, specific, informed, and explicit consent are now subject to Article 5 prohibition. Application begins on December 2, 2026, and the fine upon violation is set at 15 million EUR or 7 percent of total worldwide turnover, whichever is higher - lower than the 35 million EUR ceiling for other Article 5 violations. Korean companies operating image editing apps or character generation services that incorporate deepfake synthesis features must address this provision within this year.

The effective response when facing Article 5 is blocking at the design stage rather than ex-post remediation. Embedding prohibited items in advance into every system architecture, algorithm, data collection method, and user interface incurs far less cost than removing features after launch. As there are not a few passages where the statutory wording reads abstractly, I advise incorporating a procedure into early planning stages to consult lawyers or consultants well-versed in the AI Act and AI ethics to review both your system's intended purpose and reasonably foreseeable misuse.

Reference Materials

Dates August 1, 2024: AI Act entry into force. February 2, 2025: Article 5 fully applicable; fine imposition authority commences. May 7, 2026: Provisional agreement on the Digital Omnibus. June 29, 2026: Final approval by the Council. December 2, 2026: Prohibition on NCII and CSAM generation becomes applicable.

Fines General Article 5 violations: 35 million EUR or 7 percent of total worldwide annual turnover, whichever is higher (Article 99(3)). NCII and CSAM generation violations: 15

million EUR or 7 percent of total worldwide annual turnover, whichever is higher.

Prohibited Practices at a Glance (a) Subliminal manipulation. (b) Exploitation of vulnerabilities. (c) Social scoring, across both public and private sectors. (d) Crime risk assessment based solely on profiling; assisting human evaluation based on objective facts is exempt. (e) Construction of facial recognition databases via untargeted scraping of facial images. (f) Emotion recognition in workplaces and educational institutions; medical or safety purposes are exempt. (g) Categorization inferring race, political opinions, trade union membership, religion, or sexual orientation based on biometric data. (h) Real-time remote biometric identification in public spaces for law enforcement purposes; three narrow exceptions apply. Generation of NCII and CSAM: Newly added effective December 2026.

Our System Assessment Checklist Is our AI functionality designed in a manner that obscures user judgment? Are there features targeting vulnerable groups such as children, persons with disabilities, or the elderly? Are there features inferring emotion or state from the behavior, facial expressions, or voice of employees or students? Does it aggregate an individual's online behavior to score trustworthiness or risk levels and utilize the results in another context? Does it collect facial images in bulk from the internet or CCTV to build a database? Do body image synthesis or editing features operate without consent verification procedures? Have we designated an Authorised Representative within the EU, and addressed dual obligations alongside the domestic Framework Act on AI?

Article 5 is not a provision resolved merely by maintaining documentation and following procedures, unlike high-risk system obligations covered in other chapters. For high-risk systems, registering them, undergoing assessment, and attaching markings open a path to remain on the market. The moment a system hits Article 5, no such path exists. Crossing the line means it is over. That facial expression analysis feature in the Pangyo conference room was ultimately removed from hiring at the German branch. It is still being used for domestic recruitment. How much longer we can sustain the current situation, where the exact same code shifts between lawful and unlawful every time it crosses a border, is another question entirely.

Part VI

Intersection with Korean Law

Chapter 16. Dual Regulation of the Framework Act on AI and the EU AI Act

Let me start with the story of a hiring solution startup in Pangyo. This company developed an AI that analyzes resumes and interview videos to rank job applicants. While supplying its product to the HR team of a major domestic conglomerate, it was also preparing to deploy the exact same product to a subsidiary of an automotive parts company located in Hungary. The CEO asked me this: Since they passed the Framework Act on AI review in Korea, wouldn't they be able to sell it in Europe just as it is? I told him right then and there that the answer was no. Although both laws target the same activity - recruitment - for regulation, the required documentation, procedures, and even the evaluating entities are completely different. This company now finds itself in a position where it must satisfy two sets of laws simultaneously. Chapter 16 addresses precisely this space - the reality in which Korean companies actually stand between domestic and EU laws.

16.1 Two Layers of the Amended Framework Act on AI Taking Effect in 2026

Korea's Framework Act on the Development of Artificial Intelligence and Creation of Reliability (hereinafter referred to as the Framework Act on AI) saw its main body of law take effect on January 22, 2026. What came into force at this point was the core structure of regulation. Obligations such as safety measures for high-impact AI business operators, labeling obligations for generative AI, and the obligation to designate a domestic representative were activated at that time. However, as the Enforcement Decree was amended and took effect on July 21, 2026, another layer was added. The center of gravity of this Enforcement Decree amendment lies not in regulation, but in promotion. The AI Product and Service Verification System and preferential measures for public procurement were specifically detailed in this Enforcement Decree.

I summarize this structure as follows: the main law set up the regulatory framework in January 2026, and half a year later, the Enforcement Decree added muscle to foster the industry on top of that framework. A single law essentially houses both a policing hand and a helping hand, and this duality is the primary key to understanding the Framework Act on AI. The EU AI Act is closer to a mandatory regulatory instrument tightly woven with penalties and market entry requirements, whereas the Korean Framework Act on AI is closer to laying down a minimal safety net and then incentivizing companies through procurement and certification. If you miss this difference in tone when reading the provisions, you will end up misjudging domestic regulations by expecting EU-style strictness.

The government announced that it would operate a grace period of at least one year during the initial phase of implementation. This means providing recommendations for corrective action and room for improvement first, rather than immediately imposing administrative fines. The EU virtually lacks such a grace period concept. Starting August 2, 2026, the AI Office will directly exercise its power to impose fines for violations of high-risk AI and GPAI obligations. In the same year, both laws will enter full application side by side, yet one begins with guidance while the other begins with fines. This difference in tone creates an illusion for Korean companies. I believe cases will arise where practices that passed without issue domestically are carried straight into the EU, resulting in receiving notices of financial penalties right from the start.

16.2 High-Impact AI and High-Risk AI: Similar Concepts, Different Mesh Sizes

Both laws speak the same language in that they adopt a risk-based approach. However, the way they cast their nets is different.

Article 2, Item 4 of the Framework Act on AI defines high-impact artificial intelligence as artificial intelligence that has a significant impact on human life, bodily safety, or fundamental rights, or is likely to cause such risks, while strictly limiting its scope to cases where it is used in 11 domains specified by Presidential Decree. These include energy supply, drinking water production processes, provision of healthcare and operation of related systems, development and use of medical devices and digital health products, management and operation of nuclear power facilities, criminal investigation and arrest, decisions significantly affecting rights and obligations such as hiring and loan screening, operation and management of transport vehicles, facilities, and systems, public services, student evaluations in education, and other domains specified by Presidential Decree. Since ten items are enumerated and the last one is a delegation provision to the Presidential Decree, they are commonly referred to as the 11 domains.

The EU AI Act opens two doors under the single concept of high-risk. One is AI used as a safety component of existing product safety legislation listed in Annex I (such as medical devices, automobiles, and toys), and the other consists of specific use cases specified in Annex III - namely biometric identification, critical infrastructure, education, employment, essential private and public services, law enforcement, migration, asylum and border control management, and administration of justice and democratic processes. A striking point emerges here. EU Annex III explicitly brings in not only law enforcement, but also migration, asylum, and border control management, as well as the administration of justice and democratic processes. While Korea's 11 domains include criminal investigation, there are no items directly targeting migration, judicial administration, or election-related systems. The mesh size of the net itself is fundamentally different. A Korean company developing AI for screening immigrants or assisting judicial trials must separately check the possibility of its product being classified as a high-risk system when selling to the EU, even if it does not fall under high-impact domains under Korean law.

The EU Act contains another mechanism that does not exist in domestic law: the filter provision of Article 6(3). Even if an AI system is listed in Annex III, it opens a pathway to

escape high-risk classification if it poses no significant risk of harm to health, safety, or fundamental rights, or does not materially influence the outcome of decision-making. While this filter offers breathing room for companies, it presents a point of divided interpretation for enforcement authorities. I read this provision as follows: being listed in Annex III does not automatically make a system high-risk; rather, there is an additional door through which companies must prove the lack of material influence themselves. Domestic law lacks such an exception filter. If a system falls into one of the 11 domains, it becomes a high-impact AI system by default. While Korea offers no exit door, its criteria for judgment remain clear, whereas the EU leaves the door open but imposes the burden of proof on companies to pass through it.

The weight of obligations also differs. Providers of EU high-risk AI must establish a risk management system (Article 9), data governance (Article 10), technical documentation (Article 11), record-keeping (Article 12), human oversight design (Article 14), and accuracy, robustness, and cybersecurity (Article 15), while completing third-party conformity assessments, CE marking, and registration in the EU database before entering the market. It is a mandatory ex-ante certification system. Korea's Framework Act on AI imposes obligations on high-impact AI operators to take measures to ensure safety and reliability, but the method relies heavily on best-effort measures such as self risk management, documentation, and user notification, maintaining a phased structure that escalates from recommendations for corrective action and improvement orders to administrative fines in case of violations. This means it is closer to ex-post verification and requests for improvement rather than ex-ante certification. If the EU's approach is to inspect documents at the door, Korea's approach is to let products in first and ask them to fix issues if problems arise. Which approach is right is a matter of policy judgment, but from a practitioner's perspective, the conclusion is clear: you must establish separate schedules - an ex-ante certification timeline for products targeting the EU and an ex-post compliance system for products targeting the domestic market.

Obligations for Annex III high-risk AI systems fully apply from August 2, 2026, while those embedded in Annex I products will be fully implemented from August 2, 2027. Korea's high-impact AI obligations are already active as of January 22, 2026. In terms of timing alone, Korea is ahead by more than half a year, but I emphasize once again that being ahead does not mean being more stringent.

16.3 Deepfakes: Visible to the Human Eye or Readable by Machines?

Both laws address deepfakes head-on, but they differ in where they require labeling methods.

Article 31 of the Framework Act on AI structures transparency obligations into three components: prior notification that high-impact or generative AI is operating, labeling indicating that content is AI-generated, and separate notification and labeling for deepfake outputs. The key is that the labeling must be verifiable through the user's senses. The law requires methods that can be easily verified visually, auditorily, or using software, and that can be clearly recognized taking into account the age, physical, or social conditions of the primary users. I view putting a watermark in a corner of the screen or displaying a notification caption at the beginning of playback as falling under this requirement. This represents human-perceivable labeling - a method where people perceive and recognize it with their eyes or ears.

Article 50 of the EU AI Act takes a different approach. Article 50(2) requires providers of systems generating synthetic audio, image, video, or text content to mark their outputs in a machine-readable format (machine-readable marks). Rather than captions visible to the human eye, this refers to watermarking or metadata signatures that detection software can read inside the file. This marking is even subject to criteria that it must be effective, interoperable, robust, and reliable. Furthermore, Article 50(4) separately requires deployers generating deepfakes to disclose that the image, audio, or video has been artificially generated or manipulated. In summary, the EU imposes a dual requirement of machine-readable marking and human-facing disclosure, while Korea puts weight on human-perceivable labeling.

This difference dictates the direction of technology investment. If looking solely at the domestic market, screen captions or voice announcements are sufficient to fulfill the obligation. However, to export the same content to the EU market, you must separately integrate watermarking technologies or content provenance standards (think of frameworks like C2PA). Taking the AI avatar advertising video produced by the Pangyo startup as an example, a single line of AI generation caption at the bottom of the screen suffices for the domestic version, but uploading the exact same video to a European entity's channel means embedding a machine-readable signature directly into the video file itself. It is not a matter of adding one more caption, but rather integrating a separate module into the content production pipeline.

The timelines for implementation are also misaligned. Korea's Article 31 obligation is already active as of January 22, 2026. EU Article 50 applies from August 2, 2026, with a grace period granted until December 2, 2026, for labeling and detection obligations of generative AI systems already on the market before that date, while content created prior to that date is not subject to retroactive labeling obligations. As of this moment, Korean companies should already be fulfilling domestic labeling obligations and treating EU labeling obligations as being in a preparation phase.

16.4 Domestic Representatives and EU Authorised Representatives: Thresholds vs. No Thresholds

Both laws share the concept of requiring offshore business operators to appoint a representative within their jurisdiction. However, they diverge decisively on who bears that obligation.

Article 36 of the Framework Act on AI and Article 29 of its Enforcement Decree require the designation of a domestic representative only for AI business operators without an address or place of business in Korea that exceed a certain size threshold. The size criteria are: annual revenue of 1 trillion KRW or more in the preceding year, annual revenue of 10 billion KRW or more in the AI service sector in the preceding year, or an average of 1 million or more daily domestic users during the preceding 3 months as of the end of the previous year. Crossing any one of these three requires appointing a representative, but conversely, overseas operators below these thresholds are exempt from the domestic representative obligation. The role of the domestic representative is also distinct: representing the operator in submitting safety assurance results for large-scale AI, requesting verification of high-impact AI, and supporting the implementation of safety and reliability assurance measures for high-impact AI.

The EU side has no such revenue or user count thresholds at all. Article 22 requires providers of high-risk AI systems, and Article 54 requires providers of general-purpose AI (GPAI) models, to designate an Authorised Representative within the EU via a written mandate prior to placing their products on the EU market. Regardless of whether a company is small or large, or what its revenue is, the moment it releases a high-risk AI or GPAI onto the EU market, it must appoint a representative without exception. The sole exception is open-source GPAI models that do not pose systemic risks. The representative verifies and retains compliance documentation for 10 years, cooperates with market surveillance authorities and the AI Office upon request, and handles procedures such as EU database registration. If the provider violates the law, the representative also holds the authority to terminate the mandate and immediately notify the authorities.

Looking at this difference through the eyes of the Pangyo startup: if this company attempts to launch its hiring AI as a high-risk system in the European market, it must appoint an EU Authorised Representative to enter the market even if it is an early-stage startup with revenue under 10 billion KRW. Conversely, if a European AI enterprise enters Korea to sell its hiring AI, it can conduct business without appointing a domestic

representative as long as its domestic service revenue or user count does not exceed the Enforcement Decree thresholds. Between Korea, which casts its net based on scale, and the EU, which casts its net based solely on the nature of the activity, I see the EU burden weighing relatively heavier on smaller Korean enterprises.

Please also take note of the fine structure. For violations of prohibited practices, the EU imposes up to 7% of total worldwide annual turnover or 35 million EUR, whichever is higher; for violations of other high-risk or GPAI obligations, it imposes up to 3% or 15 million EUR, whichever is higher. This enforcement power takes effect from February 2, 2025 for prohibited practices, and from August 2, 2026 for the remaining obligations. Without an Authorised Representative, there is no contact point to address this fine risk, and the outcome is effectively no different from being pushed out of the EU market.

16.5 AI Product and Service Verification System and Procurement Preferences

This section addresses incentives rather than regulations. The amended Enforcement Decree that took effect on July 21, 2026 detailed the AI Product and Service Verification System. When a company applies for verification with the Korea Software Industry Association (KOSA), an affiliate of the National Information Society Agency, the Telecommunications Technology Association (TTA) technically evaluates the actual usage level of the artificial intelligence system, and KOSA issues a verification certificate to products and services that pass the evaluation. AI products and services that receive this verification certificate will enjoy preferential treatment in the public procurement market starting in August 2026.

The EU AI Act has no such concept of procurement preference. The EU chose an approach of locking market entry behind conformity assessments and CE marking, and products that cross that threshold compete in the market without additional preferences. Conversely, Korea leaves market entry relatively open provided minimum safety requirements are met, guiding companies to voluntarily undergo verification procedures through the clear incentive of public procurement. The picture of the EU locking the door with a whip versus Korea lining companies up with a carrot becomes distinct here.

I advise practitioners on this verification system as follows: if you are an AI company doing business with domestic public institutions, approach obtaining this verification certificate not as regulatory compliance, but as a sales strategy. Conversely, if it is a product line targeted at the EU market, this verification certificate cannot replace the EU's CE marking or conformity assessment. The domestic verification certificate is a marker meaningful only within the domestic procurement market; in the EU market, you must separately go through the Annex III procedures again from scratch. I have actually seen cases where companies mistook the two as equivalent and presented domestic verification certificates to EU buyers.

Reference Materials

Dual Regulation Comparison Table

Item / Domestic Framework Act on AI / EU AI Act Regulatory Nature / Main law on 2026-01-22 for regulatory framework, Enforcement Decree on 2026-07-21 for promotion / Mandatory ex-ante certification, immediate fines upon violation Core Concept /
--

High-Impact Artificial Intelligence (Article 2, Item 4; 11 domains in Presidential Decree) / High-Risk AI (Article 6; Annexes I & III) Exception Filter / None; determined upon falling within the 11 domains / Article 6(3); exemption possible if no material influence Domain Scope / Energy, drinking water, healthcare, nuclear energy, criminal investigation, hiring/loans, transport, public services, education, etc. / In addition to the above, includes migration, asylum, border control management, administration of justice, and democratic processes Deepfake Labeling / Article 31; notification/labeling perceivable by sight or sound / Article 50; machine-readable watermarking + separate disclosure Representative Designation Criteria / Article 36 & Enforcement Decree Article 29; choose one among 1 trillion KRW revenue, 10 billion KRW AI revenue, or 1 million users / Articles 22 & 54; all high-risk & GPAI without thresholds Verification & Procurement System / AI Product and Service Verification System (issued by KOSA, evaluated by TTA), procurement preference from 2026-08 / No such system; only CE marking and market competition exist

Checklist for Teams Preparing Concurrent Launch in Korea and the EU

First, have you separately indicated whether your product falls under the domestic 11 high-impact domains and EU Annex III items respectively? Second, have you reviewed whether there is room to escape high-risk classification under the EU Article 6(3) filter and documented the rationale for that judgment? Third, have you embedded human-perceivable labeling for domestic use and machine-readable watermarking for EU use into generated content, respectively? Fourth, have you calculated revenue and user count thresholds to confirm whether a domestic representative must be designated? Fifth, if expanding to the EU, have you completed the designation of an EU Authorised Representative via a written mandate, regardless of revenue scale? Sixth, are you separately managing the domestic verification certificate and the EU conformity assessment as distinct procedures?

That CEO in Pangyo ultimately pinned two compliance calendars side by side on his wall. One was for domestic high-impact verification and procurement preference application schedules, and the other was for EU Authorised Representative mandates and Annex III conformity assessment schedules. Same product, same company, yet two calendars. I am still watching to see which calendar this company will fill first next quarter.

Chapter 17. Intersection with Data Protection

I received a call from a medical imaging AI company based in Pangyo. They were planning to supply software that assists in diagnosing lung cancer to several hospitals in Germany, but their local representative sent them a questionnaire asking: Where were the chest X-ray images used for training obtained? Was patient consent obtained? If consent was not obtained, on what legal basis were they processed? The CEO asked me why data privacy issues kept coming up when they thought they only needed to prepare for compliance with the AI Act. I replied as follows: the AI Act and data protection regulations are not two separate matters that can be divided in the first place. The moment you record a person with a camera, collect medical records, or train a model on conversations, it is already a data protection issue, and the AI Act is layered on top of it.

17.1 Simultaneous Application of the GDPR and the Personal Information Protection Act

The AI Act (Regulation (EU) 2024/1689) entered into force on August 1, 2024. Article 2(7) of this Regulation explicitly provides that Union law on the protection of personal data, privacy, and confidentiality of communications applies to personal data processed in connection with the rights and obligations laid down under the AI Act, and that the AI Act does not affect the GDPR (Regulation (EU) 2016/679) or the Law Enforcement Directive (Directive (EU) 2016/680). Put simply, the AI Act does not displace the GDPR; rather, it sits on top of it.

The focus of the two laws diverges. The GDPR examines how an individual's data is collected, processed, and transferred, and whether that person's rights are safeguarded. The AI Act looks at the safety, performance, and risk management of the system itself. However, because both address transparency, accountability, data quality, and human oversight, there are many areas of overlap. A prime example is automated decision-making. If an AI system makes decisions about an individual automatically, Article 22 of the GDPR applies, and if that system is classified as a high-risk AI system, separate obligations under the AI Act apply as well. In effect, both laws address the exact same act simultaneously.

This is where Korean companies often miss a key point. A Korean company entering the EU is actually caught in a three-layered net. The first is the GDPR. Due to the extraterritorial applicability provision in Article 3, even if a company has no establishment in the EU, the GDPR applies directly if it offers goods or services to individuals in the EU or monitors their behavior. The second is Korea's Personal Information Protection Act (PIPA). As a Korean company, it must naturally comply with domestic laws, and provisions such as those governing cross-border transfers apply equally to data flowing out to the EU. The third is the AI Act. Under Article 2(1)(c), if the output of an AI system is used within the EU, it falls under its scope of application, and since provisions related to personal data are embedded within the AI Act itself, companies end up going through a three-fold regulatory filter. This was my answer to the CEO of the Pangyo company: data protection is not separate from the AI Act, but is already embedded within it.

The penalty structure is also twofold. The AI Act is implemented in phases. Starting February 2, 2025, violating prohibited AI practices under Article 5 carries fines of up to €35 million or 7% of total worldwide annual turnover, whichever is higher. From August 2,

2026, obligations concerning high-risk AI systems and general-purpose AI models will carry fines of up to €15 million or 3% of turnover, at which point the AI Office's enforcement authority to impose fines also comes into effect. If a GDPR violation overlaps, a separate fine of up to €20 million or 4% of worldwide turnover may be added. This means that violating both laws entails separate fines for each, meaning a single data processing mistake could result in two separate lines of penalties on the ledger.

There is also a welcome aspect for practical compliance. Article 27 of the AI Act requires deployers of high-risk AI systems to conduct a Fundamental Rights Impact Assessment (FRIA), and specifies that this assessment can complement the Data Protection Impact Assessment (DPIA) under Article 35 of the GDPR. This means that a company already performing a DPIA can simply add the items required by the AI Act to that framework and consolidate them into a single assessment. In other words, there is no need to produce two separate sets of documentation - a small relief for compliance officers.

17.2 Cross-Border Data Transfers and Adequacy Decisions

The AI Act contains no new provisions regarding cross-border data transfers. When personal data leaves the EU via an AI system, Chapter V of the GDPR governing cross-border transfers applies entirely and fully. Therefore, there is only one question a Korean enterprise must ask: does a path exist in our current data flows to satisfy Chapter V of the GDPR?

The easiest path is an adequacy decision. If the European Commission recognizes that a non-EU country provides an adequate level of data protection equivalent to that of the EU, transfers of data to that country are permitted without requiring Standard Contractual Clauses or separate safeguards. South Korea received this adequacy decision on December 17, 2021. However, a crucial distinction must be made here. One often sees materials claiming that South Korea received full approval for personal data transfers from the EU based on this decision, but that is an inaccurate description. This decision covers only the scope supervised by the Personal Information Protection Commission (PIPC). Personal credit information under the supervision of the Financial Services Commission - that is, financial data processed pursuant to the Credit Information Use and Protection Act - is excluded from this adequacy decision. For banks, credit card companies, or fintech firms to transfer credit information of EU customers to South Korea, they must still put in place separate safeguards such as Standard Contractual Clauses (SCCs). Non-financial companies, such as medical imaging AI developers, fully benefit from the adequacy decision, but companies building fintech or credit-scoring AI must verify this exception.

Traces of cross-border transfers also appear indirectly in the AI Act. Article 2(4) provides that where public authorities of a third country or international organizations use AI systems within the framework of law enforcement or judicial cooperation, the AI Act does not apply if that authority provides adequate safeguards regarding the protection of individuals' fundamental rights. This borrows the exact conceptual terminology used by the GDPR. This signifies that the language used by both laws to address cross-border transfers originates from the same root, and in practice, applying the evaluation criteria of Chapter V of the GDPR will not deviate significantly from the required standards.

The case where data is transferred onwards from South Korea to a third country is also easily overlooked. Thanks to the adequacy decision, the path from the EU to South Korea is smooth. However, if a Korean company transfers that data onwards to a development

team in Vietnam or a cloud server in the United States, that second transfer is independent of South Korea's adequacy decision. The company must re-evaluate whether that country has received a separate adequacy decision from the EU, or whether safeguards like SCCs are required. Because AI development frequently involves data passing through servers and personnel across multiple countries, it is safer to map out this second transfer segment in a table.

17.3 Legal Basis for Processing Personal Data in Training Datasets

Article 10 of the AI Act imposes data governance obligations on providers building high-risk AI systems. It requires that datasets used for training, validation, and testing be relevant, sufficiently representative, free of errors, and complete. However, this requirement in no way implies that a lawful basis for processing under the GDPR is not required when such data contains personal data. Data quality requirements and legal basis requirements are distinct criteria that stand side-by-side. No matter how clean and representative data may be, if there was no lawful basis to collect the individual's information in the first place, a GDPR violation remains.

Article 6(1) of the GDPR lists six legal bases for processing personal data, and in AI model development, point (f), legitimate interest, is commonly relied upon in practice. The European Data Protection Board (EDPB), in Opinion 28/2024 adopted on December 17, 2024, outlined how to apply this basis to the development and deployment stages of AI models. Issued upon request by the Irish Data Protection Commission (DPC), this opinion acknowledges that legitimate interest can serve as a legal basis for AI model development, but requires a three-step test for its assessment: balancing whether the interest pursued is lawful, clear, and real (not speculative); whether the processing is strictly necessary to achieve that interest; and whether that interest overrides the rights and freedoms of the data subject. Citing conversational agents assisting users or cybersecurity enhancements as examples, the EDPB asserts that while such services may proceed under legitimate interest, companies must prove that the processing is strictly necessary and that a balancing test of rights was genuinely performed.

The picture in South Korea has a different texture. Korea's PIPA places consent at the core of processing bases. However, Article 28-2 allows for the processing of pseudonymized information without data subject consent for purposes of compiling statistics, scientific research, and archiving in the public interest. In August 2025, the PIPC released guidelines on processing personal data in generative AI, establishing directions for applying these pseudonymization provisions to AI training data. Thus, while Europe seeks to justify the processing of personal data in AI training datasets through the balancing mechanism of legitimate interest, South Korea does so through two avenues: consent and pseudonymization. Because the requirements for these two bases do not overlap perfectly, whether data collected in Europe based on legitimate interest can be reused in South Korea under domestic standards must be independently reviewed by each company.

This area remains in flux. It has been only about half a year since the EDPB opinion was issued, and even less time since the PIPC guidelines were published. Definitive answers to questions such as how to apply the balancing test to large-scale training data collected via web scraping, or whether pseudonymization provisions extend to the AI fine-tuning phase, are yet to be established. At this point, my view is as follows: while the broad trajectory splits between balancing tests in Europe and consent/pseudonymization in Korea, there remains room for detailed criteria to be refined several more times. This is not to suggest postponing decisions, but rather to document the legal basis established today and re-align it as new guidelines are released.

Article 59 of the AI Act opens one additional exception. If an AI system is developed within a regulatory sandbox for reasons of substantial public interest, such as public safety or public health, personal data originally collected for other purposes may be reused for training. However, the conditions are strict: achieving the purpose must be impossible using anonymized or synthetic data; the data must be retained in a segregated environment; it cannot be shared outside the sandbox; and it must be deleted upon completion of participation. The sheer density of these conditions signals that this exception can hardly serve as a routine commercial basis.

Reference Materials

Dates December 17, 2021: Adoption of the EU GDPR adequacy decision for South Korea (excluding personal credit information supervised by the Financial Services Commission)
August 1, 2024: Entry into force of the AI Act
December 17, 2024: Adoption of EDPB Opinion 28/2024
February 2, 2025: Enforcement of prohibited AI practices under Article 5 of the AI Act (up to €35 million or 7% of turnover)
August 2025: Release of PIPC guidelines on processing personal data in generative AI
August 2, 2025: Implementation of obligations for general-purpose AI providers
August 2, 2026: Full enforcement of obligations for high-risk AI systems; launch of the AI Office's fine-imposing authority (up to €15 million or 3% of turnover)

Checklist Does the training dataset contain personal data, and if so, is the legal basis under Article 6 of the GDPR documented? Are special categories of personal data (subject to Article 9) intermingled? Have transfer pathways - from the EU to Korea, and from Korea to third countries - been verified separately? For financial sector data, is reliance placed on SCCs rather than the adequacy decision? Is there room to integrate the DPIA and FRIA into a single assessment? Does the EU authorized representative also handle data protection document retention and inquiry responses?

Part VII

Execution: Building a Compliance Program

Chapter 18. 90-Day Implementation Plan

At an HR software company in Pangyo, a compliance team lead was summoned to the CEO's office. A European client had inquired about their AI Act compliance status before signing a contract, and the CEO asked the team lead: What do we need to prepare within 90 days, and where do we even begin? In the team lead's hand was a legal text containing over a hundred articles plus annexes, with no clear sense of which page to open first.

The answer to this question lies in order. The AI Act does not hit a company all at once. Different articles have different enforcement dates, and the weight of the fines varies as well. While 90 days is too short to comply with every single article in the law, it is plenty of time to establish the order of what to tackle first. This chapter divides that sequence into four steps: building an inventory, classifying systems, measuring gaps, and setting priorities.

18.1 Building an AI Inventory: Department by Department, Down to Invisible AI

To be clear, nowhere in the AI Act does an article explicitly state that you must create an inventory. Yet the reason this task takes the top spot in the 90-day plan is simple: whether it is high-risk classification under Article 6, technical documentation under Article 11, or AI Office notifications under Article 52, you cannot even get started if you don't know what AI is running where inside your company. It is not something written in the articles, but a practical step you must go through to comply with them. I call this a *de facto* obligation.

Where companies most frequently slip up during an inventory audit is not visible AI products, but invisible AI. Everyone knows the AI services the company advertises and sells. The problem is what sits right next to them: the automated recommendation feature in commercial software used by hiring teams to screen resumes, response recommendation engines attached to customer support chatbots, targeting algorithms in advertising optimization tools used by marketing, and anomaly detection features in finance. The company didn't build these directly, nor does it label them as AI, yet they generally fall under the AI Act's definition. The audit cannot end with the engineering team alone. You must sweep department by department - HR, marketing, customer support, finance, and legal.

In practice, asking around yields mixed responses. If you ask the marketing team whether they use AI, their immediate answer is no. But if you follow up by asking which ad platforms they use, the story changes. Features that automatically rank ad impressions or group customers by similar tendencies are often already embedded in those platforms. The marketing team doesn't call them AI; they just call them ad tools. You must reframe the question from "Do you use AI?" to "Do you have features that automatically rank, recommend, or filter?" to get proper answers. Legal also needs a separate set of questions. Contracts frequently contain clauses stating that services provided by third parties include AI functionality, but those clauses stay buried in legal's drawers while the operational teams using the services remain unaware that AI is built into them.

The audit methodology doesn't need to be grand. Start by asking each department head to list every tool they use in their work that features automatic decision-making, recommendation, or filtering functions. You must include not only in-house developed systems, but also externally procured SaaS tools where AI features came bundled inside. Cases where companies didn't realize a tool was AI-based upon adoption, only to discover

it later, are far from rare in practice.

For each item, there are five pieces of information to record: what it is used for, what data it processes, whether it was built in-house or procured externally, whether it has touchpoints with the European market or data originating from Europe, and whether the compute scale used for training is known. The last item applies only to companies training their own models, but if you fail to record this now, you will have to restart the audit from scratch later when determining whether a model qualifies as a general-purpose AI model.

This inventory serves two roles in practice. First, it feeds into the next section: classification and role determination. Second, it forms the foundation for external registrations required by law, such as registering high-risk systems in the EU database under Article 49. Registration is essentially an inventory submitted outside the company, while the one you build now is its internal version.

18.2 System-by-System Classification and Role Determination

Once all items are listed in the inventory, evaluate each one against the four tiers established in Chapter 2: prohibited practices, high-risk, subject to transparency obligations, or minimal risk. A tool that ranks job candidates is highly likely to fall under the employment domain of Annex III, while ad targeting tools usually stay in minimal risk. Chatbots that generate deepfakes or simulate human conversation fall under transparency obligations, which is Article 50's domain.

If your company trains general-purpose AI models in-house, one additional check applies. You must determine whether the model qualifies as a general-purpose AI (GPAI) model that performs a wide range of tasks after self-supervised training on large volumes of data, and whether its cumulative training compute exceeds (10^{25}) floating-point operations (FLOPs), potentially classifying it as a model with systemic risk. The measurement methodology for this threshold is not yet fully settled. Because this remains uncertain, if your company runs in-house models and sits on the fence regarding this threshold, it is safer to err on the side of caution and treat it as close to systemic risk.

Once you assign tiers, overlay the role matrix from Chapter 1 onto each item. Are we a provider, a deployer, or both for this system? If you simply use an externally procured hiring tool for internal operations, you are a deployer. If you pass an in-house developed model to downstream providers in the European market, you become a provider of that model. There is one point where practitioners often get confused here: if you take a general-purpose model and fine-tune or structurally alter it, depending on the extent of that modification, a company that was originally a deployer can become a new provider of the modified model. A task you thought was finished with a few lines of fine-tuning can legally create new provider status. I have seen many companies miss this exact point. If you touched the model, you must first evaluate whether that modification shifts your status to a provider.

Let's examine a real-world scenario where role determination gets confusing. A fintech firm in Yeouido takes an open-source language model released overseas and uses it to polish credit evaluation phrasing. If they had used the original model as-is, this firm would be a deployer, and the foreign company that released the model would be the provider. But if this firm retrains the model on domestic loan underwriting data and modifies it to assign credit risk grades, the story changes. If the scope of modification is substantial, this fintech firm itself becomes the provider of that modified model, bearing

the full provider obligations for a high-risk system that falls under Annex III's credit evaluation domain. Simply using open source does not free you from obligations; what you modified and by how much determines your legal status.

The deliverable of this section is a single table: a matrix where every system and model listed in the inventory is assigned one risk tier and one legal role. This table serves as the input for the next section's gap analysis.

One additional note: the European Commission was originally expected to issue guidelines on Article 6 high-risk classification criteria by February 2026, but the actual draft came out later on May 19, 2026, with the public consultation closing on June 23. This was nearly three months behind schedule. Because detailed interpretations of the classification criteria have not yet set in stone, I advise you to take a conservative approach for systems sitting on ambiguous boundaries until the final version of these guidelines is released.

18.3 Gap Analysis

Once the inventory and classification table are in place, lay that table side by side with the legal articles and measure what is missing. This is called a gap analysis. This term is not explicitly written in the legal articles either. Article 56 establishes a Code of Practice for general-purpose AI model providers, stipulating that providers who do not adhere to this code must demonstrate compliance through other means. The true substance of those "other means" is comparing your current measures against what the code requires. Following the Code of Practice leaves you with a much lighter burden of proof.

For systems classified as high-risk, the list of articles to cross-check is long. Do you have a risk management system under Article 9? Are you maintaining data governance under Article 10 - that is, checking training data quality and bias? Have you authored technical documentation under Article 11 covering system architecture and development processes? Check each requirement one by one, down to automatic logging under Article 12, human oversight by design under Article 14, and ensuring accuracy, robustness, and cybersecurity under Article 15. Anything missing is your gap.

If you are a provider of general-purpose AI models, check whether you have these four elements in place: technical documentation, information to pass to downstream providers integrating the model, a copyright compliance policy, and a publicly available summary of training content. If classified as a model with systemic risk, model evaluations, incident reporting, and cybersecurity are added on top.

The output of a gap analysis is a gap list and a roadmap to bridge those gaps. Categorize and record items requiring technical improvements, items requiring new procedures, items where organizing documentation suffices, and items requiring new hires or training. Documented records of actually implementing measures specified in the Code of Practice can serve as mitigating factors when calculating fines under Article 101. There is a substantial difference between merely identifying gaps and leaving them unattended, versus keeping a record of finding and closing those gaps when standing before fines later on.

Let's return to the candidate ranking tool mentioned earlier. It was picked up in the inventory, classified under Annex III's employment domain, and the company was designated as a deployer. Now cross-check this tool against Articles 9 through 15. There is no risk management system. There are no records of auditing resume data for bias. The

technical documentation rests with the vendor who sold the tool, not in the company's hands. A procedure allowing humans to override final decisions exists, but it is not documented. These four lines become the exact first four lines of your gap list. Deployer obligations are lighter than provider obligations, but they are not non-existent; you must begin by requesting a copy of technical documentation and bias audit results from the vendor.

18.4 Setting Priorities: By Dates and Fine Severity

You need a benchmark to align the inventory, classification, and gap list gathered so far into a single line. I use two axes: when enforcement begins, and how much you pay if you violate it. Intersecting these two axes produces the sequence naturally.

Before establishing this order, there is one thing to clarify. The NotebookLM material that initiated this research assumed that high-risk AI system obligations would fully apply starting August 2, 2026. Under the law's original timeline, that is correct. However, recent developments confirmed via search present a different picture. The Digital Omnibus - passed by the European Parliament on June 16, approved by the Council on June 29, and signed on July 8 - postponed the enforcement date for Annex III standalone high-risk system obligations by 16 months, moving it from August 2, 2026, to December 2, 2027. Obligations for Annex I embedded high-risk systems were pushed back by a year, from August 2, 2027, to August 2, 2028. This amendment enters into force three days after publication in the Official Journal, which, as of this writing, is still pending. While publication is expected by late July at the latest, if publication is unexpectedly delayed, the original timeline of full application on August 2 will revive. Because this detail remains unconfirmed, I advise checking the Official Journal once more in early August.

This delay significantly shifts priorities. The order is as follows:

Priority 1 is Article 5, prohibited practices. It has been in effect since February 2, 2025, and carries the heaviest fines in the law: up to 35 million EUR or 7 percent of global annual turnover. On top of that, the Digital Omnibus added new items to the prohibited list, such as non-consensual generation of synthetic sexual imagery. Because this provision is already in force, you must check right now - in the very first week of your 90 days - whether any of your company's AI crosses this line.

Priority 2 is Article 50, transparency obligations. This covers disclosures identifying chatbots as such, and labeling deepfakes and synthetic content. Unaffected by the Omnibus, these obligations will take effect as originally scheduled on August 2, 2026. On that exact same day, the AI Office's general fine enforcement power also kicks in. In effect, the obligations and the enforcement authority march out on the same day. Given the short window left between now and enforcement, you should allocate the largest share of actual workload in your 90-day plan to this area.

Priority 3 is general-purpose AI model obligations. Having taken effect on August 2, 2025, if you haven't addressed them yet, you are already behind schedule. This is precisely why you must finish your review within 90 days.

Priority 4 is Annex III high-risk system obligations. Originally, this would have taken precedence over Priority 2, but the Omnibus bought us time by pushing it to December 2027. Buying time does not mean exemption, however. Take it to mean that while you should press ahead with the roadmap from your gap analysis, you now have breathing room to adjust your pace.

Priority 5 is Annex I embedded high-risk system obligations, which offer the most generous timeframe, extending to August 2028.

When assigning these five priorities, I do not look at dates alone. When an item with an imminent deadline but light fines clashes with an item with a distant deadline but heavy fines, I place more weight on the fine side. This is why I did not drop Annex III obligations below Priority 5 even after being pushed to December 2027. Although you bought time, the fine remains in the tier of up to 15 million EUR or 3 percent, carrying the same weight as transparency obligations. My recommended allocation is to maintain its place at Priority 4 on your roadmap, while concentrating the actual working hours spent during the 90 days on Priority 2 and Priority 3.

Overlay the appointment of an EU authorized representative onto these priorities. If you introduce high-risk systems or general-purpose AI models into Europe as a third-country provider, you must appoint an authorized representative under Articles 22 and 54. As soon as your inventory and roadmap are ready, hand over real operational data to this representative so that when authorities inquire, the representative is not standing empty-handed.

Breaking down the 90 days into three months looks like this:

Month 1 (Days 1 to 30). Launch the department-by-department inventory audit across all divisions simultaneously: engineering, HR, marketing, customer support, finance, and legal. In that same week, check for potential matches with Article 5 prohibited practices. Because this provision is already in effect, it is the most urgent.

Month 2 (Days 31 to 60). Map items listed in the inventory against the four risk tiers and role matrix to finalize classification. Execute actual implementation of Article 50 transparency obligations during this month, applying chatbot disclosures and deepfake

labeling technologies. Items classified as high-risk or general-purpose AI models move into gap analysis.

Month 3 (Days 61 to 90). Wrap up the gap analysis and finalize the gap list. Report the priority roadmap structured around dates and fine severity to executive management, securing approval for budget and headcount. If an EU authorized representative has been designated, hand over the inventory and roadmap. In early August 2026, reconfirm whether the Digital Omnibus has actually been published in the Official Journal, checking if the schedule remains original or modified.

Reference summary for quick lookup:

Enforcement dates: Prohibited practices: February 2, 2025. General-purpose AI model obligations: August 2, 2025 (models placed on the market before August 2, 2025 have a grace period until August 2, 2027). Transparency obligations (Article 50) and AI Office fine enforcement authority: August 2, 2026. Annex III high-risk system obligations: originally August 2, 2026, but moved to December 2, 2027 via the Digital Omnibus (pending Official Journal publication verification). Annex I high-risk (product-embedded) obligations: originally August 2, 2027, but moved to August 2, 2028.

Fines: Violations of prohibited practices: up to 35 million EUR or 7 percent of global annual turnover. Most other obligations (high-risk, transparency, general-purpose AI models in general): up to 15 million EUR or 3 percent. Supplying incorrect information: up to 7.5 million EUR or 1 percent. Based on global turnover, whichever is higher.

Departmental inventory checks: Engineering: in-house models and API integrations; HR: hiring and performance evaluation tools; Marketing: targeting and recommendation engines; Customer Support: chatbots and response recommendations; Finance: anomaly detection; Legal: contractual AI clauses on record.

Ninety days later, that team lead in Pangyo walked into the CEO's office holding a single inventory sheet: 42 systems total - three high-risk candidates, five subject to transparency obligations, and the rest minimal risk. Alongside the table were five neatly listed priority items. What the CEO had asked for was not flawless perfection from day one, but knowing where to start, and this table was that answer. Ninety days is not enough to comply with the entire law. It is plenty of time to know what to comply with first.

Chapter 19. Documentation Systems

A startup in Pangyo decided to launch a high-risk AI system in Europe. They conducted risk management, established data governance, and completed testing. However, the audit lawyer asked in the conference room: "Where is the documentation to prove it?" The team members looked at one another. They had done the work, but left no records behind.

The position this company faces before the AI Act is precisely the topic of this chapter. Supervisory authorities cannot see inside a company's mind. The only thing authorities can hold in their hands is documentation. No matter how meticulously risk was managed, if that process was not documented, before the authorities, the company stands in the same position as if it had not managed the risk at all. Documentation is not a secondary administrative task, but the sole channel for proving one's legal existence.

19.1 What, in What Format, and for How Long Should Records Be Retained?

What to retain depends on one's role. For a provider of a high-risk system, technical documentation under Article 11 and Quality Management System (QMS) records under Article 17 are required. Article 17 requires a documented system covering AI development governance, quality management, testing and validation, risk management, data governance, technical documentation specifications, post-market monitoring, and serious incident reporting. For a provider of a general-purpose AI (GPAI) model, technical documentation under Article 53 and a publicly available summary of training content must be retained. For a deployer, supervisory records and automatically generated logs under Article 26 must be managed.

Regarding format, there are cases with prescribed templates and cases without. As for the publicly available summary of training content, the AI Office issued a standard template on July 24, 2025, so that template can be used directly. For technical documentation where Annex XI or Annex IV specifies required items, one needs to fill out those items one by one. Where a template exists, it is safer to follow that template. With self-created forms, it is difficult even for oneself to later identify what has been omitted.

The fact that retention periods vary by article is the point most frequently missed in practice. Automatically generated logs of high-risk systems under Article 12 must be kept for at least six months, or longer if personal data laws require a longer period. Under Article 18, providers must retain technical documentation, quality management system records, notified body decisions, and copies of the EU declaration of conformity for 10 years after the system has been placed on the market or put into service. An authorized representative in the Union acting on behalf of a provider headquartered in a third country must likewise keep this technical documentation available to authorities for the same 10-year period. Unless this difference - six months versus 10 years - is mapped out in a table by document type, standards will falter whenever personnel change.

19.2 Audit Response: How to Prove Consistency Between Documentation and Actual Operation

Having documentation and passing an audit are two different matters. What is truly put to the test during an audit is how closely the documented contents align with how the system actually operates.

A common discrepancy looks like this: the technical documentation states that the system makes decisions through a specific procedure, but after several updates, the actual system is operating in a different way. The documentation sits in a drawer exactly as first drafted, while only the system has moved forward. When authorities cross-check documentation against logs, or documentation against actual outputs, this discrepancy is quickly exposed. And inconsistent documentation is worse than having no documentation at all, because it leaves the impression that the company submitted false information to the authorities. Quality management system records required under Article 17 prove their value at precisely this point. By demonstrating the procedures and results actually applied across the entire process of design, development, deployment, operation, and modification, they serve as evidence proving that the documentation matches actual operation.

Therefore, I recommend treating documentation not as a one-time task, but as a procedure of maintenance. This means embedding the documentation update step directly within the development process for modifying the system. Adding just a single item to the deployment approval checklist - verifying whether the relevant technical documentation has been updated - can prevent a gap from widening between documentation and reality. Although it requires a little extra effort, companies whose documentation matches their systems receive completely different treatment in an audit room compared to those that do not.

19.3 Procedures for Responding to Information Requests from EU Authorities

The AI Act grants authorities broad powers to request information. Upon a reasoned request from a national competent authority, providers of high-risk systems must provide all information and documentation necessary to demonstrate compliance with the requirements set out in Section 2 (Article 21(1)). This documentation must be provided in an official language of the EU designated by the authority that can be easily understood by that authority. If logs are under the provider's control, access to them must also be granted (Article 21(2)). Market surveillance authorities can go a step further and request access to the training, validation, and testing datasets themselves via API or other remote technical means (Article 74(12)). National authorities supervising the protection of fundamental rights also have the power to request documentation in an accessible language and format regarding high-risk systems under Annex III (Article 77(1)).

For general-purpose AI models, the point of contact is different. The AI Office may request all information necessary to assess compliance, including technical documentation pursuant to Article 53(1)(a) (Article 91), and this documentation must contain the items set out in Annex XI. It must be provided in an up-to-date form within the period specified by the AI Office, and if necessary, access to the model itself may be required (Article 92). Providers of models with systemic risk exceeding 10^{25} FLOPs must notify the AI Office of this fact without delay and at the latest within two weeks (Article 52(1)). Similarly, there is an obligation to report serious incidents or cybersecurity breaches to the AI Office and relevant authorities (Article 55(1)(c)). Throughout all these procedures, trade secrets, intellectual property rights, and confidential information are protected (Article 78). If documents submitted in response to a request contain sensitive content, confidentiality must be explicitly requested and the documents marked as confidential.

If this channel is not designated within the company in advance, when a request actually arrives, the legal, engineering, and business departments will pass responsibility to one another and miss the deadline. You must map out in advance which channel requests will come through, who receives and forwards them to which department, and who compiles and submits the documents in which language and format. For a third-country provider, since the authorized representative in the Union serves as the primary contact with authorities, required documents must be handed over to the representative in advance along with response protocols. If the representative lacks the documents, there is no way to comply with the request. The AI Office has stated that it will adopt a cooperative,

phased, and proportionate approach, encouraging entities to reach out first and seek support if they experience compliance difficulties. However, I read that cooperative stance as applying only to companies that respond faithfully. Inadequate or uncooperative responses to information requests are themselves treated as separate violations.

Dates must also be kept in mind. The AI Act entered into force on August 1, 2024, and obligations for general-purpose AI model providers apply from August 2, 2025, from which point the AI Office can issue decisions. Obligations for standalone high-risk systems under Annex III apply fully from August 2, 2026, and the AI Office's power to impose fines takes full effect on the same date. Fines vary by type of violation. Violations of general obligations are subject to fines of up to 3 percent of total worldwide annual turnover for the preceding financial year or 15 million euros, whichever is higher, while violations of prohibited practices carry fines of up to 7 percent or 35 million euros, whichever is higher. Failure to comply with an information request or providing misleading information is also included in the basis for calculating these fines.

One issue remains open. How far specific expressions such as "reasoned request," "language easily understood," or "all necessary information" extend in practice leaves room for differing interpretations depending on the case. The European Commission has an obligation to issue guidelines pursuant to Article 96, and on November 19, 2025, guidelines on the scope of obligations for general-purpose AI model providers (C(2025) 7719 final) were released. However, even these guidelines include a caveat that the specific characteristics of individual cases must be taken into account. In the end, there is no choice but to interpret rules in light of one's own circumstances and preserve the grounds for that interpretation.

Reference Materials

Document Types and Retention Periods. Automatically generated logs for high-risk systems: at least 6 months (Article 12). Technical documentation, Quality Management System records, notified body decisions, and EU declaration of conformity: 10 years after placing on the market (Article 18). Technical documentation retained by an authorized representative in the Union for a third-country provider: also 10 years.

Legal Grounds for Information Requests. For high-risk systems: Article 21 (information and documentation), Article 74(12) (access to datasets), Article 77(1) (access by fundamental rights authorities). For general-purpose AI models: Article 91 (request for

information), Article 92 (access for model evaluation), Article 53(1)(a) (technical documentation), Article 52(1) (notification of systemic risk, within 2 weeks), Article 55(1)(c) (serious incident reporting). Confidentiality protection: Article 78.

Key Dates. Entry into force: August 1, 2024. Application of obligations for general-purpose AI models: August 2, 2025. Full application for high-risk systems and activation of fining powers: August 2, 2026. Fines: up to 3 percent or 15 million euros; for prohibited practices, up to 7 percent or 35 million euros, whichever is higher.

Internal Preparation Checklist. Prepare a retention period schedule by document type. Embed a document update step into the system modification process. Designate in advance the entry channel and responsible department for information requests. Transfer documents to the authorized representative in the Union in advance. Mark confidential information prior to submission.

Audits do not arrive without notice. The clock starts ticking the moment a single request letter arrives. In that moment, the difference between a company that rummages through drawers and one that opens an already organized folder decides the outcome.

Chapter 20. Contracts and the Supply Chain

Suppose a startup in Pangyo is building a conversational chatbot. Having no funds to train its own model, it fetches a European Big Tech company's general-purpose AI model (GPAI model) via API, layers its own logic on top, and releases it as a single AI system. The contract is usually a set of standard terms of service pre-made by the Big Tech firm - a bundle of clauses agreed to with a single click. Would these standard terms alone suffice to fulfill the compliance obligations under the European Union Artificial Intelligence Act (Regulation (EU) 2024/1689, hereinafter the AI Act)? I do not think so. The AI Act explicitly stipulates in its provisions what must be exchanged, when, and in what format between upstream model providers and downstream system providers, and those provisions take precedence over standard terms of service. It is the contract that must be rewritten to match those provisions.

Chapter 19 covered how a company should respond when authorities request information. Chapter 20 examines the preceding stage - what clauses companies should include in contracts within the supply chain. A contract is a tool for regulatory compliance, not something that replaces regulations.

20.1 Provisions to Include in AI Supplier Contracts

Article 25(4) of the AI Act stipulates that a provider of a high-risk AI system and a third party supplying AI systems, tools, services, components, or processes integrated into that system must enter into a written contract. The contract must specify, in accordance with the state of the art, the necessary information, capabilities, technical access, and other support required to enable the downstream provider to fully comply with its obligations under the AI Act. The format is written. A single email is insufficient. However, this obligation does not apply to tools, services, or components released under free and open-source licenses. The legal team of a medical device company in Busan once overlooked the fact that general-purpose AI models themselves are excluded from this exception even if they are open source.

Let's divide the items to be included in the contract into three categories.

Information Provision Clauses. Explicitly specify what to hand over regarding the components supplied by the upstream provider. Providers of high-risk AI systems bear obligations regarding risk management systems (Article 9), data governance (Article 10), technical documentation (Article 11), automatic logging (Article 12), human oversight (Article 14), and accuracy, robustness, and cybersecurity (Article 15). To prove compliance with these, one must know the data sources, limitations, and logging mechanisms of the upstream components. You must go through these six articles one by one to nail down which documents will be provided by when. If you vaguely write "provide necessary support," it leaves room for the upstream provider to argue in a dispute that they provided the minimum required.

Cooperation Clauses. Define how quickly the upstream provider will respond when audits, incident investigations, or requests for information from authorities occur. As seen in 19.3, requests for information from authorities must be responded to within a designated period, but if a substantial portion of that information is in the hands of the upstream provider, the downstream provider alone cannot meet the deadline. It is better to explicitly state the cooperation timeframe in business days and even specify remedies for non-compliance. This is a clause particularly often omitted from contracts.

Change Notification Clauses. This is a clause requiring prior notification to the downstream provider if the upstream model is updated or safety measures change. The AI Act imposes an obligation on GPAI model providers to keep technical documentation

updated throughout the model lifecycle (Article 53(1)). If the fact of an update is not automatically communicated to downstream customers, this obligation exists only on paper. I have seen a case where a change in a single model forced an entire rewrite of a downstream system's risk management plan. By establishing deadlines for change notifications and the level of detail of such notifications in advance in the contract, you can buy time to respond.

The Model Contractual Clauses for AI (MCC-AI), published in March 2025 by the EU Public Procurement Community, were designed for public procurement, but their framework for information provision and cooperation clauses is worth referencing. However, as it explicitly states that it does not cover intellectual property rights, data protection, or indemnification on its own, it is not something to simply copy and paste.

20.2 Information to Receive from Upstream Model Providers

Conversely, if a Korean company is in the position of a downstream system provider, what must it secure from the upstream model provider for the contract to be complete?

The primary pillar placed at the forefront is Article 53(1)(b) and Annex XII. A provider of a GPAI model must prepare and keep updated information and documentation so that downstream providers can integrate the model into their systems. What Annex XII requires includes a general description of the model, intended tasks, integration methods with other systems, release date and distribution methods, and interaction with hardware and software. Without this information, downstream providers cannot complete Annex IV technical documentation, and failing to complete it blocks the conformity assessment itself.

Three additional items must be secured here. A policy to comply with copyright law, which includes policies adhering to the text and data mining exception under Directive 2019/790 pursuant to Article 53(1)(c) and identification of reservations of rights. A publicly available summary of training content, which is prepared and made public in accordance with the guidance and standard template released by the European Commission on July 24, 2025, as required by Article 53(1)(d). Support required for integration, which refers to the information, capabilities, and technical access required by Article 25(4). The very provision discussed in 20.1 as items to be requested by downstream providers reappears here in reverse as items to be received.

If the model carries systemic risk, another layer of information is added. Article 51 classifies models trained using cumulative amount of computation greater than 10^{25} floating-point operations (10^{25} FLOPs) as models with systemic risk, and Article 55 requires their providers to evaluate models, report serious incidents and cybersecurity breaches, and notify the AI Office within two weeks when reaching the threshold. Downstream providers must secure these evaluation results and incident histories by contract to reflect them in their own risk management plans. A situation where an upstream model is quietly classified as a model with systemic risk while the downstream provider remains unaware - I consider this the most dangerous blind spot.

There is an exception for open-source models. Models under free and open-source licenses may be exempted from technical documentation and downstream provider information obligations under certain conditions, but to do so, model parameters

(including weights), architecture information, and usage information must be disclosed. The disclosure method simply shifts from contracts to public channels.

Obligations for GPAI model providers began applying on August 2, 2025, while models released prior to that date have a grace period until August 2, 2027. If a company is entering into a new contract, it must first check whether the counterpart's model falls within this grace period.

20.3 Information to Provide to Downstream Customers

Now let's suppose a Korean company is an upstream provider that develops and sells GPAI models. The items in 20.2 are now reversed into obligations to provide. You must be prepared to hand over Annex XII information, copyright policies, and publicly available summaries of training content to downstream customers.

The problem arises here. Among the information required by Annex XII is how and on what data the model was trained, which constitutes a core technical asset and trade secret of the company. While Article 78(1) of the AI Act stipulates that authorities must protect trade secrets, intellectual property rights, and data confidentiality, this provision is directed at authorities and is not a license to withhold information from downstream customers. The minimum items of Annex XII are not of a nature that can be pared down under the pretext of trade secrets.

Therefore, in practice, information is divided into two layers. One layer is the information specified in Annex XII that is strictly necessary for downstream customers to fulfill their obligations. This cannot be omitted on the grounds of trade secrets. The other layer consists of detailed implementations supporting those items, such as specific code of training algorithms or internal hyperparameters. This can be delivered at an abstracted level explaining what types of data were used and how it was trained, while still satisfying the requirements of Annex XII. I recommend creating a template based on this distinction and attaching it to the contract.

There is also the method of overlaying a Non-Disclosure Agreement (NDA). If more detailed information needs to be provided, you can designate that information as confidential and restrict its purpose of use. However, an NDA is a tool to limit the usage of information provided, not a tool to reduce the scope of information that must be provided. There is another lever here. Article 89(2) of the AI Act recognizes the right of downstream providers to lodge a complaint with the AI Office for infringements by upstream GPAI model providers. If an upstream provider neglects information provision, it may find itself standing before regulatory authorities rather than facing a simple contract dispute.

20.4 Liability Allocation and the Limits of Indemnification Clauses

No matter how meticulously a contract is drafted, there is a line that cannot be crossed. The AI Act defines the roles of provider, deployer, importer, and distributor in Article 3 and imposes different obligations on each role. The obligations of risk management, data governance, technical documentation, and conformity assessment borne by high-risk AI system providers are linked to the objectives of public safety and fundamental rights protection, and thus cannot be transferred to others through a single line in a contract. I refer to these as non-delegable obligations. While it is possible to set indemnification ratios for damages by contract, a contract cannot alter whom regulatory authorities will penalize.

The AI Office and national competent authorities hold the power to impose fines directly on violating parties and order corrective measures. No matter how robust an indemnification clause a downstream provider has secured, the AI Office will not look at that clause; it will target the provider bearing responsibility under the AI Act. An indemnification clause is useful only when settling accounts between parties later; it cannot serve as a shield to block enforcement by authorities.

This limitation becomes starkly clear when looking at the scale of fines. Infringements of prohibited practices (Article 5) are subject to fines of up to 35 million euros or 7 percent of total worldwide annual turnover, whichever is higher, while non-compliance with other obligations carries fines of up to 15 million euros or 3 percent of turnover, whichever is higher. If personal data is handled, GDPR violation fines of up to 20 million euros or 4 percent of turnover can overlap. While liability caps in commercial contracts are usually set at a multiple of the contract value, these fines calculated as percentages of turnover easily surpass those caps. Non-financial losses, such as market withdrawal orders or reputational damage, are areas that indemnification clauses are not even designed to target in the first place.

The legal landscape has also shifted. On February 11, 2025, the European Commission announced its intention to withdraw the draft AI Liability Directive, and it was officially repealed upon publication in the Official Journal on October 6, 2025. The ostensible reason was a lack of consensus among stakeholders. Instead, the revised Product Liability Directive ((EU) 2024/2853) applies to products placed on the market after December 9, 2026, explicitly treating software and AI systems as products. If damage arises from a defect in an AI system, a framework opens up starting in late 2026 where manufacturers

bear strict liability without proof of fault. Even if a defect persists due to a failure to apply updates or security patches in a timely manner, it is deemed a defect. Administrative fines under the AI Act and civil liability for damages under the Product Liability Directive can attach simultaneously to a single incident.

When a Korean company headquartered in Seoul places an AI system on the EU market or its output is used within the EU (Article 2(1)(c) of the AI Act), these four sections from 20.1 to 20.4 of Chapter 20 are interconnected within a single set of contracts. If the provisions of 20.1 are inadequate, the information to be received in 20.2 will not come through in the first place, and failing to receive that information renders the information to be handed over to downstream customers in 20.3 inadequate as well. And as 20.4 notes, no matter how well all these provisions are drafted, individual companies retain ultimate responsibility before regulatory authorities. A contract is not a shield, but a map. However, that map does not erase the enforcement boundaries drawn by the AI Office.

I suggest opening up the supplier contract you currently hold. Check whether article numbers are specified in the information provision clause, whether cooperation deadlines are fixed in business days, and whether a change notification clause even exists. If not, that contract still remains in the world before the AI Act.

Reference Materials

Contract Checklist 1. Are information provision items corresponding to Articles 9, 10, 11, 12, 14, and 15 specified alongside their respective article numbers? 2. Does the cooperation clause state deadlines for responding to audits, incident investigations, and requests from authorities in business days? 3. Does the change notification clause specify the timing of notification and the required level of detail? 4. Is there a procedure to receive Annex XII items, copyright policies, and links to publicly available summaries of training content from upstream GPAI model providers? 5. Is there a clause to be notified if the model is classified as a model with systemic risk (exceeding 10^{25} FLOPs)? 6. Is there a benchmark document establishing criteria for separating information to provide to downstream customers from trade secrets to be protected? 7. Is the indemnification liability cap designed taking into account AI Act fines (3 percent to 7 percent of turnover) and GDPR fines (4 percent of turnover)? 8. Have defect and strict liability clauses been reviewed in preparation for the application of the Product Liability Directive after December 9, 2026?

Key Dates August 1, 2024: AI Act enters into force. February 2, 2025: Prohibited practices regulations become fully applicable. March 2025: EU Public Procurement Community releases Model Contractual Clauses for AI (MCC-AI). July 24, 2025: Guidance and standard template for publicly available summaries of training content released. August 2, 2025: Obligations for GPAI model providers take effect; AI Office authority to make GPAI decisions takes effect. October 6, 2025: Withdrawal of AI Liability Directive published in the Official Journal. August 2, 2026: Obligations for high-risk AI systems become fully applicable; AI Office fine imposition powers take full effect. December 9, 2026: Revised Product Liability Directive ((EU) 2024/2853) starts applying. August 2, 2027: End of grace period for GPAI models released prior to August 2, 2025 to comply with obligations.

Chapter 21. Organization and People

This scene played out in a real meeting at a software company in Pangyo. Deciding to introduce an AI employment screening system at a European subsidiary, the head of legal insisted, "The development team needs to draft the technical documentation." The lead developer countered, "We only write code; employee training and regulatory notifications belong to legal." The head of business listened in silence throughout the meeting before finally asking: "So, which of the three of us is ultimately responsible for this?" The fact that there is no set answer to this question forms the starting point of this chapter. The AI Act does not designate specific departments or job titles. Yet how you structure your people and organization turns out to be a decisive variable for potential fines.

21.1 AI Literacy Obligations: Article 4 and Article 26 Are Distinct Provisions

Article 4 of the AI Act requires providers and deployers to ensure that their staff and persons handling AI on their behalf possess an adequate level of AI literacy. This provision has been in effect since February 2, 2025. The risk classification of the AI used by a company does not matter. Even if a system is not high-risk - even if a company merely deploys a single chatbot - Article 4 applies.

A common misconception arises here. During the drafting stage of Article 4, discussions leaned toward imposing specific obligations on enterprises. However, the Commission shifted the center of gravity toward policy tasks for member states and the Commission itself. Consequently, what Article 4 itself demands from companies reads more like a policy declaration, while concrete obligations emerge from Article 26. Article 26 requires deployers of high-risk AI systems to train their staff to perform human oversight. This training obligation was originally scheduled for full application on August 2, 2026, for standalone high-risk systems under Annex III. As noted in Chapter 13, the provisional agreement on the Digital Omnibus in May 2026 opened the possibility of pushing this back to December 2, 2027. In my assessment, this point must not be confused. Although the high-risk training obligation (Article 26) may be delayed, the general literacy requirement (Article 4) binding all providers and deployers is already in force and falls outside the scope of Digital Omnibus discussions. If the company in Pangyo assumes it can relax upon hearing about the August 2 postponement simply because it only uses a chatbot, it is conflating Article 4 with Article 26.

The law does not specify a list of training topics. Given that practical precedents compiled by the AI Office and upcoming guidelines on Articles 8 through 15 and Article 25 will address training alongside human oversight and risk management, it is clear that this literacy training will not end with a single formal certificate of completion. I recommend building a training program centered on four core elements: how the system operates, its limits and potential biases, how to read performance metrics, and how to spot moments requiring intervention. Keep completion records. Managing these records as part of the quality management system under Article 17 or technical documentation under Article 11 eliminates the need to dig for them during a future audit. The penalty structure remains ambiguous. Article 4 is not explicitly listed among the fine brackets enumerated in Article 99. It is accurate to view that no fixed penalty tier exists solely for violating Article 4. In practice, market surveillance authorities in member states are likely to apply the general obligation penalty tier (up to 15 million euros or 3 percent of global annual turnover) or

cite it as background context for human oversight failure when an incident occurs.

21.2 Creating Roles Outside the Org Chart

The AI Act does not explicitly mandate that any company set up an AI governance team. Yet hardly any company can expect a single person to handle risk management systems (Article 9), data governance (Article 10), technical documentation (Article 11), quality management systems (Article 17), and human oversight (Article 14) concurrently. While individual articles do not explicitly demand an organization, taken together, compliance becomes impossible without one. I view this as a blank space left by the law: the structure is left flexible, but the substance is mandatory.

A practice known as the "90-Day Kickoff Plan" circulates widely in the field. It consists of three steps: building an AI system inventory, defining whether your role for each system is provider or deployer, and conducting a gap analysis to spot deficiencies. You must first decide who will lead these three steps. I recommend establishing an AI governance lead directly under the CEO - whether titled Chief AI Officer or AI Compliance Lead does not matter - supported by a structure where representatives from Legal, Development, and Business meet regularly under that lead. If any single department takes full charge, the other two get sidelined, and sidelined departments end up rushing through their assigned documentation. I have seen such failures play out repeatedly.

The AI Office focuses precisely on this aspect. It has signaled that it will examine whether risk management and quality management systems exist merely on paper or operate effectively within the organization. If a company drafts an impressive policy document without allocating personnel, budget, or authority to execute it, that paper becomes damning evidence during an audit. Having a policy without a designated owner is read as a vulnerability in risk management itself.

21.3 Legal Focuses on Law, Development on Technology, Business on the Market, and the Table Where All Three Decide Together

The legal department's role is clear. It includes mapping overlaps between the EU AI Act and domestic AI framework laws, inserting AI Act clauses into vendor and customer contracts, appointing and managing authorized representatives within the EU, and responding to information requests from the AI Office or member state authorities. For companies handling general-purpose AI (GPAI) models, verifying compliance with copyright policy also falls to legal.

The development team embeds legal compliance into code right from the design phase. This involves implementing risk management systems and data governance inside production pipelines, writing and updating technical documentation, attaching automated logging functions, and integrating machine-readable marking of AI-generated content into systems as required by Article 50. The belief that developers need not know the legal text is dangerous. Code written without knowing the legal requirements cannot be fixed later, no matter how much legal attempts to patch it.

The business department looks at the market. It first establishes whether the company acts as a provider, deployer, or importer within the supply chain, identifies system risk classifications, and manages launch timing and strategy. Explaining system intended purposes and limitations to downstream deployers - and providing training when necessary - also belongs to business. If a GPAI model presents systemic risks, the company must notify the AI Office within two weeks. Managing this tight deadline falls to the business team because it occupies the front line to spot such developments first.

Even with these divisions, decisions remain that do not fall neatly into a single bucket. Determining whether fine-tuning a GPAI model creates a new "provider" status, or whether your system falls under credit scoring or insurance risk assessment in Annex III, cannot be made by legal, development, or business in isolation. These determinations emerge only when analyzing legal interpretations, technical modifications, and market implications together at one table. Bringing all three into a single joint body represents substance rather than mere procedure. Rather than dispersing after merely drafting meeting minutes, the meeting gains real value only when it specifies who will execute and who will verify the conclusions reached.

Key points to verify in this chapter:

Article 4 AI Literacy: In effect for all providers and deployers since February 2, 2025. Applies regardless of risk classification.

Article 26 Human Oversight Training: Targets deployers of high-risk AI systems. Originally set for August 2, 2026; Digital Omnibus introduces potential delay to December 2, 2027. Do not confuse with Article 4.

Penalty Brackets: Article 99 contains no dedicated provision for Article 4. Violations will likely be subsumed under general obligations (up to 15 million euros or 3 percent of global annual turnover) or treated as background context in incident investigations.

Governance Structure: The law does not mandate a specific structure, but establishing a lead reporting directly to the CEO alongside regular meetings with representatives from Legal, Development, and Business is recommended.

Departmental Roles: Legal handles regulations, contracts, authorized representatives, and authority communications. Development oversees design, technical documentation, logging, and machine-readable marking. Business handles role determination, risk classification, launch strategies, downstream training, and notifications.

Cross-Functional Decisions: Changes to provider status and Annex III scope must be evaluated jointly by all three departments, establishing clear owners for execution and verification.

Let us return to the meeting room in Pangyo. There is no single correct answer to the question asked by the head of business - which of the three of us is responsible? Yet if that meeting ends without anyone answering the question, that company is likely already missing a requirement, whether under Article 4 or Article 26.

Chapter 22. Regulatory Sandboxes and Prior Consultation

The CEO of an AI healthcare startup in Pangyo asked me a question. He wanted to sell software that reads retinal images to screen for diabetic retinopathy to German hospitals, and asked if there was any way to validate it with real patient data before receiving official certification. I answered as follows: starting August 2, 2026, there will be an answer - the AI regulatory sandboxes operated by each Member State.

22.1 Member State Sandbox Requirements

Article 57 of the AI Act (Regulation (EU) 2024/1689) explicitly mandates that Member States establish at least one national AI regulatory sandbox and have it operationally active by August 2, 2026. The key point is that the subject of this obligation is the Member State. This means it is an infrastructure that governments must create, rather than something companies apply to establish. This obligation may be fulfilled by establishing a sandbox jointly with national competent authorities of other Member States, or substituted by participating in an already operational sandbox of another Member State, provided that such participation guarantees an equivalent level of access for national citizens. Setting up additional sandboxes at regional or local levels is not precluded, and the European Data Protection Supervisor (EDPS) may establish a separate sandbox for EU institutions, acting in a capacity equivalent to national competent authorities. The European Commission may provide technical support, advice, and tools, but does not assume the obligation of establishment itself.

Activities taking place within the sandbox follow a sandbox plan agreed upon in advance between the provider and the competent authority. The core is developing, training, testing, and validating the system for a specified duration in accordance with the plan, and testing under real-world conditions may also be conducted under the authority's supervision. The competent authority acts as a supervisor and advisor assisting regulatory compliance. It offers guidance and support in identifying and mitigating risks - specifically those related to fundamental rights, health, and safety - and, upon request by the provider, issues written proof of activities conducted in the sandbox alongside an exit report. These two documents carry substantial weight during subsequent conformity assessments or market surveillance phases, as the law explicitly requires market surveillance authorities and notified bodies to take this report into positive account and expedite procedures within a reasonable scope.

The most notable provision is the immunity from administrative fines. Sandbox participants cannot escape civil liability if they cause damage to third parties during testing. That much is obvious. However, administrative fines for violations of the AI Act will not be imposed, provided that the participant adhered to the plan and conditions of participation and followed the guidance of the competent authority in good faith. The same applies to violations of other EU or national laws, provided the relevant authority actively supervised and provided guidance. I interpret this provision as a declaration that regulatory authorities will not treat companies antagonistically from the outset.

Nevertheless, if significant risks to health, safety, or fundamental rights emerge, the authority may suspend or permanently terminate sandbox activities at any time, and that decision is communicated to the AI Office. This signifies that the sandbox is not a safe haven from regulation.

When national competent authorities establish a sandbox, they must inform the AI Office and the AI Board, and the AI Office publishes and maintains up-to-date lists and plans of Member States' sandboxes on its website. Common principles regarding detailed operation - such as eligibility and selection criteria, application procedures, and conditions of participation - are established by the Commission through implementing acts pursuant to Article 58. These implementing acts serve as a safeguard against Member States operating sandboxes under fragmented standards across the EU, establishing free access for SMEs and startups as a fundamental principle.

From the perspective of South Korean companies, the calculation is relatively clear. For companies subject to extraterritorial application under Article 2(1)(c) of the AI Act simply because their output is used in the EU, it is safer to go through a sandbox before placing high-risk AI systems on the market. Considering the fine structure - reaching up to 35 million euros or 7 percent of global annual turnover for violations - having a path to avoid the entire fine solely by following guidance in good faith is no small incentive. Holding an exit report and written proof also shortens the path to CE marking. However, an authorized representative in the EU must serve as the point of contact with authorities during this process, and separate liability insurance must be arranged to cover potential third-party damages occurring during testing.

Certain uncertainties remain. Because the implementing acts have not yet been finalized, the specific shape of eligibility criteria and application procedures may vary across Member States. Furthermore, assessing how strictly authorities will evaluate the requirement of following guidance in good faith post hoc, or to what extent non-EU enterprises will actually receive priority access comparable to EU-based SMEs, will require a body of early cases in each country. I maintain cautious optimism, but generally recommend observing precedents a few months later rather than being among the very first applicants.

22.2 Measures for SMEs

Article 62 of the AI Act requires Member States and the AI Office to take into specific account the circumstances of small and medium-sized enterprises (SMEs), and startups in particular, starting from the design phase of regulations. A question posed to me by the CEO of a medical device startup in Busan finds its answer right in this article: namely, whether they must pay the exact same conformity assessment fees as large corporations. The answer is no. Article 62(1) stipulates that fees for conformity assessments of high-risk systems must be reduced proportionally to metrics such as company size and market size, and these criteria are to be periodically reviewed upon request by the AI Board.

In terms of sandbox access as well, SMEs and startups take front-row priority. As long as they fulfill eligibility requirements and selection criteria, they are granted a right of prioritized access free of charge, with only exceptional costs charged separately. Additionally, Member States must establish and maintain standardized application templates and a single information platform across the EU, organize awareness-raising activities and training targeted at SMEs, and evaluate and disseminate best practices in AI-related public procurement. Article 63 goes one step further, allowing microenterprises (enterprises employing fewer than 10 persons with an annual turnover or balance sheet total not exceeding 2 million euros, without partner or linked enterprises) to fulfill certain quality management system requirements in a manner involving a lower administrative burden.

An intermediate category known as "small mid-caps" also emerges. This refers to enterprises employing fewer than 750 persons and having an annual turnover not exceeding 150 million euros or an annual balance sheet total not exceeding 129 million euros, extending certain documentation relief benefits previously reserved for microenterprises and SMEs to companies of this scale. Technical documentation requirements are reduced, and fine ceilings for violations are calculated based on the lower amount of 3 percent or 15 million euros, rather than the higher amount. However, there is one major exception: if an entity engages in prohibited practices such as manipulative AI practices or social scoring prohibited under Article 5, the general ceiling of up to 35 million euros or 7 percent of global annual turnover applies even to small mid-caps. I place particular emphasis on this point during consultations, as leniency based on size applies to the administrative and procedural burden, not to the gravity of prohibited conduct itself.

Practical considerations for South Korean companies can be summarized into three steps. First, verify which category your company falls into. Evaluating employee headcount, annual turnover, and balance sheet total to determine whether you qualify as an SME, startup, microenterprise, or small mid-cap dictates which documentation needs to be submitted to national competent authorities. Second, actively apply for fee reductions and priority access rather than leaving them unclaimed; even if regulations exist, they are not applied automatically without an application. Third, if you qualify as a microenterprise or small mid-cap, negotiate in advance to narrow the scope of compliance documentation relying on the documentation simplification provisions under Article 63. However, whether phrases implying "within the EU" apply directly to non-EU enterprises has not been extensively confirmed by precedents, so I recommend confirming this in writing with the national competent authority of each Member State through an EU authorized representative in advance.

22.3 Communication Channels with the AI Office

If sandboxes and fee reductions form the structural framework of the system, communication channels serve as the doorway through which enterprises actually enter it. Under Article 58, the Commission is set to develop a single dedicated interface allowing stakeholders to contact competent authorities and seek non-binding guidance. Separately, in October 2025, the AI Office officially launched the AI Act Service Desk and the Single Information Platform. Under this structure, submitting a question via an online form prompts responses from an expert team working in close cooperation with the AI Office; national competent authority information and the status of sandboxes across Member States can also be verified on this platform. On November 19, 2025, guidelines (C(2025) 7719 final) narrowing and explaining the scope of obligations for providers of general-purpose AI (GPAI) models were approved and published alongside it.

Cases where competent authorities or the AI Office initiate contact are also grounded in law. Providers of high-risk systems must provide documentation demonstrating compliance upon reasonable request by competent authorities under Article 21, and such documentation must be drafted in an official language of the EU that can be easily understood by the authority. Access to automatically generated logs (Article 21(2)) or market surveillance authorities' access to training, validation, and testing datasets (Article 74(12)) may also be requested. For providers of general-purpose AI (GPAI) models, the dynamic is slightly different. The AI Office acts directly as the contact point to demand information requests under Article 91 and model evaluation access under Article 92; providers of models presenting systemic risk - that is, models trained with cumulative computation exceeding 10^{25} FLOPs - must notify the AI Office no later than two weeks after meeting the threshold. In the event of serious incidents or cybersecurity breaches, a prompt reporting obligation also applies. Regardless of the communication channel, trade secrets, intellectual property rights, and confidential business information are protected under Article 78.

The AI Office positions itself as a collaborative, phased, and proportionate enforcer. The spirit is that if enterprises encounter compliance difficulties, they should reach out first rather than worry immediately about fines; conversely, for providers that do not adhere to codes of practice or whose compliance remains opaque, the Office notes that it may request additional information and even demand access for model evaluation. At this juncture, I consistently offer the same advice to South Korean enterprises: establish an internal response team operating jointly across legal, engineering, and business

departments, and keep technical documentation, QMS records, logs, and summaries of GPAI training content organized for immediate retrieval. Requests for information usually arrive with tight deadlines. Searching for documents from scratch at that moment is already too late.

Reference Materials

Key Deadlines August 2, 2026: Deadline for Member State AI regulatory sandboxes to become operational; full application of high-risk AI system obligations; full activation of the AI Office's power to impose fines August 2, 2025: Application of obligations for providers of general-purpose AI (GPAI) models begins February 2, 2025: Full enforcement of prohibited practices under Article 5 October 2025: Launch of the AI Act Service Desk and Single Information Platform November 19, 2025: Approval of GPAI Guidelines C(2025) 7719 final

Checklist Items 1. Has your company verified whether it qualifies as an SME, startup, microenterprise, or small mid-cap based on employee headcount, annual turnover, and balance sheet criteria? 2. Have you verified the sandbox plans and application portals of your target Member States via the AI Act Service Desk? 3. Is your EU authorized representative prepared to serve as the point of contact with competent authorities? 4. Have you reviewed liability insurance options to cover potential third-party damages occurring during sandbox participation? 5. Are technical documentation, QMS records, and logs maintained in a state of continuous organization to respond immediately to information requests?

The door established by law stands open. Whether to knock on that door first or follow after watching other companies enter remains, ultimately, a matter of each enterprise's risk appetite.

Epilogue. Should Regulation Be Seen as a Cost or Used as a Barrier to Entry?

If you are the company officer who read this book to this point, you are probably exhausted. There are many things to comply with. You must label, prepare documentation, appoint an authorized representative, undergo assessment, and register. To sell anything in Europe, you have to clear all of these hurdles. So viewing regulation as a cost is only natural. It is a burden that consumes money and time without increasing revenue. That is how many companies view the EU AI Act.

In this final chapter, I want to offer a different perspective. It is a view that sees regulation not as a cost, but as a barrier to entry. To be precise, it is a perspective focused on finding the position where a barrier for others becomes a door for you.

Heavy regulation means it is heavy for everyone. Competitors from other countries targeting the European market stand before the exact same wall. A company that crosses that wall before others is already inside a market that hesitant competitors cannot yet enter. Completing regulatory compliance early becomes an advantageous position in itself. Looking at it from another angle, the fact that Europe imposed such heavy regulation on artificial intelligence means that only the AI systems clearing those regulations will remain in the European market. Entering that passing minority is the real competition.

I believe this perspective holds special meaning for Korean companies. Korean firms have grown over a long time by passing strict regulations. They raised product quality while exporting to markets with strict safety rules, and that quality became their competitive strength. AI regulation can be placed in that exact same spot. The data governance and risk management systems built while meeting European requirements won't just be used in Europe; they make a company's entire AI setup stronger. Regulatory compliance becomes a core capability of the firm.

Of course, this doesn't mean the regulation becomes lighter. As emphasized throughout this book, obligations are heavy, deadlines are approaching, and penalties are steep. Using regulation as a barrier to entry doesn't mean denying its weight; it means bearing that weight first. The company that prepares first while others delay will end up standing inside the narrow door created by regulation.

Closing this book, I return to that initial question. The Pangyo company in the prologue didn't know it was subject to the regulation. A company that read this book now knows. It knows through which path it is caught, under which classification it falls, and what it must do by when. A company with this knowledge and one without stand in completely

different positions facing the same regulation. One receives a call on August 2; the other finishes preparations before then.

Should regulation be seen as a cost or used as a barrier to entry? The answer to this question isn't in the legal provisions. It lies within a company's choice of when and how to face those provisions. The wall is already built. All that remains is deciding when to cross it. I hope to see you on the side that crossed first.

Appendix

Appendix 1. Integrated Implementation Schedule Timeline

As of July 23, 2026. The dates in this timeline are subject to change. Always reconfirm the latest status before undertaking actual preparations.

EU Artificial Intelligence Act

August 1, 2024 Entry into force of the Artificial Intelligence Act.

February 2, 2025 Application of Article 5 (Prohibited Practices) begins. Application of Article 4 (AI Literacy Obligations) begins.

August 2, 2025 Application of obligations for providers of general-purpose AI models (Articles 53 and 55) begins. Governance rules apply. Legal basis for imposing fines under Article 99 enters into force.

June 29, 2026 Final Council approval of the Digital Omnibus revision. Scheduled to enter into force on the 3rd day following publication in the Official Journal.

August 2, 2026 Full application of Article 50 (Transparency Obligations). Enforcement powers regarding general-purpose AI models activated.

December 2, 2026 Prohibition on generating NCII and CSAM under Article 5 enters into force. Grace period ends for labeling obligations on existing generative AI systems placed on the market before August 2, 2026.

February 2, 2027 Deadline for equipping interoperable solutions for general-purpose AI watermark detection.

August 2, 2027 Compliance deadline for existing general-purpose AI models placed on the market before August 2, 2025.

December 2, 2027 Application of obligations for Annex III stand-alone high-risk AI systems.

August 2, 2028 Application of obligations for Annex I embedded high-risk AI systems.

Korea Framework Act on Artificial Intelligence

January 21, 2025 Enactment and promulgation of the Framework Act on Artificial Intelligence.

January 22, 2026 Full implementation of the main Act. Transparency obligations (Article 31), safety assurance obligations (Article 32), duties of high-impact AI business operators (Article 34), domestic agent requirements (Article 36), etc.

July 21, 2026 Implementation of the revised Enforcement Decree. Promotion measures including the AI product and service verification system, preferential treatment in public procurement, and usage cost support.

What Has Been Postponed and What Has Not

Postponed Annex III stand-alone high-risk AI systems (December 2, 2027), Annex I embedded high-risk AI systems (August 2, 2028).

Not Postponed Article 50 transparency obligations (remains August 2, 2026). Prohibition on NCII and CSAM is a newly added provision, not a postponement (December 2, 2026).

Appendix 2. System Classification Flowchart

Follow the answers step by step to determine the location of your system. Detailed explanations can be found in the chapter numbers listed next to each question.

Step 1 Check Extraterritorial Application (Chapter 1)

Question Is our AI placed on the EU market, used by someone in the EU, or is its output used in the EU?

Yes to any of the three -> Subject to application. Proceed to Step 2.

No to all three -> Currently out of scope. Reconfirm whether there is truly no contact point.

Step 2 Determine Your Role (Chapter 1)

Question What is our role in this system?

Develop and place on the market -> Provider.

Deploy and use -> Deployer.

Import into the market -> Importer. / Supply in the distribution chain -> Distributor.

Multiple roles may overlap for a single system.

Step 3 Check Prohibited Line (Chapter 15)

Question Does this system touch upon one of the eight prohibited practices under Article 5? (Emotion recognition in workplace/educational settings, social scoring, exploitation of vulnerabilities, untargeted scraping of facial images, generation of NCII/CSAM, etc.)

Yes -> Prohibited. Cannot be used in the EU. Stop here.

No -> Proceed to Step 4.

Step 4 Risk Tier Classification (Chapter 2)

Question A Does it interact with humans (chatbots) or generate synthetic content/deepfakes?

Yes -> Subject to transparency. Article 50 obligations (Chapters 4, 5, 6). August 2, 2026.

Question B Do you develop or substantially modify and release a general-purpose AI model? (Chapter 7)

Yes -> General-purpose AI model provider. Article 53 obligations (Chapter 8). If exceeding 10^{25} FLOPs, additional Article 55 obligations apply (Chapter 9).

Question C Is it integrated as a safety component into an Annex I product (medical devices, automobiles, machinery, etc.) and subject to third-party conformity assessment? (Chapter 11)

Yes -> Annex I high-risk system. No filter applies. August 2, 2028. Chapters 12 and 13 obligations.

Question D Does it fall under one of the eight areas in Annex III (biometrics, critical infrastructure, education, employment, essential services, law enforcement, migration, administration of justice)? (Chapter 11)

Yes -> Candidate for Annex III high-risk system. Proceed to Step 5 filter. December 2, 2027.

Falls under none of the above -> Minimal risk. No special obligations.

Step 5 Annex III Filter (Chapter 11)

Question Does it meet one of the four exceptions under Article 6(3) (narrow procedural task, improving human outcome, pattern detection, preparatory task)?

Yes, and it does not perform profiling of natural persons -> Eligible for high-risk exemption. Subject

to documenting technical grounds + EU DB registration.

Performs profiling -> Exemption disqualified. Confirmed as high-risk system.

No -> Confirmed as Annex III high-risk system. Chapters 12 and 13 obligations.

Step 6 Cross-Check Korean Law Compliance (Chapters 16 & 17)

Even after completing the EU classification, cross-check in parallel against the Korean Framework Act on AI (High-Impact AI under Article 2, Labelling under Article 31, Domestic Representative under Article 36) and data protection laws (GDPR & PIPA). Compliance with one side does not grant exemption from the other.

Appendix 3. Self-Assessment Checklist

Designed to print and keep on your desk. Answer Yes or No to each item, and review the corresponding chapter for any item marked No.

A. Scope and Roles (Chapter 1)

Have you checked whether your AI connects to the European market, users, or outputs?
Have you defined whether your role for each system is a provider, deployer, importer, or distributor?
As a third-country provider placing a high-risk system or general-purpose model on the EU market, have you appointed an authorized representative in the EU (Articles 22 and 54)?
Have you actually transferred technical documentation and required information to the representative?

B. Prohibited Practices (Chapter 15)

Have you evaluated your active AI systems against the eight prohibited practices under Article 5?
Are you refraining from using emotion recognition tools in workplace or educational settings?
Do you refrain from scoring individuals to impose disadvantages in unrelated contexts?
Does your generative model have technical safeguards to block non-consensual intimate imagery (NCII) and child sexual abuse material (CSAM), and can you prove they work?

C. Transparency (Chapters 4, 5, and 6)

Is an AI notice displayed on the initial screen of chatbots or voicebots interacting with EU users?
Do generative outputs carry machine-readable markings that survive capture and re-uploading?
Do you visibly disclose deepfakes when creating and distributing them?
Is your marking system multi-layered (C2PA metadata + invisible watermarking)?
Do you have a plan to ensure watermark detection interoperability by February 2, 2027?

D. General-Purpose AI Models (Chapters 7, 8, 9, and 10)

Have you evaluated whether your model possesses significant generality to meet the general-purpose AI model threshold?
Have you prepared technical documentation (Annex XI), downstream information (Annex XII), a copyright policy, and a summary of training content?
Have you verified whether cumulative training compute exceeds 10^{25} FLOPs using two calculation methods?
If seeking an open-source exemption, have you met the four freedoms and free-of-charge requirements while fulfilling the remaining two obligations (copyright policy and training content summary)?

E. High-Risk AI Systems (Chapters 11, 12, 13, and 14)

Have you determined whether your product qualifies as high-risk under Annex I or Annex III?
If relying on derogations or filtering criteria, have you documented the justification for that assessment?
Have you established risk management (Article 9), data governance (Article 10), technical documentation and logging (Articles 11 and 12), human oversight (Article 14), and accuracy, robustness, and cybersecurity (Article 15)?
Do you have plans for conformity assessment, CE marking, and EU database registration (Articles 43,

48, and 49)?

If you are a deployer, have you verified whether worker notification and a fundamental rights impact assessment (Article 27) apply to you?

Have you initiated data governance now during the grace period?

F. Korean Law and Data Privacy (Chapters 16 and 17)

Have you checked whether your system qualifies as high-impact under the Korean AI Framework Act and verified marking obligations (Article 31)?

If you are a foreign business entering Korea, have you verified domestic representative requirements (Article 36), or EU representative requirements if expanding to the EU?

Have you verified that processing EU users' personal data triggers both GDPR and Korea's PIPA?

Have you established legal bases for processing personal data in training sets under both EU and Korean law?

G. Execution Framework (Chapters 18, 19, 20, and 21)

Have you created an AI inventory across departments, including shadow AI?

Does your documentation match actual system operations, and have you tied documentation updates to system updates?

Have you included clauses for information provision, cooperation, and change notification in vendor contracts?

Have you implemented and documented AI literacy (Article 4) tailored to specific roles?

Have you designated points of contact and procedures in advance to respond to regulatory requests for information?

Appendix 4. Sample Disclosure Statements for Article 50 (Korean / English)

These are starting points, not final plug-and-play text. Refine actual wording to fit the service context and legal counsel. Disclosures must be presented no later than the time of first interaction or exposure, in a clear and distinguishable manner, and in accordance with accessibility requirements.

1. Chatbots / Interactive Agents (Article 50(1))

Korean

This session is conducted by artificial intelligence. If you would like to connect with a human agent, please request it below.

English

You are chatting with an AI system. If you would like to speak with a human agent, please request it below.

Short Form

Korean AI Assistant.

English This is an AI assistant.

2. Voice Agents (Article 50(1))

Korean

Hello. This is an automated AI voice assistant. All interactions from this point forward will be handled by artificial intelligence.

English

Hello. This is an automated AI voice service. Your call will be handled by an AI system.

3. Labeling Generative Content (Article 50(2), Human-Perceptible Disclosures)

Korean

This content was generated by artificial intelligence.

English

This content was generated by AI.

Caption for Images and Videos

Korean AI-generated image

English AI-generated image

Note: Separate from this visual disclosure, machine-readable markings (such as C2PA) must also be embedded into the file to satisfy Article 50(2). See Chapter 5.

4. Deepfake Disclosures (Article 50(4))

Korean

This video contains scenes artificially generated or manipulated using artificial intelligence.

English

This video contains images or audio that have been artificially generated or manipulated using AI.

Mitigated Disclosure for Artistic / Satirical Works (In a Manner That Does Not Impair Enjoyment)

Korean AI synthesis technology was used in this work. (Indicated in end credits, etc.)

English This work uses AI-synthesised elements. (in end credits or equivalent)

5. Public Interest Text (Article 50(4))

Korean

This text was generated by artificial intelligence and published following editorial review.

English

This text was generated by AI and published under human editorial review.

Note: This exception cannot be invoked without substantive human review and editorial responsibility. See Chapter 4.

6. Emotion Recognition and Biometric Categorisation Disclosures (Article 50(3))

Korean

This service uses artificial intelligence to analyse your emotional state. (or to categorise biometric data.)

English

This service uses AI to analyse your emotional state. (or to categorise biometric data.)

Note: Emotion recognition in workplaces and educational institutions is subject to prohibition rather than disclosure (Article 5). See Chapter 15.

Appendix 5. Technical Documentation Table of Contents Template

Sample table of contents for items required by Annex IV for high-risk systems and Annex XI for general-purpose models. Cross-reference the original text of the respective annexes when drafting. The document must correspond to the actual system and must be updated alongside system modifications (Chapter 19).

A. Technical Documentation for High-Risk AI Systems (Based on Annex IV)

1. General Information on the System

- 1.1 Intended Purpose, Provider, and Version
- 1.2 Other Systems and Hardware with Which the System Interacts
- 1.3 Description of the User Interface
- 1.4 Instructions for Use for Deployers

2. System Design and Development

- 2.1 Development Methods and Stages
- 2.2 Design Specifications and System Architecture
- 2.3 Data Requirements (Origin, Scope, and Characteristics of Training, Validation, and Testing Data)
- 2.4 Human Oversight Measures and Their Design
- 2.5 Pre-determined Changes and Methods for Maintaining Continuous Conformity

3. System Performance and Limitations

- 3.1 Accuracy, Robustness, and Cybersecurity Metrics and Test Results
- 3.2 Foreseeable Malfunctions and Risks
- 3.3 Potential Performance Discrepancies Across Specific Groups

4. Risk Management

- 4.1 Description of the Risk Management System (Article 9)
- 4.2 Identified Risks and Mitigation Measures

5. Change History and Maintenance

- 5.1 Record of Changes by Version
- 5.2 Post-Market Monitoring Plan (Article 72)

6. Conformity

- 6.1 Applied Harmonised Standards and Common Specifications
- 6.2 Copy of the EU Declaration of Conformity (Article 47)

B. Technical Documentation for General-Purpose AI Models (Based on Annex XI)

1. General Information on the Model

- 1.1 Model Name, Provider, Release Date, and Version
- 1.2 Architecture and Number of Parameters
- 1.3 Input and Output Formats
- 1.4 License

2. Development Information

2.1 Training Methods and Process

2.2 Types and Provenance of Training and Testing Data, and Data Processing Methods

2.3 Computational Resources and Training Compute (FLOP Estimates and Rationale)

2.4 Known Capabilities and Limitations

3. Information for Downstream Providers (Annex XII)

3.1 Technical Information Required for Model Integration

3.2 Capabilities, Limitations, and Conditions of Use

4. Copyright and Training Content

4.1 Policy to Comply with Union Copyright Law (Article 4(3) of Directive 2019/790)

4.2 Publicly Disclosed Summary of Training Content (Article 53(1)(d), AI Office Standard Template)

5. Systemic Risks (Where Applicable, Article 55)

5.1 Model Evaluation and Adversarial Testing Results

5.2 Systemic Risk Assessment and Mitigation Measures

5.3 Serious Incident Response Framework

5.4 Cybersecurity Measures

Appendix 6. Sample AI Provider Contractual Clauses

This is a draft version. The actual contract should be finalized following legal advice. Refer to Chapter 20 for the background of each clause.

1. Definitions Clause

In this Agreement, the AI Act refers to Regulation (EU) 2024/1689 and its amendments and secondary legislation. Terms such as provider, deployer, high-risk, and general-purpose AI model shall have the meanings set forth in the AI Act.

2. Confirmation of Roles and Status Clause

The parties hereby confirm their respective statuses under the AI Act (provider, deployer, importer, distributor) with respect to the AI system or model supplied under this Agreement as follows: (To be specified in a table.)

The supplier shall notify the customer in writing as to whether the relevant system constitutes a high-risk AI system or poses a systemic risk, along with the grounds therefor.

3. Provision of Information Clause

The supplier shall provide the customer, prior to commencement of performance under the contract and promptly upon request, with information necessary for the customer to fulfill its obligations under the AI Act (including Annex XII and Article 13 information).

In the case of general-purpose AI models, the supplier shall provide information regarding technical documentation, capabilities and limitations, integration conditions, copyright compliance, and a summary of training content.

4. Notice of Modification Clause

The supplier shall give prior notice to the customer in the event of material changes to the performance, architecture, or training data of the supplied system or model. If a modification made without notice causes the customer to breach its obligations, the supplier shall bear the liability therefor.

5. Cooperation Clause

In the event of information requests, investigations, audits, or serious incident responses by competent authorities, the parties shall cooperate with each other in good faith and provide the necessary materials within a reasonable timeframe.

6. Confidentiality Clause

Any information provided under this Agreement that constitutes a trade secret shall not be used for purposes other than intended or disclosed to a third party without the prior consent of the other party; provided, however, that disclosures made pursuant to applicable laws, regulations, or lawful requests of competent authorities shall be excepted.

7. EU Authorised Representative Clause (If Applicable)

If the supplier is a provider located in a third country, the supplier shall appoint an EU authorised representative pursuant to Article 22 or Article 54 of the AI Act, notify the customer of the representative's details, and provide technical documentation and information to enable the

representative to perform its tasks.

8. Limitation of Liability and Indemnification Clause

The allocation of damages between the parties shall be governed by this Agreement. Provided, however, that regulatory obligations directly imposed on each party by the AI Act shall remain fully binding on each respective party, notwithstanding any allocation of liability or indemnification set forth in this Agreement.

9. Governing Law and Jurisdiction Clause

(To be determined based on the case. Review the impact of the choice of governing law and jurisdiction on actual enforcement in disputes related to the European market.)

Appendix 7. Terminology Comparative Table (Korean · English · Legal Provisions)

EU Artificial Intelligence Act Regulation (EU) 2024/1689, Korean Framework Act on Artificial Intelligence enforced as of July 21, 2026. Articles verified on July 23, 2026.

EU Artificial Intelligence Act

Provider / provider / Article 3(3)

Deployer / deployer / Article 3(4)

Authorised representative / authorised representative / Article 3(5) (High-risk Article 22, General-purpose models Article 54)

Importer / importer / Article 3(6)

Distributor / distributor / Article 3(7)

Operator / operator / Article 3(8)

Placing on the market / placing on the market / Article 3(9)

Making available on the market / making available on the market / Article 3(10)

Putting into service / putting into service / Article 3(11)

AI literacy / AI literacy / Article 4

Prohibited AI practices / prohibited AI practices / Article 5

High-risk AI system / high-risk AI system / Article 6 (Annexes I & III)

Risk management system / risk management system / Article 9

Data and data governance / data and data governance / Article 10

Technical documentation / technical documentation / Article 11 (Annex IV)

Record-keeping, logs / record-keeping, logs / Article 12

Transparency and provision of information / transparency and provision of information / Article 13

Human oversight / human oversight / Article 14

Accuracy, robustness and cybersecurity / accuracy, robustness and cybersecurity / Article 15

Obligations of high-risk providers / obligations of providers / Article 16

Quality management system / quality management system / Article 17

Responsibilities along the AI value chain / responsibilities along the AI value chain / Article 25

Obligations of deployers / obligations of deployers / Article 26

Fundamental rights impact assessment (FRIA) / fundamental rights impact assessment (FRIA) / Article 27

Harmonised standards / harmonised standards / Article 40

Common specifications / common specifications / Article 41

Conformity assessment / conformity assessment / Article 43

EU declaration of conformity / EU declaration of conformity / Article 47

CE marking / CE marking / Article 48

Registration / registration / Article 49

Transparency obligations for certain AI systems / transparency obligations for certain AI systems / Article 50

Deepfake / deepfake / Article 50(4)

Codes of practice on labelling / codes of practice on labelling / Article 50(7)

General-purpose AI model / general-purpose AI model / Article 3(63) (Classification Article 51)

Systemic risk / systemic risk / Article 3(65) (Article 51)

GPAI model with systemic risk / GPAI model with systemic risk / Article 51

Obligations for providers of GPAI models / obligations for providers of GPAI models / Article 53

Summary of content used for training / summary of content used for training / Article 53(1)(d)
Authorised representatives of GPAI providers / authorised representatives of GPAI providers / Article 54
Obligations for providers of GPAI models with systemic risk / obligations for GPAI with systemic risk / Article 55
Codes of practice / codes of practice / Article 56
AI regulatory sandbox / AI regulatory sandbox / Article 57
Testing in real world conditions / testing in real world conditions / Article 60
AI Office / AI Office / Article 3(47)
European Artificial Intelligence Board / European Artificial Intelligence Board / Article 65
National competent authorities / national competent authorities / Article 70
EU database for high-risk AI systems / EU database for high-risk AI systems / Article 71
Post-market monitoring / post-market monitoring / Article 72
Reporting of serious incidents / reporting of serious incidents / Article 73
Market surveillance / market surveillance / Article 74
Right to lodge a complaint / right to lodge a complaint / Article 85
Right to explanation / right to explanation / Article 86
Penalties / penalties / Article 99
Fines for GPAI providers / fines for GPAI providers / Article 101
Entry into force and application / entry into force and application / Article 113

Korean Framework Act on Artificial Intelligence

Artificial intelligence / artificial intelligence / Article 2(1)
AI system / AI system / Article 2(2)
AI technology / AI technology / Article 2(3)
High-impact AI / high-impact AI / Article 2(4)
Generative AI / generative AI / Article 2(5)
AI business operator / AI business operator / Article 2(7)
AI development business operator / AI development business operator / Article 2(7)(a)
AI use business operator / AI use business operator / Article 2(7)(b)
User / user / Article 2(8)
Affected person / affected person / Article 2(9)
Training data / training data / Article 2(12)
AI ethics principles / AI ethics principles / Article 27
Transparency obligation / transparency obligation / Article 31
Safety obligation / safety obligation / Article 32
Confirmation of high-impact AI / confirmation of high-impact AI / Article 33
Obligations of high-impact AI operators / obligations of high-impact AI operators / Article 34
High-impact AI impact assessment / high-impact AI impact assessment / Article 35
Designation of domestic representative / designation of domestic representative / Article 36
Fact-finding investigation / fact-finding investigation / Article 40
Criminal penalties / criminal penalties / Article 42
Administrative fines / administrative fines / Article 43

Appendix 8. EU Authorities and Domestic Support Agencies

Because contact details change frequently, official agencies and their contact channels are listed instead of specific telephone numbers. Prior to making actual inquiries, verify the latest contact channels on each agency's official website.

EU Side

European Commission AI Office (European Commission, AI Office)

Hub for supervising general-purpose AI models and publishing guidelines, codes of practice, and templates.

Official portal digital-strategy.ec.europa.eu (Shaping Europe's digital future)

AI Act Service Desk

Official Commission helpdesk providing legal provisions, annexes, and FAQs.

Official portal ai-act-service-desk.ec.europa.eu

European Artificial Intelligence Board

Coordination body for Member States. Discloses its activities through the Commission portal.

National Competent Authorities of Member States (national competent authorities)

Market surveillance and notifying authorities designated by each Member State. The authorities of the target Member State of expansion must be checked individually. They differ by Member State, and the list is updated on the Commission portal.

Official EU Legal Portal (EUR-Lex)

Source for verifying official texts and Official Journal publications of Regulation (EU) 2024/1689 and its amendments.

Official portal eur-lex.europa.eu

Korean Side

Ministry of Science and ICT (MSIT)

Competent ministry for the Framework Act on AI. Contact channel for high-impact determination, domestic representative notifications, and submission of safety implementation results.

Official portal msit.go.kr

Personal Information Protection Commission (PIPC)

Competent authority for the Personal Information Protection Act; issues AI-related personal data guidelines. Contact channel for Korea-EU adequacy decisions.

Official portal pipc.go.kr

Korea Internet & Security Agency (KISA)

Practical support for personal data protection and cybersecurity.

Korea Trade-Investment Promotion Agency (KOTRA)

Support for companies expanding into Europe and local regulatory information.

Ministry of SMEs and Startups (MSS) and Affiliated Agencies

Support for SMEs in responding to overseas regulations and certifications.

Practical Recommendations

First check the national competent authorities of the target Member State for expansion and

whether a regulatory sandbox (Chapter 22) is operational. Primary EU sources must always be cross-referenced with the original texts on the official portals listed above. Use secondary commentary (such as law firm newsletters) for general direction, but confirm dates and legal provisions using primary sources.

Appendix 9. Summary of Key Articles

When legal citations are required, always cross-reference the official text of Regulation (EU) 2024/1689 on EUR-Lex (and its 2026 amendments). Below is a summary of key articles frequently referenced in practice, not the literal text of the articles themselves.

EU AI Act

Article 2 (Scope)

Applies to providers placing AI on the market in Europe, deployers within Europe, and providers and deployers in third countries whose outputs are used in Europe. Legal basis for extraterritorial application. (Chapter 1)

Article 3 (Definitions)

Definitions of provider, deployer, importer, distributor, authorized representative, placing on the market, putting into service, general-purpose AI model, systemic risk, etc. (Chapters 1 & 7)

Article 4 (AI Literacy)

Obligation to ensure that personnel operating AI possess adequate AI literacy. Applicable from February 2, 2025. (Chapter 21)

Article 5 (Prohibited Practices)

Prohibition of subliminal manipulation, exploitation of vulnerabilities, social scoring, untargeted scraping of facial images, emotion recognition in workplace/educational institutions, biometric categorisation inferring sensitive data, real-time remote biometric identification, etc. 2026 amendments add prohibitions on generating NCII and CSAM (December 2, 2026). (Chapters 15 & 5)

Article 6 (Classification Rules for High-Risk AI Systems)

Classification into high-risk systems embedded in Annex I products and standalone systems under Annex III. Paragraph 3 serves as a filter for Annex III (excluded if there is no significant risk, with exceptions barred in cases of profiling). (Chapter 11)

Articles 9 to 15 (Requirements for High-Risk AI Providers)

Risk management, data governance, technical documentation, record-keeping, transparency and provision of information, human oversight, accuracy, robustness, and cybersecurity. (Chapter 12)

Article 22 & Article 54 (Authorized Representative)

Third-country providers must appoint an authorized representative in the EU before placing high-risk systems (Article 22) or general-purpose AI models (Article 54) on the European market. The representative must retain technical documentation (10 years for general-purpose AI models) and cooperate with authorities. (Chapter 1)

Article 26 & Article 27 (Deployer Obligations & Fundamental Rights Impact Assessment)

Compliance with instructions for use, assignment of human oversight, retention of logs (at least 6 months), worker notification. Deployers that are public authorities or providers of public services must perform a Fundamental Rights Impact Assessment (FRIA) for Annex III systems. (Chapter 13)

Article 43, Article 48, & Article 49 (Conformity Assessment, CE Marking, Registration)

First area of Annex III may involve notified bodies; internal control for the rest. CE marking and registration in the EU database. (Chapter 12)

Article 50 (Transparency Obligations)

Paragraph 1: chatbot notification (providers); Paragraph 2: machine-readable marking of synthetic content (providers); Paragraph 3: emotion recognition and biometric categorisation disclosure (deployers); Paragraph 4: disclosure of deepfakes and public interest text (deployers); Paragraph 5: timing and format. Paragraph 7: codes of practice. Applicable from August 2, 2026. (Chapters 4 & 6)

Article 51, Article 53, & Article 55 (General-Purpose AI Models)

Article 51: presumption of systemic risk (10^{25} FLOPs). Article 53: basic obligations (technical documentation, downstream information, copyright policy, summary of training content). Article 55: additional obligations for systemic risk (evaluations, red-teaming, incident reporting, cybersecurity). (Chapters 7, 8, & 9)

Article 99 (Fines)

Infringements of prohibitions: up to €35 million or 7% of total worldwide turnover; other obligations: up to €15 million or 3%; incorrect information: up to €7.5 million or 1%. Whichever is higher. (Chapter 2)

Article 113 (Entry into Force and Application)

Base provision for phased application dates. (Chapter 3)

Related EU Laws

Directive (EU) 2019/790 Article 4(3)

Reservation of rights by copyright holders regarding text and data mining. Legal basis for copyright policies on general-purpose AI models. (Chapter 8)

Korea's Framework Act on AI

Article 2, Item 4 (High-Impact AI)

Artificial intelligence used in 11 specified domains that significantly impacts life, physical safety, or fundamental rights. (Chapter 16)

Article 31 (Obligation to Ensure Transparency)

Prior notification for high-impact/generative AI use, labeling of generative AI outputs, clear disclosure for deepfake recognition. (Chapter 16)

Article 32 (Obligation to Ensure Safety) & Article 34 (Responsibilities of High-Impact AI Operators)

Risk management for systems meeting or exceeding compute thresholds; risk management, explanation, and document retention for high-impact AI operators. (Chapter 16)

Article 36 (Domestic Representative)

Designation and reporting of a domestic representative by overseas operators exceeding revenue or user thresholds. (Chapter 16)

The EU AI Act

ISBN |

Author | Kim Kyung-jin

Publisher | Kim Kyung-jin Law Office

Publication Registration | 2025. 3. 10. (No. 2025-000015)

Address | 304 Baegil Bldg., 91 Jeonnong-ro, Dongdaemun-gu, Seoul, Republic of Korea

Tel | +82-2-6338-1905

Email | kimkj008@gmail.com

© Kim Kyung-jin 2026

This book is the intellectual property of the author. Unauthorized reproduction or distribution is prohibited.

Note) Parts of this book's content and editing process were assisted by artificial intelligence tools.

Support This AI Library Book

If this book has been helpful, please consider supporting future translations, public editions, and open AI knowledge projects via PayPal.



PayPal

paypal.me/kimkjj

Scan the QR code or open the link above.