

CERTIFICATION COURSE COMPANION

EU AI ACT COMPLIANCE & AI ETHICS

The complete bilingual guide to Europe's Artificial Intelligence Act — risk classification, high-risk obligations, generative AI governance, enforcement, ethics, and the Spanish legal framework.

7 MODULES · 28 TOPICS · 100 ASSESSMENT QUESTIONS

ENGLISH · ESPAÑOL

EU AI Act Compliance & AI Ethics **Certification**

Course Companion Ebook · Ebook Complementario del Curso

Bilingual Edition — English / Español

EU AI Act Compliance & AI Ethics Certification — Course Companion

Published by the Spanish Compliance Institute · spanishcomplianceinstitute.com

© Spanish Compliance Institute. All rights reserved. No part of this publication may be reproduced, distributed or transmitted in any form without the prior written permission of the publisher, except for brief quotations in reviews and certain other non-commercial uses permitted by copyright law. This publication provides general educational information about the EU Artificial Intelligence Act and related frameworks; it does not constitute legal advice. Organizations should consult qualified counsel for advice on specific compliance situations.

Todos los derechos reservados. Esta publicación ofrece información educativa general sobre el Reglamento de IA de la UE y marcos relacionados; no constituye asesoramiento jurídico.

Contents

Índice

- **About This Ebook**

Sobre este ebook · How to use it · Cómo usarlo

- **Introduction**

Introducción

- 1 Introduction to EU AI Regulation**

Introducción a la Regulación de IA de la UE

1.1 Origins and Legislative Structure of the EU AI Act · 1.2 Key Definitions and Regulatory Terminology · 1.3 EU Digital Strategy and Alignment with Other Laws · 1.4 Obligations for Providers, Deployers, Importers, and Distributors · Module Assessment

- 2 EU AI Act Risk Categories**

Categorías de Riesgo de la Ley de IA de la UE

2.1 Prohibited AI Practices and Unacceptable Risk · 2.2 High-Risk AI Use Cases and Classification Criteria · 2.3 Transparency Rules for Limited-Risk AI · 2.4 Minimal-Risk AI and Voluntary Best Practices · Module Assessment

- 3 High-Risk AI System Requirements**

Requisitos para Sistemas de IA de Alto Riesgo

3.1 Governance and Lifecycle Risk Management · 3.2 Data Quality, Training Practices, and Bias Control · 3.3 Documentation, Conformity Assessment, and CE Marking · 3.4 Human Oversight, Monitoring, and System Reliability · Module Assessment

- 4 Governance of General-Purpose and Generative AI**

Gobernanza de IA de Propósito General y Generativa

4.1 General-Purpose AI Models and Regulatory Scope · 4.2 Systemic GPAI Obligations and Testing Requirements · 4.3 Transparency, Watermarking, and Synthetic Content Rules · 4.4 Responsibilities for Fine-Tuning and Downstream Deployment · Module Assessment

- 5 Enforcement, Risk, and Liability Frameworks**

Ejecución, Riesgo y Marcos de Responsabilidad

5.1 Market Surveillance, Audits, and Regulatory Powers · 5.2 Penalties, Fines, and Non-Compliance Consequences · 5.3 Civil and Product Liability in AI-Driven Systems · 5.4 Contractual Controls and Third-Party Risk Management · Module Assessment

- 6 Ethical AI Principles and Governance Structures**

Principios Éticos de IA y Estructuras de Gobernanza

6.1 Fairness, Non-Discrimination, and Human Rights Alignment · 6.2 Transparency, Explainability, and Trustworthiness · 6.3 Ethical Risk Assessment and Mitigation Techniques · 6.4 Organizational AI Governance, Oversight, and Accountability · Module Assessment

- 7 Spanish AI Legal and Regulatory Framework**

Marco Legal y Regulatorio de IA en España

7.1 Spain's Implementation of the EU AI Act · 7.2 GDPR, LOPDGDD, and National Data Protection Duties · 7.3 Spanish Labor, Workplace, and Equality Regulations · 7.4 National AI Strategy, Public Sector Rules, and Enforcement Bodies · Module Assessment

- **Final Certification Exam**

Examen Final de Certificación · 30 questions

- **Conclusion**

Conclusión

- A Appendix A — Answer Key**

Apéndice A — Clave de Respuestas

B **Appendix B — Glossary of Key Terms**
Apéndice B — Glosario de Términos Clave

About This Ebook

Sobre este ebook

This ebook is the complete companion text for the EU AI Act Compliance & AI Ethics Certification, a self-paced course from the Spanish Compliance Institute. It converts the full course curriculum — seven modules, twenty-eight topics, real Spanish and EU case studies, and one hundred assessment questions — into a structured reference you can read cover to cover or return to whenever a compliance question arises at work.

Este ebook es el texto complementario completo de la Certificación en Cumplimiento del EU AI Act y Ética en IA, un curso a ritmo propio del Spanish Compliance Institute. Convierte el plan de estudios completo del curso — siete módulos, veintiocho temas, casos reales de España y la UE, y cien preguntas de evaluación — en una referencia estructurada que puedes leer de principio a fin o consultar cuando surja una duda de cumplimiento en el trabajo.

7 MODULES · MÓDULOS	28 TOPICS · TEMAS	100 QUESTIONS · PREGUNTAS	EN·ES BILINGUAL · BILINGÜE
-------------------------------	-----------------------------	--	--------------------------------------

How to Use This Book

Cómo usar este libro

- 1

Read in your language · Lee en tu idioma
Every chapter runs in two parallel columns: English on the left, Spanish on the right. Switch between them freely — the content is mirrored paragraph for paragraph.

Cada capítulo se presenta en dos columnas paralelas: inglés a la izquierda, español a la derecha. Cambia libremente entre ambas — el contenido está reflejado párrafo a párrafo.
- 2

Learn from cases · Aprende con casos
Tinted panels contain real-world case studies, decision-point scenarios and applied examples drawn from Spanish and EU practice. Treat each one as a rehearsal for decisions you may face.

Los paneles sombreados contienen casos reales, escenarios de decisión y ejemplos aplicados de la práctica española y europea. Trata cada uno como un ensayo de decisiones que podrías enfrentar.
- 3

Pause and reflect · Pausa y reflexiona
Feynman reflection boxes ask you to write down your intuitions before a module and revisit them after. The gap between the two is where learning becomes visible.

Las cajas de reflexión Feynman te piden anotar tus intuiciones antes de un módulo y revisarlas después. La distancia entre ambas es donde el aprendizaje se hace visible.

4**Test yourself · Ponte a prueba**

Each module closes with a ten-question assessment, and the book ends with a thirty-question final exam. Answers are collected in Appendix A — resist checking them until you have committed.

Cada módulo cierra con una evaluación de diez preguntas, y el libro termina con un examen final de treinta preguntas. Las respuestas están en el Apéndice A — evita consultarlas antes de responder.

Why the EU AI Act Matters Now

Por qué el EU AI Act importa ahora

ENGLISH

Artificial intelligence has moved from the margins of European business into its decision-making core. Systems now shortlist job candidates, price insurance, allocate shifts, approve loans and support medical judgments. The EU AI Act — the world's first comprehensive AI law — answers a simple question with far-reaching consequences: who is accountable when an algorithm shapes a human life?

This book teaches the Act the way practitioners need it: as a working system, not a legal abstraction. You will follow its logic from origins and definitions, through the four-tier risk pyramid, into the heavy machinery of high-risk obligations, the new rules for general-purpose and generative models, the enforcement and liability apparatus that gives the law its teeth, the ethical principles beneath it all — and finally into Spain, where the LOPDGDD, the Workers' Statute, the AEPD and AESIA turn European rules into daily operational reality.

By the final page you will be able to classify a system, name the obligations that follow, identify who carries them, and defend those judgments to a regulator, an executive or a works council.

ESPAÑOL

La inteligencia artificial ha pasado de los márgenes de la empresa europea al núcleo de su toma de decisiones. Los sistemas ya preseleccionan candidatos, calculan primas de seguros, asignan turnos, aprueban préstamos y apoyan juicios médicos. El EU AI Act — la primera ley integral de IA del mundo — responde a una pregunta simple con consecuencias profundas: ¿quién es responsable cuando un algoritmo moldea una vida humana?

Este libro enseña la Ley como la necesitan los profesionales: como un sistema operativo, no como una abstracción legal. Seguirás su lógica desde los orígenes y definiciones, pasando por la pirámide de riesgo de cuatro niveles, la maquinaria de las obligaciones de alto riesgo, las nuevas reglas para modelos de propósito general y generativos, el aparato de ejecución y responsabilidad que da fuerza a la ley, los principios éticos que la sustentan — y finalmente España, donde la LOPDGDD, el Estatuto de los Trabajadores, la AEPD y AESIA convierten las reglas europeas en realidad operativa diaria.

Al llegar a la última página podrás clasificar un sistema, nombrar las obligaciones que le corresponden, identificar quién las asume y defender esos juicios ante un regulador, un directivo o un comité de empresa.

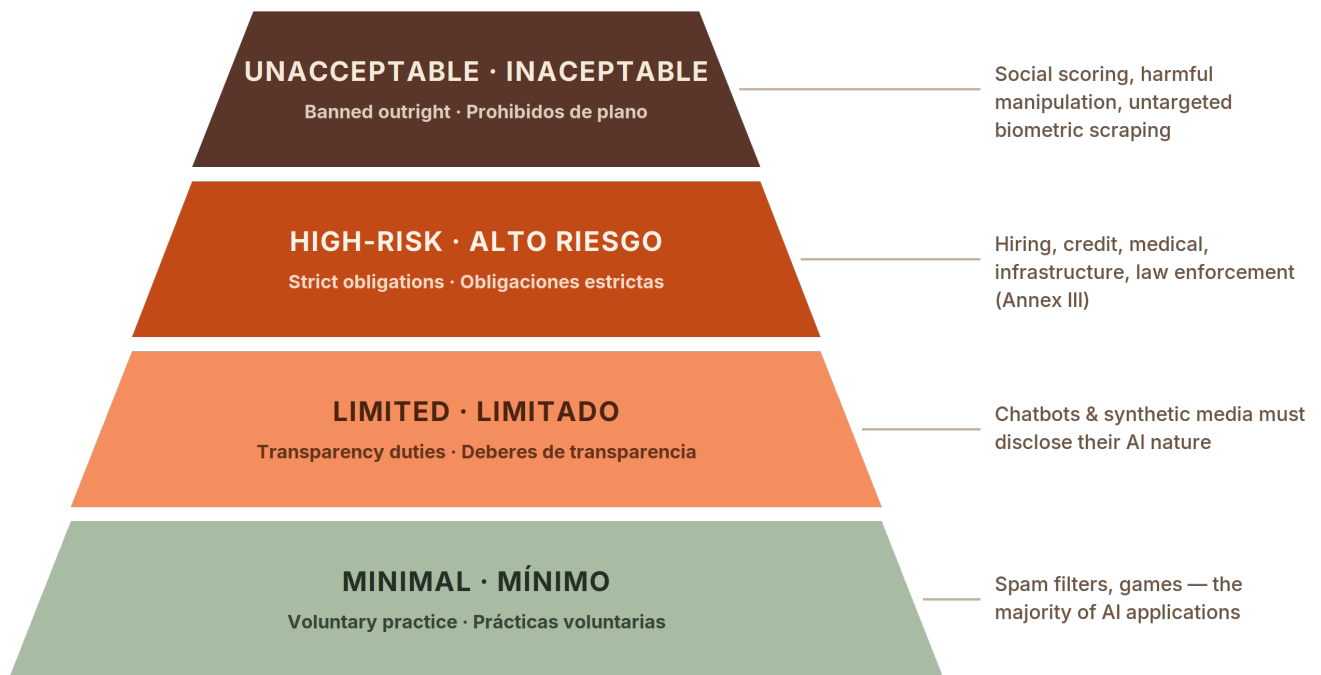


Figure 1 — The four risk tiers of the EU AI Act · Los cuatro niveles de riesgo del EU AI Act

1

MODULE ONE · MÓDULO 1

Introduction to EU AI Regulation

Introducción a la Regulación de IA de la UE

1.1 Origins and Legislative Structure of the EU AI Act

Orígenes y Estructura Legislativa del EU AI Act

1.2 Key Definitions and Regulatory Terminology

Definiciones Clave y Terminología Regulatoria

1.3 EU Digital Strategy and Alignment with Other Laws

Estrategia Digital de la UE y Alineación con Otras Leyes

1.4 Obligations for Providers, Deployers, Importers, and Distributors

Obligaciones para Proveedores, Desplegadores, Importadores y Distribuidores

Module 1 — Introduction to EU AI Regulation

Módulo 1 — Introducción a la Regulación de IA de la UE

ENGLISH

ESPAÑOL

WHY THIS MODULE · POR QUÉ ESTE MÓDULO

The regulatory roots of the European Union's AI Act can be traced back to the very same principles used to oversee consumer product safety. That unexpected influence reveals how seriously Europe takes AI's impact on society. AI is not treated as a distant innovation, but as a system capable of transforming access to services, opportunities and protections. This unusual starting point is the first key to understanding why the Act took its current shape.

The importance of this topic goes far beyond legal curiosity. Without a solid understanding of how the Act was built, it becomes difficult to interpret any of its terms, obligations or enforcement mechanisms. Spain's legal framework rests directly on this foundation, which means the clarity you gain here will support every compliance decision you make later. In this module you will discover the structure, reasoning and regulatory philosophy that guide the Act's design.

With every concept you master, you strengthen your confidence to operate in a complex regulatory environment. Let this moment set the pace of your growth and remind you that mastering the fundamentals is one of the most strategic steps toward becoming an informed, trusted voice in AI governance.

Las raíces regulatorias de la Ley de IA de la Unión Europea pueden rastrearse hasta los mismos principios utilizados para supervisar la seguridad de los productos de consumo. Esta influencia inesperada revela cuán seriamente Europa considera el impacto de la IA en la sociedad. No se la aborda como una innovación lejana, sino como un sistema capaz de transformar el acceso a servicios, oportunidades y protecciones. Este punto de partida poco común es la primera clave para comprender por qué la Ley adoptó su forma actual.

La importancia de este tema va mucho más allá de la curiosidad legal. Sin una comprensión sólida de cómo se construyó la Ley, resulta difícil interpretar cualquiera de sus términos, obligaciones o mecanismos de aplicación. El marco jurídico de España se apoya directamente en esta base, lo que significa que la claridad que adquieras aquí respaldará cada decisión de cumplimiento que tomes más adelante. En este módulo descubrirás la estructura, el razonamiento y la filosofía regulatoria que guían el diseño de la Ley.

Con cada concepto que comprendes, fortaleces tu confianza para desenvolverte en un entorno regulatorio complejo. Permite que este momento marque el ritmo de tu crecimiento y te recuerde que dominar los fundamentos es uno de los pasos más estratégicos para convertirte en una voz informada y confiable en la gobernanza de la IA. ## Feynman Before Section / Método Feynman – Antes de la Sección

PAUSE & REFLECT — BEFORE YOU BEGIN · REFLEXIÓN INICIAL DE FEYNMAN

Before you continue, pause for a moment. Test your intuitive understanding of AI regulation with a few simple questions. Start with this one: when you hear the phrase "AI regulation", what do you think it primarily tries to protect — people, organizations, or innovation? Don't overthink it; write down the first idea that comes to mind.

Now consider two more questions. Why do you think the EU built a dedicated framework for AI instead of relying only on existing laws? And when defining AI in a legal context, what matters more: how the technology works, or how it affects real life? Take a moment to write down your answers. Keep them — you will return to them later, and the progress you notice may surprise you.

Antes de continuar, hagamos una pausa. Quiero que evalúes tu comprensión intuitiva de la regulación de la IA respondiendo algunas preguntas sencillas. Comienza con esta: cuando escuchas la expresión "regulación de la IA", ¿qué crees que intenta proteger principalmente: a las personas, a las organizaciones o a la innovación? No lo pienses demasiado; anota la primera idea que te venga a la mente.

Ahora considera las siguientes dos preguntas. ¿Por qué crees que la UE construyó un marco específico para la IA en lugar de apoyarse únicamente en leyes existentes? Y al definir la IA en un contexto legal, ¿qué consideras más importante: cómo funciona la tecnología o cómo afecta a la vida real? Tómame un momento para escribir tus respuestas. Guárdalas; volverás a ellas más adelante y el progreso que notes puede sorprenderte. ## Topic 1: Origins and Legislative Structure of the EU AI Act / Tema 1: Orígenes y Estructura Legislativa de la Ley de IA de la UE

1.1 Origins and Legislative Structure of the EU AI Act

Orígenes y Estructura Legislativa del EU AI Act

To appreciate the EU AI Act, you need to understand the environment that shaped it. As AI systems began to influence hiring decisions, credit approvals, public services and medical assessments, policymakers identified the need for a unified approach. Existing laws were never designed for automated decision-making at this scale, and leaving AI unregulated risked widening inequalities and eroding public trust. Europe responded by treating AI regulation as a long-term investment in safety, fairness and democratic stability.

The Act's development began with broad consultations across academic, technical and legal communities. Researchers documented incidents of biased algorithms, flawed predictions and opaque automated decisions affecting millions of people. These findings informed the European Commission's first proposal in 2021. Spain took an active part in the negotiations, emphasizing the need for strong labor protections and transparent AI systems in both the public and private sectors.

Para apreciar la Ley de IA de la UE, es necesario comprender el entorno que la moldeó. A medida que los sistemas de IA comenzaron a influir en decisiones de contratación, aprobaciones de crédito, servicios públicos y evaluaciones médicas, los responsables políticos identificaron la necesidad de un enfoque unificado. Las leyes existentes nunca fueron diseñadas para la toma de decisiones automatizada a esta escala, y dejar la IA sin regular implicaba el riesgo de ampliar desigualdades y debilitar la confianza pública. Europa respondió tratando la regulación de la IA como una inversión a largo plazo en seguridad, equidad y estabilidad democrática.

El desarrollo de la Ley comenzó con amplias consultas en comunidades académicas, técnicas y jurídicas. Investigadores documentaron incidentes de algoritmos sesgados, predicciones defectuosas y decisiones automatizadas opacas que afectaban a millones de personas. Estos hallazgos informaron la primera propuesta de la Comisión Europea en 2021. España participó activamente en las negociaciones,

The structure of the Act reflects a deliberate hierarchy. It begins by anchoring every requirement in fundamental rights — the backbone of European law — before introducing classifications, obligations and penalties.

One of the most interesting design decisions is the use of mechanisms already proven in product safety legislation. Rather than inventing new compliance models, the EU incorporated familiar elements: conformity assessments, market surveillance, technical documentation and robust enforcement pathways.

What ultimately makes the Act unique is its flexibility. It avoids tying itself to specific technologies and focuses instead on purpose, context and impact. This structure ensures the rules stay relevant whether AI is used in a retail chatbot, a power grid or an immigration system.

enfaticando la necesidad de una fuerte protección laboral y de sistemas de IA transparentes tanto en el sector público como en el privado.

La estructura de la Ley refleja una jerarquía deliberada. Comienza anclando cada requisito en los derechos fundamentales, que constituyen la columna vertebral del derecho europeo, antes de introducir clasificaciones, obligaciones y sanciones.

Una de las decisiones de diseño más interesantes es el uso de mecanismos ya probados en la legislación de seguridad de productos. En lugar de inventar nuevos modelos de cumplimiento, la UE incorporó elementos conocidos: evaluaciones de conformidad, vigilancia del mercado, documentación técnica y vías sólidas de aplicación.

Lo que finalmente hace única a la Ley es su flexibilidad. Evita limitarse a tecnologías específicas y se centra en el propósito, el contexto y el impacto. Esta estructura garantiza que la normativa siga siendo relevante tanto si la IA se utiliza en un chatbot minorista como en una red eléctrica o un sistema de inmigración. ## Topic 2: Key Definitions and Regulatory Terminology + Reflection Story / Tema 2: Definiciones Clave y Terminología Regulatoria + Historia de Reflexión

1.2 Key Definitions and Regulatory Terminology

Definiciones Clave y Terminología Regulatoria

Clear terminology is one of the pillars of EU AI Act compliance. Definitions are not academic; they determine whether an organization must comply, which responsibilities apply, and how regulators assess risk. The definition of "AI system" is intentionally broad and covers technologies that generate outputs such as predictions, recommendations or decisions.

The legal roles defined in the Act carry significant weight. A "provider" is the entity that places an AI system on the EU market under its own name. A "deployer" uses AI within its operations. An "importer" brings systems in from outside the EU, and a "distributor" handles their commercial transfer. Each designation carries specific obligations — from documentation to oversight — and many organizations misread these roles without realizing it.

La terminología clara es uno de los pilares del cumplimiento de la Ley de IA de la UE. Las definiciones no son académicas; determinan si una organización debe cumplir, qué responsabilidades se aplican y cómo los reguladores evalúan el riesgo. La definición de "sistema de IA" es intencionalmente amplia e incluye tecnologías que generan resultados como predicciones, recomendaciones o decisiones.

Los roles legales definidos en la Ley tienen un peso significativo. Un "proveedor" es la entidad que introduce un sistema de IA en el mercado de la UE bajo su propio nombre. Un "deployer" utiliza la IA dentro de sus operaciones. Un "importador" introduce sistemas desde fuera de la UE, y un "distribuidor" gestiona su transferencia comercial. Cada designación conlleva obligaciones específicas, desde documentación hasta supervisión, y muchas

The mismatch between legal definitions and operational reality is one of the most common compliance mistakes. Some companies believe the original developer carries all responsibility. Yet if a deployer modifies the system's logic, retrains a model or changes its intended purpose, it can legally become a provider — and the regulatory obligations shift immediately. Understanding this distinction prevents unintentional infringements.

organizaciones malinterpretan estos roles sin darse cuenta.

La desalineación entre las definiciones legales y la realidad operativa es uno de los errores de cumplimiento más comunes. Algunas empresas creen que el desarrollador original asume toda la responsabilidad. Sin embargo, si un deployer modifica la lógica del sistema, reentrena un modelo o cambia su propósito previsto, puede convertirse legalmente en proveedor, y las obligaciones regulatorias cambian de inmediato. Comprender esta distinción evita infracciones involuntarias.

CASE IN PRACTICE · CASO PRÁCTICO

An HR technology company in Barcelona acquired an automated candidate-scoring tool from a US vendor. Believing they were acting purely as deployers, they integrated the tool and adjusted certain weightings to fit Spanish hiring norms. In legal terms, that modification made them providers of the system, taking on responsibility for documentation, data governance and risk mitigation. When questions arose about fairness in hiring, they had no compliance records to present.

Once they recognized their correct legal role, the organization rebuilt its compliance posture. They documented model performance, implemented oversight steps and clarified the system's intended purpose.

Una empresa de tecnología de recursos humanos en Barcelona adquirió una herramienta de puntuación automatizada de un proveedor estadounidense. Creyendo que actuaban únicamente como deployers, integraron la herramienta y ajustaron ciertos pesos para adaptarla a las normas de contratación en España. En términos legales, esta modificación los convirtió en proveedores del sistema, asumiendo la responsabilidad por documentación, gobernanza de datos y mitigación de riesgos. Cuando surgieron dudas sobre la equidad en la contratación, no contaban con registros de cumplimiento que presentar.

Una vez que reconocieron su rol legal correcto, la organización reconstruyó su postura de cumplimiento. Documentaron el rendimiento del modelo, implementaron pasos de supervisión y aclararon el propósito previsto del sistema. ## Topic 3: EU Digital Strategy and Alignment With Other Laws + Scenario-Based Mini Assessment / Tema 3: Estrategia Digital de la UE y Alineación con Otras Leyes + Mini Evaluación Basada en Escenarios

1.3 EU Digital Strategy and Alignment with Other Laws

Estrategia Digital de la UE y Alineación con Otras Leyes

When the EU AI Act is placed within Europe's broader digital strategy, its purpose becomes even clearer. Europe understood that AI, data protection, cybersecurity and digital markets could not be governed in isolation. The Act therefore aligns with

Cuando se sitúa la Ley de IA de la UE dentro de la estrategia digital europea más amplia, su propósito resulta aún más claro. Europa comprendió que la IA, la protección de datos, la ciberseguridad y los mercados digitales no podían gobernarse de forma aislada. Por

the GDPR, the Digital Services Act, the Digital Markets Act and sector-specific safety rules. Together they form a synchronized framework that ensures technology develops responsibly and rights are respected in Spain and across the EU. This alignment prevents regulatory gaps in data handling, AI model deployment or the integration of foreign tools.

The most important connection is with the GDPR. AI systems typically depend on personal data, and GDPR rules on lawful processing, transparency and accountability remain fully in force. The AI Act does not replace the GDPR; it reinforces it with additional requirements for high-risk models, documentation and human oversight. In Spain, the LOPDGDD adds another layer, strengthening workers' rights and employer responsibilities.

There is also strong alignment with sector regulations. AI-based medical tools, for instance, must comply with both the AI Act and the Medical Devices Regulation. Energy systems using predictive AI must continue to respect critical-infrastructure safety rules.

To see how this alignment works in practice, consider a Spanish fintech startup developing an AI-based credit-scoring tool. Even before it is classified as high-risk under the AI Act, the company must comply with the GDPR on fairness, transparency and automated decisions. The AI Act then adds layers of documentation, testing and oversight.

ello, la Ley se alinea con el GDPR, la Ley de Servicios Digitales, la Ley de Mercados Digitales y normas sectoriales de seguridad. En conjunto, forman un marco sincronizado que garantiza el desarrollo responsable de la tecnología y el respeto de los derechos en España y en toda la UE. Esta alineación evita vacíos regulatorios en la gestión de datos, el despliegue de modelos de IA o la integración de herramientas extranjeras.

La conexión más relevante es con el GDPR. Los sistemas de IA suelen depender de datos personales, y las normas del GDPR sobre tratamiento lícito, transparencia y responsabilidad siguen plenamente vigentes. La Ley de IA no sustituye al GDPR; lo refuerza mediante requisitos adicionales para modelos de alto riesgo, documentación y supervisión humana. En España, la LOPDGDD añade otra capa, reforzando los derechos de los trabajadores y las responsabilidades del empleador.

También existe una fuerte alineación con regulaciones sectoriales. Por ejemplo, las herramientas médicas basadas en IA deben cumplir tanto con la Ley de IA como con el Reglamento de Productos Sanitarios. Los sistemas energéticos que utilizan IA predictiva deben seguir respetando las normas de seguridad de infraestructuras críticas.

Para comprender cómo funciona esta alineación en la práctica, considera una startup fintech española que desarrolla una herramienta de puntuación crediticia basada en IA. Incluso antes de clasificarla como de alto riesgo según la Ley de IA, la empresa debe cumplir con el GDPR en materia de equidad, transparencia y decisiones automatizadas. La Ley de IA añade luego capas de documentación, pruebas y supervisión.

DECISION POINT · PUNTO DE DECISIÓN

A Madrid-based logistics company adopts an AI scheduling tool developed in the United States and integrates it without reviewing alignment with the GDPR or the AI Act. Three options emerge: A) Use the tool immediately because the vendor claims it is "fully compliant". B) Conduct only a GDPR review, assuming AI Act obligations apply only to developers. C) Assess both GDPR and AI Act requirements, verifying documentation and risk controls before deployment.

Una empresa logística con sede en Madrid adopta una herramienta de programación basada en IA desarrollada en Estados Unidos e la integra sin revisar la alineación con el GDPR ni con la Ley de IA. Surgen tres opciones: A) Usar la herramienta de inmediato porque el proveedor afirma que es "totalmente compatible". B) Realizar solo una revisión del GDPR, asumiendo que las obligaciones de la Ley de IA solo aplican a los desarrolladores. C) Evaluar tanto los requisitos del GDPR como de la Ley de IA, verificando documentación y controles de riesgo antes del despliegue. Correct Action: C — porque los deployers en España deben garantizar el cumplimiento legal, independientemente del origen del sistema o de quién lo haya desarrollado. ## Topic 4: Obligations for Providers, Deployers, Importers, and Distributors + Practical Example / Tema 4: Obligaciones de Proveedores, Deployers, Importadores y Distribuidores + Ejemplo Práctico

Correct answer: C — Deployers in Spain must ensure legal compliance regardless of where the system originated or who developed it.

Respuesta correcta: C — los deployers en España deben garantizar el cumplimiento legal, independientemente del origen del sistema o de quién lo haya desarrollado.

1.4 Obligations for Providers, Deployers, Importers, and Distributors

Obligaciones para Proveedores, Desplegadores, Importadores y Distribuidores

The roles defined in the AI Act are not abstract; they determine accountability. Providers must ensure the system meets every legal requirement before reaching the market. Deployers answer for how AI is used, monitored and supervised inside the organization. Importers must verify that third-country systems meet EU standards before bringing them into Europe. Distributors must avoid supplying systems that appear non-compliant. Each role sustains a chain of responsibility designed to leave no gap where harm could occur.

Providers shoulder the most demanding obligations. They must maintain technical documentation, guarantee data quality, carry out risk management, implement cybersecurity safeguards and build clear human-oversight instructions into the system.

Los roles definidos en la Ley de IA no son abstractos; determinan la rendición de cuentas. Los proveedores deben garantizar que el sistema cumpla todos los requisitos legales antes de llegar al mercado. Los deployers son responsables de cómo se utiliza, monitorea y supervisa la IA dentro de la organización. Los importadores deben verificar que los sistemas de terceros países cumplan los estándares de la UE antes de introducirlos en Europa. Los distribuidores deben evitar suministrar sistemas que parezcan no conformes. Cada rol sostiene una cadena de responsabilidad diseñada para evitar vacíos donde pueda producirse daño.

Los proveedores asumen las obligaciones más exigentes. Deben mantener documentación técnica, garantizar la calidad de los datos, realizar gestión de

Importers and distributors often underestimate their responsibilities. An importer bringing AI tools from the United States into Spain cannot simply assume the foreign provider is compliant. It must verify the conformity documentation, confirm CE marking where required, and make sure the system's instructions align with EU expectations.

riesgos, implementar salvaguardas de ciberseguridad e incorporar instrucciones claras de supervisión humana en el sistema.

Importadores y distribuidores suelen subestimar sus responsabilidades. Un importador que introduce herramientas de IA desde Estados Unidos a España no puede asumir automáticamente que el proveedor extranjero cumple la normativa. Debe verificar la documentación de conformidad, confirmar la existencia del marcado CE cuando sea requerido y asegurarse de que las instrucciones del sistema se alineen con las expectativas de la UE.

APPLIED EXAMPLE · EJEMPLO APLICADO

Consider a US supermarket chain expanding into Spain with an AI-based theft-prevention system. In the United States, the system faces minimal oversight requirements. Once introduced into Spain, however, the importer must verify whether the AI uses biometric categorization, whether it qualifies as high-risk, and whether the provider's documentation complies with EU and Spanish labor and privacy law.

These obligations keep AI use legal, fair and transparent. They also protect organizations from severe penalties, reputational damage and operational disruption.

Considera una cadena de supermercados estadounidense que se expande a España con un sistema de prevención de robos basado en IA. En Estados Unidos, el sistema tiene requisitos mínimos de supervisión. Sin embargo, al introducirse en España, el importador debe verificar si la IA utiliza categorización biométrica, si se considera de alto riesgo y si la documentación del proveedor cumple con la legislación laboral y de privacidad de la UE y de España.

Estas obligaciones garantizan que el uso de la IA se mantenga legal, equitativo y transparente. También protegen a las organizaciones frente a sanciones severas, daños reputacionales e interrupciones operativas. ## Feynman After Section / Método Feynman – Después de la Sección

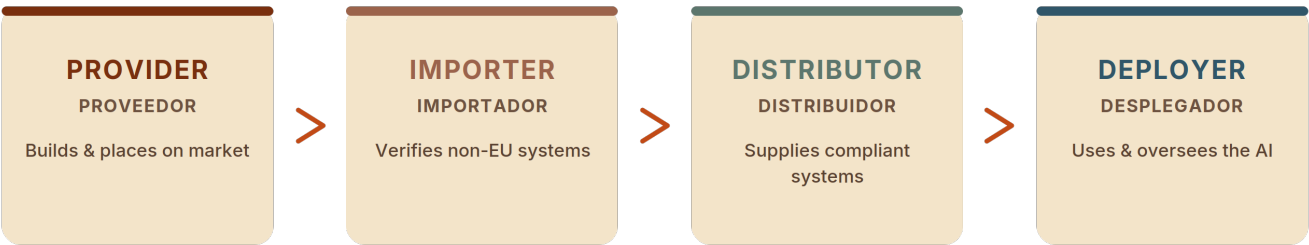


Figure 2 — The chain of responsibility · La cadena de responsabilidad

PAUSE & REFLECT — LOOKING BACK · REFLEXIÓN FINAL DE FEYNMAN

Let's return to the opening questions with your new understanding. The first asked what AI regulation tries to protect. You can now see that it protects people, organizations and innovation alike by creating structured responsibilities instead of letting risks accumulate. That balance is central to the European approach.

The next questions explored why the EU created a dedicated AI law and whether function or impact matters more. The answer is now clear: existing laws could not address AI's complex, dynamic nature — and impact is what ultimately defines risk. The Act focuses on what AI systems do in society, not just on how they are programmed.

Volvamos a las preguntas iniciales con tu nueva comprensión. La primera preguntaba qué intenta proteger la regulación de la IA. Ahora puedes ver que protege a las personas, a las organizaciones y a la innovación al crear responsabilidades estructuradas en lugar de permitir que los riesgos se acumulen. Este equilibrio es central en el enfoque europeo.

Las siguientes preguntas exploraban por qué la UE creó una ley específica para la IA y si la función o el impacto es más importante. Ahora la respuesta es clara: las leyes existentes no podían abordar la naturaleza compleja y dinámica de la IA, y el impacto es lo que finalmente define el riesgo. La Ley se centra en lo que los sistemas de IA hacen en la sociedad, no solo en cómo están programados. ## Module Wrap-Up + Practical Life Implementation / Cierre del Módulo + Aplicación Práctica en la Vida Real

Module Wrap-Up & Practical Implementation

Cierre del Módulo e Implementación Práctica

We have explored the origins, definitions, strategic connections and legal responsibilities built into the EU AI Act. Each element lays a foundation for understanding how AI must be governed in Spain and across the European Union. This matters because AI is no longer optional: it shapes decisions, workplaces and public services at an accelerating pace. The clarity you have gained here prepares you for the more advanced modules ahead.

Remember the reflection story and the scenario assessment. Both showed how small misunderstandings — misassigning a role, skipping compliance checks — can produce significant legal and ethical consequences. They also showed that taking the right approach restores confidence, strengthens governance and protects organizational reputation. As you move forward, apply this knowledge to your own environment. Identify your organization's role — provider, deployer, importer or distributor — and assess whether the corresponding responsibilities are being met.

Hemos explorado los orígenes, las definiciones, las conexiones estratégicas y las responsabilidades legales integradas en la Ley de IA de la UE. Cada elemento construye una base para comprender cómo debe gobernarse la IA en España y en toda la Unión Europea. Esto es fundamental porque la IA ya no es opcional: influye en decisiones, entornos laborales y servicios públicos a una velocidad creciente. La claridad adquirida aquí te prepara para los módulos más avanzados que siguen.

Recuerda la historia de reflexión y la evaluación basada en escenarios. Ambas demostraron cómo pequeños malentendidos, como asignar incorrectamente un rol o omitir verificaciones de cumplimiento, pueden generar consecuencias legales y éticas significativas. También mostraron que adoptar el enfoque correcto restaura la confianza, fortalece la gobernanza y protege la reputación organizacional. A medida que avances, aplica este conocimiento a tu entorno real. Identifica el rol de tu organización —proveedor, deployer, importador o

distribuidor— y evalúa si las responsabilidades correspondientes se están cumpliendo.

MODULE 1 ASSESSMENT · 10 QUESTIONS

Test Your Understanding

Evalúa tu comprensión — Respuestas en el Apéndice A

1. What was the primary motivation behind creating the EU AI Act?

¿Cuál fue la motivación principal detrás de la creación del Reglamento de IA de la UE?

- A.** To promote unrestricted AI innovation / Promover la innovación en IA sin restricciones
- B.** To replace the GDPR with AI-specific legislation / Reemplazar el GDPR con una legislación específica para IA
- C.** To establish a harmonized legal framework for AI risks / Establecer un marco legal armonizado para los riesgos de la IA
- D.** To regulate only military AI systems / Regular únicamente los sistemas de IA militares

2. Which EU institution formally proposed the AI Act?

¿Qué institución de la UE propuso formalmente el Reglamento de IA?

- A.** European Parliament / Parlamento Europeo
- B.** European Commission / Comisión Europea
- C.** European Data Protection Board / Comité Europeo de Protección de Datos
- D.** Court of Justice of the EU / Tribunal de Justicia de la Unión Europea

3. Under the EU AI Act, the term "provider" primarily refers to:

Dentro del Reglamento de IA de la UE, el término "proveedor" se refiere principalmente a:

- A.** Any user interacting with an AI system / Cualquier usuario que interactúe con un sistema de IA
- B.** An entity importing AI into the EU / Una entidad que importa IA a la UE
- C.** An entity that develops or places an AI system on the market / Una entidad que desarrolla o pone un sistema de IA en el mercado
- D.** A regulatory authority supervising AI systems / Una autoridad reguladora que supervisa los sistemas de IA

4. Which existing EU law does the AI Act align with most closely in terms of risk-based regulation?

¿Con qué ley existente de la UE se alinea más el Reglamento de IA en términos de regulación basada en riesgos?

- A.** Digital Services Act / Ley de Servicios Digitales
- B.** GDPR / GDPR
- C.** ePrivacy Directive / Directiva de ePrivacidad
- D.** Competition Law / Ley de Competencia

5. Which role is responsible for ensuring an AI system is used according to its intended purpose?

¿Qué rol es responsable de garantizar que un sistema de IA se utilice según su propósito previsto?

- A.** Provider / Proveedor
- B.** Importer / Importador
- C.** Distributor / Distribuidor
- D.** Deployer / Desplegador

6. The EU AI Act is best described as:

El Reglamento de IA de la UE se describe mejor como:

- A.** A voluntary ethical guideline / Una guía ética voluntaria
- B.** A sector-specific directive / Una directiva sectorial específica
- C.** A binding horizontal regulation / Un reglamento horizontal vinculante
- D.** A national-level policy framework / Un marco de políticas a nivel nacional

7. The EU AI Act applies only to AI systems developed within the European Union.

El Reglamento de IA de la UE se aplica únicamente a los sistemas de IA desarrollados dentro de la Unión Europea.

- A.** True / Verdadero
- B.** False / Falso
- C.** Partially true / Parcialmente verdadero
- D.** Not defined / No definido

8. Distributors under the EU AI Act have no compliance obligations once an AI system is placed on the market.

Los distribuidores bajo el Reglamento de IA de la UE no tienen obligaciones de cumplimiento una vez que un sistema de IA se coloca en el mercado.

- A.** True / Verdadero
- B.** False / Falso
- C.** Only for high-risk AI / Solo para IA de alto riesgo
- D.** Only outside the EU / Solo fuera de la UE

9. The EU AI Act follows a _____ approach to regulating AI systems.

El Reglamento de IA de la UE sigue un enfoque _____ para regular los sistemas de inteligencia artificial.

- A.** Technology-neutral / Neutral en tecnología
- B.** Risk-based / Basado en riesgos
- C.** Market-driven / Impulsado por el mercado
- D.** Innovation-first / Innovación primero

10. The EU AI Act is part of the EU's broader _____ strategy aimed at regulating digital technologies.

El Reglamento de IA de la UE forma parte de la estrategia más amplia de la UE en _____ destinada a regular las tecnologías digitales.

- A.** Industrial / Industrial
- B.** Economic / Económica
- C.** Digital / Digital
- D.** Cybersecurity / Ciberseguridad

2

MODULE TWO · MÓDULO 2

EU AI Act Risk Categories

Categorías de Riesgo de la Ley de IA de la UE

2.1 Prohibited AI Practices and Unacceptable Risk

Prácticas de IA Prohibidas y Riesgo Inaceptable

2.2 High-Risk AI Use Cases and Classification Criteria

Casos de Uso de IA de Alto Riesgo y Criterios de Clasificación

2.3 Transparency Rules for Limited-Risk AI

Reglas de Transparencia para IA de Riesgo Limitado

2.4 Minimal-Risk AI and Voluntary Best Practices

IA de Riesgo Mínimo y Mejores Prácticas Voluntarias

Module 2 — EU AI Act Risk Categories

Módulo 2 — Categorías de Riesgo de la Ley de IA de la UE

ENGLISH

ESPAÑOL

WHY THIS MODULE · POR QUÉ ESTE MÓDULO

If an AI system were going to make decisions about your job, your credit score or your medical care, how much confidence would you want to feel in its fairness and safety? That simple question led the EU to design one of the AI Act's defining features: the risk classification system. Once you understand this system, the rest of the law stops feeling abstract and becomes deeply practical.

Risk categories matter because they determine obligations, responsibilities, penalties — and even whether an AI system may be permitted in the EU at all. Every tool, whether used in Spain, Italy or the US before deployment, must be placed in the correct risk tier. Today you will explore prohibited AI systems, high-risk systems, limited-risk systems with transparency rules, and minimal-risk systems — each one a pillar of safe, lawful AI use in Europe.

Understanding risk categories gives you the clarity, control and strategic insight that separate solid AI governance from accidental non-compliance. Let that spark your confidence: mastering the risk tiers is the first key step toward mastering the EU AI landscape.

Si un sistema de IA tomará decisiones sobre tu empleo, tu puntaje crediticio o tu atención médica, ¿qué nivel de confianza querrías sentir en su equidad y seguridad? Esa simple pregunta llevó a la UE a diseñar una de las características más definitorias de la Ley de IA: el sistema de clasificación de riesgos. Una vez que entiendes este sistema, el resto de la ley deja de sentirse abstracto y comienza a ser profundamente práctico.

Las categorías de riesgo importan porque determinan obligaciones, responsabilidades, sanciones e incluso si un sistema de IA puede permitirse en la UE. Cada herramienta —ya sea utilizada en España, Italia o EE. UU. antes de su implementación— debe ubicarse en el nivel de riesgo correcto. Hoy explorarás los sistemas de IA prohibidos, de alto riesgo, de riesgo limitado con reglas de transparencia y de riesgo mínimo, cada uno formando un pilar del uso seguro y legal de la IA en Europa.

Comprender las categorías de riesgo te brinda claridad, control e insight estratégico que separa una gobernanza sólida de IA de un incumplimiento accidental. Deja que esto encienda tu confianza, porque entender los niveles de riesgo es el primer paso clave para dominar el panorama de IA de la UE. ## Topic 1: Prohibited AI Practices and Unacceptable Risk / Tema 1: Prácticas de IA Prohibidas y Riesgo Inaceptable

2.1 Prohibited AI Practices and Unacceptable Risk

Prácticas de IA Prohibidas y Riesgo Inaceptable

The EU AI Act opens with its strictest category: unacceptable-risk systems. These are not systems that can be corrected or improved; they are systems the EU has decided cannot be permitted under any circumstances, because they fundamentally threaten safety, dignity or fundamental rights. Their design or purpose is incompatible with European values, and Spain reinforces these prohibitions through strong constitutional protections.

La Ley de IA de la UE comienza con la categoría más estricta: sistemas de riesgo inaceptable. No son sistemas que puedan corregirse o mejorarse; son aquellos que la UE ha decidido que no pueden permitirse bajo ninguna circunstancia porque amenazan de forma fundamental la seguridad, la dignidad o los derechos fundamentales. Su diseño o propósito es incompatible con los valores europeos, y

Prohibited systems include AI designed to manipulate human behavior in harmful ways, exploit vulnerable populations, or socially score people based on personal traits or past actions. Real-time biometric identification in public spaces also falls into this category, with very narrow exceptions such as investigations of serious crimes or disappearances. These rules exist because past abuses showed how quickly such tools can become instruments of discrimination or mass surveillance.

Another prohibited practice involves AI systems that distort human decision-making through subliminal techniques. The EU considers these systems inherently unsafe because they bypass rational thought. In Spain, where labor rights and dignity are strongly protected, this rule has direct implications for HR, marketing and public-facing technologies.

It is worth stressing that unacceptable-risk systems are not banned out of hostility to innovation. They are banned because their potential harm outweighs any possible benefit. The risk-benefit equation is so unfavorable that regulation alone cannot make them safe. This proactive stance reveals the EU's commitment to protecting society from irreversible harm.

Understanding prohibited AI gives you clarity about which lines must never be crossed. Organizations operating in Spain must make sure their tools — imported, built in-house or purchased — do not fall into these categories. The consequence of misclassification is not just a fine; it puts public trust, legal accountability and reputation at risk.

España refuerza estas prohibiciones mediante sólidas protecciones constitucionales.

Los sistemas prohibidos incluyen IA diseñada para manipular comportamientos humanos de manera dañina, explotar poblaciones vulnerables o puntuar socialmente a personas basándose en rasgos personales o acciones pasadas. La identificación biométrica en tiempo real en espacios públicos también se encuentra en esta categoría, con excepciones muy limitadas, como investigaciones por delitos graves o desapariciones. Estas reglas existen porque abusos anteriores demostraron lo rápido que tales herramientas pueden convertirse en instrumentos de discriminación o vigilancia masiva.

Otra práctica prohibida involucra sistemas de IA que distorsionan la toma de decisiones humanas mediante técnicas subliminales. La UE considera estos sistemas inherentemente inseguros porque evaden el pensamiento racional. En España, donde los derechos laborales y la dignidad están fuertemente protegidos, esta regla tiene implicaciones directas en recursos humanos, marketing y tecnologías de interacción pública.

Vale la pena destacar que los sistemas de riesgo inaceptable no están prohibidos por ser anti-innovación. Están prohibidos porque su potencial daño supera cualquier beneficio posible. En otras palabras, la ecuación riesgo-beneficio es tan desfavorable que solo la regulación no puede hacerlos seguros. Este enfoque proactivo revela el compromiso de la UE con proteger la sociedad de daños irreversibles.

Comprender la IA prohibida te da claridad sobre qué líneas no deben cruzarse. Las organizaciones que operan en España deben asegurarse de que sus herramientas —ya sean importadas, desarrolladas internamente o adquiridas— no entren en estas categorías. La consecuencia de una clasificación errónea no es solo una multa; pone en riesgo la confianza pública, la responsabilidad legal y la reputación. ## Topic 2: High-Risk AI Use Cases and Classification Criteria + Reflection Story / Tema 2: Casos de IA de Alto Riesgo y Criterios de Clasificación + Historia de Reflexión

2.2 High-Risk AI Use Cases and Classification Criteria

Casos de Uso de IA de Alto Riesgo y Criterios de Clasificación

High-risk AI sits one level below the unacceptable category, yet its obligations are stricter than in any other tier. It covers systems used in critical sectors such as employment, credit, health, transport, public services and security — areas where errors or bias can cause significant harm. Providers and deployers must therefore meet strict requirements before these systems can be used.

High-risk classification does not depend on the technology alone, but on its intended use. A machine-learning model that suggests playlists is harmless; the same architecture used to assess loan eligibility is high-risk. This purpose-based approach ensures the law captures what really matters: the consequences of AI decisions for real people.

A system is considered high-risk if it appears in Annex III of the Act or forms part of a regulated product requiring conformity assessments. In Spain, this includes AI affecting working conditions, public-administration decisions and safety systems. Classification triggers obligations such as documentation, risk management, oversight, cybersecurity and continuous monitoring.

La IA de alto riesgo se encuentra un nivel por debajo de la categoría inaceptable, pero sus obligaciones son más estrictas que en cualquier otra categoría. Involucra sistemas utilizados en sectores críticos como empleo, crédito, salud, transporte, servicios públicos y seguridad. Aquí, errores o sesgos pueden causar daños significativos. Por ello, proveedores y deployers deben cumplir requisitos estrictos antes de que estos sistemas puedan usarse.

La clasificación de alto riesgo no depende solo de la tecnología, sino de su uso previsto. Por ejemplo, un modelo de aprendizaje automático que sugiere listas de reproducción es inofensivo, pero la misma arquitectura usada para evaluar la elegibilidad de préstamos se considera de alto riesgo. Este enfoque basado en propósito garantiza que la ley capture lo que realmente importa: las consecuencias de las decisiones de la IA sobre personas reales.

Un sistema se considera de alto riesgo si aparece en el Anexo III de la Ley o forma parte de un producto regulado que requiere evaluaciones de conformidad. En España, esto incluye IA que afecta condiciones laborales, decisiones de la administración pública y sistemas de seguridad. Esta clasificación activa obligaciones como documentación, gestión de riesgos, supervisión, ciberseguridad y monitoreo continuo.

CASE IN PRACTICE · CASO PRÁCTICO

In northern Spain, a mid-sized employer rolled out an AI-based interview-ranking tool, assuming it was harmless because it merely "sorted candidates". They treated it as a simple productivity tool. Months later, applicants complained about biased outcomes. Regulators determined the system should have been classified as high-risk because it influenced hiring decisions — one of the clearest Annex III categories. The company lacked documentation, oversight mechanisms and fairness testing. After correct classification, they rebuilt the system using transparent scoring, human review and validated datasets. Complaints stopped and trust improved markedly.

En el norte de España, un empleador mediano implementó una herramienta de ranking de entrevistas basada en IA, suponiendo que era inofensiva porque solo "ordenaba candidatos". La trataron como una simple herramienta de productividad. Meses después, los postulantes se quejaron de resultados sesgados. Los reguladores determinaron que debía haberse clasificado como alto riesgo porque influía en decisiones de contratación —una de las categorías más claras del Anexo III. La empresa carecía de documentación, mecanismos de supervisión y pruebas de equidad. Tras la correcta clasificación, reconstruyeron el sistema usando puntuación transparente, revisión humana y conjuntos de datos validados. Las quejas cesaron y la confianza mejoró notablemente.

Danger rarely comes from intentional misuse; it more commonly arises from misunderstanding what qualifies as high-risk. Correct classification turns vague risks into structured actions, guiding organizations toward responsible, lawful AI adoption.

El peligro rara vez proviene de un mal uso intencional; más comúnmente surge de malinterpretar qué califica como alto riesgo. Una clasificación correcta transforma riesgos vagos en acciones estructuradas, guiando a las organizaciones hacia una adopción responsable y legal de la IA. ## Topic 3: Transparency Rules for Limited-Risk AI + Regular Mini Assessment / Tema 3: Reglas de Transparencia para IA de Riesgo Limitado + Mini Evaluación

2.3 Transparency Rules for Limited-Risk AI

Reglas de Transparencia para IA de Riesgo Limitado

Limited-risk AI covers tools that interact directly with people but lack the gravity of high-risk systems — chatbots, content generators, synthetic voices and emotion-detection tools in non-critical settings. For these systems the law focuses on transparency, ensuring users know they are interacting with AI and can make informed decisions.

Transparency obligations require organizations to disclose AI involvement clearly and understandably. A chatbot must tell the user it is not human. A generative text tool must indicate that its output was machine-produced. Spain's data-protection culture supports these rules, emphasizing informed consent and personal autonomy.

La IA de riesgo limitado incluye herramientas que interactúan directamente con personas, pero que no tienen la gravedad de los sistemas de alto riesgo. Ejemplos: chatbots, generadores de contenido, voces sintéticas y herramientas de detección de emociones en entornos no críticos. Para estos sistemas, la ley se centra en la transparencia, garantizando que los usuarios sepan que están interactuando con IA y puedan tomar decisiones informadas.

Las obligaciones de transparencia requieren que las organizaciones divulguen la participación de la IA de forma clara y comprensible. Por ejemplo, un chatbot debe informar al usuario que no es un humano. Una herramienta de texto generativo debe indicar que su salida fue producida por una máquina. La cultura de

protección de datos en España respalda estas reglas, enfatizando el consentimiento informado y la autonomía personal.

THINK ABOUT IT · PARA REFLEXIONAR

A customer-service center in Madrid deploys a voice chatbot for simple queries. A customer begins explaining a sensitive matter, believing they are speaking with a human. Now consider: should the system announce its AI nature immediately, or only if the user asks? How do the customer's expectations change once they know they are talking to AI? Reflect on whether failing to disclose AI involvement could affect trust, transparency or the user's comfort.

Un centro de atención al cliente en Madrid implementa un chatbot de voz para consultas simples. Un cliente comienza a explicar un asunto sensible, creyendo que habla con un humano. Ahora considera: ¿el sistema debe anunciar su naturaleza de IA inmediatamente o solo si el usuario lo solicita? ¿Cómo cambia la expectativa del cliente al saber que está interactuando con IA? Reflexiona sobre si no divulgar la participación de IA podría afectar la confianza, la transparencia o la comodidad del usuario.

Transparency is simple in concept but powerful in impact. It preserves fairness by making sure users know the nature of the system they are dealing with. It also prevents discomfort or confusion, especially for vulnerable people. In Spain, where digital literacy and consumer rights matter more every year, getting transparency right builds trust and reduces the risk of complaints or misuse.

Understanding limited-risk transparency rules strengthens your ability to design responsible interactions. It ensures people are informed participants rather than unwitting subjects of algorithmic influence, reinforcing consumer rights in the Spanish market.

La transparencia es simple en concepto pero poderosa en impacto. Preserva la equidad al asegurar que los usuarios conozcan la naturaleza del sistema con el que interactúan. También previene incomodidades o confusión, especialmente para personas vulnerables. En España, donde la alfabetización digital y los derechos del consumidor son cada vez más importantes, la transparencia correcta fomenta la confianza y reduce riesgos de quejas o mal uso.

Comprender las reglas de transparencia de IA de riesgo limitado fortalece tu capacidad de diseñar interacciones responsables. Asegura que las personas sean participantes informados y no sujetos inadvertidos de la influencia algorítmica, reforzando los derechos del consumidor en el mercado español. ## Topic 4: Minimal-Risk AI and Voluntary Best Practices + Practical Example / Tema 4: IA de Riesgo Mínimo y Buenas Prácticas Voluntarias + Ejemplo Práctico

2.4 Minimal-Risk AI and Voluntary Best Practices

IA de Riesgo Mínimo y Mejores Prácticas Voluntarias

Minimal-risk AI systems make up the majority of AI applications — video-game engines, spam filters, personalized recommendations. They are permitted without specific legal obligations because they pose little or no risk to fundamental rights or safety. Voluntary best practices are nevertheless

Los sistemas de IA de riesgo mínimo constituyen la mayoría de las aplicaciones de IA. Ejemplos: motores de videojuegos, filtros de spam y recomendaciones personalizadas. Están permitidos sin obligaciones legales específicas porque representan poco o ningún riesgo para los derechos fundamentales o la

recommended, particularly across Spain's public and private sectors.

Voluntary best practices include maintaining basic transparency, monitoring for unexpected outcomes and documenting updates. Even low-risk systems can create problems when they scale or interact with higher-risk systems if basic governance is ignored. Spain's digital strategy strongly encourages responsible innovation even where the law is permissive.

seguridad. Sin embargo, se recomienda adoptar buenas prácticas voluntarias, especialmente en sectores públicos y privados en España.

Las buenas prácticas voluntarias incluyen: mantener transparencia básica, monitorear resultados inesperados y documentar actualizaciones. Aunque estos sistemas son de bajo riesgo, ignorar la gobernanza básica puede generar problemas al escalar o interactuar con sistemas de mayor riesgo. La estrategia digital de España fomenta fuertemente la innovación responsable, incluso cuando la ley es permisiva.

APPLIED EXAMPLE · EJEMPLO APLICADO

Consider a US supermarket chain expanding into Spain, using a low-risk recommendation engine to suggest products to customers. No high-risk classification applies, but voluntary practices still matter. The team documents how the algorithm selects products, monitors customer feedback and checks that no discriminatory or misleading patterns emerge. In a future audit, this documentation demonstrates maturity and foresight — qualities increasingly valued in Spain.

Considera un supermercado estadounidense expandiéndose a España, usando un motor de recomendaciones de bajo riesgo para sugerir productos a clientes. No existe clasificación de alto riesgo, pero las prácticas voluntarias siguen siendo importantes. El equipo documenta cómo el algoritmo selecciona los productos, monitorea la retroalimentación de clientes y asegura que no aparezcan patrones discriminatorios o engañosos. Ante una futura auditoría, esta documentación demuestra madurez y previsión, cualidades cada vez más valoradas en España.

Minimal risk does not mean "no responsibility". It simply means organizations have more flexibility to manage systems in ways that strengthen trust and reduce long-term risk. Adopting voluntary governance also makes compliance easier if the system's functions evolve later.

Ultimately, understanding minimal-risk AI completes the risk spectrum. It helps organizations distinguish between what demands strict oversight and what benefits from good practice. With that clarity, businesses in Spain can innovate confidently without unnecessary regulatory friction.

Riesgo mínimo no significa "sin responsabilidad". Solo implica que las organizaciones tienen más flexibilidad para gestionar sistemas de manera que fortalezcan la confianza y reduzcan riesgos a largo plazo. Adoptar gobernanza voluntaria facilita el cumplimiento si las funciones del sistema evolucionan después.

En definitiva, comprender la IA de riesgo mínimo completa el espectro de riesgos. Ayuda a las organizaciones a distinguir entre lo que requiere supervisión estricta y lo que se beneficia de buenas prácticas. Con esta claridad, los negocios en España pueden innovar con confianza sin fricciones regulatorias innecesarias. ## Module Wrap-Up + Practical Life Implementation / Cierre del Módulo e Implementación Práctica

Module Wrap-Up & Practical Implementation

Cierre del Módulo e Implementación Práctica

You have now explored all four AI risk tiers — from fully prohibited systems to minimal-risk tools. This hierarchy forms the backbone of the EU AI Act and guides every compliance decision that follows. Knowing which tier a system belongs to lets you anticipate obligations instead of reacting to surprises.

Remember the reflection story and the assessment scenario. Both showed how easily organizations drift off course by misclassifying AI or neglecting transparency — and how the right actions quickly restore fairness, trust and compliance. The lesson is simple: classification is not just labeling, it is responsibility.

As you continue the course, take one immediate step: review an AI system your organization uses and determine its risk category. That small action can prevent major compliance problems. Remember — clarity is power, and you now have the knowledge to navigate AI risk classification with confidence.

Ahora has explorado los cuatro niveles de riesgo de la IA —desde sistemas completamente prohibidos hasta herramientas de riesgo mínimo. Esta jerarquía forma la columna vertebral de la Ley de IA de la UE y guía cada decisión de cumplimiento futura. Saber en qué nivel se encuentra un sistema te permite anticipar obligaciones en lugar de reaccionar a sorpresas.

Recuerda la historia de reflexión y el escenario de evaluación. Ambos mostraron lo fácil que es que las organizaciones se desvíen al mal clasificar la IA o ignorar la transparencia. Pero también demostraron cómo las acciones correctas restauran rápidamente equidad, confianza y cumplimiento. La lección es simple: la clasificación no es solo un etiquetado, es una responsabilidad.

A medida que avances en el curso, toma un paso inmediato: revisa un sistema de IA que tu organización utilice y determina su categoría de riesgo. Esta pequeña acción puede prevenir grandes problemas de cumplimiento. Recuerda: la claridad es poder, y ahora posees el conocimiento para navegar la clasificación de riesgos de IA con confianza.

MODULE 2 ASSESSMENT · 10 QUESTIONS

Test Your Understanding

Evalúa tu comprensión — Respuestas en el Apéndice A

1. Which AI practice is classified as "unacceptable risk" under the EU AI Act?

¿Qué práctica de IA se clasifica como "riesgo inaceptable" según el Reglamento de IA de la UE?

- A. AI-based chatbots / Chatbots basados en IA
- B. Emotion recognition in public spaces for law enforcement / Reconocimiento de emociones en espacios públicos para la aplicación de la ley
- C. Recommendation systems / Sistemas de recomendación
- D. AI-powered spell checkers / Correctores ortográficos impulsados por IA

2. High-risk AI systems are identified primarily on the basis of:

Los sistemas de IA de alto riesgo se identifican principalmente en función de:

- A. The size of the AI model / El tamaño del modelo de IA
- B. The country of development / El país de desarrollo
- C. Their intended use and potential impact on fundamental rights / Su uso previsto y el posible impacto en los derechos fundamentales
- D. Whether they use personal data / Si utilizan datos personales

3. Which sector is explicitly listed for high-risk AI use cases?

¿Qué sector está explícitamente listado para casos de uso de IA de alto riesgo?

- A.** Online advertising / Publicidad en línea
- B.** Health diagnostics / Diagnóstico de salud
- C.** Entertainment media / Medios de entretenimiento
- D.** Video games / Videojuegos

4. Transparency obligations apply mainly to which AI risk category?

Las obligaciones de transparencia se aplican principalmente a qué categoría de riesgo de IA:

- A.** Prohibited-risk AI / IA de riesgo prohibido
- B.** High-risk AI / IA de alto riesgo
- C.** Limited-risk AI / IA de riesgo limitado
- D.** Minimal-risk AI / IA de riesgo mínimo

5. Minimal-risk AI systems are:

Los sistemas de IA de riesgo mínimo son:

- A.** Completely prohibited / Completamente prohibidos
- B.** Subject to mandatory conformity assessments / Sujetos a evaluaciones de conformidad obligatorias
- C.** Regulated under criminal law / Regulados bajo la ley penal
- D.** Permitted with voluntary best practices / Permitidos con mejores prácticas voluntarias

6. Which AI application is typically classified as limited-risk AI?

¿Qué aplicación de IA típicamente se clasifica como IA de riesgo limitado?

- A.** Social scoring systems / Sistemas de puntuación social
- B.** Biometric identification / Identificación biométrica
- C.** Conversational AI systems / Sistemas de IA conversacional
- D.** Critical-infrastructure control systems / Sistemas de control de infraestructura crítica

7. All AI systems are subject to the same compliance requirements under the EU AI Act.

Todos los sistemas de IA están sujetos a los mismos requisitos de cumplimiento según el Reglamento de IA de la UE.

- A.** True / Verdadero
- B.** False / Falso
- C.** Only in Spain / Solo en España
- D.** Only for the public sector / Solo para el sector público

8. High-risk AI systems are automatically prohibited under the EU AI Act.

Los sistemas de IA de alto riesgo están automáticamente prohibidos según el Reglamento de IA de la UE.

- A.** True / Verdadero
- B.** False / Falso
- C.** Only biometric systems / Solo sistemas biométricos
- D.** Only for private companies / Solo para empresas privadas

9. AI systems posing unacceptable risk are _____ under the EU AI Act.

Los sistemas de IA que representan un riesgo inaceptable están _____ según el Reglamento de IA de la UE.

- A.** Regulated / Regulados
- B.** Certified / Certificados
- C.** Restricted / Restringidos
- D.** Prohibited / Prohibidos

10. Transparency duties apply mainly to _____-risk AI systems.

Los deberes de transparencia se aplican principalmente a los sistemas de IA de _____ riesgo.

- A.** Minimal / Mínimo
- B.** Limited / Limitado
- C.** High / Alto
- D.** Unacceptable / Inaceptable

3

MODULE THREE · MÓDULO 3

High-Risk AI System Requirements

Requisitos para Sistemas de IA de Alto Riesgo

3.1 Governance and Lifecycle Risk Management

Gobernanza y Gestión de Riesgos en el Ciclo de Vida

3.2 Data Quality, Training Practices, and Bias Control

Calidad de Datos, Prácticas de Entrenamiento y Control de Sesgos

3.3 Documentation, Conformity Assessment, and CE Marking

Documentación, Evaluación de Conformidad y Mercado CE

3.4 Human Oversight, Monitoring, and System Reliability

Supervisión Humana, Monitoreo y Confiabilidad del Sistema

Module 3 — High-Risk AI System Requirements

Módulo 3 — Requisitos para Sistemas de IA de Alto Riesgo

ENGLISH

WHY THIS MODULE · POR QUÉ ESTE MÓDULO

You would never guess how much damage a single undocumented change caused in one European healthcare system. A minor update, barely noticed by the technical team, produced inaccurate patient prioritizations that affected thousands. That incident became one of the triggers for stricter high-risk AI requirements across Europe — and it shows exactly why this module matters: high-risk AI is not risky because it is complex, but because it touches real people's lives.

High-risk systems demand the greatest level of discipline, transparency and control. These are the systems that influence hiring, loan approvals, educational paths, medical decisions, transport safety and public services throughout Spain. Today you will explore governance, data quality, documentation, oversight and reliability — the pillars that keep high-risk AI legal and trustworthy.

Master this, and everything else becomes more manageable. Let it strengthen your confidence: learning to control high-risk AI is one of the most valuable skills you can bring to any organization in Spain.

ESPAÑOL

Nunca imaginarías cuánto daño causó un solo cambio no documentado en un sistema de salud europeo. Una actualización menor, casi desapercibida por el equipo técnico, provocó priorizaciones inexactas de pacientes que afectaron a miles. Ese incidente se convirtió en uno de los detonantes para impulsar requisitos más estrictos para la IA de alto riesgo en Europa. Y demuestra exactamente por qué este módulo es importante: la IA de alto riesgo no es riesgosa por su complejidad, sino porque afecta la vida de personas reales.

Los sistemas de alto riesgo requieren el mayor nivel de disciplina, transparencia y control. Son los sistemas que influyen en contratación, aprobación de préstamos, trayectorias educativas, decisiones médicas, seguridad en transporte y servicios públicos en toda España. Hoy explorarás gobernanza, calidad de datos, documentación, supervisión y confiabilidad: los pilares que mantienen la IA de alto riesgo legal y confiable.

Si dominas esto, todo lo demás se vuelve más manejable. Refuerza tu confianza, porque aprender a controlar la IA de alto riesgo es una de las habilidades más valiosas que puedes aportar a cualquier organización en España. ## Feynman Before Section / Reflexión Inicial de Feynman

PAUSE & REFLECT — BEFORE YOU BEGIN · REFLEXIÓN INICIAL DE FEYNMAN

Before you begin, take a moment to reflect. Start with this question: what is the first thing that comes to mind when you hear "high-risk AI"? Fear, regulation, opportunity or responsibility? Write down whichever answer feels true right now.

Now consider two more questions. Why do you think the EU insists on such strict documentation and oversight for these systems? And which do you believe matters more for reducing risk: technical accuracy or human oversight? Note your instinctive thoughts and keep them close — we will revisit them later and see how your understanding evolves.

Antes de comenzar, tómate un momento para reflexionar. Comienza con esta pregunta: ¿qué es lo primero que se te viene a la mente al escuchar "IA de alto riesgo"? ¿Es miedo, regulación, oportunidad o responsabilidad? Escribe la respuesta que sientas como verdadera ahora mismo.

Ahora considera dos preguntas más. ¿Por qué crees que la UE insiste en documentación y supervisión tan estrictas para estos sistemas? ¿Y qué crees que es más importante para reducir riesgos: la precisión técnica o la supervisión humana? Apunta tus pensamientos instintivos. Manténlos cerca; los revisaremos más adelante y veremos cómo evoluciona tu comprensión. ## Topic 1: Governance and Lifecycle Risk Management / Tema 1: Gobernanza y Gestión de Riesgos en el Ciclo de Vida

3.1 Governance and Lifecycle Risk Management

Gobernanza y Gestión de Riesgos en el Ciclo de Vida

Governance is the backbone of high-risk AI. It ensures every stage — from design through deployment to monitoring — follows clear processes. In the EU, and especially in Spain, governance is not optional; it is the mechanism that aligns AI with legal obligations, organizational ethics and public expectations. Without strong governance, even well-designed systems can cause harm or drift toward misuse.

Lifecycle risk management means treating AI as a living system, not a static product. Before deployment, risks must be identified, assessed and mitigated. During deployment, systems require continuous monitoring. After deployment, organizations must be ready to update models, correct errors or even suspend AI use if necessary. Spanish regulators expect this ongoing diligence, particularly in sectors such as employment, health and public administration.

Effective governance requires cross-functional teams: technical experts, legal advisors, compliance officers, HR leaders and domain specialists. Collaboration matters because high-risk AI touches many areas at once. Working together, these teams create checks and balances no single department could achieve.

La gobernanza es la columna vertebral de la IA de alto riesgo. Asegura que cada etapa —desde el diseño hasta el despliegue y monitoreo— siga procesos claros. En la UE y especialmente en España, la gobernanza no es opcional; es el mecanismo que alinea la IA con obligaciones legales, ética organizacional y expectativas públicas. Sin una gobernanza sólida, incluso sistemas bien diseñados pueden causar daño o desviarse hacia un uso indebido.

La gestión de riesgos del ciclo de vida significa tratar la IA como un sistema vivo, no como un producto estático. Antes del despliegue, se deben identificar, evaluar y mitigar los riesgos. Durante el despliegue, los sistemas requieren monitoreo continuo. Después del despliegue, las organizaciones deben estar preparadas para actualizar modelos, corregir errores o incluso suspender el uso de IA si es necesario. Los reguladores españoles esperan esta diligencia continua, especialmente en sectores como empleo, salud y administración pública.

La gobernanza efectiva requiere equipos multifuncionales: expertos técnicos, asesores legales,

A solid governance model also includes escalation paths. When something fails — an anomaly, a bias alert, an unexplainable decision — there must be a clear route for reporting and resolving the problem. In Spain, embedding these paths into internal compliance systems and workplace protocols is recommended to ensure transparency with workers.

Ultimately, governance turns risk from something unpredictable into something manageable. With defined responsibilities, regular audits and structured oversight, organizations create environments where high-risk AI can operate safely and reliably.

oficiales de cumplimiento, líderes de RR.HH. y especialistas en dominio. La colaboración importa porque la IA de alto riesgo impacta múltiples áreas simultáneamente. Cuando estos equipos trabajan juntos, crean controles y equilibrios que un solo departamento no podría lograr.

Un modelo de gobernanza sólido también incluye rutas de escalamiento. Cuando algo falla —una anomalía, alerta de sesgo o decisión inexplicable— debe existir un camino claro para reportar y resolver el problema. En España, se recomienda integrar estas rutas en los sistemas internos de cumplimiento y protocolos laborales para garantizar transparencia con los trabajadores.

En última instancia, la gobernanza transforma el riesgo de algo impredecible a algo manejable. Con responsabilidades definidas, auditorías regulares y supervisión estructurada, las organizaciones crean entornos donde la IA de alto riesgo puede operar de manera segura y confiable. ## Topic 2: Data Quality, Training Practices, and Bias Control + Reflection Story / Tema 2: Calidad de Datos, Prácticas de Entrenamiento y Control de Sesgos + Historia de Reflexión

3.2 Data Quality, Training Practices, and Bias Control

Calidad de Datos, Prácticas de Entrenamiento y Control de Sesgos

High-risk AI is only as trustworthy as the data behind it. Poor-quality data produces poor-quality — sometimes harmful — decisions. The EU AI Act requires providers and deployers to ensure data is relevant, accurate, representative and free of discriminatory distortions. In Spain, where equality and labor rights are strongly protected, data governance becomes even more critical.

Training data must reflect the diversity of the population the AI system will serve. When datasets are skewed — missing key demographic groups or carrying historical bias — the system amplifies those flaws. Organizations must therefore document data sources, assess representativeness and apply corrective measures such as rebalancing or data augmentation.

Bias control is an ongoing process, not a one-time check. High-risk systems must be tested, validated and monitored continuously for unfair outcomes. Spanish regulators expect clear documentation

La IA de alto riesgo solo es confiable en la medida en que lo son los datos que la respaldan. Datos de baja calidad generan decisiones de baja calidad, a veces perjudiciales. La Ley de IA de la UE exige a proveedores y operadores garantizar que los datos sean relevantes, precisos, representativos y libres de distorsiones discriminatorias. En España, donde la igualdad y los derechos laborales están fuertemente protegidos, la gobernanza de datos se vuelve aún más crítica.

Los datos de entrenamiento deben reflejar la diversidad de la población a la que servirá el sistema de IA. Cuando los conjuntos de datos están sesgados —faltan grupos demográficos clave o contienen sesgos históricos— el sistema amplifica estas deficiencias. Por ello, las organizaciones deben documentar las fuentes de datos, evaluar su representatividad e implementar medidas correctivas, como reequilibrio o aumento de datos.

showing how systems were tested, why certain variables were included or excluded, and how fairness thresholds were chosen.

El control de sesgos es un proceso continuo, no un chequeo único. Los sistemas de alto riesgo deben ser probados, validados y monitoreados permanentemente para detectar resultados injustos. Los reguladores españoles esperan documentación clara que muestre cómo se probaron los sistemas, por qué se incluyeron o excluyeron ciertas variables y cómo se eligieron los umbrales de equidad.

CASE IN PRACTICE · CASO PRÁCTICO

A school internship-matching algorithm in Valencia offers a striking example. Designed to recommend students for placements, it leaned heavily on historical data — data that carried inherited bias: students from certain neighborhoods had historically received fewer opportunities. When the system launched, those patterns simply repeated themselves. Schools and parents complained. After an investigation, the team retrained the model on balanced data, removed harmful variables and added fairness constraints. Placement equity improved markedly and the system regained trust.

Un algoritmo de asignación de prácticas escolares en Valencia ofrece un ejemplo llamativo. Diseñado para recomendar estudiantes para pasantías, dependía fuertemente de datos históricos. Pero esos datos tenían sesgos heredados: estudiantes de ciertos barrios históricamente recibían menos oportunidades. Al lanzar el sistema, esos patrones se reprodujeron. Escuelas y padres se quejaron. Tras una investigación, el equipo reentrenó el modelo con datos equilibrados, eliminó variables dañinas y agregó restricciones de equidad. La equidad en la asignación mejoró notablemente y el sistema recuperó confianza.

Bias is not always intentional. It often creeps in silently through historical patterns. But with strong data-quality practices and continuous testing, organizations can transform biased systems into fair, reliable tools.

El sesgo no siempre es intencional. A menudo surge de forma silenciosa a través de patrones históricos. Pero con prácticas sólidas de calidad de datos y pruebas continuas, las organizaciones pueden transformar sistemas sesgados en herramientas justas y confiables. ## Topic 3: Documentation, Conformity Assessment, and CE Marking + Scenario-Based Mini Assessment / Tema 3: Documentación, Evaluación de Conformidad y Marcado CE + Mini Evaluación Basada en Escenarios

3.3 Documentation, Conformity Assessment, and CE Marking

Documentación, Evaluación de Conformidad y Marcado CE

Documentation is one of the most demanding parts of high-risk AI — and one of the most powerful. It forces organizations to translate their process into a clear narrative: how the system was built, how it works, what data it uses, what risks exist and how they are mitigated. Regulators in Spain rely heavily on documentation during inspections, so accuracy is essential.

La documentación es una de las partes más exigentes de la IA de alto riesgo, pero también una de las más poderosas. Obliga a las organizaciones a traducir su proceso en una narrativa clara: cómo se construyó el sistema, cómo funciona, qué datos usa, qué riesgos existen y cómo se mitigan. Los reguladores en España dependen en gran medida de la documentación durante las inspecciones, por lo que la precisión es esencial.

Conformity assessments verify that high-risk AI meets every legal requirement before entering the market. Providers must perform internal audits or work with external notified bodies where required. These assessments examine everything from data quality to oversight instructions. Once approved, the system may carry the CE marking, signaling compliance across the EU.

The CE marking is not a decorative label. It represents legal responsibility, requiring organizations to maintain post-market monitoring, incident reporting and transparency. If the system fails to perform safely, the CE marking can be withdrawn — halting deployment immediately. For public-sector systems in Spain, this process is especially critical given public-accountability standards.

Documentation also strengthens internal understanding. Teams that document properly gain clarity about assumptions, design decisions and model limitations. That transparency becomes essential when updating models or responding to audits, incidents or workers' inquiries.

Las evaluaciones de conformidad aseguran que la IA de alto riesgo cumpla con todos los requisitos legales antes de entrar al mercado. Los proveedores deben realizar auditorías internas o trabajar con organismos notificados externos cuando sea necesario. Estas evaluaciones examinan todo, desde la calidad de los datos hasta las instrucciones de supervisión. Una vez aprobado, el sistema puede recibir el marcado CE, que indica cumplimiento en toda la UE.

El marcado CE no es una etiqueta decorativa. Representa responsabilidad legal, requiriendo que las organizaciones mantengan monitoreo post-mercado, reporte de incidentes y transparencia. Si el sistema no funciona de manera segura, el marcado CE puede retirarse, deteniendo inmediatamente el despliegue. Para sistemas del sector público en España, este proceso es especialmente crítico debido a los estándares de responsabilidad pública.

La documentación también fortalece la comprensión interna. Cuando los equipos documentan correctamente, obtienen claridad sobre suposiciones, decisiones de diseño y limitaciones del modelo. Esta transparencia se vuelve esencial al actualizar modelos o responder a auditorías, incidentes o consultas de los trabajadores.

DECISION POINT • PUNTO DE DECISIÓN

A regional bank in Spain wants to deploy a high-risk AI credit-assessment system. The provider claims the system passed conformity assessments in another EU country. What should the deployer do? A) Trust the existing certification and deploy immediately. B) Request full documentation, review the conformity evidence and verify the system meets Spain's contextual requirements. C) Skip the documentation because oversight only applies to providers.

Un banco regional en España desea desplegar un sistema de IA de evaluación crediticia de alto riesgo. El proveedor afirma que el sistema pasó evaluaciones de conformidad en otro país de la UE. ¿Qué debería hacer el operador? A) Confiar en la certificación existente y desplegar de inmediato. B) Solicitar documentación completa, revisar la evidencia de conformidad y verificar que el sistema cumpla con los requisitos contextuales de España. C) Omitir la documentación porque la supervisión solo aplica a proveedores.

Correct answer: B — Deployers must verify documentation and ensure the system suits the Spanish legal and demographic context.

Respuesta correcta: B — los operadores deben verificar la documentación y asegurar que el sistema sea adecuado para el contexto legal y demográfico español.

3.4 Human Oversight, Monitoring, and System Reliability

Supervisión Humana, Monitoreo y Confiabilidad del Sistema

Human oversight is one of the most important requirements for high-risk AI. The EU insists that humans retain meaningful control over decisions — not as passive observers, but as active supervisors able to intervene, correct errors or reject AI recommendations. This keeps accountability with people, never with algorithms.

Oversight must be designed into the system itself. Providers must produce clear instructions explaining when human intervention is required and what actions supervisors may take. Deployers in Spain must ensure staff are trained to understand the system's limits and to recognize anomalies or unsafe behavior.

Monitoring keeps high-risk systems aligned with expectations over time. Performance must be tracked, errors logged and unexpected patterns investigated. Reliability is never static: AI models change, environments evolve and new risks emerge. Continuous monitoring lets organizations catch problems early.

La supervisión humana es uno de los requisitos más importantes para la IA de alto riesgo. La UE insiste en que los humanos mantengan control significativo sobre las decisiones, no como observadores pasivos, sino como supervisores activos capaces de intervenir, corregir errores o rechazar recomendaciones de IA. Esto garantiza que la responsabilidad siempre recaiga en personas, no en algoritmos.

La supervisión debe estar diseñada dentro del propio sistema. Los proveedores deben crear instrucciones claras que expliquen cuándo se requiere intervención humana y qué acciones pueden tomar los supervisores. Los operadores en España deben asegurar que el personal reciba capacitación para entender los límites del sistema y poder identificar anomalías o comportamientos inseguros.

El monitoreo mantiene los sistemas de alto riesgo alineados con las expectativas a lo largo del tiempo. Se debe rastrear el desempeño, registrar errores e investigar patrones inesperados. La confiabilidad no es estática: los modelos de IA cambian, los entornos evolucionan y surgen nuevos riesgos. El monitoreo continuo permite a las organizaciones detectar problemas tempranamente.

APPLIED EXAMPLE · EJEMPLO APLICADO

Consider a US medical-technology company introducing an AI diagnostic assistant into Spanish hospitals. However accurate the system, Spanish deployers must ensure clinicians remain the decision-makers. Oversight protocols require doctors to review AI suggestions, confirm diagnoses and log discrepancies. Continuous monitoring reveals shifts in model accuracy — perhaps caused by demographic variation. This combination of oversight and monitoring keeps the tool trustworthy and safe within the Spanish healthcare system.

Considera una empresa estadounidense de tecnología médica que introduce un asistente diagnóstico de IA en hospitales españoles. Aunque el sistema es preciso, los operadores españoles deben asegurarse de que los clínicos sigan siendo los tomadores de decisiones. Los protocolos de supervisión requieren que los médicos revisen las sugerencias de la IA, confirmen diagnósticos y registren discrepancias. El monitoreo continuo revela cambios en la precisión del modelo, tal vez debido a variaciones demográficas. Esta combinación de supervisión y monitoreo garantiza que la herramienta siga siendo confiable y segura en el sistema de salud español.

Human oversight and monitoring are not burdens; they are safeguards protecting the integrity of

La supervisión humana y el monitoreo no son cargas, sino salvaguardas que protegen la integridad de las

decisions. When organizations embrace them, systems become more reliable, more stable and better aligned with Spain's ethical expectations.

decisiones. Cuando las organizaciones los adoptan, los sistemas se vuelven más confiables, estables y alineados con las expectativas éticas de España. ## Feynman After Section / Reflexión Final de Feynman

PAUSE & REFLECT — LOOKING BACK · REFLEXIÓN FINAL DE FEYNMAN

Return to your first reflections. When you first thought about high-risk AI, your mind may have focused on technological complexity. Now you can see the real challenge lies in governance, data quality, oversight and accountability. Those elements — not accuracy alone — determine whether an AI system is safe, lawful and fair.

You also weighed whether documentation or human oversight matters more. Now you know they work together: documentation creates clarity, and oversight ensures ongoing accountability. Your first instincts should now feel more solid and better grounded — that is what deepening understanding feels like.

Volvamos a tus reflexiones iniciales. Cuando pensaste por primera vez en IA de alto riesgo, tu mente tal vez se centró en la complejidad tecnológica. Ahora puedes ver que el verdadero desafío es la gobernanza, calidad de datos, supervisión y responsabilidad. Estos elementos, no solo la precisión, determinan si un sistema de IA es seguro, legal y justo.

También consideraste si la documentación o la supervisión humana importan más. Ahora sabes que funcionan en conjunto: la documentación genera claridad y la supervisión asegura responsabilidad continua. Tus pensamientos iniciales ahora se sienten más sólidos y fundamentados: así es como se profundiza la comprensión. ## Module Wrap-Up + Practical Life Implementation / Cierre del Módulo + Implementación Práctica

Module Wrap-Up & Practical Implementation

Cierre del Módulo e Implementación Práctica

You have explored the essential pillars of high-risk AI: governance, data practices, documentation, conformity, oversight and reliability. These requirements form the core of EU compliance and determine how AI enters workplaces, hospitals, banks and public institutions across Spain. Mastering them gives you the tools to evaluate any high-risk system with confidence.

Remember the reflection story and the scenario assessment. Both showed how organizations run into trouble not through intent to harm, but through incomplete understanding of their responsibilities — and how transparency, analysis and oversight quickly restore fairness and trust.

Now pick a high-risk or potentially high-risk AI system in your organization and map its lifecycle — from data to oversight. Identify one improvement you can implement this month. You now have the knowledge

Has explorado los pilares esenciales de la IA de alto riesgo: gobernanza, prácticas de datos, documentación, conformidad, supervisión y confiabilidad. Estos requisitos forman el núcleo del cumplimiento en la UE y determinan cómo la IA ingresa a lugares de trabajo, hospitales, bancos e instituciones públicas en España. Dominarlos te brinda las herramientas para evaluar cualquier sistema de alto riesgo con confianza.

Recuerda la historia de reflexión y la evaluación basada en escenarios. Ambos mostraron cómo las organizaciones pueden tener problemas no por intención de daño, sino por no comprender completamente sus responsabilidades. Pero también demostraron cómo la transparencia, el análisis y la supervisión pueden restaurar rápidamente la equidad y la confianza.

to lead high-risk AI responsibly and reliably in Spain's evolving digital landscape.

Selecciona un sistema de IA de alto riesgo o potencialmente de alto riesgo en tu organización y mapea su ciclo de vida —desde los datos hasta la supervisión—. Identifica una mejora que puedas implementar este mes. Ahora tienes el conocimiento para liderar la IA de alto riesgo de manera responsable y confiable en el paisaje digital en evolución de España.

MODULE 3 ASSESSMENT · 10 QUESTIONS

Test Your Understanding

Evalúa tu comprensión — Respuestas en el Apéndice A

1. Lifecycle risk management for high-risk AI focuses on:

La gestión de riesgos a lo largo del ciclo de vida de la IA de alto riesgo se centra en:

- A. Marketing performance / Desempeño de marketing
- B. Continuous risk identification and mitigation / Identificación y mitigación continua de riesgos
- C. Increasing processing speed / Incrementar la velocidad de procesamiento
- D. Reducing system costs / Reducir costos del sistema

2. Data used to train high-risk AI systems must be:

Los datos utilizados para entrenar sistemas de IA de alto riesgo deben ser:

- A. Exclusively proprietary / Exclusivamente propietarios
- B. Randomly collected / Recogidos aleatoriamente
- C. Relevant, representative and free of bias / Relevantes, representativos y libres de sesgos
- D. Publicly accessible / De acceso público

3. Which document demonstrates compliance before placing a high-risk AI system on the EU market?

¿Qué documento demuestra el cumplimiento antes de colocar un sistema de IA de alto riesgo en el mercado de la UE?

- A. Privacy notice / Aviso de privacidad
- B. Ethics charter / Carta ética
- C. Risk statement / Declaración de riesgos
- D. Technical documentation / Documentación técnica

4. CE marking on AI systems indicates:

La marca CE para sistemas de IA indica:

- A. Ethical approval / Aprobación ética
- B. Market popularity / Popularidad en el mercado
- C. Compliance with EU requirements / Cumplimiento con los requisitos de la UE
- D. Open-source licensing / Licencia de código abierto

5. Human-oversight mechanisms are required in order to:

Los mecanismos de supervisión humana son necesarios para:

- A. Replace automated decision-making / Reemplazar la toma de decisiones automatizada
- B. Eliminate AI errors entirely / Eliminar completamente los errores de IA
- C. Enable intervention and override of AI outputs / Permitir la intervención y anulación de los resultados de la IA
- D. Reduce development time / Reducir el tiempo de desarrollo

6. Post-market monitoring ensures:

La monitorización posterior a la comercialización asegura:

- A. Marketing efficiency / Eficiencia en marketing
- B. Ongoing compliance and risk control / Cumplimiento continuo y control de riesgos
- C. Data monetization / Monetización de datos
- D. Automated model retraining / Automatización del reentrenamiento del modelo

7. High-risk AI systems require conformity assessments before deployment.

Los sistemas de IA de alto riesgo requieren evaluaciones de conformidad antes de su despliegue.

- A. True / Verdadero
- B. False / Falso
- C. Only for public authorities / Solo para autoridades públicas
- D. Only for importers / Solo para importadores

8. Bias-control measures are optional for high-risk AI systems.

Las medidas de control de sesgos son opcionales para los sistemas de IA de alto riesgo.

- A. True / Verdadero
- B. False / Falso
- C. Only for biometric AI / Solo para IA biométrica
- D. Only under the GDPR / Solo bajo el GDPR

9. High-risk AI systems must maintain detailed _____ throughout their lifecycle.

Los sistemas de IA de alto riesgo deben mantener detallada _____ a lo largo de su ciclo de vida.

- A.** Marketing plans / Planes de marketing
- B.** User manuals / Manuales de usuario
- C.** Technical documentation / Documentación técnica
- D.** Financial records / Registros financieros

10. Human oversight enables _____ intervention in AI decision-making.

La supervisión humana permite la intervención _____ en la toma de decisiones de la IA.

- A.** Automatic / Automática
- B.** Manual / Manual
- C.** Random / Aleatoria
- D.** Deferred / Diferida

4

MODULE FOUR · MÓDULO 4

Governance of General-Purpose and Generative AI

Gobernanza de IA de Propósito General y Generativa

4.1 General-Purpose AI Models and Regulatory Scope

Modelos de IA de Propósito General y Alcance Regulatorio

4.2 Systemic GPAI Obligations and Testing Requirements

Obligaciones Sistémicas de GPAI y Requisitos de Pruebas

4.3 Transparency, Watermarking, and Synthetic Content Rules

Transparencia, Marcas de Agua y Reglas de Contenido Sintético

4.4 Responsibilities for Fine-Tuning and Downstream Deployment

Responsabilidades en Ajuste Fino y Despliegue Posterior

Module 4 — Governance of General-Purpose and Generative AI

Módulo 4 — Gobernanza de IA de Propósito General y Generativa

ENGLISH

ESPAÑOL

WHY THIS MODULE · POR QUÉ ESTE MÓDULO

Most of today's most influential AI systems were not designed for a single task. They were built to learn almost anything: writing, predicting, reasoning, coding, generating images, analyzing data. These general-purpose AI models — GPAI — now sit at the center of Europe's regulatory attention. Why? Because when a model becomes the foundation for thousands of applications across Spain and the EU, a single failure can ripple through entire sectors.

GPAI governance matters because these systems amplify the impact of everything connected to them. A misalignment in one foundation model can instantly affect users, hospitals, businesses and governments. Today you will explore how the EU regulates general-purpose and generative AI: its scope, systemic obligations, transparency rules and the responsibilities that arise when organizations fine-tune or deploy these models in real environments.

As generative AI transforms how Spain works, learns and creates, understanding its governance is more than useful — it positions you as someone able to guide your organization through the most transformative technological shift of our era.

La mayoría de los sistemas de IA más influyentes no fueron diseñados para una sola tarea. Fueron contruidos para aprender casi cualquier cosa: escribir, predecir, razonar, codificar, generar imágenes o analizar datos. Estos modelos de IA de propósito general, o GPAI, ahora están en el centro de la atención regulatoria en Europa. ¿Por qué? Porque cuando un modelo se convierte en la base de miles de aplicaciones en España y la UE, un solo fallo puede repercutir en sectores completos.

La gobernanza de GPAI importa porque estos sistemas amplifican el impacto de todo lo conectado a ellos. Un desalineamiento en un modelo fundamental puede afectar instantáneamente a usuarios, hospitales, empresas y gobiernos. Hoy explorarás cómo la UE regula la IA de propósito general y generativa: su alcance, obligaciones sistémicas, reglas de transparencia y responsabilidades que surgen cuando las organizaciones ajustan o implementan estos modelos en entornos reales.

A medida que la IA generativa transforma cómo España trabaja, aprende y crea, comprender la gobernanza no solo es útil, sino que te posiciona como alguien capaz de guiar a tu organización en el cambio tecnológico más transformador de nuestra era. ## Topic 1: General-Purpose AI Models and Regulatory Scope / Tema 1: Modelos de IA de Propósito General y Alcance Regulatorio

4.1 General-Purpose AI Models and Regulatory Scope

Modelos de IA de Propósito General y Alcance Regulatorio

General-purpose AI, or GPAI, covers systems designed to perform many tasks rather than one specialized function. These models can summarize documents, classify content, generate text, create images, analyze patterns and power countless downstream applications. Because their reach is

La IA de propósito general, o GPAI, abarca sistemas diseñados para realizar múltiples tareas en lugar de una función especializada. Estos modelos pueden resumir documentos, clasificar contenido, generar texto, crear imágenes, analizar patrones y potenciar innumerables aplicaciones posteriores. Debido a su

broad and not tied to one use case, they pose unique regulatory challenges.

The EU AI Act identifies GPAI as models that can be integrated into many other systems — from customer-service chatbots to decision-support tools in healthcare or public administration. Generative AI falls within this framework because of its multi-purpose potential: creating synthetic images, drafting documents, assisting with code.

GPAI regulation differs from high-risk regulation. Instead of focusing solely on the end-use context, the EU also assesses the model's inherent scale, its capabilities, its training methods and its possible societal impact. This dual perspective allows rules for both providers and deployers, securing safety across the whole ecosystem.

Spain's digital strategy aligns closely with the EU framework. National authorities recognize that GPAI models influence education, labor markets, public services and creative industries. Misuse, opacity or bias in one foundation model can affect millions, which makes governance an essential safeguard.

Understanding scope is the first step. GPAI governance recognizes that responsibility must exist at every level: from the creators who build the models, to the organizations that adapt them, to the teams that deploy them in real contexts.

alcance amplio y no limitado a un solo caso de uso, presentan desafíos regulatorios únicos.

La Ley de IA de la UE identifica a GPAI como modelos que pueden integrarse en muchos otros sistemas, desde chatbots en atención al cliente hasta herramientas de apoyo a decisiones en salud o administración pública. La IA generativa entra dentro de este marco debido a su potencial de uso múltiple: creación de imágenes sintéticas, redacción de documentos o asistencia en programación.

La regulación de GPAI difiere de la IA de alto riesgo. En lugar de centrarse únicamente en el contexto de uso final, la UE también evalúa la escala inherente del modelo, sus capacidades, métodos de entrenamiento y el posible impacto social. Esta perspectiva dual permite establecer normas tanto para proveedores como para desplegados, garantizando seguridad en todo el ecosistema.

La estrategia digital de España se alinea estrechamente con el marco de la UE. Las autoridades nacionales reconocen que los modelos GPAI influyen en educación, mercados laborales, servicios públicos e industrias creativas. El uso indebido, la opacidad o los sesgos en un modelo fundamental pueden afectar a millones, por lo que la gobernanza es una salvaguarda esencial.

Comprender el alcance es el primer paso. La gobernanza de GPAI reconoce que la responsabilidad debe existir en todos los niveles: desde los creadores que construyen los modelos, hasta las organizaciones que los adaptan y los equipos que los implementan en contextos reales. ## Topic 2: Systemic GPAI Obligations and Testing Requirements + Reflection Story / Tema 2: Obligaciones Sistémicas de GPAI y Requisitos de Pruebas + Historia de Reflexión

4.2 Systemic GPAI Obligations and Testing Requirements

Obligaciones Sistémicas de GPAI y Requisitos de Pruebas

Some GPAI systems are deemed "systemic models" — models whose capabilities or scale give them widespread influence. They carry special obligations because their behavior can propagate across applications, sectors and borders. Providers must run rigorous testing, assess systemic risks and maintain extensive documentation of model behavior, limitations and potential harms.

Algunos sistemas GPAI se consideran "modelos sistémicos", lo que significa que poseen capacidades o escala que permiten una influencia generalizada. Estos modelos tienen obligaciones especiales porque su comportamiento puede propagarse a través de aplicaciones, sectores y fronteras. Los proveedores deben realizar pruebas rigurosas, evaluar riesgos sistémicos y mantener documentación extensa sobre

Testing requirements include evaluating robustness, cybersecurity, bias, misuse risks and error patterns. Providers must anticipate not only how the system works, but how it might fail. Spain places particular emphasis on these evaluations because GPAI often enters public-sector workflows, educational settings and workplace decision-making.

Systemic GPAI obligations also demand continuous monitoring of real-world performance. These systems must be watched constantly for emerging risks — hallucinations, misinformation generation, unfair outputs, or unexpected behavior that could affect individuals or communities.

el comportamiento del modelo, sus limitaciones y posibles daños.

Los requisitos de prueba incluyen evaluar robustez, ciberseguridad, sesgos, riesgos de uso indebido y patrones de error. Los proveedores deben anticipar no solo cómo funciona el sistema, sino también cómo podría fallar. España pone especial énfasis en estas evaluaciones porque GPAI a menudo entra en flujos de trabajo del sector público, entornos educativos y procesos de toma de decisiones laborales.

Las obligaciones sistémicas de GPAI también exigen monitoreo continuo del desempeño en el mundo real. Estos sistemas deben observarse constantemente para detectar riesgos emergentes, incluyendo alucinaciones, generación de desinformación, resultados injustos o comportamientos inesperados que puedan impactar a individuos o comunidades.

CASE IN PRACTICE · CASO PRÁCTICO

A public education platform in Madrid integrated a generative AI model to help students practice their writing. At first the tool seemed harmless. After a few months, however, teachers noticed the generated essays contained subtle cultural biases and inaccuracies about Spanish history. Parents raised concerns about misinformation. On investigation, the platform found the GPAI provider had not thoroughly tested the model's cultural knowledge. Working with the provider, they ran targeted evaluations, added safeguards and retrained components with Spanish educational material. The platform regained trust — not by abandoning AI, but by strengthening governance and testing.

Ejemplo: una plataforma educativa pública en Madrid integró un modelo de IA generativa para ayudar a los estudiantes a practicar la escritura. Inicialmente, la herramienta parecía inofensiva. Pero tras unos meses, los docentes notaron que los ensayos generados presentaban sesgos culturales sutiles e imprecisiones históricas españolas. Los padres manifestaron preocupación por la desinformación. Al investigar, la plataforma descubrió que el proveedor de GPAI no había probado exhaustivamente el conocimiento cultural del modelo. Tras colaborar con el proveedor, realizaron evaluaciones específicas, añadieron salvaguardas y reentrenaron componentes con material educativo español. La plataforma recuperó la confianza, no abandonando la IA, sino fortaleciendo la gobernanza y las pruebas.

When a model influences thousands of users, governance failures amplify fast. Testing and monitoring are not optional extras — they are what keeps reliability and social trust intact.

Quando un modelo influye en miles de usuarios, los errores de gobernanza se amplifican rápidamente. Las pruebas y el monitoreo no son opcionales: son necesarios para mantener la fiabilidad y la confianza social. ## Topic 3: Transparency, Watermarking, and Synthetic Content Rules + REGULAR MINI ASSESSMENT / Tema 3: Transparencia, Marcas de Agua y Reglas de Contenido Sintético + Mini Evaluación Regular

4.3 Transparency, Watermarking, and Synthetic Content Rules

Transparencia, Marcas de Agua y Reglas de Contenido Sintético

Transparency for GPAI systems focuses primarily on helping users understand when outputs are synthetic and when they were made by humans. Watermarking, labels and disclosure requirements support that clarity. These rules prevent confusion, reduce misinformation and protect public trust — particularly in journalism, education, political communication and online platforms.

Generative AI models must include mechanisms that allow deployers to identify synthetic content. Watermarks may be invisible to the user but detectable by platforms and regulators. Labels must be clear enough for users to understand they are engaging with AI-generated content. Spain places special emphasis on this because of concerns about deepfakes, electoral misinformation and deceptive content.

These transparency rules reinforce autonomy. Users deserve to know when an image, voice or document was produced by AI rather than a human. Transparency also helps organizations preserve credibility by preventing misunderstandings about AI involvement.

La transparencia en sistemas GPAI se centra principalmente en ayudar a los usuarios a entender cuándo los resultados generados son sintéticos y cuándo fueron creados por humanos. Las marcas de agua, etiquetas y requisitos de divulgación apoyan esta claridad. Estas reglas previenen confusión, reducen la desinformación y protegen la confianza pública, especialmente en periodismo, educación, comunicación política y plataformas en línea.

Los modelos de IA generativa deben incluir mecanismos que permitan a los desplegados identificar contenido sintético. Las marcas de agua pueden ser invisibles para el usuario, pero detectables por plataformas y reguladores. Las etiquetas deben ser claras para que los usuarios comprendan que están interactuando con contenido generado por IA. España pone especial énfasis en esto debido a preocupaciones sobre deepfakes, desinformación electoral y contenido engañoso.

Estas reglas de transparencia refuerzan la autonomía. Los usuarios merecen saber cuándo una imagen, voz o documento fue producido por IA y no por un humano. La transparencia también ayuda a las organizaciones a mantener credibilidad evitando malentendidos sobre la participación de la IA.

THINK ABOUT IT · PARA REFLEXIONAR

A cultural museum in Barcelona uses generative AI to create synthetic historical recreations for visitors. The visuals are realistic, and many assume the material is archival. Ask yourself: could failing to label the content as synthetic damage trust? Could visitors misread history or form false beliefs? What responsibilities does the museum carry to disclose AI involvement? Reflect on how transparency protects accuracy, accountability and cultural integrity.

Ejemplo: un museo cultural en Barcelona utiliza IA generativa para crear recreaciones históricas sintéticas para los visitantes. El contenido visual es realista y muchos asumen que el material es archivístico. Pregúntate: ¿no etiquetar el contenido como sintético podría afectar la confianza? ¿Podrían los visitantes interpretar mal la historia o formar creencias falsas? ¿Qué responsabilidades tiene el museo para divulgar la participación de IA? Reflexiona sobre cómo la transparencia protege la precisión, la responsabilidad y la integridad cultural.

Transparency rules are not merely regulatory requirements — they are ethical commitments. They ensure AI-generated content enhances communication rather than distorting it. When

Las reglas de transparencia no son solo requisitos regulatorios: son compromisos éticos. Aseguran que el contenido generado por IA mejore la comunicación y no la distorsione. Cuando las organizaciones aplican

organizations apply these rules consistently, users feel informed and respected.

estas reglas de manera consistente, los usuarios se sienten informados y respetados. ## Topic 4: Responsibilities for Fine-Tuning and Downstream Deployment + Practical Example / Tema 4: Responsabilidades en Ajuste Fino y Despliegue Posterior + Ejemplo Práctico

4.4 Responsibilities for Fine-Tuning and Downstream Deployment

Responsabilidades en Ajuste Fino y Despliegue Posterior

Fine-tuning turns a general-purpose model into something specialized. But it also introduces new responsibilities. When an organization adapts a GPAI model, it becomes partly accountable for the resulting behavior. The EU AI Act treats fine-tuning as a modification that can change regulatory obligations — especially if the adapted model becomes high-risk or systemic.

Downstream deployers — organizations that use GPAI models in real workflows — must make sure outputs are used responsibly. They must establish oversight processes, document how the model is applied and communicate its limitations to users. Spain's regulatory culture reinforces these responsibilities, particularly in sectors serving vulnerable populations or making sensitive decisions.

Organizations must also assess whether fine-tuning introduces new biases, safety issues or misuse risks. Deploying a model without an impact analysis exposes the organization to regulatory scrutiny and potential harm to users.

El ajuste fino transforma un modelo de propósito general en algo especializado. Pero también introduce nuevas responsabilidades. Cuando una organización adapta un modelo GPAI, se vuelve parcialmente responsable del comportamiento resultante. La Ley de IA de la UE considera el ajuste fino como una modificación que puede cambiar las obligaciones regulatorias, especialmente si el modelo adaptado se vuelve de alto riesgo o sistémico.

Los desplegados posteriores, organizaciones que usan modelos GPAI en flujos de trabajo reales, deben garantizar que los resultados se usen de manera responsable. Deben establecer procesos de supervisión, documentar cómo se aplica el modelo y comunicar las limitaciones a los usuarios. La cultura regulatoria española refuerza estas responsabilidades, especialmente en sectores con poblaciones vulnerables o decisiones sensibles.

Las organizaciones también deben evaluar si el ajuste fino introduce nuevos sesgos, problemas de seguridad o riesgos de uso indebido. Desplegar un modelo sin un análisis de impacto expone a la organización a escrutinio regulatorio y posibles daños a los usuarios.

APPLIED EXAMPLE · EJEMPLO APLICADO

A US health-tech startup expands into Spain. They fine-tune a generative model to draft medical admission notes. Although the base model performed well in English, the tuned version struggled with Spanish medical terminology and misspelled drug names. Staff noticed inconsistencies that created operational friction. By running additional training cycles on Spanish datasets and establishing human review, the startup earned regulatory approval and dramatically improved output quality. The example shows how fine-tuning and deployment must adapt to local expectations, data and oversight norms.

Ejemplo: una startup estadounidense de salud se expande a España. Ajustan finamente un modelo generativo para redactar notas de ingreso médico. Aunque el modelo base funcionaba bien en inglés, la versión ajustada tuvo dificultades con terminología médica española y nombres de medicamentos mal escritos. El personal notó inconsistencias que generaron tensión operativa. Implementando ciclos adicionales de entrenamiento con datasets españoles y estableciendo una revisión humana, la startup obtuvo aprobación regulatoria y mejoró significativamente la calidad de resultados. Este ejemplo muestra cómo ajuste fino y despliegue deben adaptarse a expectativas locales, datos y normas de supervisión.

Fine-tuning and deployment responsibilities ensure AI stays safe not just at launch, but throughout its lifecycle. Organizations that honor these obligations protect themselves, their users and their reputation.

Las responsabilidades de ajuste fino y despliegue aseguran que la IA sea segura no solo al inicio, sino durante todo su ciclo de vida. Las organizaciones que cumplen estas obligaciones protegen a sí mismas, a sus usuarios y su reputación. ## Module Wrap-Up + Practical Life Implementation / Cierre del Módulo + Implementación Práctica

Module Wrap-Up & Practical Implementation

Cierre del Módulo e Implementación Práctica

You have explored the governance of general-purpose and generative AI — from regulatory scope through systemic obligations, transparency rules and downstream deployment responsibilities. These requirements ensure GPAI technologies contribute positively to Spain's digital transformation without violating rights or spreading misinformation.

Recall the reflection story and the scenario. Both illustrated how generative AI can shape learning, culture and trust — and how quickly problems surface when transparency or testing is neglected. They also showed how thoughtful, consistent governance practices restore confidence.

Identify one GPAI or generative tool your organization uses. Ask: are the transparency measures clear? Is oversight defined? Have downstream impacts been analyzed? Applying these questions strengthens your

Has explorado la gobernanza de la IA de propósito general y generativa: desde el alcance regulatorio hasta obligaciones sistémicas, reglas de transparencia y responsabilidades en despliegue posterior. Estos requisitos aseguran que las tecnologías GPAI contribuyan positivamente a la transformación digital de España sin vulnerar derechos ni difundir desinformación.

Recuerda la historia de reflexión y el escenario. Ambos ilustraron cómo la IA generativa puede influir en el aprendizaje, la cultura y la confianza, y cuán rápido surgen problemas cuando se descuidan la transparencia o las pruebas. También mostraron cómo las prácticas de gobernanza pueden restaurar la confianza si se aplican de manera reflexiva y consistente.

leadership in responsible AI adoption in a fast-moving landscape.

Identifica una herramienta GPAI o generativa que tu organización utilice. Pregúntate: ¿las medidas de transparencia son claras? ¿Está definida la supervisión? ¿Se han analizado los impactos posteriores? Aplicar estas preguntas fortalece tu liderazgo en la adopción responsable de IA en un panorama que evoluciona rápidamente.

MODULE 4 ASSESSMENT · 10 QUESTIONS

Test Your Understanding

Evalúa tu comprensión — Respuestas en el Apéndice A

1. General-purpose AI models are designed to:

Los modelos de IA de propósito general están diseñados para:

- A.** Perform one specific task / Realizar una tarea específica
- B.** Operate only in regulated sectors / Operar únicamente en sectores regulados
- C.** Be adapted to multiple downstream tasks / Ser adaptados a múltiples tareas derivadas
- D.** Replace human intelligence / Reemplazar la inteligencia humana

2. Systemic general-purpose AI (GPAI) refers to models that:

La IA de propósito general sistémica (GPAI) se refiere a modelos que:

- A.** Are open source / Son de código abierto
- B.** Have limited computing power / Tienen potencia computacional limitada
- C.** Pose significant systemic risks / Plantean riesgos sistémicos significativos
- D.** Are used only by governments / Son utilizados únicamente por gobiernos

3. Which obligation applies specifically to systemic GPAI models?

¿Qué obligación se aplica específicamente a los modelos GPAI sistémicos?

- A.** Product labeling / Etiquetado del producto
- B.** Advanced testing and risk evaluation / Pruebas y evaluación avanzadas de riesgos
- C.** Mandatory public code disclosure / Divulgación pública obligatoria del código
- D.** National registration only / Registro nacional únicamente

4. Watermarking requirements aim to:

Los requisitos de marca de agua tienen como objetivo:

- A.** Improve AI accuracy / Mejorar la precisión de la IA
- B.** Protect intellectual property / Proteger la propiedad intelectual
- C.** Identify AI-generated content / Identificar contenido generado por IA
- D.** Reduce computational costs / Reducir los costos computacionales

5. After fine-tuning, compliance responsibility falls primarily on:

La responsabilidad del cumplimiento después del ajuste fino recae principalmente en:

- A.** The original model developer / El desarrollador original del modelo
- B.** The cloud provider / El proveedor en la nube
- C.** The downstream deployer / El implementador posterior (downstream deployer)
- D.** The end user / El usuario final

6. Transparency obligations for generative AI require disclosure when:

Las obligaciones de transparencia para la IA generativa requieren divulgación cuando:

- A.** AI improves efficiency / La IA mejora la eficiencia
- B.** AI content replaces human work / El contenido de IA reemplaza trabajo humano
- C.** Users interact with AI-generated content / Los usuarios interactúan con contenido generado por IA
- D.** AI systems are proprietary / Los sistemas de IA son propietarios

7. All GPAI models are classified as high-risk by default.

Todos los modelos GPAI se clasifican como de alto riesgo por defecto.

- A.** True / Verdadero
- B.** False / Falso
- C.** Only systemic GPAI / Solo GPAI sistémica
- D.** Only generative AI / Solo IA generativa

8. Synthetic content must be clearly identifiable under the EU AI Act.

El contenido sintético debe ser claramente identificable según el Reglamento de IA de la UE.

- A.** True / Verdadero
- B.** False / Falso
- C.** Only for images / Solo para imágenes
- D.** Only for audio / Solo para audio

9. GPAI obligations increase when models pose _____ risks.

Las obligaciones de GPAI aumentan cuando los modelos presentan riesgos _____.

- A.** Economic / Económicos
- B.** Systemic / Sistémicos
- C.** Minimal / Mínimos
- D.** Commercial / Comerciales

10. Watermarking helps users identify _____ content.

La marca de agua ayuda a los usuarios a identificar contenido _____.

- A.** Licensed / Licenciado
- B.** Encrypted / Encriptado
- C.** AI-generated / Generado por IA
- D.** Personalized / Personalizado

5

MODULE FIVE · MÓDULO 5

Enforcement, Risk, and Liability Frameworks

Ejecución, Riesgo y Marcos de Responsabilidad

5.1 Market Surveillance, Audits, and Regulatory Powers

Supervisión del Mercado, Auditorías y Poderes Regulatorios

5.2 Penalties, Fines, and Non-Compliance Consequences

Sanciones, Multas y Consecuencias por Incumplimiento

5.3 Civil and Product Liability in AI-Driven Systems

Responsabilidad Civil y de Producto en Sistemas Impulsados por IA

5.4 Contractual Controls and Third-Party Risk Management

Controles Contractuales y Gestión de Riesgos de Terceros

Module 5 — Enforcement, Risk, and Liability Frameworks

Módulo 5 — Ejecución, Riesgo y Marcos de Responsabilidad

ENGLISH

ESPAÑOL

WHY THIS MODULE · POR QUÉ ESTE MÓDULO

A small administrative oversight is what pushed EU regulators to tighten enforcement around AI. It was not a massive technical failure — it was a quiet documentation gap. A regulator asked a company for the paperwork behind an AI tool influencing thousands of decisions, and the organization could not produce a single verifiable file. The lesson was unmistakable: without strong enforcement, even the best-designed framework collapses.

This module explores the part of the EU AI Act that turns rules into reality. Enforcement, audits, penalties, liability and contractual controls determine how organizations protect themselves, their customers and their staff. Spain's enforcement culture adds an extra layer of clarity and seriousness about how sanctions and responsibilities are applied. Understanding these mechanisms lets you see AI governance not as theory, but as an operating system that keeps organizations accountable.

This is your moment to build confidence in navigating the legal and practical consequences of AI use. The better you understand enforcement and liability, the better prepared you are to build systems and organizations that withstand scrutiny with strength and credibility.

Un pequeño descuido administrativo impulsó a los reguladores de la UE a reforzar la ejecución en torno a la IA: no se trató de un fallo técnico masivo, sino de un vacío documental silencioso. Un regulador solicitó a una empresa la documentación de una herramienta de IA que influía en miles de decisiones, y la organización no pudo proporcionar un solo archivo verificable. Esta situación dejó algo claro: sin una ejecución sólida, incluso el marco mejor diseñado se desmorona.

Este módulo explora la parte de la Ley de IA de la UE que convierte las reglas en realidad. La ejecución, auditorías, sanciones, responsabilidad y controles contractuales determinan cómo las organizaciones se protegen a sí mismas, a sus clientes y a su personal. La cultura de ejecución de España añade un nivel adicional de claridad y seriedad sobre cómo se aplican las sanciones y responsabilidades. Comprender estos mecanismos permite ver la gobernanza de IA no como teoría, sino como un sistema operativo que mantiene a las organizaciones responsables.

Este es tu momento para construir confianza en la navegación de las consecuencias legales y prácticas del uso de IA. Cuanto más comprendas la ejecución y la responsabilidad, mejor preparado estarás para construir sistemas y organizaciones que resistan el escrutinio con solidez y credibilidad. ## Topic 1: Market Surveillance, Audits, and Regulatory Powers / Supervisión del Mercado, Auditorías y Poderes Regulatorios

5.1 Market Surveillance, Audits, and Regulatory Powers

Supervisión del Mercado, Auditorías y Poderes Regulatorios

Enforcement begins with market surveillance — a structured system ensuring that AI placed on the EU market meets legal requirements. Authorities inspect systems, request documentation and evaluate compliance across sectors. Spain participates actively through the AEPD, labor authorities and sector

La ejecución comienza con la supervisión del mercado, un sistema estructurado que asegura que la IA introducida en el mercado de la UE cumpla con los requisitos legales. Las autoridades inspeccionan sistemas, solicitan documentación y evalúan el cumplimiento en diversos sectores. España participa activamente a través de la AEPD, autoridades

regulators overseeing AI in health, finance, transport and public administration.

Regulators hold broad powers to inspect AI systems both before and after deployment. They can demand technical documentation, request logs, interview staff and test AI behavior in controlled scenarios. These powers allow them to detect risks, verify claims and challenge organizations that fall short of compliance expectations.

Audits are a central enforcement mechanism. They assess whether systems remain aligned with legal, ethical and operational requirements. For high-risk AI, audits verify that data-quality standards, oversight protocols and monitoring practices remain active. Spanish regulators expect thorough internal traceability so that design decisions, operations and outcomes can be clearly linked.

Regulators may also carry out targeted inspections based on complaints, incident reports or observed anomalies. If a Spanish public agency receives repeated concerns about unfair automated decisions, for instance, authorities can open a formal investigation demanding evidence from providers and deployers alike. Accountability is never left uncovered.

These enforcement powers serve a stabilizing purpose: they build trust across the AI ecosystem by making sure organizations cannot hide failures or cut standards. Strong enforcement does not limit innovation; it strengthens the foundation that makes innovation safe, ethical and sustainable.

laborales y reguladores sectoriales que supervisan IA en salud, finanzas, transporte y administración pública.

Los reguladores tienen amplios poderes para inspeccionar sistemas de IA tanto antes como después de su despliegue. Pueden exigir documentación técnica, solicitar registros, entrevistar al personal y probar el comportamiento de la IA en escenarios controlados. Estos poderes permiten detectar riesgos, verificar afirmaciones y desafiar a organizaciones que no cumplan con las expectativas de cumplimiento.

Las auditorías son un mecanismo central de ejecución. Evalúan si los sistemas permanecen alineados con los requisitos legales, éticos y operativos. Para IA de alto riesgo, las auditorías verifican que los estándares de calidad de datos, los protocolos de supervisión y las prácticas de monitoreo sigan activos. Los reguladores españoles esperan trazas internas exhaustivas para poder vincular decisiones de diseño, operativas y resultados de manera clara.

Los reguladores también pueden realizar inspecciones específicas basadas en denuncias, informes de incidentes o anomalías observadas. Por ejemplo, si una agencia pública española recibe preocupaciones repetidas sobre decisiones automatizadas injustas, las autoridades pueden iniciar una investigación formal que exija evidencia tanto a proveedores como a deployers. Esto garantiza que la responsabilidad nunca quede sin cubrir.

Estos poderes de ejecución cumplen un propósito estabilizador: crean confianza en todo el ecosistema de IA al asegurar que las organizaciones no puedan ocultar fallas o reducir estándares. Una ejecución sólida no limita la innovación; fortalece la base que hace que la innovación sea segura, ética y sostenible.

Topic 2: Penalties, Fines, and Non-Compliance Consequences + Reflection Story / Sanciones, Multas y Consecuencias por Incumplimiento + Historia de Reflexión

5.2 Penalties, Fines, and Non-Compliance Consequences

Sanciones, Multas y Consecuencias por Incumplimiento

The EU AI Act introduces one of the most robust penalty structures in global technology regulation. Fines depend on the severity of the infringement, the

La Ley de IA de la UE introduce una de las estructuras de sanción más robustas en la regulación tecnológica global. Las multas dependen de la gravedad de la

type of system involved and the potential harm. The penalty structure resembles the GDPR's, adapted to AI's unique risks. Spain applies these sanctions rigorously, especially in cases affecting the public or workers.

The heaviest penalties apply to prohibited systems, breaches of high-risk obligations and violations affecting fundamental rights. Fines can reach tens of millions of euros or a percentage of global turnover — whichever is higher. The message is unambiguous: AI misuse is not a minor compliance issue; it is a serious legal and ethical infringement.

Non-compliance also carries operational consequences. Regulators can suspend deployments, order corrective action, withdraw CE marking or restrict access to systems until problems are fixed. The reputational cost can be just as damaging — particularly in Spain, where public trust is central to organizational credibility.

infracción, el tipo de sistema involucrado y el daño potencial. La estructura de sanciones se asemeja al GDPR, pero adaptada a los riesgos únicos de la IA. España aplica estas sanciones rigurosamente, especialmente en casos de impacto público o sobre trabajadores.

Las sanciones más severas se aplican a sistemas prohibidos, incumplimientos de obligaciones de alto riesgo y violaciones que afecten derechos fundamentales. Las multas pueden alcanzar decenas de millones de euros o un porcentaje del volumen de negocios global, lo que sea mayor. Estas sanciones envían un mensaje claro: el mal uso de IA no es un problema menor de cumplimiento; es una infracción legal y ética significativa.

El incumplimiento también conlleva consecuencias operativas. Los reguladores pueden suspender despliegues, exigir acciones correctivas, retirar el marcado CE o restringir el acceso a sistemas hasta que se solucionen los problemas. El costo reputacional puede ser igual de dañino, especialmente en España, donde la confianza pública es clave para la credibilidad organizacional.

CASE IN PRACTICE · CASO PRÁCTICO

A logistics company in Seville integrated an AI tool to manage worker assignments. Believing the system was low-risk, they skipped transparency measures and never documented oversight procedures. Workers began reporting inconsistent schedules and suspected algorithmic discrimination. When regulators opened an investigation, the company could not produce records or explain the model's decisions. They were fined, forced to halt the system and required to rebuild governance from scratch. Months later — after implementing documentation, fairness testing and oversight protocols — the company regained trust and resumed operations.

Una empresa logística en Sevilla integró una herramienta de IA para gestionar asignaciones de trabajadores. Creyendo que el sistema era de bajo riesgo, omitieron medidas de transparencia y no documentaron los procedimientos de supervisión. Los trabajadores comenzaron a reportar horarios inconsistentes y sospechas de discriminación algorítmica. Cuando los reguladores iniciaron una investigación, la empresa no pudo presentar registros ni explicar las decisiones del modelo. Recibieron sanciones, debieron detener el sistema y reconstruir la gobernanza desde cero. Meses después, tras implementar documentación, pruebas de equidad y protocolos de supervisión, la empresa recuperó la confianza y reanudó operaciones.

The story shows how non-compliance usually stems from assumptions rather than malice. Penalties are not merely punitive; they push organizations toward disciplined, responsible AI use. When compliance is woven into daily practice, risk drops dramatically.

Esta historia muestra cómo el incumplimiento suele originarse en suposiciones más que en malicia. Las sanciones no son solo punitivas; impulsan a las organizaciones hacia un uso disciplinado y responsable de la IA. Cuando el cumplimiento se integra en la práctica diaria, los riesgos disminuyen

5.3 Civil and Product Liability in AI-Driven Systems

Responsabilidad Civil y de Producto en Sistemas Impulsados por IA

Civil liability frameworks determine who is responsible when AI causes harm. The EU recognized that traditional liability rules were never designed for autonomous systems, so it introduced updated guidance to clarify responsibility. In Spain, civil liability is closely tied to consumer protection and labor law, so organizations must demonstrate diligence and oversight.

Product liability applies to AI systems deemed defective or unsafe. If a system behaves unpredictably, generates harmful outputs or fails to perform its intended function, both providers and deployers can be held responsible — including cases where AI decisions cause financial loss, discriminatory treatment or safety risks.

Civil liability becomes complicated when multiple actors are involved: foundation-model creators, fine-tuners, deployers, vendors. The EU AI Act clarifies that liability follows contribution to risk. If a deployer misconfigures a system or ignores oversight rules, it can bear responsibility even when the model itself is compliant.

Organizations reduce liability by documenting decisions, training staff, implementing oversight, monitoring system performance and keeping clear records. Spanish authorities expect organizations to show what actions they took — not just what the system was programmed to do.

Los marcos de responsabilidad civil determinan quién es responsable cuando la IA causa daño. La UE reconoce que las normas tradicionales de responsabilidad no estaban diseñadas para sistemas autónomos, por lo que introdujo directrices actualizadas para aclarar la responsabilidad. En España, la responsabilidad civil está estrechamente ligada a la protección del consumidor y a las leyes laborales, por lo que las organizaciones deben demostrar diligencia y supervisión.

La responsabilidad de producto se aplica a sistemas de IA considerados defectuosos o inseguros. Si un sistema se comporta de forma impredecible, genera resultados dañinos o no cumple su función prevista, tanto proveedores como deployers pueden ser responsables. Esto incluye casos en que decisiones de IA provoquen pérdidas financieras, trato discriminatorio o riesgos de seguridad.

La responsabilidad civil se complica cuando intervienen múltiples actores: creadores de modelos fundamentales, ajustadores, deployers y vendedores. La Ley de IA de la UE aclara que la responsabilidad sigue la contribución al riesgo. Si un deployer configura incorrectamente un sistema o ignora las normas de supervisión, puede asumir responsabilidad incluso si el modelo es conforme.

Las organizaciones reducen la responsabilidad documentando decisiones, capacitando al personal, implementando supervisión, monitoreando el desempeño del sistema y manteniendo registros claros. Las autoridades españolas esperan que las organizaciones demuestren qué acciones tomaron, no solo lo que el sistema estaba programado para hacer.

DECISION POINT · PUNTO DE DECISIÓN

A Spanish insurance provider uses an AI tool to estimate claim payouts. One day a claimant receives a drastically incorrect assessment that causes financial harm. Who is liable? A) Only the AI provider, because it built the model. B) Only the deployer, because it used the model. C) Potentially both, depending on whether documentation, configuration and oversight duties were met.

Un proveedor de seguros español utiliza una herramienta de IA para estimar pagos de reclamaciones. Un día, un reclamante recibe una evaluación drásticamente incorrecta que provoca perjuicio económico. ¿Quién es responsable? A) Solo el proveedor de IA, porque construyó el modelo. B) Solo el deployer, porque utilizó el modelo. C) Potencialmente ambos, dependiendo de si la documentación, configuración y supervisión se cumplieron.

Correct answer: C — Liability depends on shared contribution to risk and on whether each actor met its obligations.
Respuesta correcta: C — la responsabilidad depende de la contribución compartida al riesgo y del cumplimiento de las obligaciones de cada actor.

5.4 Contractual Controls and Third-Party Risk Management

Controles Contractuales y Gestión de Riesgos de Terceros

Contractual controls determine how organizations manage risk when using third-party AI systems. Contracts should define responsibilities, documentation access, data-usage rights, model limitations and obligations around updates or incident notification. Without these controls, organizations expose themselves to regulatory uncertainty and legal vulnerability.

Third-party risk management acknowledges that many organizations in Spain rely on imported or externally developed AI tools. Deployers must evaluate whether vendors provide sufficient documentation, meet transparency requirements, follow cybersecurity standards and align with EU legislation. A vendor's compliance claims are not enough; organizations must verify them.

Contracts should include terms on performance monitoring, audit rights, liability allocation and obligations when the system drifts or behaves unexpectedly. Strong contractual governance ensures every party understands its role in maintaining safety and compliance.

Los controles contractuales determinan cómo las organizaciones gestionan el riesgo al usar sistemas de IA de terceros. Los contratos deben definir responsabilidades, acceso a documentación, derechos de uso de datos, limitaciones del modelo y obligaciones de actualización o notificación de incidentes. Sin estos controles, las organizaciones se exponen a incertidumbre regulatoria y vulnerabilidad legal.

La gestión de riesgos de terceros reconoce que muchas organizaciones en España dependen de herramientas de IA importadas o desarrolladas externamente. Los deployers deben evaluar si los proveedores ofrecen suficiente documentación, cumplen con requisitos de transparencia, siguen estándares de ciberseguridad y se alinean con la legislación de la UE. Las afirmaciones de cumplimiento de un proveedor no son suficientes; las organizaciones deben verificarlas.

Los contratos deben incluir términos sobre monitoreo de desempeño, derechos de auditoría, asignación de responsabilidad y obligaciones ante desviaciones o comportamientos inesperados del sistema. Una gobernanza contractual sólida asegura que cada parte comprenda su rol en el mantenimiento de la seguridad y cumplimiento.

APPLIED EXAMPLE · EJEMPLO APLICADO

A US financial-services firm expands into Spain offering an automated credit-risk assessment tool. The Spanish deployer demands contractual guarantees on data provenance, fairness testing, transparency features and model-update protocols. When the system's outputs begin drifting due to local demographic variation, the deployer invokes its contractual audit rights to investigate performance. The collaboration leads to retraining on Spain-specific data — reducing errors and aligning the system with national fairness expectations.

Una empresa de servicios financieros estadounidense se expande a España ofreciendo una herramienta automatizada de evaluación de riesgo crediticio. El deployer español exige garantías contractuales sobre la procedencia de los datos, pruebas de equidad, características de transparencia y protocolos de actualización del modelo. Cuando los resultados del sistema empiezan a desviarse debido a variaciones demográficas locales, el deployer utiliza los derechos de auditoría contractuales para investigar el desempeño. Esta colaboración conduce a un reentrenamiento con datos específicos de España, reduciendo errores y alineando el sistema con las expectativas nacionales de equidad.

The example shows that contracts are not just legal shields; they are tools for operational reliability. With solid agreements and diligent vendor oversight, organizations prevent governance gaps and strengthen the entire AI supply chain.

Este ejemplo demuestra que los contratos no son solo protecciones legales; son herramientas para la fiabilidad operativa. Con acuerdos sólidos y supervisión diligente de proveedores, las organizaciones previenen vacíos de gobernanza y fortalecen toda la cadena de suministro de IA. ## Module Wrap-Up + Practical Life Implementation / Cierre del Módulo + Implementación Práctica

Module Wrap-Up & Practical Implementation

Cierre del Módulo e Implementación Práctica

You now have a clear picture of how enforcement, penalties, liability and contractual controls work together to shape AI governance in Spain and the EU. These mechanisms ensure organizations cannot rely on good intentions alone: they must demonstrate accountability, maintain documentation and align their systems with legal expectations.

The story and the scenario showed that the consequences of weak governance are rarely immediate; they accumulate quietly until something critical fails. They also showed that organizations can recover quickly by establishing discipline, transparency and accountability.

Pick one AI tool your organization uses. Review your documentation, contractual terms, incident-logging processes and compliance posture. By applying these

Ahora tienes un entendimiento claro de cómo la ejecución, sanciones, responsabilidad y controles contractuales trabajan juntos para conformar la gobernanza de IA en España y la UE. Estos mecanismos aseguran que las organizaciones no puedan depender únicamente de buenas intenciones: deben demostrar responsabilidad, mantener documentación y alinear sus sistemas con las expectativas legales.

La historia y el escenario demostraron que las consecuencias de una gobernanza deficiente rara vez son inmediatas; se acumulan silenciosamente hasta que algo crítico falla. Pero también mostraron que las organizaciones pueden recuperarse rápidamente al establecer disciplina, transparencia y responsabilidad.

practices you move from merely using AI to governing it — protecting your organization, your users and your reputation.

Selecciona una herramienta de IA que tu organización utilice. Revisa tu documentación, términos contractuales, procesos de registro de incidentes y postura de cumplimiento. Al aplicar estas prácticas, pasas de simplemente usar IA a gobernarla— protegiendo tu organización, tus usuarios y tu reputación.

MODULE 5 ASSESSMENT · 10 QUESTIONS

Test Your Understanding

Evalúa tu comprensión — Respuestas en el Apéndice A

1. Which authority is responsible for market surveillance under the EU AI Act?

¿Qué autoridad es responsable de la vigilancia del mercado bajo el Reglamento de IA de la UE?

- A. National supervisory authorities / Autoridades supervisoras nacionales
- B. Private auditors / Auditores privados
- C. Industry associations / Asociaciones industriales
- D. AI providers / Proveedores de IA

2. Maximum fines under the EU AI Act are calculated based on:

Las multas máximas bajo el Reglamento de IA de la UE se calculan en función de:

- A. Company size / Tamaño de la empresa
- B. Net profit / Beneficio neto
- C. Global annual turnover / Facturación anual global
- D. National revenue / Ingresos nacionales

3. Non-compliance with prohibited AI practices can result in:

El incumplimiento de prácticas de IA prohibidas puede resultar en:

- A. Warning letters / Cartas de advertencia
- B. Mandatory retraining / Reentrenamiento obligatorio
- C. Administrative fines / Multas administrativas
- D. Criminal penalties only / Sanciones penales únicamente

4. AI-related civil liability mainly refers to:

La responsabilidad civil relacionada con la IA se refiere principalmente a:

- A. Ethical misconduct / Conducta ética indebida
- B. Market competition / Competencia en el mercado
- C. Harm caused by AI systems / Daños causados por sistemas de IA
- D. Contractual breaches only / Incumplimientos contractuales únicamente

5. Contractual controls are mainly used to:

Los controles contractuales se utilizan principalmente para:

- A. Avoid AI regulation / Evitar la regulación de IA
- B. Allocate responsibilities between parties / Asignar responsabilidades entre las partes
- C. Reduce development costs / Reducir costos de desarrollo
- D. Eliminate liability / Eliminar la responsabilidad

6. Third-party risk management focuses on:

La gestión de riesgos de terceros se enfoca en:

- A. Vendor accountability / Responsabilidad del proveedor
- B. User behavior / Comportamiento del usuario
- C. Market expansion / Expansión de mercado
- D. Data localization / Localización de datos

7. EU AI Act penalties apply only to companies headquartered in the EU.

Las sanciones del Reglamento de IA de la UE se aplican únicamente a empresas con sede en la UE.

- A. True / Verdadero
- B. False / Falso
- C. Only public entities / Solo a entidades públicas
- D. Only providers / Solo a proveedores

8. Audits can be initiated by regulators to verify AI compliance.

Los auditores pueden ser iniciados por reguladores para verificar el cumplimiento de la IA.

- A. True / Verdadero
- B. False / Falso
- C. Only voluntarily / Solo de manera voluntaria
- D. Only after an incident / Solo después de un incidente

9. Fines can reach a percentage of a company's _____ turnover.

Las multas pueden alcanzar un porcentaje de la facturación _____ de la empresa.

- A.** National / Nacional
- B.** Monthly / Mensual
- C.** Global / Global
- D.** Project / De proyecto

10. Liability frameworks address AI-related harm _____ to individuals.

Los marcos de responsabilidad abordan los daños _____ relacionados con la IA causados a los individuos.

- A.** Benefits / Beneficios
- B.** Innovation / Innovación
- C.** Caused / Causados
- D.** Automation / Automatización

6

MODULE SIX · MÓDULO 6

Ethical AI Principles and Governance Structures

Principios Éticos de IA y Estructuras de Gobernanza

6.1 Fairness, Non-Discrimination, and Human Rights Alignment

Equidad, No Discriminación y Alineación con Derechos Humanos

6.2 Transparency, Explainability, and Trustworthiness

Transparencia, Explicabilidad y Confiabilidad

6.3 Ethical Risk Assessment and Mitigation Techniques

Evaluación de Riesgos Éticos y Técnicas de Mitigación

6.4 Organizational AI Governance, Oversight, and Accountability

Gobernanza Organizacional de IA, Supervisión y Responsabilidad

Module 6 — Ethical AI Principles and Governance Structures

Módulo 6 — Principios Éticos de IA y Estructuras de Gobernanza

ENGLISH

WHY THIS MODULE · POR QUÉ ESTE MÓDULO

Organizations that excel at ethical AI don't just avoid harm — they unlock greater innovation, better decision-making and deeper trust from the people they serve. Ethical principles are not ornamental; they are operational tools that shape how AI behaves, how teams interact with it and how society receives it. The shift is especially visible in Spain, where public institutions and private companies increasingly treat ethical alignment as a competitive advantage.

Fairness, transparency, explainability, ethical risk assessment — and the governance structures that keep responsibility intact. These principles transform AI from a technical product into a system that reflects values, safeguards rights and strengthens social trust. Ethical frameworks matter because AI does not merely make predictions; it influences people's opportunities, freedoms and dignity.

This is where understanding moves from compliance to leadership: when ethical principles guide your decisions, you strengthen your influence, protect your organization and help build AI systems that genuinely serve society.

ESPAÑOL

Las organizaciones que sobresalen en IA ética no solo evitan daños, sino que desbloquean mayor innovación, mejor toma de decisiones y una confianza más profunda de quienes sirven. Los principios éticos no son un adorno; son herramientas operativas que moldean cómo se comporta la IA, cómo los equipos interactúan con ella y cómo la sociedad la recibe. Este cambio es especialmente visible en España, donde las instituciones públicas y las empresas privadas consideran cada vez más la alineación ética como una ventaja competitiva.

Equidad, transparencia, explicabilidad, evaluación de riesgos éticos y las estructuras de gobernanza que aseguran que la responsabilidad permanezca intacta. Estos principios transforman la IA de un producto técnico en un sistema que refleja valores, salvaguarda derechos y fortalece la confianza social. Los marcos éticos son importantes porque la IA no solo hace predicciones, sino que influye en las oportunidades, libertades y dignidad de las personas.

Comprender desde el cumplimiento hasta el liderazgo: cuando los principios éticos guían tus decisiones, fortaleces tu influencia, proteges tu organización y contribuyes a construir sistemas de IA que realmente sirvan a la sociedad. ## Topic 1: Fairness, Non-Discrimination, and Human Rights Alignment / Tema 1: Equidad, No Discriminación y Alineación con Derechos Humanos

6.1 Fairness, Non-Discrimination, and Human Rights Alignment

Equidad, No Discriminación y Alineación con Derechos Humanos

Fairness is one of the most closely watched ethical principles in the EU — for good reason. AI systems can unintentionally reproduce or amplify biases present in historical data, organizational practices or social patterns. In Spain, anti-discrimination law intersects with AI regulation, creating an expectation

La equidad es uno de los principios éticos más monitoreados en la UE, y por buenas razones. Los sistemas de IA pueden reproducir o amplificar involuntariamente sesgos presentes en datos históricos, prácticas organizacionales o patrones sociales. En España, las leyes contra la discriminación se intersectan con la regulación de IA, creando expectativas de que los sistemas deben tratar a las

that systems treat people equitably across gender, ethnicity, disability, age and socioeconomic factors.

Human-rights alignment serves as the foundation for fairness. AI must respect individual dignity, autonomy and equal access to opportunity. When systems influence hiring, credit, insurance, education or public services, fairness becomes not just a technical aspiration but a legal and moral requirement. Fairness means ensuring AI decisions do not create structural disadvantage.

Organizations must assess how models behave across demographic groups, understand where disparities arise and document the steps taken to mitigate them. Fairness monitoring is continuous, not one-off. A system that performs equitably today can drift tomorrow through new data or shifting patterns; ethical oversight must adapt accordingly.

Spain's regulatory expectations emphasize fairness especially in labor contexts. Automated decision tools used in hiring, promotion, scheduling or evaluation must demonstrate that they do not degrade workers' rights. Ethical alignment protects both organizational integrity and social wellbeing.

When fairness is embedded across the AI lifecycle — from data selection through model testing to deployment — trust builds far more easily. Fairness is not one principle among many; it is the principle that makes every other ethical standard truly meaningful.

personas de manera equitativa respecto a género, etnia, discapacidad, edad y factores socioeconómicos.

La alineación con los derechos humanos sirve como base para la equidad. La IA debe respetar la dignidad individual, la autonomía y el acceso equitativo a oportunidades. Cuando los sistemas influyen en contratación, crédito, seguros, educación o servicios públicos, la equidad se convierte no solo en una aspiración técnica, sino en un requisito legal y moral. Equidad significa garantizar que las decisiones de IA no creen desventajas estructurales.

Las organizaciones deben evaluar cómo se comportan los modelos en distintos grupos demográficos, entender dónde surgen disparidades y documentar los pasos tomados para mitigarlas. El monitoreo de equidad es continuo, no único. Un sistema que funciona de manera equitativa hoy puede desviarse mañana debido a nuevos datos o patrones cambiantes. La supervisión ética debe adaptarse en consecuencia.

Las expectativas regulatorias de España enfatizan la equidad especialmente en contextos laborales. Las herramientas de decisión automatizada utilizadas en contratación, promociones, programación o evaluación deben demostrar que no degradan los derechos de los trabajadores. La alineación ética protege tanto la integridad organizacional como el bienestar social.

Cuando la equidad se integra a lo largo del ciclo de vida de la IA —desde la selección de datos, pruebas del modelo hasta el despliegue— la confianza se construye con mayor facilidad. La equidad no es un principio más entre muchos; es el principio que garantiza que todos los demás estándares éticos sean realmente significativos. ## Topic 2: Transparency, Explainability, and Trustworthiness + Reflection Story / Tema 2: Transparencia, Explicabilidad y Confiabilidad + Historia de Reflexión

6.2 Transparency, Explainability, and Trustworthiness

Transparencia, Explicabilidad y Confiabilidad

Transparency ensures users understand when AI is influencing an interaction, recommendation or decision. This matters because people make better choices when they know whether they are dealing with a system or a human. Transparency also creates

La transparencia asegura que los usuarios comprendan cuándo la IA está influyendo en una interacción, recomendación o decisión. Esto importa porque las personas toman mejores decisiones cuando saben si interactúan con un sistema o con un

accountability: it prevents organizations from hiding behind algorithmic processes.

Explainability goes a step further. It provides insight into why the system produced a particular result. Even when the full technical process is complex, organizations must offer clear explanations that help users interpret outcomes. Spain's legal environment supports this expectation, especially in sectors where decisions affect rights, such as employment or public services.

Trustworthiness emerges when transparency and explainability work hand in hand. Users trust systems that behave predictably, provide clarity and offer meaningful routes for appeal or correction. Without these, even high-performing systems face skepticism.

humano. La transparencia también genera responsabilidad: evita que las organizaciones se escuden detrás de procesos algorítmicos.

La explicabilidad va un paso más allá. Brinda información sobre por qué el sistema produjo un resultado particular. Incluso si el proceso técnico completo es complejo, las organizaciones deben ofrecer explicaciones claras que ayuden a los usuarios a interpretar los resultados. El entorno legal de España respalda esta expectativa, especialmente en sectores donde las decisiones afectan derechos, como empleo o servicios públicos.

La confiabilidad surge cuando transparencia y explicabilidad trabajan de la mano. Los usuarios confían en sistemas que se comportan de manera predecible, ofrecen claridad y brindan vías significativas para apelación o corrección. Sin estos elementos, incluso sistemas de buen desempeño enfrentan escepticismo.

CASE IN PRACTICE · CASO PRÁCTICO

A Madrid-based insurer integrated an AI model to support claims assessment. Customers began receiving automated decisions without understanding how the model reached its conclusions. Complaints rose and trust fell. When the company introduced transparency disclosures, simplified explanations and a clear appeal process, trust began to recover. Interestingly, the system's performance never changed — only the way the organization communicated about it. Better clarity became a strategic advantage.

Ejemplo: una aseguradora con sede en Madrid integró un modelo de IA para apoyar la evaluación de reclamaciones. Los clientes comenzaron a recibir decisiones automáticas sin entender cómo el modelo llegaba a sus conclusiones. Aumentaron las quejas y disminuyó la confianza. Cuando la empresa introdujo divulgaciones de transparencia, explicaciones simplificadas y un proceso de apelación claro, la confianza comenzó a recuperarse. Curiosamente, el desempeño del sistema no cambió; solo cambió la forma en que la organización se comunicaba al respecto. La claridad mejorada se convirtió en una ventaja estratégica.

Transparency is not a technical add-on. It is a relational practice. When organizations offer honest, understandable explanations, they invite trust and create systems people are willing to work with — not resist.

La transparencia no es un complemento técnico. Es una práctica relacional. Cuando las organizaciones ofrecen explicaciones honestas y comprensibles, invitan a la confianza y crean sistemas con los que las personas están dispuestas a colaborar, no a resistirse.

Topic 3: Ethical Risk Assessment and Mitigation Techniques + Scenario-Based Mini Assessment / Tema 3: Evaluación de Riesgos Éticos y Técnicas de Mitigación + Mini Evaluación por Escenario

6.3 Ethical Risk Assessment and Mitigation Techniques

Evaluación de Riesgos Éticos y Técnicas de Mitigación

Ethical risk assessment identifies how AI could cause harm, intentionally or not. Ethical risks include bias, privacy violations, manipulation, misinformation, over-reliance on automation and exploitation of vulnerabilities. Unlike traditional risk models focused on financial or operational damage, ethical risk assessment examines how dignity, fairness and human rights may be affected.

Mitigation requires foresight. Organizations map potential harms, consider worst-case scenarios, analyze how marginalized groups could be affected and design countermeasures. These may include model retraining, data filtering, human review loops, user education or workflow redesign. Ethical risk mitigation is iterative and adaptive.

Many organizations in Spain now build ethical review into procurement and design stages. Before adopting an external AI tool, leaders ask: does the system respect rights? Could it misrepresent people? Does it create asymmetry between the organization and those it serves? These questions stop ethical problems before they start.

Monitoring mechanisms are essential. Even a system launched with strong safeguards can reveal new concerns in real-world use. Organizations maintain alerts, user-feedback loops and oversight dashboards tracking how the model behaves over time.

La evaluación de riesgos éticos identifica cómo la IA puede causar daño, de manera intencional o no. Los riesgos éticos incluyen sesgos, violaciones de privacidad, manipulación, desinformación, exceso de dependencia de la automatización y explotación de vulnerabilidades. A diferencia de los modelos de riesgo tradicionales enfocados en daños financieros u operativos, la evaluación de riesgos éticos examina cómo la dignidad, equidad y derechos humanos pueden verse afectados.

La mitigación requiere previsión. Las organizaciones evalúan daños potenciales, consideran escenarios de peor caso, analizan cómo los grupos marginados pueden verse afectados y diseñan contramedidas. Esto puede incluir reentrenamiento de modelos, filtrado de datos, ciclos de revisión humana, educación del usuario o rediseño de flujos de trabajo. La mitigación de riesgos éticos es iterativa y adaptativa.

Muchas organizaciones en España ahora integran revisiones éticas en las etapas de adquisición y diseño. Antes de adoptar una herramienta de IA externa, los líderes preguntan: ¿El sistema respeta los derechos? ¿Podría malrepresentar a las personas? ¿Crea asimetría entre la organización y quienes sirve? Estas preguntas evitan que surjan problemas éticos más adelante.

Los mecanismos de monitoreo son esenciales. Incluso si un sistema se lanza con fuertes salvaguardas, el comportamiento en el mundo real puede revelar nuevas preocupaciones. Las organizaciones mantienen alertas, bucles de retroalimentación de usuarios y tableros de supervisión que rastrean cómo se comporta el modelo con el tiempo.

DECISION POINT · PUNTO DE DECISIÓN

A Spanish university deploys an AI tool to help professors predict which students may struggle academically. After several weeks, students from a particular socioeconomic background are flagged disproportionately. What should the university do? A) Assume the tool is correct and keep using it. B) Investigate data representation, review the model's logic and assess whether contextual factors are producing unfair signals. C) Disable the tool permanently without further analysis.

Ejemplo: una universidad española despliega una herramienta de IA para ayudar a los profesores a predecir qué estudiantes podrían tener dificultades académicas. Tras varias semanas, los estudiantes de un determinado estrato socioeconómico son señalados de manera desproporcionada. ¿Qué debería hacer la universidad? A) Asumir que la herramienta es correcta y seguir utilizándola. B) Investigar la representación de los datos, revisar la lógica del modelo y evaluar si factores contextuales generan señales injustas. C) Deshabilitar la herramienta permanentemente sin más análisis.

Correct answer: B — Ethical risk assessment demands examination, not assumptions.

Respuesta correcta: B — la evaluación ética del riesgo requiere examen, no suposiciones.

6.4 Organizational AI Governance, Oversight, and Accountability

Gobernanza Organizacional de IA, Supervisión y Responsabilidad

Ethical AI cannot function without governance structures that define responsibilities, processes and escalation paths. Governance ensures decisions are documented, oversight is active and ethical principles are applied consistently across departments. Without it, even good intentions crumble under operational pressure.

Oversight assigns humans to monitor AI systems proactively. Oversight teams evaluate outputs, detect anomalies, escalate concerns and intervene when systems behave unexpectedly. This protects users and strengthens organizational credibility. Spain's regulatory culture strongly supports human oversight in every context touching rights or public trust.

Accountability structures clarify who owns each risk. Providers, deployers, fine-tuners and oversight teams each play a role. Accountability is not about assigning blame — it creates the clarity that enables more responsible system management. When everyone knows their role, ethical alignment becomes achievable.

La IA ética no puede funcionar sin estructuras de gobernanza que definan responsabilidades, procesos y vías de escalamiento. La gobernanza asegura que las decisiones se documenten, la supervisión sea activa y los principios éticos se apliquen de manera consistente en todos los departamentos. Sin gobernanza, incluso las buenas intenciones se desmoronan bajo presión operativa.

La supervisión asigna humanos para monitorear los sistemas de IA de manera proactiva. Los equipos de supervisión evalúan resultados, detectan anomalías, escalan preocupaciones e intervienen cuando los sistemas se comportan de manera inesperada. Esto protege a los usuarios y fortalece la credibilidad organizacional. La cultura regulatoria de España respalda fuertemente la supervisión humana en todos los contextos que afectan derechos o confianza pública.

Las estructuras de responsabilidad aclaran quién posee cada riesgo. Proveedores, implementadores, afinadores y equipos de supervisión tienen cada uno un rol. La responsabilidad no solo asigna culpa: crea claridad que permite una gestión más responsable del sistema. Cuando todos conocen su rol, la alineación ética se vuelve alcanzable.

APPLIED EXAMPLE · EJEMPLO APLICADO

A US healthcare-technology company expanded into Spain with a decision-support system for nurses. Before deployment, they established an AI governance committee, created oversight guidelines, defined clear documentation processes and trained staff on proper escalation protocols. When nurses noticed inconsistent recommendations under specific conditions, they followed the oversight path. Engineers traced the issue to a data-mapping error introduced during localization and fixed it quickly. Because governance was in place, patient safety was protected and regulatory confidence strengthened.

Ejemplo: una empresa estadounidense de tecnología sanitaria se expandió a España con un sistema de apoyo a decisiones para enfermeras. Antes del despliegue, establecieron un comité de gobernanza de IA, crearon pautas de supervisión, definieron procesos claros de documentación y capacitaron al personal en protocolos de escalamiento adecuados. Cuando las enfermeras notaron recomendaciones inconsistentes en condiciones específicas, siguieron la vía de supervisión. Los ingenieros descubrieron un error de mapeo de datos introducido durante la localización y lo corrigieron rápidamente. Gracias a la gobernanza, se protegió la seguridad del paciente y se fortaleció la confianza regulatoria.

Problems never disappear — but they become manageable. Governance transforms ethical risk from a growing threat into a structured, solvable challenge.

Los problemas no desaparecen, pero se vuelven manejables. La gobernanza transforma el riesgo ético de una amenaza creciente en un desafío estructurado y resoluble. ## Module Wrap-Up + Practical Life Implementation / Cierre del Módulo + Implementación Práctica

Module Wrap-Up & Practical Implementation

Cierre del Módulo e Implementación Práctica

You have explored the ethical principles and governance structures that form the backbone of responsible AI in Spain and across the EU. Fairness, transparency, risk assessment, oversight and accountability work together to ensure AI strengthens society rather than undermining it. Ethical principles are not abstract; they shape everyday interactions and decisions.

The reflection story and the scenario revealed something essential: ethical problems usually surface subtly before becoming visible. But when organizations embed strong ethical governance, they spot issues early, respond quickly and preserve trust even under pressure. Ethical AI is resilient AI.

Choose one AI system you interact with and examine it through an ethical lens. Ask whether fairness controls exist, whether transparency is clear, whether oversight is active and whether accountability is documented. Every step you take strengthens not

Ahora has explorado los principios éticos y las estructuras de gobernanza que forman la columna vertebral de la IA responsable en España y en toda la UE. Equidad, transparencia, evaluación de riesgos, supervisión y responsabilidad trabajan juntos para garantizar que la IA fortalezca la sociedad en lugar de socavarla. Los principios éticos no son abstractos; moldean las interacciones y decisiones cotidianas.

La historia de reflexión y el escenario revelaron algo esencial: los problemas éticos suelen aparecer de manera sutil antes de volverse visibles. Pero cuando las organizaciones integran una gobernanza ética sólida, identifican problemas temprano, responden rápidamente y mantienen la confianza incluso durante desafíos. La IA ética es IA resiliente.

Elige un sistema de IA con el que interactúes y examínalo a través de un lente ético. Pregúntate si existen controles de equidad, si la transparencia es clara, si la supervisión está activa y si la

only your organization's compliance, but its relationship with the people it serves.

responsabilidad está documentada. Cada paso que tomes fortalece no solo el cumplimiento de tu organización, sino su relación con las personas a las que sirve.

MODULE 6 ASSESSMENT · 10 QUESTIONS

Test Your Understanding

Evalúa tu comprensión — Respuestas en el Apéndice A

1. Fairness in AI primarily aims to prevent:

La equidad en la IA tiene como objetivo principal prevenir:

- A. Automation / La automatización
- B. Bias and discrimination / Sesgos y discriminación
- C. Slower innovation / La ralentización de la innovación
- D. Data sharing / El intercambio de datos

2. Explainability supports:

La explicabilidad (explainability) apoya:

- A. Model accuracy / La precisión del modelo
- B. User trust / La confianza del usuario
- C. Faster deployment / El despliegue más rápido
- D. Cost reduction / La reducción de costos

3. Ethical risk assessments are used to:

Las evaluaciones de riesgo ético se utilizan para:

- A. Replace legal compliance / Reemplazar el cumplimiento legal
- B. Identify and mitigate ethical harms / Identificar y mitigar daños éticos
- C. Increase AI autonomy / Aumentar la autonomía de la IA
- D. Reduce documentation / Reducir la documentación

4. Human-rights alignment ensures AI respects:

La alineación con los derechos humanos asegura que la IA respete:

- A. Business goals / Objetivos empresariales
- B. Organizational culture / Cultura organizacional
- C. Fundamental rights / Derechos fundamentales
- D. Market competition / Competencia en el mercado

5. Organizational AI governance assigns:

La gobernanza organizacional de la IA asigna:

- A. Random responsibility / Responsabilidad aleatoria
- B. Informal oversight / Supervisión informal
- C. Clear roles and responsibilities / Roles y responsabilidades claros
- D. External control only / Control externo únicamente

6. Trustworthy AI requires systems to be:

La IA confiable requiere que los sistemas sean:

- A. Profitable / Rentables
- B. Transparent and accountable / Transparentes y responsables
- C. Fully autonomous / Totalmente autónomos
- D. Open source / De código abierto

7. Ethical AI governance goes beyond legal compliance.

La gobernanza ética de la IA va más allá del cumplimiento legal.

- A. True / Verdadero
- B. False / Falso
- C. Only for high-risk AI / Solo para IA de alto riesgo
- D. Only in Spain / Solo en España

8. Transparency means publicly disclosing all source code.

Transparencia significa revelar todo el código fuente públicamente.

- A. True / Verdadero
- B. False / Falso
- C. Only for GPAI / Solo para GPAI
- D. Only under GDPR / Solo bajo GDPR

9. Ethical AI emphasizes _____ and accountability.

La IA ética enfatiza la _____ y la responsabilidad.

- A. Automation / Automatización
- B. Profitability / Rentabilidad
- C. Trustworthiness / Confiabilidad (trustworthiness)
- D. Efficiency / Eficiencia

10. Oversight mechanisms help ensure _____ use of AI systems.

Los mecanismos de supervisión ayudan a asegurar un uso _____ de los sistemas de IA.

- A. Legal / Legal
- B. Responsible / Responsable
- C. Experimental / Experimental
- D. Unlimited / Ilimitado

7

MODULE SEVEN · MÓDULO 7

Spanish AI Legal and Regulatory Framework

Marco Legal y Regulatorio de IA en España

7.1 Spain's Implementation of the EU AI Act

Implementación Española de la AI Act de la UE

7.2 GDPR, LOPDGDD, and National Data Protection Duties

GDPR, LOPDGDD y Obligaciones Nacionales de Protección de Datos

7.3 Spanish Labor, Workplace, and Equality Regulations

Regulación Laboral, de Igualdad y Uso de IA en el Trabajo

7.4 National AI Strategy, Public Sector Rules, and Enforcement Bodies

Estrategia Nacional de IA, Normas del Sector Público y Órganos de Supervisión

Module 7 — Spanish AI Legal and Regulatory Framework

Módulo 7 — Marco Legal y Regulatorio de IA en España

ENGLISH

ESPAÑOL

WHY THIS MODULE · POR QUÉ ESTE MÓDULO

Spain is one of the EU member states most proactively preparing for the AI Act. While the regulation sets a common European baseline, Spain adds layers of its own: data-protection rules, labor standards, equality obligations and public-sector governance requirements. Together they create a compliance environment where AI must satisfy not only European expectations but also national values grounded in transparency, fairness and social responsibility.

In this module you will explore Spain's implementation of the AI Act, national data-protection rules such as the GDPR and LOPDGDD, labor and equality regulation, and the structure of Spain's AI governance ecosystem. Understanding these frameworks is essential because they determine how organizations operate AI responsibly and lawfully. Spain's regulatory approach reflects its commitment to protecting fundamental rights and ensuring AI supports — rather than undermines — social and institutional systems.

Strengthen your grasp of how AI is governed in the Spanish national context. When you understand how Spain interprets and applies AI rules, you gain the clarity needed to design systems that meet expectations today and withstand legal scrutiny tomorrow.

España es uno de los estados miembros de la UE más proactivos en prepararse para la IA Act. Mientras que la regulación establece una base europea común, España añade sus propias capas: normas de protección de datos, estándares laborales, obligaciones de igualdad y requisitos de gobernanza en el sector público. Estas capas crean un entorno de cumplimiento donde la IA debe cumplir no solo con las expectativas europeas, sino también con valores nacionales basados en transparencia, equidad y responsabilidad social.

En este módulo, explorarás la implementación española de la IA Act, las normas nacionales de protección de datos como GDPR y LOPDGDD, la regulación laboral y de igualdad, y la estructura del ecosistema de gobernanza de IA en España. Comprender estos marcos es esencial porque determinan cómo las organizaciones operan la IA de manera responsable y legal. El enfoque regulatorio español refleja su compromiso con la protección de derechos fundamentales y asegura que la IA apoye—y no socave—los sistemas sociales e institucionales.

Fortalece tu comprensión de cómo se gobierna la IA en el contexto nacional español. Cuando entiendes cómo España interpreta y aplica las reglas de IA, obtienes la claridad necesaria para diseñar sistemas que cumplan expectativas hoy y resistan el escrutinio legal mañana. ### Topic 1: Spain's Implementation of the EU AI Act / Tema 1: Implementación Española de la IA Act de la UE

7.1 Spain's Implementation of the EU AI Act

Implementación Española de la IA Act de la UE

Spain adopts the EU AI Act as directly applicable regulation, but it also designates national authorities responsible for supervision and enforcement. Spain is building governance structures that coordinate market surveillance, oversight of high-risk systems and evaluation of AI used in public administration. The

España adopta la IA Act de la UE como regulación directamente aplicable, pero también define autoridades nacionales responsables de supervisión y cumplimiento. España está preparando estructuras de gobernanza que coordinan la vigilancia del mercado, supervisión de sistemas de alto riesgo y evaluación de IA utilizada en la administración pública. Esto crea un marco nacional que refleja las prioridades sociales

result is a national framework reflecting Spanish social priorities while remaining aligned with EU rules.

The Ministry for Digital Transformation plays a central role in coordinating AI policy, supported by the Spanish Agency for the Supervision of Artificial Intelligence (AESIA). AESIA is set to become one of Europe's first dedicated national AI supervisory agencies. It will review AI use, conduct enforcement actions, collaborate with European-level bodies and guide organizations toward lawful, ethical AI deployment.

Spain places strong emphasis on transparency and human rights within its implementation of the AI Act. Public institutions must document how automated systems influence decisions, provide explanations to citizens and maintain clear lines of accountability. These expectations reinforce Spain's administrative culture of openness and fairness.

High-risk AI systems deployed in Spain must meet EU standards, but deployers also face specific duties around labor rights, non-discrimination and documentation. Spain expects organizations to build transparency into communication with employees and users — especially where AI influences workplace evaluations, public benefits or social services.

Spain's implementation model demonstrates commitment not only to legal compliance, but to building a trustworthy, people-centered digital ecosystem. By integrating EU rules with national protections, Spain positions itself as a leader in ethical, socially responsible AI adoption.

españolas, manteniéndose alineado con las normas de la UE.

El Ministerio de Transformación Digital tiene un rol central en la coordinación de políticas de IA, apoyado por la Agencia Española de Supervisión de la Inteligencia Artificial (AESIA). AESIA se convertirá en una de las primeras agencias nacionales de supervisión de IA en Europa. Revisará el uso de IA, realizará acciones de cumplimiento, colaborará con organismos a nivel europeo y guiará a las organizaciones hacia un despliegue legal y ético de los sistemas de IA.

España pone un fuerte énfasis en transparencia y derechos humanos dentro de su implementación de la AI Act. Las instituciones públicas deben documentar cómo los sistemas automatizados influyen en las decisiones, proporcionar explicaciones a los ciudadanos y mantener líneas claras de rendición de cuentas. Estas expectativas refuerzan la cultura administrativa española de apertura y equidad.

Los sistemas de IA de alto riesgo desplegados en España deben cumplir los estándares de la UE, pero los implementadores también enfrentan deberes específicos respecto a derechos laborales, no discriminación y documentación. España espera que las organizaciones integren la transparencia en la comunicación con empleados y usuarios, especialmente cuando la IA influye en evaluaciones laborales, beneficios públicos o servicios sociales.

El modelo de implementación español demuestra un compromiso no solo con el cumplimiento legal, sino con la construcción de un ecosistema digital confiable y centrado en las personas. Al integrar las normas de la UE con protecciones nacionales, España se posiciona como líder en adopción ética y socialmente responsable de IA. #### Topic 2: GDPR, LOPDGDD, and National Data Protection Duties + Reflection Story / Tema 2: GDPR, LOPDGDD y Obligaciones Nacionales de Protección de Datos + Historia de Reflexión

7.2 GDPR, LOPDGDD, and National Data Protection Duties

GDPR, LOPDGDD y Obligaciones Nacionales de Protección de Datos

Spain implements the GDPR through its national data-protection law, the LOPDGDD, which reinforces and extends GDPR obligations. For AI systems, this means organizations must secure a legal basis for

España integra GDPR mediante su ley nacional de protección de datos, la LOPDGDD, que refuerza y amplía las obligaciones del GDPR. Para sistemas de IA, esto significa que las organizaciones deben

processing, protect special categories of data and follow the rules on automated decision-making. Spain emphasizes clarity, transparency and minimization: collect only what is necessary and explain plainly how data is used.

The AEPD — Spain's data-protection authority — plays a key role in supervising AI-related data processing. It monitors compliance, issues guidance, conducts investigations and can impose sanctions. Many AI failures trace back to inadequate data governance, which makes the AEPD's role critical in the Spanish landscape. AI development requires strict adherence to GDPR principles such as purpose limitation, fairness and accountability.

Both the GDPR and the LOPDGDD require organizations to provide meaningful information about automated decisions — including the logic used, the significance of the results and their possible implications. Individuals must be able to request human review of AI-driven decisions. In Spain, these rights apply especially in employment, healthcare, credit scoring and public administration.

garantizar bases legales para el procesamiento, proteger categorías especiales de datos y seguir normas para la toma de decisiones automatizada. España enfatiza claridad, transparencia y minimización: recopilar solo lo necesario y explicar claramente cómo se usan los datos.

La AEPD (Agencia Española de Protección de Datos) desempeña un papel clave en la supervisión del procesamiento de datos relacionado con IA. La AEPD monitorea el cumplimiento, emite guías, realiza investigaciones y puede imponer sanciones. Muchos fallos de IA provienen de una gobernanza inadecuada de datos, haciendo que el rol de la AEPD sea crítico en el panorama español. El desarrollo de IA requiere adherencia estricta a principios de GDPR como finalidad, equidad y responsabilidad.

Tanto GDPR como LOPDGDD requieren que las organizaciones ofrezcan información significativa sobre decisiones automatizadas, incluyendo la lógica utilizada, la relevancia de los resultados y posibles implicaciones. Los individuos deben poder solicitar revisión humana de decisiones impulsadas por IA. En España, estos derechos aplican especialmente en empleo, salud, scoring crediticio y administración pública.

CASE IN PRACTICE · CASO PRÁCTICO

A telecommunications company in Spain deployed an automated retention-risk model to predict customer cancellations. It collected extensive behavioral data without proper consent and never told users profiling was taking place. Customer complaints prompted an AEPD investigation. The company had to redesign its consent practices, minimize the data collected and implement transparency notices. Once these problems were fixed, customer trust actually rose. The system stayed in use — now aligned with legal expectations and user comfort.

Una compañía de telecomunicaciones en España desplegó un modelo automatizado de riesgo de retención para predecir cancelaciones de clientes. Recopiló extensos datos de comportamiento sin consentimiento adecuado y no informó a los usuarios que se realizaba perfilado. Ante quejas de los clientes, la AEPD inició una investigación. La compañía tuvo que rediseñar prácticas de consentimiento, minimizar datos recopilados e implementar avisos de transparencia. Después de corregir estos problemas, la confianza de los clientes aumentó. El sistema permaneció en uso, pero ahora alineado con expectativas legales y comodidad de los usuarios.

The reflection shows how solid data-protection practice supports — rather than obstructs — AI adoption. The GDPR and LOPDGDD create boundaries that protect people and strengthen an organization's ethical footing.

Esta reflexión muestra cómo las prácticas sólidas de protección de datos apoyan—y no obstaculizan—la adopción de IA. GDPR y LOPDGDD crean límites que protegen a las personas y fortalecen la base ética de las organizaciones. ### Topic 3: Spanish Labor, Workplace, and Equality Regulations + Scenario-

7.3 Spanish Labor, Workplace, and Equality Regulations

Regulación Laboral, de Igualdad y Uso de IA en el Trabajo

Spanish labor law strongly shapes how AI may be used in the workplace. Automated tools involved in recruitment, scheduling, evaluation or task assignment must respect workers' rights and avoid discriminatory outcomes. Employers must notify workers when AI contributes to decision-making and guarantee active human oversight. Workers retain the right to understand decisions that affect them.

The Workers' Statute — Spain's foundational labor law — requires transparency and prohibits unfair algorithmic decisions. New provisions under national and EU rules address algorithmic management in workplaces, covering gig platforms, traditional employers and public institutions alike. Algorithms may not erode labor protections or create opaque evaluation systems.

Equality regulation reinforces these duties. Spain's Equality Law requires organizations to prevent discrimination by gender, disability, ethnicity, age, religion, sexual orientation and socioeconomic status. AI systems used in hiring or performance evaluation must undergo fairness testing and trigger corrective action when disparities appear.

Collective bargaining agreements in Spain increasingly include algorithmic-transparency provisions. Worker representatives can request documentation on how automated systems influence workloads, schedules or evaluations. These agreements ensure technological change respects labor rights and human dignity.

Las leyes laborales españolas influyen fuertemente en cómo se puede utilizar IA en el lugar de trabajo. Herramientas automatizadas involucradas en reclutamiento, programación, evaluación o asignación de tareas deben respetar los derechos de los trabajadores y evitar resultados discriminatorios. Los empleadores deben notificar a los trabajadores cuando la IA contribuye a la toma de decisiones y garantizar supervisión humana activa. Los trabajadores conservan el derecho a entender decisiones que los afecten.

El Estatuto de los Trabajadores, ley laboral fundamental de España, exige transparencia y prohíbe decisiones algorítmicas injustas. Nuevas disposiciones bajo normas nacionales y de la UE abordan la gestión algorítmica en los lugares de trabajo, incluyendo plataformas de economía colaborativa, empleadores tradicionales e instituciones públicas. Los algoritmos no pueden disminuir las protecciones laborales ni crear sistemas de evaluación opacos.

Las regulaciones de igualdad refuerzan estas responsabilidades. La Ley de Igualdad exige que las organizaciones prevengan discriminación por género, discapacidad, etnia, edad, religión, orientación sexual y estatus socioeconómico. Los sistemas de IA usados en contratación o evaluación de desempeño deben someterse a pruebas de equidad y tomar medidas correctivas ante disparidades.

Los convenios colectivos en España incluyen cada vez más normas para la transparencia algorítmica. Los representantes de los trabajadores pueden solicitar documentación sobre cómo los sistemas automatizados influyen en cargas de trabajo, horarios o evaluaciones. Estos convenios aseguran que el cambio tecnológico respete derechos laborales y dignidad humana.

DECISION POINT · PUNTO DE DECISIÓN

A Spanish retail company uses AI to assign shifts based on predicted productivity. After a few weeks, employees over 50 receive fewer night shifts, reducing their income. What must the organization do? A) Assume the results reflect natural performance differences. B) Run a fairness review, analyze whether age bias exists and adjust the model or policies accordingly. C) Ignore concerns unless formal complaints are filed.

Una empresa minorista española utiliza IA para asignar turnos según productividad predicha. Después de unas semanas, empleados mayores de 50 años reciben menos turnos nocturnos, reduciendo sus ingresos. ¿Qué debe hacer la organización? A) Asumir que los resultados reflejan diferencias naturales de desempeño. B) Realizar revisión de equidad, analizar si existe sesgo por edad y ajustar modelo o políticas según corresponda. C) Ignorar preocupaciones a menos que se presenten quejas formales.

Correct answer: B — Spain requires corrective action to prevent discriminatory outcomes in the workplace.

Respuesta correcta: B — España exige acción correctiva para prevenir resultados discriminatorios en el trabajo.

7.4 National AI Strategy, Public Sector Rules, and Enforcement Bodies

Estrategia Nacional de IA, Normas del Sector Público y Órganos de Supervisión

Spain's National AI Strategy, España Digital 2026 and subsequent policy updates set out the country's vision for responsible AI development. They highlight ethical innovation, economic competitiveness, digital inclusion and respect for rights — emphasizing transparency, accountability and people-centered AI across public and private sectors.

Public-sector AI use is tightly regulated. Government agencies must document automated decision systems, offer explanations to citizens and maintain appeal mechanisms. Administrative-transparency laws require institutions to disclose when decisions are influenced by algorithms. Spanish public institutions are expected to set the standard for trustworthy AI.

Multiple authorities contribute to supervision. The AEPD handles data protection. AESIA oversees AI governance and market surveillance. Sector regulators — such as the CNMV in finance or the Ministry of Health — assess AI use within their domains. These bodies collaborate to ensure AI respects safety, rights and fairness.

Spain also invests heavily in education, research and innovation through AI research centers and partnerships with universities and regional governments. The goal is an ecosystem where ethical,

La Estrategia Nacional de IA de España, España Digital 2026, y actualizaciones de políticas posteriores, delinean la visión del país para un desarrollo responsable de IA. Estas estrategias destacan innovación ética, competitividad económica, inclusión digital y respeto por derechos. Enfatizan transparencia, responsabilidad y IA centrada en las personas, tanto en sector público como privado.

El uso de IA en el sector público está regulado estrictamente. Las agencias gubernamentales deben documentar sistemas automatizados de decisión, ofrecer explicaciones a los ciudadanos y mantener mecanismos de apelación. Las leyes de transparencia administrativa requieren que las instituciones divulguen cuándo las decisiones son influenciadas por algoritmos. Las instituciones públicas españolas deben marcar el estándar para IA confiable.

Múltiples autoridades contribuyen a la supervisión. La AEPD maneja protección de datos. AESIA supervisa gobernanza de IA y vigilancia del mercado. Reguladores sectoriales—como CNMV en finanzas o Ministerio de Salud—evalúan el uso de IA en sus dominios. Estos organismos colaboran para asegurar que la IA respete seguridad, derechos y equidad.

España también invierte fuertemente en educación, investigación e innovación a través de centros de investigación en IA y asociaciones con universidades

technically sound AI drives social and economic development without sacrificing legal protections.

y gobiernos regionales. El objetivo es construir un ecosistema donde la IA ética y técnicamente sólida impulse desarrollo social y económico sin sacrificar protecciones legales.

APPLIED EXAMPLE · EJEMPLO APLICADO

A US company providing AI traffic-management solutions expanded into Spain. Before deployment, it had to submit documentation to local authorities, explain the algorithm's criteria and demonstrate transparency and public-accountability mechanisms. Because Spanish public-sector rules protect citizens' rights, the company added real-time explanation features and built oversight dashboards for municipal staff. The project succeeded because it aligned with Spain's expectations of administrative transparency.

Una empresa estadounidense proveedora de soluciones de IA para gestión de tráfico se expandió a España. Antes del despliegue, debía presentar documentación a autoridades locales, explicar criterios del algoritmo y demostrar mecanismos de transparencia y rendición de cuentas pública. Debido a que las normas del sector público español protegen derechos de los ciudadanos, la empresa añadió funciones de explicación en tiempo real y creó paneles de supervisión para el personal municipal. El proyecto tuvo éxito porque se alineó con expectativas de transparencia administrativa españolas. ### Module Wrap-Up + Practical Life Implementation / Cierre del Módulo + Implementación Práctica

Module Wrap-Up & Practical Implementation

Cierre del Módulo e Implementación Práctica

You have now seen how Spain integrates EU requirements with its own legal commitments to rights, equality, transparency and oversight. Spain's regulatory approach reinforces the idea that AI systems must serve people and protect their dignity. These frameworks guide every step of AI deployment — from data collection through decision-making to public accountability.

The reflection story and the workplace scenario showed that legal compliance and ethical responsibility operate together. Mistakes can begin quietly — unclear consent, biased datasets, opaque workplace tools — but Spanish regulation provides correction paths that restore trust and fairness. Strong compliance improves both organizational and social outcomes.

Identify an AI system that influences people inside your organization and assess whether it meets Spanish legal expectations. Ask whether data practices align with the GDPR and LOPDGDD, whether workplace uses respect equality rules and whether

Ahora has visto cómo España integra los requisitos de la UE con sus propios compromisos legales respecto a derechos, igualdad, transparencia y supervisión. El enfoque regulatorio español refuerza la idea de que los sistemas de IA deben servir a las personas y proteger su dignidad. Estos marcos guían cada paso del despliegue de IA, desde la recopilación de datos hasta la toma de decisiones y la rendición de cuentas pública.

La historia de reflexión y el escenario laboral mostraron que el cumplimiento legal y la responsabilidad ética operan de manera conjunta. Los errores pueden comenzar de forma silenciosa— consentimiento poco claro, conjuntos de datos sesgados, herramientas laborales opacas—pero las regulaciones españolas proporcionan caminos de corrección que restauran confianza y equidad. Un cumplimiento sólido mejora tanto resultados organizacionales como sociales.

Identifica un sistema de IA que influya en personas dentro de tu organización y evalúa si cumple las

transparency measures are meaningful. Every improvement strengthens ethical leadership and prepares your organization for an AI-driven future in Spain.

expectativas legales españolas. Pregúntate si las prácticas de datos se alinean con GDPR y LOPDGDD, si los usos laborales respetan reglas de igualdad y si las medidas de transparencia son significativas. Cada mejora fortalece el liderazgo ético y prepara a tu organización para un futuro impulsado por IA en España.

MODULE 7 ASSESSMENT · 10 QUESTIONS

Test Your Understanding

Evalúa tu comprensión — Respuestas en el Apéndice A

1. Spain's implementation of the EU AI Act is coordinated at which level?

La implementación en España del Reglamento de IA de la UE se coordina a nivel:

- A. Municipal / Municipal
- B. Autonomous-community only / Solo de la comunidad autónoma
- C. National / Nacional
- D. Private sector / Sector privado

2. The LOPDGDD complements which EU regulation?

La LOPDGDD complementa cuál regulación de la UE:

- A. Digital Services Act / Ley de Servicios Digitales
- B. EU AI Act / Reglamento de IA de la UE
- C. GDPR / GDPR
- D. Data Governance Act / Ley de Gobernanza de Datos

3. Spanish labor laws affect the use of AI in:

Las leyes laborales españolas afectan el uso de la IA en:

- A. Marketing analytics / Análisis de marketing
- B. Workplace monitoring / Supervisión en el lugar de trabajo
- C. Online gaming / Juegos en línea
- D. Product design / Diseño de productos

4. Spain's equality regulations seek to prevent:

Las regulaciones de igualdad en España buscan prevenir:

- A. Data leaks / Filtraciones de datos
- B. Algorithmic discrimination / Discriminación algorítmica
- C. Market dominance / Dominio del mercado
- D. AI automation / Automatización de IA

5. Spain's National AI Strategy promotes:

La Estrategia Nacional de IA de España promueve:

- A. Unregulated innovation / Innovación no regulada
- B. Ethical, human-centric AI / IA ética y centrada en las personas
- C. Military AI systems / Sistemas de IA militares
- D. Full automation / Automatización completa

6. Public-sector AI systems in Spain are subject to:

Los sistemas de IA del sector público en España están sujetos a:

- A. Voluntary guidelines only / Solo directrices voluntarias
- B. No oversight / Sin supervisión
- C. Reinforced transparency obligations / Obligaciones de transparencia reforzadas
- D. Private certification / Certificación privada

7. Spanish authorities can impose penalties under the EU AI Act.

Las autoridades españolas pueden imponer sanciones bajo el Reglamento de IA de la UE.

- A. True / Verdadero
- B. False / Falso
- C. Only EU bodies / Solo organismos de la UE
- D. Only courts / Solo tribunales

8. GDPR obligations are replaced by the EU AI Act in Spain.

Las obligaciones del GDPR son reemplazadas por el Reglamento de IA de la UE en España.

- A. True / Verdadero
- B. False / Falso
- C. Only for AI data / Solo para datos de IA
- D. Only for the public sector / Solo para el sector público

9. Spain's data protection authority is the _____.

La autoridad de protección de datos de España es la _____.

- A. CNMC / CNMC
- B. AEPD / AEPD
- C. INCIBE / INCIBE
- D. AESIA / AESIA

10. Spain established _____ as its dedicated AI supervisory authority.

España estableció a _____ como autoridad supervisora dedicada a la IA.

- A.** CNMV / CNMV
- B.** AEPD / AEPD
- C.** AESIA / AESIA
- D.** SEPE / SEPE



FINAL CERTIFICATION EXAM · EXAMEN FINAL

Final Certification Exam

Examen Final de Certificación — 30 preguntas

- **Thirty questions spanning all seven modules**
Treinta preguntas que abarcan los siete módulos · Answers in Appendix A / Respuestas en el Apéndice A

Final Certification Exam

Examen Final de Certificación · Answers in Appendix A / Respuestas en el Apéndice A

1. The main objective of the EU AI Act is to:

El objetivo principal del Reglamento de IA de la UE es:

- A.** Promote unrestricted AI innovation / Promover la innovación en IA sin restricciones
- B.** Harmonize legal requirements for AI risk management / Armonizar los requisitos legales para la gestión de riesgos de IA
- C.** Replace the GDPR / Reemplazar el GDPR
- D.** Regulate only European AI companies / Regular solo a empresas europeas de IA

2. High-risk AI systems are classified based on:

Los sistemas de IA de alto riesgo se clasifican en función de:

- A.** Model size / Tamaño del modelo
- B.** Intended use and impact on fundamental rights / Uso previsto e impacto en los derechos fundamentales
- C.** Country of origin / País de origen
- D.** Development cost / Costo de desarrollo

3. Which role ensures an AI system is used according to its intended purpose?

¿Qué rol asegura que un sistema de IA se utilice según su propósito previsto?

- A.** Provider / Proveedor
- B.** Distributor / Distribuidor
- C.** Deployer / Implementador (Deployer)
- D.** Importer / Importador

4. CE marking on high-risk AI indicates:

La marca CE en IA de alto riesgo indica:

- A.** Ethical approval / Aprobación ética
- B.** Compliance with EU requirements / Cumplimiento con los requisitos de la UE
- C.** Market preference / Preferencia de mercado
- D.** Public endorsement / Respaldo público

5. Transparency obligations apply mainly to:

Las obligaciones de transparencia se aplican principalmente a:

- A.** High-risk AI / IA de alto riesgo
- B.** Minimal-risk AI / IA de riesgo mínimo
- C.** Limited-risk AI / IA de riesgo limitado
- D.** Prohibited AI / IA prohibida

6. Systemic general-purpose AI requires:

La IA de propósito general sistémica requiere:

- A.** Advanced risk testing / Pruebas avanzadas de riesgo
- B.** Voluntary best practice / Mejores prácticas voluntarias
- C.** No monitoring / Sin monitoreo
- D.** CE marking only / Solo marca CE

7. All AI systems face the same compliance requirements under the EU AI Act.

Todos los sistemas de IA están sujetos a los mismos requisitos de cumplimiento según el Reglamento de IA de la UE.

- A.** True / Verdadero
- B.** False / Falso
- C.** Only in Spain / Solo en España
- D.** Only for public authorities / Solo para autoridades públicas

8. Spanish authorities can impose penalties under the EU AI Act.

Las autoridades españolas pueden imponer sanciones bajo el Reglamento de IA de la UE.

- A.** True / Verdadero
- B.** False / Falso
- C.** Only EU authorities / Solo autoridades de la UE
- D.** Only courts / Solo tribunales

9. The EU AI Act uses a _____ approach to regulating AI systems.

El Reglamento de IA de la UE utiliza un enfoque _____ para regular los sistemas de IA.

- A.** Risk-based / Basado en riesgos
- B.** Technology-specific / Específico de tecnología
- C.** Market-driven / Impulsado por el mercado
- D.** Voluntary / Voluntario

10. Spain's data protection authority is called _____.

La autoridad de protección de datos de España se llama _____.

- A.** CNMV / CNMV
- B.** AEPD / AEPD
- C.** AESIA / AESIA
- D.** SEPE / SEPE

11. Which AI practice is considered prohibited under the EU AI Act?

¿Qué práctica de IA se considera prohibida según el Reglamento de IA de la UE?

- A.** Emotion recognition for law enforcement / Reconocimiento de emociones para aplicación de la ley
- B.** Conversational AI / IA conversacional
- C.** Recommendation systems / Sistemas de recomendación
- D.** Grammar checkers / Correctores gramaticales

12. Bias-control measures in high-risk AI systems are:

Las medidas de control de sesgos en sistemas de IA de alto riesgo son:

- A.** Mandatory / Obligatorias
- B.** Optional / Opcionales
- C.** Only for public AI / Solo para IA pública
- D.** Voluntary in Spain / Voluntarias en España

13. Watermarking AI-generated content helps to:

La marca de agua en contenido generado por IA ayuda a:

- A.** Improve AI efficiency / Mejorar la eficiencia de la IA
- B.** Identify AI-generated content / Identificar contenido generado por IA
- C.** Reduce AI training costs / Reducir costos de entrenamiento de IA
- D.** Promote open-source AI / Promover IA de código abierto

14. Human oversight in AI systems ensures:

La supervisión humana en sistemas de IA asegura:

- A.** Replacement of AI decisions / Reemplazo de decisiones de IA
- B.** The ability to intervene and override outputs / Capacidad de intervenir y anular resultados
- C.** Full automation / Automatización completa
- D.** Faster processing / Procesamiento más rápido

15. AI non-compliance penalties are calculated based on:

Las sanciones por incumplimiento de IA se calculan en función de:

- A.** Net profit / Beneficio neto
- B.** Company size / Tamaño de la empresa
- C.** Global annual turnover / Facturación anual global
- D.** Local revenue / Ingresos locales

16. Transparency, explainability and trustworthiness are pillars of:

La transparencia, explicabilidad y confiabilidad son pilares de:

- A.** AI technical standards / Normas técnicas de IA
- B.** Ethical AI principles / Principios éticos de IA
- C.** Market surveillance / Vigilancia de mercado
- D.** Minimal-risk AI / IA de riesgo mínimo

17. Spanish labor laws regulate the use of AI in:

Las leyes laborales españolas regulan el uso de IA en:

- A.** Marketing analytics / Análisis de marketing
- B.** Workplace monitoring / Supervisión en el lugar de trabajo
- C.** Entertainment AI / IA para entretenimiento
- D.** Video games / Videojuegos

18. Required documentation for high-risk AI includes:

La documentación requerida para IA de alto riesgo incluye:

- A.** A marketing plan / Plan de marketing
- B.** Technical documentation / Documentación técnica
- C.** User satisfaction reports / Informes de satisfacción del usuario
- D.** Competitor analysis / Análisis de la competencia

19. Limited-risk AI is subject to:

La IA de riesgo limitado está sujeta a:

- A.** Mandatory fines / Multas obligatorias
- B.** Voluntary transparency obligations / Obligaciones de transparencia voluntarias
- C.** CE marking / Marca CE
- D.** Prohibition / Prohibición

20. Contractual controls in AI systems aim to:

Los controles contractuales en sistemas de IA buscan:

- A.** Reduce AI accuracy / Reducir la precisión de la IA
- B.** Allocate responsibilities between parties / Asignar responsabilidades entre las partes
- C.** Avoid compliance / Evitar el cumplimiento
- D.** Automate decisions / Automatizar decisiones

21. Ethical AI governance focuses solely on legal compliance.

La gobernanza ética de la IA se centra únicamente en el cumplimiento legal.

- A.** True / Verdadero
- B.** False / Falso
- C.** Only for high-risk AI / Solo para IA de alto riesgo
- D.** Only in Spain / Solo en España

22. High-risk AI systems require post-market monitoring.

Los sistemas de IA de alto riesgo requieren monitoreo posterior a la comercialización.

- A.** True / Verdadero
- B.** False / Falso
- C.** Only in Spain / Solo en España
- D.** Only voluntary / Solo voluntario

23. AI systems posing unacceptable risk are _____ under the EU AI Act.

Los sistemas de IA que representan un riesgo inaceptable están _____ según el Reglamento de IA de la UE.

- A.** Restricted / Restringidos
- B.** Certified / Certificados
- C.** Prohibited / Prohibidos
- D.** Recommended / Recomendados

24. Oversight mechanisms ensure _____ use of AI systems.

Los mecanismos de supervisión aseguran un uso _____ de los sistemas de IA.

- A.** Responsible / Responsable
- B.** Unrestricted / Sin restricciones
- C.** Experimental / Experimental
- D.** Commercial / Comercial

25. General-purpose AI models are designed to:

Los modelos de IA de propósito general están diseñados para:

- A.** Solve one specific task / Resolver una tarea específica
- B.** Adapt to many downstream tasks / Adaptarse a múltiples tareas derivadas
- C.** Replace all human jobs / Reemplazar todos los trabajos humanos
- D.** Avoid regulation / Evitar regulación

26. Civil liability in AI mainly addresses:

La responsabilidad civil en IA aborda principalmente:

- A.** Harm caused to individuals / Daños causados a individuos
- B.** Marketing errors / Errores de marketing
- C.** Competitive behavior / Comportamiento competitivo
- D.** Ethical conflicts only / Conflictos éticos únicamente

27. Spain's National AI Strategy promotes:

La Estrategia Nacional de IA de España promueve:

- A.** Military AI / IA militar
- B.** Ethical, human-centric AI / IA ética y centrada en las personas
- C.** Full automation / Automatización completa
- D.** Unregulated innovation / Innovación no regulada

28. High-risk AI systems must undergo:

Los sistemas de IA de alto riesgo deben someterse a:

- A.** Marketing assessments / Evaluaciones de marketing
- B.** Conformity assessments / Evaluaciones de conformidad
- C.** Voluntary disclosure / Divulgación voluntaria
- D.** Social media evaluation / Evaluación en redes sociales

29. Transparency obligations for generative AI require disclosure when:

Las obligaciones de transparencia para IA generativa requieren divulgación cuando:

- A.** AI reduces costs / La IA reduce costos
- B.** AI content is generated for users / El contenido de IA se genera para los usuarios
- C.** AI improves performance / La IA mejora el rendimiento
- D.** AI is open source / La IA es de código abierto

30. Organizational AI governance ensures:

La gobernanza organizacional de IA asegura:

- A.** Random oversight / Supervisión aleatoria
- B.** Clear roles, accountability and ethical compliance / Roles claros, responsabilidad y cumplimiento ético
- C.** External oversight only / Supervisión solo externa
- D.** Full AI autonomy / Autonomía completa de la IA

From Regulation to Practice

De la regulación a la práctica

ENGLISH

You began this book with a regulation that may have felt distant and technical. You finish it holding a complete operating model: four risk tiers that classify any system, four legal roles that assign every obligation, a documentation and oversight discipline that survives audits, and the Spanish institutional map — AEPD, AESIA, the Workers' Statute, the Equality Law — that turns European principle into local practice.

Compliance, done well, is not a defensive posture. As the cases in these pages showed again and again, the organizations that classify honestly, document thoroughly and oversee actively are the ones that keep public trust when others lose it. That is the real certification this course offers: not a document, but a way of working.

Complete your module assessments and the final exam on the course platform to earn your certificate — and keep this book close. The Act will evolve, guidance will accumulate, and the professional who understands the structure beneath the rules will always read the changes faster than the one who memorized them.

ESPAÑOL

Comenzaste este libro con un reglamento que quizá parecía distante y técnico. Lo terminas con un modelo operativo completo: cuatro niveles de riesgo que clasifican cualquier sistema, cuatro roles legales que asignan cada obligación, una disciplina de documentación y supervisión que resiste auditorías, y el mapa institucional español — AEPD, AESIA, el Estatuto de los Trabajadores, la Ley de Igualdad — que convierte el principio europeo en práctica local.

El cumplimiento, bien hecho, no es una postura defensiva. Como mostraron una y otra vez los casos de estas páginas, las organizaciones que clasifican con honestidad, documentan con rigor y supervisan activamente son las que conservan la confianza pública cuando otras la pierden. Esa es la verdadera certificación que ofrece este curso: no un documento, sino una forma de trabajar.

Completa las evaluaciones de módulo y el examen final en la plataforma del curso para obtener tu certificado — y mantén este libro cerca. La Ley evolucionará, las guías se acumularán, y el profesional que entiende la estructura detrás de las reglas siempre leerá los cambios más rápido que quien solo las memorizó.

Answer Key

Clave de Respuestas

Module 1 — Introduction to EU AI Regulation · Módulo 1

Module	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10
Answers	C	B	C	B	D	C	B	B	B	C

Module 2 — EU AI Act Risk Categories · Módulo 2

Module	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10
Answers	B	C	B	C	D	C	B	B	D	B

Module 3 — High-Risk AI System Requirements · Módulo 3

Module	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10
Answers	B	C	D	C	C	B	A	B	C	B

Module 4 — Governance of General-Purpose and Generative AI · Módulo 4

Module	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10
Answers	C	C	B	C	C	C	B	A	B	C

Module 5 — Enforcement, Risk, and Liability Frameworks · Módulo 5

Module	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10
Answers	A	C	C	C	B	A	B	A	C	C

Module 6 — Ethical AI Principles and Governance Structures · Módulo 6

Module	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10
Answers	B	B	B	C	C	B	A	B	C	B

Module 7 — Spanish AI Legal and Regulatory Framework · Módulo 7

Module	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10
Answers	C	C	B	B	B	C	A	B	B	C

Final Exam — Questions 1–15 · Examen Final 1–15

Module	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Q9	Q10	Q11	Q12	Q13	Q14	Q15
Answers	B	B	C	B	C	A	B	A	A	B	A	A	B	B	C

Final Exam — Questions 16–30 · Examen Final 16–30

Module	Q16	Q17	Q18	Q19	Q20	Q21	Q22	Q23	Q24	Q25	Q26	Q27	Q28	Q29	Q30
Answers	B	B	B	B	B	B	A	C	A	B	A	B	B	B	B

Score each module assessment out of 10 and the final exam out of 30. A score of 70% or higher on the course platform is required for certification. · Puntúa cada evaluación de módulo sobre 10 y el examen final sobre 30. Se requiere un 70% o más en la plataforma del curso para la certificación.

Glossary of Key Terms

Glosario de Términos Clave

AI system / Sistema de IA

A machine-based system that generates outputs such as predictions, recommendations or decisions influencing real or virtual environments.

Sistema basado en máquinas que genera resultados como predicciones, recomendaciones o decisiones que influyen en entornos reales o virtuales.

Provider / Proveedor

The entity that develops an AI system or places it on the EU market under its own name; carries the heaviest obligations.

La entidad que desarrolla un sistema de IA o lo introduce en el mercado de la UE bajo su propio nombre; asume las obligaciones más exigentes.

Deployer / Desplegador

The entity using an AI system within its operations; responsible for lawful use, oversight and monitoring.

La entidad que utiliza un sistema de IA en sus operaciones; responsable del uso legal, la supervisión y el monitoreo.

Importer / Importador

The entity bringing AI systems from outside the EU into the European market; must verify EU conformity first.

La entidad que introduce sistemas de IA desde fuera de la UE en el mercado europeo; debe verificar primero la conformidad con la UE.

Distributor / Distribuidor

The entity making an AI system available on the market; must not supply systems that appear non-compliant.

La entidad que comercializa un sistema de IA; no debe suministrar sistemas que parezcan no conformes.

Unacceptable risk / Riesgo inaceptable

Practices banned outright — social scoring, harmful manipulation, exploitation of vulnerable groups, untargeted biometric scraping.

Prácticas prohibidas de plano: puntuación social, manipulación dañina, explotación de grupos vulnerables, captura biométrica indiscriminada.

High-risk AI / IA de alto riesgo

Systems listed in Annex III or embedded in regulated products; subject to strict obligations before and after market entry.

Sistemas listados en el Anexo III o integrados en productos regulados; sujetos a obligaciones estrictas antes y después de entrar al mercado.

Limited risk / Riesgo limitado

Systems that interact with people (chatbots, synthetic media); must clearly disclose their AI nature.

Sistemas que interactúan con personas (chatbots, medios sintéticos); deben revelar claramente su naturaleza de IA.

Minimal risk / Riesgo mínimo

The majority of AI applications; permitted without specific obligations, voluntary best practices encouraged.

La mayoría de las aplicaciones de IA; permitidas sin obligaciones específicas, con buenas prácticas voluntarias recomendadas.

GPAI

General-purpose AI — models built for many tasks and integrable into many downstream systems; systemic GPAI carries extra testing and monitoring duties.

IA de propósito general: modelos construidos para muchas tareas e integrables en múltiples sistemas derivados; la GPAI sistémica conlleva deberes adicionales de pruebas y monitoreo.

Conformity assessment / Evaluación de conformidad

The audit process verifying a high-risk system meets all legal requirements before market entry.

El proceso de auditoría que verifica que un sistema de alto riesgo cumple todos los requisitos legales antes de entrar al mercado.

CE marking / Mercado CE

The mark signaling EU-wide conformity; carries ongoing post-market monitoring and incident-reporting duties, and can be withdrawn.

La marca que señala conformidad en toda la UE; conlleva deberes continuos de monitoreo post-mercado y notificación de incidentes, y puede retirarse.

Human oversight / Supervisión humana

Meaningful human control over AI decisions — the ability to intervene, correct or override outputs.

Control humano significativo sobre las decisiones de IA: la capacidad de intervenir, corregir o anular resultados.

Watermarking / Marca de agua

Technical mechanisms that make AI-generated content identifiable to platforms, regulators and users.

Mecanismos técnicos que hacen identificable el contenido generado por IA para plataformas, reguladores y usuarios.

GDPR / RGPD

The EU's General Data Protection Regulation; fully applicable alongside the AI Act wherever personal data is processed.

El Reglamento General de Protección de Datos de la UE; plenamente aplicable junto a la Ley de IA cuando se procesan datos personales.

LOPDGDD

Spain's national data-protection law, reinforcing and extending the GDPR — including strengthened workers' digital rights.

La ley española de protección de datos, que refuerza y amplía el RGPD — incluidos los derechos digitales reforzados de los trabajadores.

AEPD

Agencia Española de Protección de Datos — Spain's data-protection authority, supervising AI-related data processing.

Agencia Española de Protección de Datos: la autoridad española de protección de datos, que supervisa el tratamiento de datos vinculado a la IA.

AESIA

Agencia Española de Supervisión de la Inteligencia Artificial — Spain's dedicated AI supervisory agency, among the first in Europe.

Agencia Española de Supervisión de la Inteligencia Artificial: la agencia supervisora de IA de España, una de las primeras de Europa.

Annex III / Anexo III

The AI Act's list of high-risk use cases: employment, credit, education, essential services, law enforcement and more.

La lista de casos de uso de alto riesgo de la Ley de IA: empleo, crédito, educación, servicios esenciales, aplicación de la ley y más.

Workers' Statute / Estatuto de los Trabajadores

Spain's foundational labor law; requires transparency toward workers and prohibits unfair algorithmic decisions.

La ley laboral fundamental de España; exige transparencia hacia los trabajadores y prohíbe decisiones algorítmicas injustas.

Continue Your Certification

Continúa tu certificación

Complete the module assessments and the final exam on the course platform to earn your certificate from the Spanish Compliance Institute.

Completa las evaluaciones de módulo y el examen final en la plataforma del curso para obtener tu certificado del Spanish Compliance Institute.

SPANISHCOMPLIANCEINSTITUTE.COM