



Secure Journey to AI

Telefónica Tech's roadmap to adopt AI with security, governance, and compliance



telefonicatech.com

Index

Introduction	3
AI context in business environments	4
Regulatory considerations and responsibility in AI	4
Key challenges in AI adoption.....	6
Governance and compliance	6
Network and security management	6
Continuous monitoring and improvement	6
The journey towards secure AI	8
Identifying AI risks	8
Consulting for AI deployment.....	8
Identity, data, and security audit	9
Risk discovery	9
AI model security	10
Main risks of Generative AI	11
Protection against threats.....	11
Creating resilient applications.....	12
Data and identity protection (IAM).....	12
Risk mitigation in AI architecture	13
Security in high-level generative AI architecture.....	13
Best practices for protecting data in AI.....	14
360° Response	15
End-to-end visibility and security governance.....	15
Integration with Telefónica Tech's specialized SOC	16
Incident management and continuous improvement.....	16
Security as a living process.....	17
Conclusions and next steps	18
Conclusions	18
Immediate actions.....	18
Consolidating secure and responsible AI.....	19
Why Telefónica Tech	19

Introduction

Artificial Intelligence (AI) is a key driver of transformation for businesses of all sizes, boosting efficiency, personalization, and competitiveness across several industries. Its ability to analyze data, automate processes, and drive innovation makes it a cross-cutting technology with a direct impact on business operations.

In this context, Generative AI represents a qualitative leap: it expands the creative potential of technology and democratizes the use of AI throughout the company. The adoption of AI, however, increases known risks in the areas of privacy, security, regulatory compliance, and governance, and introduces new ones, such as those related to bias and model poisoning.

Telefónica Tech proposes the **Secure Journey to AI** framework to address these challenges. This comprehensive strategy enables companies to harness the potential of AI in a secure and responsible way that is aligned with their business objectives.

Our proposal is based on three key pillars:

- **Risk identification:** early detection of vulnerabilities in data, identities, infrastructure, and models.
- **Protection against threats:** implementation of robust measures to ensure resilient infrastructure, secure applications, and effective access and data control.
- **360° response:** continuous monitoring, integration with a specialized SOC, and a framework for continuous improvement in the event of incidents.

In this paper, we will analyze the business context of AI and the challenges and opportunities of its adoption with a focus on security: our strategic recommendations and principles of responsibility are designed to ensure that AI is implemented in a secure, governed, and compliant way, always aligned with business goals

With Telefónica Tech as their technology partner, companies can turn AI into a reliable driver of transformation, one that fuels innovation without compromising security or business continuity. Our Secure Journey to AI framework ensures that its adoption delivers real value, operational resilience, and a sustainable competitive edge.

Companies that integrate AI securely will transform regulation into trust and trust into competitive advantage.

AI context in business environments

The presence of AI in companies is unquestionable, and its evolution towards Generative AI has brought about a significant change. This new technology enhances operational efficiency and introduces creative and advanced automation capabilities. It allows people from different areas to access innovative tools and participate in the digital transformation of their companies.

This qualitative leap brings with it a paradox that is difficult to ignore. On the one hand, companies see AI as an unprecedented opportunity to innovate, accelerate their digital transformation, and open up new business models. On the other hand, they face risks: loss of traceability in results, amplified biases in decisions, accidental exposure of sensitive information, or attacks designed to manipulate AI data and models.

In this scenario, the question is no longer whether companies will incorporate AI into their operations, but how to navigate this journey safely and responsibly. Governing technology, anticipating risks, and aligning its use with business values and objectives will be what makes the difference between harnessing its potential or succumbing to its threats. To achieve this, companies need a comprehensive framework that provides security, regulatory compliance, and governance from the very first step.

AI is no longer a promise for the future: it is the engine that today is transforming industries and redefining how companies compete, innovate, and engage with society.

Regulatory considerations and responsibility in AI

The adoption of AI in a business environment cannot be understood in isolation from security and trust. Whenever an AI model processes sensitive information, such as personal data, medical records, or financial transactions, a regulatory and accountability framework comes into play that conditions its design and management.

The European Union's Artificial Intelligence Act establishes a comprehensive framework for classifying AI systems based on their risk and requiring transparency, traceability, and specific governance measures. Along with this, the General Data Protection Regulation (GDPR) safeguards the privacy of individuals.

Sector-specific regulations such as HIPAA in healthcare and PCI-DSS in finance also impose strict obligations to protect the confidentiality of medical information and payment and bank card transactions. These frameworks emphasize that the value of AI is only sustainable if it is based on data that is protected and treated with rigor and responsibility.

Infrastructure security is just as critical as privacy. International standards such as ISO/IEC 27001 help establish information security management systems, strengthening controls throughout the AI lifecycle. At the regulatory level, the NIS2 directive extends Cyber Security and digital resilience requirements in

essential sectors, while DORA (Digital Operational Resilience Act) establishes specific requirements for the financial sector to withstand and recover from technological incidents and cyberattacks.

Compliance with these regulations is not only a legal obligation, but also a prerequisite for building trust among customers, regulators, and society. This ensures that AI is used in a manner consistent with social values and business objectives, avoiding negative impacts on fundamental rights and on business operations and results.

Governance and compliance are not an add-on at the end, but the foundation that enables companies to adopt AI in a secure, responsible way that is compliant with regulations and aligned with their strategic objectives.



Key challenges in AI adoption

The path to successful AI adoption has its challenges. Companies face a number of recurring challenges that, if not managed strategically, can compromise both the effectiveness of projects and the trust of customers, partners, employees, and regulators.

These challenges encompass technical, organizational, and ethical dimensions and require a proactive approach that combines governance, security, and continuous adaptability.

Governance and compliance

AI governance remains one of the biggest challenges for companies. It is not enough to deploy advanced solutions: clear frameworks must be established that define responsibilities, assign specific roles, and ensure process traceability.

Training employees in the responsible use of AI is equally critical, as ignorance or lack of training can open the door to errors, biases, or misuse. A good governance model ensures that AI complies with current regulations and is also aligned with the company's values and strategic objectives.

Network and security management

Security has become a central concern as AI and Generative AI models connect to more data, applications, and infrastructure. Threats are evolving rapidly: from attacks such as prompt injection to information leaks or model poisoning attempts capable of manipulating the integrity of systems.

Given this landscape, protecting models and associated data requires redesigning architecture, strengthening access controls, and implementing advanced Cyber Security solutions.

Only then can companies reduce their attack surface and safeguard confidence in the results generated by their AI systems.

Continuous monitoring and improvement

AI is not static. Regulatory frameworks, technologies, and risks are constantly evolving, forcing companies to maintain a focus on ongoing monitoring and continuous improvement. This involves regularly auditing policies, reviewing processes, and updating controls in line with technological and regulatory changes.

Oversight should not be interpreted as a bureaucratic burden, but rather as an opportunity to strengthen the resilience of systems and adapt them to new realities. A cycle of continuous improvement turns security into a strategic asset, capable of anticipating threats and ensuring the long-term sustainability of projects.

These challenges should not be seen as obstacles to innovation, but as necessary conditions for AI to contribute to competitiveness in a sustained manner. Overcoming them allows companies to ensure resilience, mitigate risks, and strengthen the trust of customers, strategic partners, and regulators, consolidating an environment in which AI can responsibly unleash its full potential. Ataques frecuentes a la IA

Loss of information	• Privacy attacks: These include the extraction of sensitive information from training data or the context of use. ○ Data reconstruction: reconstruction of private information from responses. ○ Membership inference: determining whether an individual was part of the training set. ○ Model extraction: replicating model architecture and parameters through queries. ○ Prompt/context stealing: retrieving system instructions or confidential context.
Prompt injection	• Prompt injection. Hidden instructions in the prompt to alter model behavior, filter sensitive information, or bypass security controls. • Jailbreaks. Manipulation of an AI model to bypass its security restrictions and generate inappropriate or prohibited responses
Risk responses	• Generation of erroneous or inappropriate information. Includes false responses (hallucinations) and biased or malicious content that can affect reputation or decision-making.
Poisoning	• Data poisoning. Manipulation of training data to introduce biases, backdoors, or vulnerabilities that compromise the security or reliability of the model. • Model poisoning. Direct alteration of model parameters (e.g., in federated learning), degrading its performance or forcing manipulated outputs.

Access risk	• Unauthorized access. Misuse to extract, modify, or delete critical information.
Others	• Supply chain attacks. Attacks on the AI supply chain (datasets, libraries, pre-trained models) to introduce malicious components into business applications.

The journey towards secure AI

Identifying AI risks

As we have seen, the adoption of AI poses specific risks that, if not properly managed, can compromise the security, privacy, operational resilience, and regulatory compliance of companies.

This section addresses the first pillar of our Secure Journey to AI strategy: risk identification.

The purpose is to detect vulnerabilities and threats that could affect both the development and use of AI at an early stage, enabling safe and responsible deployment. This requires a comprehensive view that encompasses everything from technological infrastructure to data, identities, models, and applications.

Early identification of AI risks allows for the construction of models of trust and resilience.

Consulting for AI deployment

It is essential to design secure architectures from the planning stage before implementing AI solutions. Specialized consulting allows security to be incorporated as a strategic enabler rather than an add-on.

This involves:

- Evaluate the technological infrastructure to ensure that it supports AI in a secure and scalable manner.
- Design and analyze reference architectures that integrate Cyber Security controls from the outset.
- Analyze the risks related to the supply chain of models and tools.
- Apply the security by design approach, which is essential for anticipating threats before they arise, reducing costs and vulnerabilities.

A deployment that omits the review of third-party software dependencies, for example, can open the door to supply chain attacks, impacting the integrity of the models.

Integrating security at these early stages avoids cost overruns, ensures trust, and facilitates compliance with regulations such as the European AI Act or the NIS2 directive, reducing exposure to penalties, breaches, and reputational crises.

Identity, data, and security audit

Data, identities, and environment security are at the core of AI. If access, permissions, and configurations are not managed correctly, information leaks or unauthorized access can compromise both models and sensitive user information.

- The audit should include the evaluation of permissions and role segmentation applying the principle of least privilege, along with mechanisms such as multi-factor authentication (MFA) or conditional access. This ensures that each user only accesses the information that is strictly necessary, reducing the exposure surface.
- It is important to classify and label sensitive data, apply data leak prevention (DLP) policies, and establish continuous log review to identify misuse or unauthorized use. This protects both personal information and the company's intellectual property.
- It is essential to conduct a comprehensive security audit as an initial diagnosis, providing organizations with a clear and prioritized risk map before scaling up AI adoption.

Poor access design in a hospital, for example, could expose medical records to unauthorized users, with serious legal and reputational consequences.

Robust identity, data, and security controls therefore help ensure that AI systems are used appropriately, protecting both the company's intellectual property and user privacy. This is in line with regulations such as the GDPR on personal data, HIPAA in the healthcare sector, and NIS2 on critical infrastructure protection and digital resilience.

Risk discovery

Early risk identification requires the application of advanced methodologies and tools to detect vulnerabilities before they are exploited.

- Vulnerability management (VRM) is a continuous process of identifying, assessing, prioritizing, and remediating vulnerabilities in the company's systems, applications, and infrastructure.
- Attack surface analysis reveals vulnerabilities in models and applications.
- Shadow AI assessments help uncover unauthorized uses of AI tools within the company.
- Risk prioritization using criticality systems allows resources to be focused on threats with the greatest potential impact.

In sectors such as banking, identifying unauthorized use of AI applications by employees outside of company-approved environments can mean the difference between a minor incident and a massive leak of sensitive data.

Adopting a proactive discovery approach turns security into a continuous process, capable of anticipating incidents and ensuring confident and controlled AI adoption.

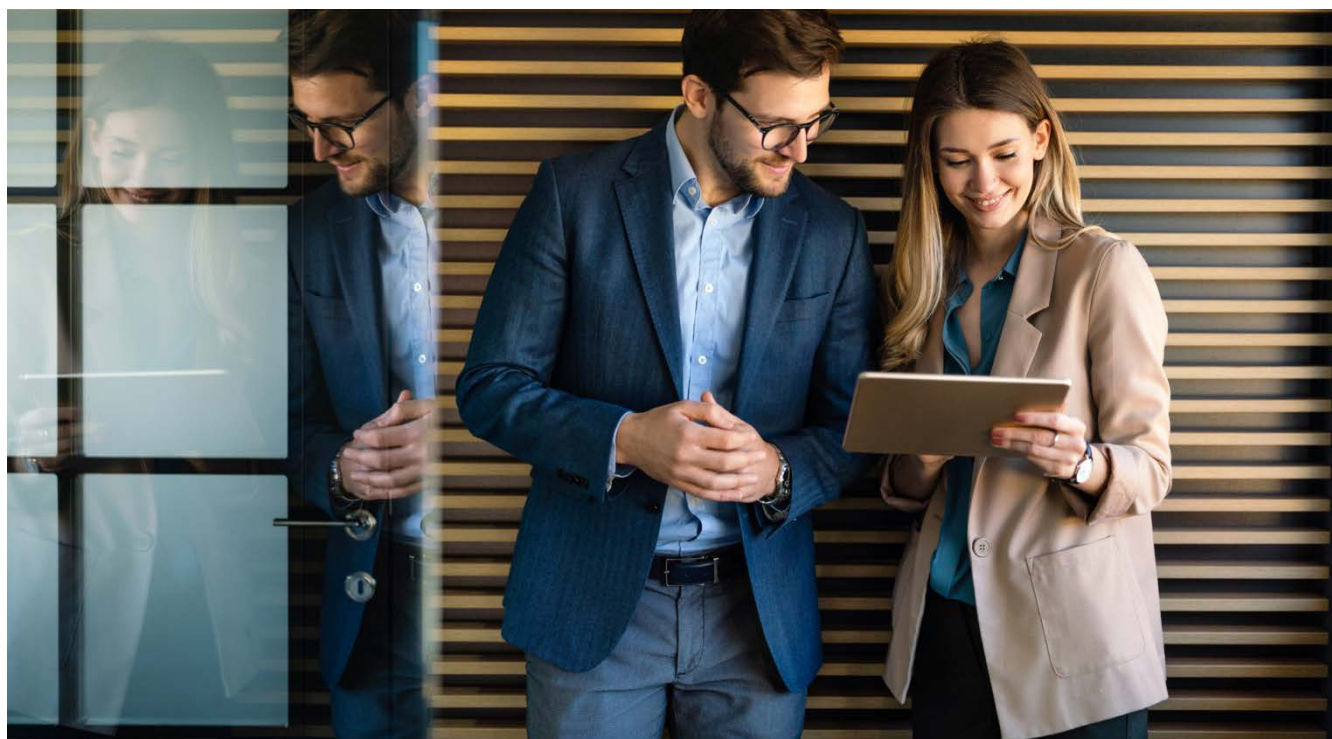
AI model security

AI models are at the core of generative applications, but they are also one of their most critical risk vectors. Protecting them requires specific testing and techniques to ensure their resilience.

- Adversarial testing allows us to check how models react to malicious inputs designed to manipulate their results.
- Analysis of APIs and integration points is essential to prevent attacks such as prompt injection, model extraction, or data poisoning.
- Offensive security methodologies geared toward generative AI and model hardening help strengthen their resistance to manipulation attempts.

A customer service model without robust controls, for instance, can be manipulated to deliver sensitive information or give fraudulent responses.

Model security protects technical performance, reputation, and business continuity: a failure at this layer can lead to loss of customer trust, regulatory non-compliance, and significant operational risks.



Main risks of Generative AI



Shadow AI.

Unauthorized use of AI tools by employees without corporate control, increasing the risk of sensitive information leaks and making it difficult to trace activities.



Sensitive data leaks.

Poorly designed prompts or insecure configurations can cause generative assistants connected to SaaS systems or internal repositories to expose confidential company information.



Hallucinations and misinformation.

AI can generate false, inconsistent, or misleading responses, damaging corporate reputation, customer trust, and affecting regulatory compliance.



Algorithmic biases.

Lack of diversity in training sets can introduce biases that lead to discriminatory, unreliable, or counterproductive decisions.



Manipulation of RAG pipelines (Retrieval Poisoning).

Attackers can introduce altered data into corporate information sources that dynamically feed the model, compromising the reliability of the system and exposing sensitive company information.



Privacy attacks.

Techniques such as training information extraction, membership inference (discovering a specific record), or prompt theft can reveal personal data, intellectual property, or corporate secrets.

These risks explain the need for a robust security and governance framework to identify and manage vulnerabilities from the outset.

Protection against threats

Once the risks have been identified, the next step in our Secure Journey to AI is to implement protective measures that safeguard models, applications, and data against cyber threats. This second pillar focuses on strengthening the resilience of the AI infrastructure, correcting vulnerabilities, and ensuring that technological environments operate under a comprehensive security framework.

Protection against threats should not be limited to the occasional correction of faults but should be conceived as a continuous and dynamic process combining good secure development practices, advanced Cyber Security controls, and integration with identity and access management.

True innovation only thrives when it is protected: shielding data, models, and infrastructure is the basis for confidently harnessing the full potential of AI.

Creating resilient applications

The development of AI applications must incorporate security from the design stage to prevent vulnerabilities that could be exploited later. The DevSecOps methodology, which integrates security into each phase of the development cycle, allows faults to be detected and corrected at an early stage, reducing the cost and impact of incidents.

- This involves having robust vulnerability management processes in place, capable of identifying and resolving errors before going into production.
- It also involves applying code analysis techniques, both static and dynamic and interactive (AST, SAST, DAST, IAST), to ensure that the software is free of critical vulnerabilities.
- Added to this are solutions such as WAAP (Web Application and API Protection) and CNAPP (Cloud-Native Application Protection Platform), designed to protect web applications and cloud-native environments from sophisticated attacks.
- The shift-left security approach also reinforces the idea of moving security controls as close as possible to the beginning of the development cycle, reducing risks before the application goes live.

Together, these practices enable the creation of resilient applications that reduce the attack surface and reinforce trust in AI, preventing it from becoming a weak point within the business ecosystem.

Data and identity protection (IAM)

Data and identities are critical assets for AI. Protection against threats involves strengthening the logical perimeter and ensuring that only authorized individuals access sensitive information.

- An advanced identity and access management (IAM) scheme includes mechanisms such as multi-factor authentication (MFA), conditional access, and granular role control, ensuring that each user only accesses the information that is strictly necessary.
- Added to this are automatic audits and early warning systems that detect anomalous behavior in real time, as well as data leak prevention (DLP) policies to prevent sensitive information from leaving the systems accidentally or intentionally.

The human dimension also plays a key role: protection against social engineering attacks, such as phishing or credential theft, is essential to reduce the risk of unauthorized access. A simple click on a fraudulent email can open the door to intrusions that compromise critical models and data.

A robust data and identity protection strategy prevents overexposure, reduces the likelihood of leaks, and ensures the integrity of models, thereby strengthening the trust of customers, partners, and users.

Risk mitigation in AI architecture

The architecture underpinning AI requires layered protection to prevent insecure configurations or accidental exposures that compromise overall security.

This requires:

- Deploying secure infrastructures, applying techniques such as microsegmentation and using specific security tools for hybrid and cloud environments. These include:
 - *CSPM (Cloud Security Posture Management)*, which monitors cloud configuration to detect deviations.
 - *CWPP (Cloud Workload Protection Platform)*, which protects cloud workloads.
 - *CIEM (Cloud Infrastructure Entitlement Management)*, which controls excessive permissions in cloud environments.
- Continuously review and correct misconfigurations in pipelines and AI models, as well as the use of AI Gateways and security policies that regulate how generative models are exposed and used.
- Protect the APIs and plugins used in integration to reduce the risks of manipulation, data extraction, or abuse in the consumption of AI services.

Security in high-level generative AI architecture

The security of a generative AI application must be understood holistically, applying specific controls at each level of its architecture. These levels range from the data that feeds the models to the final interaction with the user.

- **Model training.** Protecting training data is essential to ensuring the reliability of results. It requires robust data governance, classification and labeling of sensitive information, DLP mechanisms, and regular audits to prevent exposure or misuse of information.
- **LLM runtime.** The environment in which prompts are processed and responses are generated needs continuous protection. AI-SPM (AI Security Posture Management) capabilities are particularly important here, as they allow deviations to be monitored and corrective measures to be applied proactively, together with specific offensive assessments (such as prompt injection tests, data extraction, or output manipulation) that validate the model's resilience against advanced attacks.

- Application layer. At the point of contact with the user, it is critical to reinforce the security of APIs, plugins, and extensions. To this end, next-generation WAAP or WAF solutions are applied, combined with recurring offensive testing to ensure the resilience of the interaction against abuse or information leaks.

This three-block vision ensures that security supports the entire cycle, from training data to the end-user experience, consolidating generative AI architecture as a reliable, resilient, and compliant environment.

Best practices for protecting data in AI



Data classification and labeling

To identify the sensitivity level of information and apply specific protection policies according to its criticality.



Data leak prevention (DLP)

By implementing controls that prevent the unauthorized sending of sensitive information, both in cloud environments and on corporate devices.



Identity and access management (IAM)

Applying multi-factor authentication (MFA), granular role controls, and the principle of least privilege to limit exposure.



End-to-end encryption

That protects data both at rest and in transit to prevent unauthorized access.



Responsible use policies and training

For employees and collaborators on best practices when interacting with AI models, to reduce human error and misconfigurations.

These measures make **security an enabler** that shields critical information and strengthens the trust of customers, regulators, and partners.



360° Response

The third pillar of our Secure Journey to AI focuses on a comprehensive, 360° response. A comprehensive approach that ensures continuous visibility, incident response capability, and continuous improvement in the face of rapidly changing threats. AI introduces new and dynamic risks, and companies need a framework that not only detects vulnerabilities but also allows them to respond immediately and learn from each event.

This pillar combines advanced monitoring, integration with a specialized SOC (Security Operations Center), and systematic incident management to ensure that companies operate with confidence and transparency in a complex and changing technological environment.

In an environment where threats evolve at the speed of technology, the ability to respond quickly and transparently will be the biggest factor in building trust.

End-to-end visibility and security governance

Companies need complete visibility into what is happening at every layer, from infrastructure to models and data, in order to trust their AI systems. Without traceability and auditing, risks become invisible and trust is undermined.

- End-to-end visibility is not just a technical record of activity: it is the guarantee that every interaction can be explained, reviewed, and corrected if necessary.
- Have detailed user and prompt logs, apply simultaneous monitoring of infrastructure and applications.
- Leveraging advanced AI-based telemetry allows you to anticipate anomalies before they become incidents.

A bank that monitors the interactions of its virtual assistants in real time, for example, not only protects the integrity of its models, but can also prove to auditors and customers that each recommendation is generated under security and regulatory compliance criteria.

In this sense, visibility becomes an asset of governance and transparency, essential for complying with regulations such as the European AI Act and reinforcing the trust of customers and regulators.

Integration with Telefónica Tech's specialized SOC

Telefónica Tech's AI-specialized SOC (Security Operations Center) acts as a control tower that monitors, anticipates, and neutralizes threats in real time. This capability gives companies the assurance that their AI environments are continuously monitored, ensuring immediate detection and management of any incidents.

Beyond surveillance, the SOC offers:

- Automated responses, integrating playbooks that activate defenses in seconds.
- Intelligence against new threats, from prompt injection attacks to data poisoning attempts or context leaks.
- Human intervention when the incident requires it, coordinating the response with multidisciplinary teams capable of escalating and containing critical situations.

The value of the SOC lies not only in the technology, but also in the confidence generated by a managed and expert service. In this way, the SOC turns Cyber Security into a strategic element, ensuring business continuity, reputation, and peace of mind for companies that adopt AI.

Incident management and continuous improvement

Responding to incidents is not enough: it is essential to turn every situation into an opportunity to strengthen organizational resilience. Incident management in AI involves identifying root causes, adjusting policies, and improving processes so that every thwarted attack strengthens the overall defense.

This approach translates into:

- A post-incident cycle, where each event generates lessons learned that are incorporated into the security framework.
- Adaptive policies are adjusted based on new threats, and AI-specific metrics help evaluate the effectiveness of controls.
- AI simulation or sandbox environments allow models to be tested against adverse scenarios before putting them into production.

Consider a healthcare company that detects an attempt to manipulate its diagnostic model. It stops the threat thanks to a structured incident management process and also adapts its controls to prevent such an attack from happening again. The healthcare company therefore does not react but evolves.

An incident management and continuous improvement framework ensures that AI does not get caught in a cycle of vulnerability and reaction, but rather moves toward more robust and reliable use, reducing the recurrence of attacks and strengthening organizational confidence in its use.

The 360° Response not only strengthens resilience, it also ensures regulatory compliance.

Security as a living process

Monitor end-to-end activity	Of models, data, identities, and applications to ensure traceability and governance.
Identify anomalies	Manipulation attempts, or data leaks before they escalate into critical incidents.
Integrate with specialized SOC	And use automated playbooks to contain and mitigate incidents in an agile manner.
Learn from each incident	Adjust policies, and strengthen controls to ensure that protection evolves at the same pace as threats.
Continuous audits	That periodically validate the effectiveness of controls and measure resilience against new threats and regulatory requirements.

Security in AI is not an end goal, but a permanent cycle of anticipation, adaptation, and learning that turns trust into a sustainable competitive advantage. This continuous cycle of monitoring, response, and auditing contributes decisively to regulatory compliance. Frameworks such as the AI Act, NIS2, and DORA require traceability, governance, and periodic validation of AI systems, requirements that are naturally integrated into a 360° Response approach.

Conclusions and next steps

Conclusions

AI has established itself as a catalyst for transformation, but it introduces unprecedented risks that force companies to rethink how they adopt it. Our Secure Journey to AI roadmap addresses these new risks through three pillars: early risk identification and detection, protection against threats, and 360° incident response.

It is important to note that these pillars do not work in isolation: they constitute a cohesive framework that enables the safe and responsible use of AI in line with business objectives.

The most relevant conclusion is that security and trust are not obstacles to innovation, but rather its enablers. Only companies that integrate regulatory compliance, governance, and data protection into their strategy will be able to deploy AI sustainably and obtain real and lasting value.

Only companies that securely integrate AI into their DNA will be able to govern it effectively, comply with regulations, and turn trust into a competitive advantage.

Immediate actions

The conclusions point us in a clear direction: adopting AI requires placing security at the center and approaching it with a proactive mindset. It is not enough to react when risks arise, we must anticipate them and turn security into a competitive advantage.

Companies must take the following actions immediately in order to move forward:

- Audit current AI usage to identify and classify applications in use (including Shadow AI) to obtain an accurate view of the situation.
- Assess AI maturity to diagnose gaps in governance, security, and internal competencies, and design a feasible roadmap.
- Establish a cross-functional governance model that defines clear responsibilities, aligns areas, and provides training at all levels to build a culture of responsible AI.
- Integrate security by design by incorporating controls into data, models, and applications from the outset of projects, thereby avoiding cost overruns and building trust from the start.
- Adopt a continuous improvement approach by reviewing and updating policies and processes to adapt them to new threats and regulatory frameworks, such as the AI Act, the GDPR, or the NIS2 directive.

These immediate actions are the basis for AI adoption to maximize return on investment, minimize risk, and accelerate digital transformation in a sustainable and reliable way.

AI adoption is only sustainable when security is central and proactively addressed.

Consolidating secure and responsible AI

After identifying risks, protecting against threats, and implementing a 360° response, the great challenge for organizations is to consolidate an AI model that maintains security as its central focus and inspires long-term confidence. It is not just a matter of initiating projects, but of ensuring that their evolution remains secure, stable, and resilient, capable of adapting to changing needs and new regulatory requirements.

The future of AI depends not only on what it is capable of doing, but also on how we decide to implement, consolidate, and govern it.

Sustainability, resilience, and transparency will be the factors that make the difference, allowing AI to become a stable asset that supports business strategy. Achieving this goal requires an approach that combines innovation with security, responsibility with resilience, and transparency with regulatory compliance, so that technology has purpose and adds value.

Telefónica Tech supports companies of all sizes and sectors in this process with Secure Journey to AI, a comprehensive approach that ensures the adoption of AI translates into positive impact, sustainable value, and long-term trust.

Why Telefónica Tech

Telefónica Tech is not just a supplier on the journey towards adopting secure and responsible AI: we are a technological and strategic partner that accompanies companies every step of the way and at every milestone along the way.

Telefónica Tech translates security into the language of innovation. We support companies on their journey towards secure AI that has a positive impact on their business.

Our comprehensive approach protects key assets (data, applications, and infrastructure) while enabling innovative, responsible use that is aligned with business objectives. We provide differential value through unique capabilities that turn security into a driver of trust:

- Comprehensive asset protection. We ensure that data, applications, and AI infrastructures are always protected without compromising their strategic value, applying a cross-cutting security vision that extends from end to end.
- Three-pillar roadmap. Risk identification, threat protection, and 360° response form our comprehensive Secure Journey to AI framework, which enables companies to confidently navigate their journey to AI.
- Guaranteed regulatory compliance. We ensure that each project complies with frameworks such as the AI Act, NIS2, or DORA and anticipates their requirements, offering transparency and peace of mind to customers, employees, and regulators.
- Recognized Cyber Security capabilities. We have a robust infrastructure, advanced 24/7 monitoring services, and third-party certifications for proactive defense and management of identities and sensitive data.
- Talent and expertise in AI. More than 400 experts and 10 specialized centers combine knowledge in AI and security to create reliable and resilient environments.

Telefónica Tech is a trusted partner that helps companies integrate AI securely, efficiently, and in line with their business objectives, advancing their secure digital transformation.

Our comprehensive vision ensures both protection and value creation, strengthening customer trust and operational resilience in a competitive, changing, and regulated environment.

*The journey toward secure AI has already begun,
traveling it with confidence will make the difference
between experimenting and transforming.*

The information contained in this document is proprietary to Telefonica Cyber Security & Cloud Tech S.A together with Telefonica IoT & Big Data Tech S.A. (hereinafter "Telefonica Tech") and/or any other entity within the Telefonica Group or its licensors.

Telefónica Tech and/or any company within the Telefónica Group or Telefónica Tech's licensors reserve all intellectual property rights (including any patents or copyrights) arising from or pertaining to this document, including the rights to design, produce, reproduce, use and sell this document, except to the extent that such rights are expressly granted to third parties in writing. The information contained in this document may be subject to change at any time without prior notice.

The information contained in this document may not be copied in whole or in part, distributed, adapted or reproduced in any form without the prior written consent of Telefónica Tech.

The sole purpose of this document is to support the reader in the use of the product or service described herein. The reader agrees and undertakes to use the information contained herein for the reader's own use and not for any other use.

Telefónica Tech shall not be liable for any loss or damage arising from the use of the information contained herein or for any errors or omissions in the document or for the incorrect use of the service or product. Use of the product or service described herein shall be governed by the terms and conditions accepted by the user of this document for use.

Telefónica Tech and its brands (as well as any brand belonging to the Telefónica Group) are registered trademarks. Telefónica Tech and its subsidiaries reserve all rights therein.