

Nemko Digital Webinar Report - EU AI Act Update

Executive Summary

This webinar provided a timely and comprehensive overview of the new obligations under the EU AI Act specifically targeting general purpose AI models, which came into force on August 2, 2025. Monica Fernandez from Nemko delivered this critical update to help organizations understand their compliance responsibilities when placing general purpose AI models in the European market.

The primary objective was to clarify the regulatory landscape for AI providers, focusing on the risk-based approach of the EU AI Act and the specific requirements that general purpose AI model providers must meet to ensure legal compliance.

Main Theme and Objectives

Primary Objectives

- **Educate on new regulatory requirements** that have recently come into force or are approaching implementation
- **Enable compliance assessment** by helping organizations determine if and how regulations apply to their operations
- **Clarify specific obligations** and requirements that affected entities must meet to remain compliant
- **Address implementation timelines** including enforcement dates, grace periods, and compliance deadlines

Stakeholder-Focused Objectives

- **Support primary regulated entities** in understanding their direct regulatory responsibilities

- **Guide secondary affected parties** who may become regulated through their business activities or modifications
- **Assist organizations with integrated operations** that may trigger multiple regulatory requirements
- **Address the full industry ecosystem** by covering various types of affected stakeholders and use cases

Strategic Objectives

- **Introduce compliance resources** and guidance documents while noting their current status and applicability
- **Provide regulatory outlook** on potential future developments and evolving requirements
- **Bridge complexity gaps** by translating complex regulatory language into practical, actionable guidance
- **Enable strategic planning** for business operations within the new regulatory framework

Educational Objectives

- **Deliver comprehensive understanding** beyond simple requirement listings through context and explanation
- **Provide real-world applications** through examples and case studies that illustrate regulatory impact
- **Support ongoing compliance** by offering expert resources and professional consultation opportunities
- **Ensure immediate value** by focusing on actionable steps organizations can implement right away

Understanding the EU AI Act Framework

The foundation of the webinar rested on a comprehensive explanation of the EU AI Act's risk-based regulatory approach, which Fernandez illustrated using a pyramid

structure that clearly delineates different categories of AI systems and their corresponding regulatory requirements. This framework represents a nuanced approach to AI regulation that recognizes the varying levels of risk and societal impact associated with different types of AI applications.

At the apex of this regulatory pyramid sit prohibited AI systems, which are completely banned from the European market due to their potential for causing unacceptable harm or violating fundamental rights. These systems represent the most extreme cases where the risks are deemed to outweigh any potential benefits, establishing clear red lines for AI development and deployment.

The largest segment of the regulatory framework addresses high-risk AI systems, which encompass applications that have the potential to significantly impact safety, fundamental rights, or other critical societal interests. These systems face the most comprehensive set of regulatory requirements, including rigorous conformity assessments, quality management systems, risk management processes, and ongoing monitoring obligations. The extensive nature of these requirements reflects the EU's recognition that high-risk AI systems require robust safeguards to protect individuals and society.

Limited risk AI systems occupy a middle ground in the regulatory framework, subject to specific transparency requirements that ensure users are aware they are interacting with AI systems. These requirements are designed to enable informed decision-making by individuals who may be affected by AI-generated outputs or recommendations, while avoiding overly burdensome regulations for systems that pose moderate risks.

At the base of the pyramid, low-risk AI systems fall entirely outside the scope of the AI Act, reflecting the EU's intention to avoid stifling innovation in areas where AI applications pose minimal risks to individuals or society. This approach allows for continued development and deployment of beneficial AI applications without unnecessary regulatory overhead.

General purpose AI models represent a unique category within this framework, subject to their own specific set of requirements that reflect their distinctive characteristics and potential applications. Unlike traditional AI systems that are designed for specific use cases, general purpose AI models are characterized by their versatility and ability to be adapted for various applications, creating unique regulatory challenges that require specialized approaches.

The risk-based approach underlying the entire framework demonstrates the EU's commitment to proportionate regulation that balances innovation with protection. By tailoring requirements to the level of risk posed by different types of AI systems, the regulation aims to provide appropriate safeguards without unnecessarily hindering technological development or market competition.

Defining General Purpose AI Models: Thresholds and Examples

A significant portion of the webinar was dedicated to providing a clear and actionable definition of what constitutes a general purpose AI model under the EU AI Act. Monica Fernandez emphasized that a precise understanding of this definition is the critical first step for any organization seeking to determine its regulatory obligations. The Act defines a general purpose AI model as one that is trained on large-scale data, demonstrates the ability to perform a wide range of distinct tasks, and can be integrated into various downstream applications. This definition highlights the versatility and adaptability that distinguish these models from more specialized AI systems.

To provide a more concrete and objective measure, the regulation establishes a specific computational threshold for identifying these models. A model is considered a general purpose AI model if its training process involved a cumulative training compute greater than 10^{23} floating-point operations (FLOPs). This quantitative measure provides a clear line for providers to assess whether their models fall within the scope of the regulation. In addition to this computational threshold, the model must also be capable of generating content such as text, audio, images, or video.

Fernandez also clarified an important exclusion: models that are used exclusively for research, development, or prototyping activities before being placed on the market are not subject to these obligations. This exemption is designed to ensure that innovation and experimentation are not stifled during the early stages of a model's lifecycle.

Models with Systemic Risk

The webinar also detailed a sub-category of general purpose AI models that are identified as having "systemic risk." These models are subject to a more stringent

set of requirements due to their potential for significant impact on the EU market and society. The threshold for identifying a model with systemic risk is a training compute greater than 10^{25} FLOPs. Providers of such models have an obligation to self-identify and notify the European Commission. Furthermore, the Commission retains the authority to designate a model as having systemic risk based on a set of criteria outlined in Annex 13 of the AI Act, even if it does not meet the computational threshold.

Practical Examples

To further clarify these definitions, Fernandez presented three practical examples:

1. **Not a General Purpose AI Model:** A model trained using 10^{22} FLOPs (below the threshold) on natural language data that is not capable of performing many different tasks. This model fails to meet both the computational and functional criteria.
2. **Is a General Purpose AI Model:** A model trained using 10^{24} FLOPs (above the threshold) on a large and diverse set of natural language data from the internet and other sources. This model's training compute and its ability to handle a wide range of text-generation tasks clearly place it within the scope of the regulation.
3. **Not a General Purpose AI Model:** A model trained with 10^{24} FLOPs (above the threshold) but designed exclusively for a single, narrow task, such as transcribing speech to text. Despite meeting the computational threshold, its lack of generality and inability to perform a wide range of distinct tasks means it is not considered a general purpose AI model.

These examples provided valuable clarity, demonstrating that both the computational resources used for training and the model's functional capabilities are essential factors in determining its classification under the EU AI Act.

Obligations for General Purpose AI Model Providers

The webinar provided a detailed breakdown of the specific obligations that providers of general purpose AI models must fulfill to comply with the EU AI Act. These obligations are structured in a tiered approach, with baseline requirements for all general purpose AI models and additional requirements for those classified as having systemic risk.

Core Obligations for All General Purpose AI Models

Every provider of a general purpose AI model must adhere to three fundamental obligations that form the foundation of compliance with the regulation:

Transparency Over Training Content: Providers must provide comprehensive transparency regarding the content used to train their models. This requirement reflects the EU's commitment to ensuring that stakeholders, including downstream users and regulatory authorities, have visibility into the data sources and methodologies that shape a model's capabilities and potential biases. The transparency obligation extends beyond simple disclosure to include meaningful information that enables informed decision-making about the model's appropriate use cases and limitations.

EU Copyright Law Compliance: The second core obligation requires providers to implement and maintain a robust policy framework for complying with European Union copyright law. This requirement acknowledges the significant legal and ethical questions surrounding the use of copyrighted material in AI training datasets. Providers must demonstrate that they have established processes and safeguards to respect intellectual property rights throughout their model development and deployment processes.

Technical Documentation: The third fundamental obligation involves the creation and maintenance of comprehensive technical documentation that serves two critical purposes. First, this documentation must be available to share with the AI Office and national competent authorities upon request, enabling regulatory oversight and compliance verification. Second, the documentation must be made available to downstream providers who integrate the model into their own systems or applications, facilitating informed decision-making and appropriate risk management throughout the AI value chain.

Additional Obligations for Models with Systemic Risk

Models that meet the criteria for systemic risk face four additional obligations that reflect their potential for significant societal impact:

Model Evaluation and Adversarial Testing:

Providers of systemic risk models must conduct rigorous evaluation processes, including adversarial testing designed to identify potential vulnerabilities, biases, or harmful capabilities. This testing must be comprehensive and ongoing, reflecting the dynamic nature of AI systems and their potential for unexpected behaviors or applications. The evaluation process should encompass both technical performance metrics and broader assessments of societal impact and risk.

Risk Identification and Mitigation:

Building on the evaluation requirements, providers must establish systematic processes for identifying systemic risks associated with their models and implementing appropriate mitigation measures. This obligation requires a proactive approach to risk management that goes beyond reactive responses to identified problems. Providers must demonstrate ongoing vigilance and continuous improvement in their risk management practices.

Incident Reporting:

Providers of systemic risk models must establish robust systems for documenting and reporting serious incidents to the AI Office and national competent authorities without delay. This requirement ensures that regulatory authorities have timely access to information about significant problems or failures that could have broader implications for public safety or societal well-being. The reporting obligation extends to corrective actions taken in response to incidents, providing transparency about remediation efforts.

Cybersecurity Protection:

The final additional obligation requires providers to ensure adequate cybersecurity protection for both the model itself and its supporting infrastructure. This requirement recognizes that systemic risk models represent high-value targets for malicious actors and that security breaches could have far-reaching consequences. Providers must implement comprehensive security measures that address both technical vulnerabilities and operational risks.

These obligations reflect a comprehensive approach to managing the risks associated with powerful AI systems while enabling continued innovation and

development. The tiered structure ensures that regulatory burden is proportionate to the level of risk posed by different types of models, while still providing appropriate safeguards for all general purpose AI systems.

Downstream Modifiers: When Fine-Tuning Creates New Provider Obligations

One of the most complex and practically significant aspects of the EU AI Act's approach to general purpose AI models involves the treatment of downstream modifiers—organizations that modify, fine-tune, or adapt existing models rather than developing them from scratch. Monica Fernandez dedicated substantial attention to this topic, recognizing its critical importance for the many organizations that build upon existing foundation models rather than creating entirely new ones.

The fundamental principle underlying the regulation's approach to downstream modification is that significant changes to a model can fundamentally alter its characteristics, capabilities, and risk profile. When modifications are substantial enough, the downstream modifier effectively becomes the provider of a new general purpose AI model, with all the associated regulatory obligations. This approach ensures that regulatory oversight extends throughout the AI value chain, not just to original model developers.

The One-Third Computational Threshold

The recently published guidance has provided much-needed clarity on what constitutes a "significant" modification that would trigger provider obligations. The key threshold is whether the fine-tuning or modification process uses more than one-third of the computational resources that were used to train the original model. This quantitative measure provides a clear and objective standard that organizations can apply to assess their regulatory status.

Fernandez explained that this one-third threshold was carefully chosen to be proportional to the size of the original model, ensuring fairness across different model scales while avoiding the discouragement of improvements to smaller models. The threshold recognizes that meaningful modifications to larger models naturally require more computational resources, while still allowing for substantial fine-tuning activities that fall below the regulatory trigger point.

Alternative Thresholds for Unknown Original Compute

A practical challenge that many organizations face is that the computational resources used to train the original model may not be publicly available or may be difficult to determine accurately. To address this situation, the European Commission has provided alternative thresholds based on the standard computational benchmarks used to define general purpose AI models.

When the original model's training compute is unknown or unavailable, organizations can use these substitute thresholds:

- If the modification uses more than one-third of 10^{25} FLOPs (approximately 3.33×10^{24} FLOPs), the modifier becomes the provider of a new general purpose AI model with systemic risk.
- If the modification uses more than one-third of 10^{23} FLOPs (approximately 3.33×10^{22} FLOPs), the modifier becomes the provider of a new standard general purpose AI model.

These alternative thresholds provide a practical solution for organizations that need to assess their regulatory obligations without access to detailed information about the original model's development process.

Current Market Reality and Future Implications

Fernandez acknowledged that under current market conditions, it is relatively uncommon for downstream modifications to reach the computational thresholds that would trigger provider obligations. Most fine-tuning and adaptation activities currently performed by organizations involve significantly less computational resources than would be required to meet the one-third threshold.

However, she emphasized that this situation may change as computational resources become more accessible and affordable, and as organizations develop more sophisticated approaches to model modification and improvement. The regulatory framework is designed to be forward-looking, anticipating future developments in AI technology and market practices.

The European Commission has also indicated its intention to review and potentially update these thresholds over time to reflect changes in technology and market conditions. This adaptive approach ensures that the regulatory framework can

evolve alongside technological developments while maintaining appropriate oversight of significant modifications to AI models.

Implications for Transparency and Risk Management

When a downstream modifier becomes a provider under these rules, they assume full responsibility for the transparency, copyright compliance, and technical documentation obligations associated with their modified model. This includes providing information about both the original model and the modifications made, ensuring that downstream users and regulatory authorities have comprehensive visibility into the model's development and characteristics.

For models that reach the systemic risk threshold through modification, the new provider must also fulfill all the additional obligations associated with systemic risk models, including evaluation, testing, incident reporting, and cybersecurity requirements. This comprehensive transfer of responsibility ensures that regulatory oversight remains effective even as models evolve through downstream modification processes.

System Integration: Considering the Broader AI Ecosystem

A critical insight shared during the webinar was the recognition that general purpose AI models rarely operate in isolation. Instead, they are typically integrated into larger AI systems that serve specific purposes within particular sectors or domains. This integration creates additional layers of regulatory complexity that organizations must carefully navigate to ensure comprehensive compliance.

Monica Fernandez illustrated this concept with a practical example involving GPT-4 being used to summarize clinical guidelines within a healthcare diagnostic decision support system for doctors. In this scenario, while GPT-4 itself is a general purpose AI model subject to its own set of obligations, the broader system into which it is integrated would be classified as a high-risk AI system due to its healthcare application and potential impact on patient safety.

This dual classification means that the provider of the overall system must comply with the comprehensive requirements for high-risk AI systems, which include rigorous conformity assessments, quality management systems, risk management

processes, and ongoing monitoring obligations. Simultaneously, they must also address the general purpose AI obligations, either by ensuring their third-party model provider has met these requirements or by taking on these responsibilities themselves if they have modified the model significantly.

The complexity of this regulatory landscape requires organizations to think holistically about their AI implementations. They must consider not only the characteristics and obligations associated with individual AI components but also how these components interact within larger systems and the regulatory implications of those interactions. This systems-level thinking is essential for developing comprehensive compliance strategies that address all applicable requirements.

Furthermore, the integration context can significantly influence the overall risk profile and regulatory classification of an AI system. A general purpose AI model that might be considered relatively low-risk in isolation could become part of a high-risk system when deployed in critical applications such as healthcare, transportation, or financial services. Organizations must therefore conduct thorough assessments of their intended use cases and deployment contexts to understand their full regulatory obligations.

Critical Timeline and Implementation Requirements

The timing of the webinar was strategically chosen to address the immediate urgency surrounding the EU AI Act's implementation timeline. Monica Fernandez emphasized that the obligations for general purpose AI models officially came into force on August 2nd, 2025, making compliance an immediate requirement rather than a future consideration.

Immediate Compliance for New Models

The most pressing requirement applies to any general purpose AI models introduced to the European market after August 2nd, 2025. These models must fully comply with all applicable obligations from the moment they are placed on the market. There are no grace periods or transitional arrangements for new models, reflecting the EU's commitment to ensuring that newly developed AI systems meet the established safety and transparency standards from the outset.

Organizations that have been developing models with the intention of releasing them after this date must ensure that they have implemented all necessary compliance measures before market introduction. This includes establishing transparency documentation, implementing copyright compliance policies, creating comprehensive technical documentation, and, for systemic risk models, implementing evaluation, testing, incident reporting, and cybersecurity measures.

Grace Period for Existing Models

Recognizing the practical challenges of retrofitting existing models to meet new regulatory requirements, the EU AI Act provides a grace period for models that were already on the market before August 2nd, 2025. These existing models have until August 2nd, 2027, to achieve full compliance with the new obligations.

However, this grace period comes with important caveats and limitations. Providers of existing models are not automatically required to retrain their models or completely reconstruct their training datasets if doing so would be practically impossible or excessively burdensome. The regulation acknowledges that in some cases, the original training data may no longer be available, or the computational costs of retraining may be prohibitive.

When providers cannot fully comply with certain obligations due to these practical limitations, they must clearly document and explain these limitations in their technical documentation. This transparency requirement ensures that downstream users and regulatory authorities understand any gaps in compliance and can make informed decisions about the appropriate use of these models.

Notification Requirements for Systemic Risk Models

Providers of models that meet the criteria for systemic risk face additional time-sensitive obligations. They must notify the European Commission of their model's systemic risk status within two weeks of the August 2nd, 2025 deadline. This notification requirement applies regardless of whether the model was introduced before or after the deadline, emphasizing the particular importance of maintaining regulatory oversight over these high-impact systems.

The notification process serves multiple purposes: it enables the Commission to maintain awareness of systemic risk models operating in the EU market, facilitates

appropriate regulatory oversight and engagement, and ensures that providers acknowledge their enhanced obligations and responsibilities.

Proactive Engagement Requirements

For organizations facing compliance difficulties, Fernandez emphasized the importance of proactive engagement with the AI Office. Rather than waiting for regulatory enforcement or attempting to navigate complex requirements in isolation, providers are encouraged to reach out for guidance and support when they encounter challenges in meeting their obligations.

This proactive approach benefits both providers and regulators by enabling early identification and resolution of compliance issues, reducing the likelihood of inadvertent violations, and fostering a collaborative relationship between industry and regulatory authorities. The AI Office's willingness to engage with providers facing genuine compliance challenges reflects the EU's recognition that effective regulation requires ongoing dialogue and cooperation between all stakeholders.

The General Purpose AI Code of Practice: A Practical Compliance Framework

A significant portion of the webinar was dedicated to introducing the General Purpose AI Code of Practice, which represents a collaborative effort to translate the EU AI Act's requirements into practical, implementable guidance for industry stakeholders. Published on July 10th, 2025, this voluntary guidance framework emerged from the collective expertise of over 1,000 independent experts who worked to create feasible, tangible, and accessible guidance for providers seeking to comply with their regulatory obligations.

Development and Structure

The Code of Practice represents an unprecedented collaborative effort in AI governance, bringing together diverse perspectives from academia, industry, civil society, and regulatory bodies to create comprehensive guidance that addresses real-world implementation challenges. The involvement of such a large and diverse group of experts ensures that the guidance reflects a broad range of perspectives and practical experiences, enhancing its relevance and applicability across different organizational contexts and use cases.

The Code is structured around three main chapters that correspond directly to the core obligations established by the EU AI Act:

Transparency Chapter:

This section provides detailed guidance on how providers can fulfill their transparency obligations regarding training content. It addresses practical questions about what information must be disclosed, how it should be presented, and what level of detail is required to meet regulatory expectations. The guidance recognizes the need to balance transparency with legitimate business interests and competitive considerations.

Copyright Chapter:

The second chapter focuses on the complex intersection between AI development and intellectual property law, providing practical guidance on how providers can establish and implement policies for complying with EU copyright law. This section addresses one of the most challenging aspects of AI regulation, where technical capabilities and legal requirements intersect in complex ways.

Safety and Security Chapter:

The final chapter addresses the additional obligations that apply specifically to models with systemic risk. This section provides guidance on conducting model evaluations, implementing adversarial testing, establishing incident reporting systems, and ensuring adequate cybersecurity protections. The guidance in this chapter reflects the heightened expectations and responsibilities associated with high-impact AI systems.

Legal Status and Future Recognition

While the Code of Practice provides valuable guidance, Fernandez emphasized that it is currently undergoing legal review by the AI Office and the AI Board. This review process is critical for determining whether the Code will receive official recognition and legal validity across the European Union. If approved, the European Commission may issue an implementing act that would give the Code general legal validity throughout the EU.

This potential official recognition would transform the Code from voluntary guidance into a recognized compliance pathway, providing significant benefits for organizations that choose to follow its recommendations. Providers who demonstrate compliance with an officially recognized Code of Practice would have a clear and defensible basis for asserting that they meet their regulatory obligations.

Compliance Benefits and Flexibility

One of the most significant advantages of following the Code of Practice, particularly if it receives official recognition, is the ability to demonstrate compliance with EU AI Act obligations in a straightforward and well-documented manner. Organizations that implement the Code's recommendations would have a clear framework for organizing their compliance efforts and documenting their adherence to regulatory requirements.

However, Fernandez also emphasized that the Code of Practice remains voluntary, and organizations are free to use alternative methods to meet their obligations. This flexibility recognizes that different organizations may have different capabilities, constraints, and approaches to compliance. The regulation's focus on outcomes rather than specific processes allows for innovation and adaptation in compliance strategies.

The voluntary nature of the Code also means that organizations can choose to implement some recommendations while developing alternative approaches for others. However, Fernandez noted that opting out of specific parts of the Code would remove the compliance demonstration benefits for those particular areas, requiring organizations to develop and justify their alternative approaches.

Penalties and Enforcement

The webinar concluded with a sobering reminder of the financial consequences of non-compliance with general purpose AI model obligations. Organizations that fail to meet their regulatory requirements face penalties of up to €15 million or up to 3% of their global annual turnover, whichever is higher. These substantial penalties reflect the EU's commitment to ensuring effective compliance with the AI Act and the serious consequences that can result from regulatory violations.

The scale of these penalties underscores the importance of taking compliance obligations seriously and implementing robust systems and processes to ensure ongoing adherence to regulatory requirements. For many organizations, the potential financial impact of non-compliance far exceeds the costs of implementing appropriate compliance measures, making investment in regulatory compliance both a legal necessity and a sound business decision.

Key Takeaways and Strategic Implications

The webinar delivered several critical insights that organizations operating in the AI space must integrate into their strategic planning and operational processes. The immediate enforcement of general purpose AI model obligations represents a fundamental shift in the regulatory landscape that requires prompt and comprehensive response from affected organizations.

The most immediate takeaway is the urgent need for organizations to conduct thorough assessments of their AI models and activities to determine their regulatory status. The computational thresholds and functional criteria provided in the webinar offer clear benchmarks for this assessment, but organizations must also consider the broader context of their AI activities, including downstream modifications and system integrations that might trigger additional obligations.

The tiered approach to obligations, with baseline requirements for all general purpose AI models and enhanced requirements for those with systemic risk, provides a framework for prioritizing compliance efforts. Organizations must not only understand their current obligations but also anticipate how future developments in their AI capabilities might affect their regulatory status.

The treatment of downstream modifiers represents a particularly important consideration for the many organizations that build upon existing foundation models. The one-third computational threshold provides clarity about when modification activities trigger provider obligations, but organizations must also prepare for potential changes in this threshold as technology and market conditions evolve.

The integration of general purpose AI models into larger systems creates additional layers of complexity that require comprehensive compliance strategies addressing both model-specific and system-specific requirements. Organizations must develop

holistic approaches that consider all aspects of their AI implementations and their regulatory implications.

Value for Attendees and Industry Impact

This webinar provided exceptional value to attendees by delivering timely, practical, and actionable information at a critical moment in the implementation of EU AI regulation. The presentation's timing, just one day before the enforcement deadline, ensured that attendees received the most current information available and could take immediate action to address their compliance obligations.

The comprehensive coverage of definitions, thresholds, obligations, and implementation timelines provided attendees with a complete framework for understanding their regulatory requirements. The practical examples and case studies helped translate abstract regulatory concepts into concrete guidance that organizations can apply to their specific situations.

The introduction of the Code of Practice as a potential compliance pathway provided attendees with valuable insight into future developments in the regulatory landscape and a practical tool for organizing their compliance efforts. Understanding both the current requirements and potential future developments enables organizations to make informed decisions about their compliance strategies and investments.

For the broader AI industry, this webinar represents an important milestone in the transition from regulatory uncertainty to practical implementation. The clarity provided around definitions, thresholds, and obligations helps reduce uncertainty and enables organizations to make informed decisions about their AI development and deployment strategies.

Recommended Next Steps for Organizations

Based on the information presented in the webinar, organizations should prioritize several immediate actions to ensure compliance with their obligations:

Immediate Assessment:

Organizations should conduct comprehensive assessments of their current AI models and activities to determine which fall within the scope of the general purpose AI model obligations. This assessment should consider not only the computational resources used in training but also the functional capabilities and intended applications of their models.

Documentation Review:

Organizations should review and update their technical documentation to ensure it meets the transparency and accessibility requirements outlined in the regulation. This documentation should be prepared for potential requests from regulatory authorities and should be suitable for sharing with downstream providers.

Policy Implementation:

Organizations must establish and implement robust policies for complying with EU copyright law in their AI development processes. These policies should address both current practices and future development activities.

Risk Assessment:

Organizations with models that may have systemic risk should conduct thorough evaluations to determine their regulatory status and implement appropriate additional safeguards if required. This includes establishing systems for ongoing monitoring, evaluation, and incident reporting.

Stakeholder Engagement:

Organizations should engage proactively with regulatory authorities, industry associations, and other stakeholders to stay informed about regulatory developments and share experiences and best practices for compliance.

Continuous Monitoring:

Given the dynamic nature of both AI technology and regulatory requirements, organizations should establish systems for ongoing monitoring of their compliance status and regular review of their obligations as their AI capabilities evolve.



The webinar concluded with Monica Fernandez's invitation for attendees to engage further with Nemko through their LinkedIn presence for ongoing updates and through free consultation sessions for organizations seeking specific guidance on their use cases. This offer of continued support reflects the collaborative approach needed to navigate the complex regulatory landscape surrounding AI development and deployment.

The successful implementation of the EU AI Act's requirements for general purpose AI models will require ongoing collaboration between industry, regulators, and other stakeholders. This webinar represents an important contribution to that collaborative effort, providing the information and insights needed for effective compliance while supporting continued innovation in AI technology.