



Association des avocats en droit boursier

Paris, 9 September 2026

Artificial intelligence: the new frontiers of risk

Speech by Denis Beau

**First Deputy Governor of the Banque de France
and Delegated Chairman of the ACPR**

Ladies and Gentlemen,

I am delighted to be with you today to discuss a topic which, over the course of the past few years, has become a major focus of debate: **artificial intelligence (AI)**.

In the financial sector, the potential uses of AI are considerable. A survey conducted by the ACPR in 2025 revealed that almost all banks and insurers now have use cases in production. A recent study by the Autorité des marchés financiers (AMF) also highlighted the diversity of applications available to market players, whether in combating fraud and abusive market practices or detecting investment opportunities, or in the emergence of “augmented” advisers. Furthermore, the impacts from the proliferation of AI have only just begun. Whereas until recently AI was a tool that assisted humans, it can now take initiative and act on its own. Today, companies are developing **agents** that can interact with their environment and carry out increasingly complex tasks autonomously.

But as so often with major technological upheavals, opportunities and risks go hand in hand. Consequently, the issue from my perspective as a supervisor is knowing how to embrace innovation without renouncing the requirements for safety, transparency and resilience that underpin our approach to financial system regulation.

While not aiming to be exhaustive, this morning I would like to discuss three aspects of the risks posed by AI in the financial sector, examining them in light of the responses already undertaken by public authorities and the challenges that still lie ahead. I shall first consider the risks for customers, looking at the issue of control over AI-assisted decision-making (I); then cyber risks, which have hit the headlines recently (II); and finally, the economic and financial consequences of the rapid development of AI.

*

I/ The first set of risks concerns citizens and consumers.

Because behind the sometimes impressive performances of AI systems, we must not lose sight of their very real impact on individuals. A decision based on AI can be flawed. It can also reproduce or amplify certain biases, or can be difficult to explain when the mechanisms underpinning it are complex. In the financial sector, these risks are far from anecdotal: they can affect access to credit and insurance, and undermine customers’ ability to understand the reasons behind a decision that affects them.

Given the challenges, the European Union (EU) has adopted an ambitious framework, with the AI Act. The AI Act is based on a simple principle: the greater the risks an AI system poses to

individuals, the stricter the obligations applicable to it. In the financial sector, the systems used for creditworthiness assessment and for risk assessment and pricing in life and health insurance are classified as “high-risk”. And as such, they will have to meet stricter requirements in terms of governance, data quality, fairness, transparency and human oversight.

The ACPR will be tasked with supervising high-risk systems from December 2027 onwards. We are actively preparing for this new mission, which is a natural extension of our responsibilities in customer protection and prudential supervision. As the European framework continues to take shape, we are also contributing to making certain requirements more concrete. This is notably the aim of the **ACPR’s methodological work**: one particular example is its recent publication of a **discussion paper on algorithmic fairness in the financial sector**, which is the subject of a public consultation that will conclude at the end of September.

II/ I shall now turn to the issues of cybersecurity and sovereignty.

More and more organisations are using AI to strengthen their cybersecurity systems. But, at the same time, AI is contributing to the escalation of cyber threats, in two main ways.

First, **AI enhances attackers’ offensive capabilities** – the most advanced systems can already identify vulnerabilities in computer code, facilitate the design of malware and industrialise social engineering techniques. Second, the **AI systems used by organisations have themselves become targets** – not only end targets, but also “intermediate” targets. Since they are increasingly integrated into the critical processes of financial institutions and connected to their data and tools, they provide access to a host of sensitive resources.

These risks must be addressed, but **we are fortunately not starting with a blank slate**. The DORA regulation has already imposed a robust framework for IT risk management based on principles that are flexible enough to incorporate the risks associated with new uses of AI. The AI Act complements this approach with specific cybersecurity requirements for high-risk AI systems and for the most powerful general-purpose models.

But this framework is probably not enough to face up to the risks posed by the most advanced models: even their designers struggle to contain the dangers associated with them. The most advanced models are beginning to reveal unprecedented cyber capabilities. For instance, during the Hugging Face incident this summer, OpenAI agents managed to break out of their isolated environment. Nearly 700 of them coordinated their actions to participate in an end-to-end attack – a spectacular illustration of AI’s ability to command the entire offensive chain. **Therefore, there is an urgent need to strengthen the existing framework to better manage the dangers associated with frontier AI.**

Given the context, the ACPR is advocating at international level for the introduction of specific safeguards: a gradual and controlled roll-out of the most powerful models, access initially restricted to “trusted partners” (for example at G7 level), or the development of independent assessment capabilities, particularly in Europe. This approach is also promoted by the European Commission in its recently published Action Plan on Cybersecurity and Artificial Intelligence.

III/ I shall conclude by commenting on the economic and financial consequences of the rapid adoption of AI.

In the United States, the development of AI and data centres is having a tangible macroeconomic impact on investment and economic activity, and is generating inflationary pressures. In Europe, the effects certainly remain more modest, and are also more difficult to isolate from other ongoing structural adjustments such as climate change or the ageing of the population. Therefore, in the short term, they do not justify a shift in monetary policy. Over the longer term, we still lack the perspective required to fully assess the scale of the transformations driven by AI, and particularly its effects on productivity and employment.

Compounding these economic uncertainties are issues of financial stability. First, AI could result in major restructuring between companies and between sectors, with consequences on corporate bankruptcies in certain sectors. Second, and in the shorter term, the sheer amount of capital invested in AI and the expectations it generates are fuelling debate over the existence of a potential “bubble”, with the risk that a sharp revision of these expectations could trigger a **disorderly market correction**.

The challenge for Europe is therefore twofold: speeding up progress while managing the risks. Through the European Commission’s “AI Continent” action plan and the “Cloud and AI Development Act”, the EU is explicitly seeking to stimulate the adoption of AI solutions and mobilise significant investment in infrastructure and computing capacity. This ambition is essential to our competitiveness; at the same time, we must keep a close watch on the risks of economic and financial imbalances that could arise as this transformation gathers pace.

*

Ultimately, AI promises to fundamentally transform our societies, our businesses and our economy. **But this promise must be matched by an equally strong resolve to manage the risks it raises**, because the benefits we derive from this technological revolution will depend on our collective ability to understand its effects, to anticipate its consequences and to control its uses. Only on these terms will we be able to build an AI that is safe, responsible and trustworthy.