

AGENTIC AI IN FINANCIAL SERVICES



ONDŘEJ
KOVÁŘÍK

Former MEP, Economic and
Monetary Affairs committee

Agentic AI in European finance

Financial institutions have spent the past decade absorbing successive waves of artificial intelligence adoption. It has transformed various operations from credit-scoring models to chatbots or generative assistants. Agentic AI is different in kind. Agentic systems are able to autonomously decide and execute. They pursue a goal across multiple steps, invoke external tools and payment rails, they are able to initiate and complete transactions with minimal human involvement. That leap makes the technology commercially compelling, and very challenging from a regulatory point of view.

From assistance to execution: the value case

The specific value of agentic AI lies in closing the "last mile" that previous AI generations left open. For instance, a generative model can summarise a loan file. An agent can advance the process much more. It can assemble the file, verify documents, run affordability checks, price the loan and route only the exceptions to a human underwriter.

By performing those tasks, agents allow whole workflows to be redesigned around straight-through processing. It can generate value by following the entire process, starting with onboarding, KYC remediation, payments investigations, reconciliation, reporting and collections.

The second source of value is orchestration. Agents can synthesise unstructured inputs, such as calls or chats, with structured data to anticipate customer needs and act on them, something no single-purpose model could do. Among institutions already deploying agents at scale, roughly two-thirds expect to redesign their operating model, flattening hierarchies and shifting staff from repetitive execution to supervising workflows.

A revolution in operating models and distribution

Early experience and expectations suggest that agentic AI can reshape how financial institutions are organised or how their products are distributed. The approach to AI adoption differs across the market. Incumbents are concentrating investment in back-office efficiency. Digital-native fintechs are embedding agents throughout customer-facing journeys and product design. The strategic risk for incumbents is very often not the cost line but the customer interface.

Another aspect should be considered on the consumer side, when they start to delegate their financial decisions to personal AI agents. The agent advising on a better mortgage rate or insurance policy may become the main point of contact, not the bank's app. This can heavily impact customer preferences and brand loyalty. Banks or insurance companies would face a strategic challenge on how to address changes in customer behaviour driven by agentic AI advice. Distribution built around apps and branches could be disintermediated by an "agent layer" much as aggregators disrupted insurance.

Is the EU rulebook ready?

When it comes to a regulatory framework, Europe does not start from zero. The AI Act's high-risk regime, which captures most consequential financial-sector use cases, already mandates human oversight and auditability with the EBA and ESMA acting as AI Act supervisors within their sectors. DORA addresses operational

resilience and, since the designation of the first critical ICT third-party providers in late 2025, brings major cloud and AI infrastructure under direct oversight. GDPR governs the data agents consume; PSD2 and the forthcoming PSR govern the payments they may initiate.

Yet the framework was mainly designed for human-operated systems, and whether it fits the agentic era should be investigated. Generally speaking, there may be gaps in it. First, consent and authorisation. Payments law presumes a customer authorising each transaction, not a standing mandate delegated to an autonomous agent. The issue of liability when an agent transacts erroneously across a chain of agent providers, PSPs or merchants is not fully resolved. Second, the AI Act's classification is static, while agentic systems are dynamic and self-modifying and can change in nature over time. Third, the overlap risk falls on supervisors too. Interactions between the AI Act and sectoral financial legislation create legal uncertainty and duplicative compliance burdens. It will also be a challenge for all the supervisory authorities mandated for different rulebooks.

Early experience and expectations suggest that agentic AI can reshape how financial institutions are organized or how their products are distributed.

What is needed in the first place is clarification, such as guidance on human-oversight expectations for autonomous agents and explicit liability allocation for agent-initiated payments. The urgency also lies in closer supervisory cooperation and practice convergence so that different authorities' interpretations do not diverge and fragment the single market. Agentic AI will not wait for the guidance to arrive. Neither, on current evidence, will the businesses.



GIUSEPPE SIANI

Director General for
Financial Supervision and
Regulation – Banca d'Italia

Agentic AI in finance: autonomy, resilience and trust

AI is moving from analytical support to operational agency. In financial institutions, agentic systems can observe data, formulate plans, trigger workflows and adapt actions with limited human input. This evolution can bring material benefits: faster fraud detection, more efficient compliance, better risk analytics and more personalised services. Yet the same features that make agentic AI valuable — autonomy, speed, scalability and reliance on complex data chains — can also amplify risks.

The right question is no longer whether finance should use AI but how to deploy it responsibly, and without weakening governance, accountability and trust.

At firm level, operational risks are the first line of concern. Agentic systems embedded in credit origination, payments, trading or cyber defense may act faster than internal controls can interpret. Errors in data, prompts or model configuration can propagate through multiple processes before they are detected. A model may hallucinate, misclassify a transaction, execute an inappropriate escalation or optimise a narrow objective at the expense of resilience. The more these systems rely on cloud infrastructures, external models, data providers and software libraries, the greater the exposure to

third-party concentration, supply-chain vulnerabilities and operational disruptions. Cyber risk also changes in nature: AI can strengthen detection and response, but malicious actors can use it to scale phishing, impersonation, malware development or vulnerability discovery.

Conduct risks are equally important. Agentic tools may shape customer interactions, product recommendations and credit decisions. If the data are incomplete or biased, outcomes may be unfair, discriminatory or difficult to explain. If customers interact with automated systems without clear information about their role, transparency may suffer. If staff rely excessively on AI outputs, human judgement may be weakened precisely where it is most needed.

Prudential risks arise when AI becomes embedded in core risk management and balance-sheet decisions. Credit scoring, early-warning indicators, liquidity forecasting, capital planning and market-risk models can benefit from richer data and faster processing. But if models are insufficiently validated, poorly governed or updated without effective oversight, they may produce correlated errors or understate risks in stress. Automation can also compress reaction times.

This is the bridge from micro risk to systemic risk. Market-wide effects may emerge through common dependencies and common behaviour. Similar models trained on similar data may generate similar trading signals, risk assessments or customer responses. In benign times

**Agentic AI is not a reason
to oppose innovation.
It is a reason to make
innovation governable.**

this may look like efficiency; in stress it may become herding. Feedback loops can arise if AI systems respond to market movements by in parallel actions, accelerating outflows or disorderly trading. Disinformation generated or amplified by AI could further affect market sentiment, especially when combined with high-speed trading and social media dynamics.

Supervisors should therefore look not only at models, but at systems of control. Existing tools remain relevant: governance reviews, model validation, outsourcing oversight, cyber-resilience testing, incident reporting and stress testing already address many of the

underlying vulnerabilities. But they need to evolve. Supervisors should push intermediaries to properly manage AI inventories, critical use cases, degrees of autonomy, human intervention points, data lineage, third-party dependencies and incidents. Testing should be improved, including scenarios in which several institutions react in the same way at the same time. Auditability must cover not only code and documentation, but decision logs, escalation paths and the ability to reconstruct why an agent acted. The adoption of AI must go hand in hand with accountability, including where systems are provided by third parties, as financial institutions remain fully responsible for their use.

The supervisory objective should be proportionate and technology-neutral, but not technology-blind. High-risk uses require stronger controls, clear accountability at board and senior-management level, and credible intervention mechanisms. “Human in the loop” should not be a slogan: humans must have the expertise, authority and time to challenge, override or stop automated actions. Public authorities also need to invest in skills and SupTech, cooperate across prudential, conduct, cyber, data protection and competition domains.

Agentic AI is not a reason to oppose innovation. It is a reason to make innovation governable. In finance, autonomy can be useful only if it remains anchored in responsibility, resilience and trust. These risks are not merely legal or reputational: a system that is efficient but opaque can undermine trust in institutions and in markets. Trust is both a public good and a competitive asset; as such, it remains our top priority.



JIMMY KVARNSTRÖM

Executive Director of Markets –
Swedish Financial Supervisory
Authority (Finansinspektionen)

Autonomy can be delegated to AI agents — accountability cannot

When algorithmic trading transformed European equities, it changed speed and scale, but not responsibility. The firms using those algorithms remained responsible for their use and effects. Agentic AI raises this accountability question in a broader setting. The answer remains the same.

Most AI in finance today informs human decisions. It may score credit, flag transactions or draft analysis. AI agents do something different. They act. They pursue objectives across chains of tasks, interact with other systems and agents, and adapt as they go. They can support stronger risk monitoring, automated controls and better services. But as agents become embedded in trading, risk management, compliance and customer processes, more decisions will happen at machine speed, without a human in each loop.

This raises specific challenges for both firms and supervisors. Accountability can blur when outcomes emerge from chains of automated steps rather than identifiable decisions. Auditability

becomes harder when systems may not produce the same output every time and the basis for an action is difficult to reconstruct. When many agents interact, collective behaviour can emerge that no single firm designed. This may include correlated reactions, herding and, in markets, conduct that resembles coordination without agreement.

There is also a concentration and resilience risk. If many firms build agents on a small number of underlying models and providers, behaviour may correlate and operational dependencies may deepen — a systemic dimension that individual firms should assess and manage in their own model and provider choices. The same capabilities can also make it faster to find vulnerabilities and automate attacks, raising operational resilience risks. When agents are supplied by third parties, it is necessary to know before deployment who performs the regulated activity and who answers for it.

This is not without precedent. When algorithmic trading grew, MiFID II did not prohibit automation. It required firms to demonstrate control. Algorithms had to be tested before deployment, monitored in operation, capable of being stopped immediately, and documented so that behaviour could be reconstructed afterwards. The principle was simple. Autonomy is acceptable where control is demonstrable. This principle also applies to agentic AI. It enables rather than restrains. Firms that can demonstrate control can deploy with confidence.

For firms, these principles translate into practical safeguards, extending the logic of MiFID II's algorithmic-trading rules to a new setting. Every agent should have a clearly identified owner, meaning a function with senior accountability for its objectives, boundaries and conduct. Mandates should be explicit, with limits on what an agent may do and the resources it may commit. Testing should precede deployment, monitoring should be continuous, and intervention, including immediate shutdown, should always be possible. Logging should allow firms to reconstruct what an agent did on what inputs and on what basis. Safeguards should be proportionate. An agent drafting analysis should not be governed in the same way as an agent committing capital.

Human oversight must be meaningful rather than nominal. With autonomous agents, oversight of every action is neither possible nor the point. What matters is oversight of the system — its design, objectives, boundaries and performance — and clarity about who exercises it. A human approving

machine output faster than it can be read is not oversight. Governance should treat agents as the firm would treat anyone acting in its name, with defined authority and active supervision.

Supervision will also need to evolve. When decisions move into processes, supervisors must assess model governance, testing regimes, monitoring and controls, not only outcomes. This requires new capabilities and access to data that makes firm behaviour observable. It also requires supervisors to understand these systems from the inside, including by using them in their own work. The same technology that challenges supervision can strengthen it through sharper monitoring, faster analysis and earlier detection. This makes early dialogue between firms and supervisors important.

**Autonomy is acceptable
where control is
demonstrable.**

EU coordination would help before national practices diverge. The risks of agentic AI do not follow borders, and neither do the firms using it. Common expectations on accountability, auditability, testing, resilience and human oversight would give firms predictability and supervisors a shared baseline. This should build on technology-neutral rules and the AI Act, not duplicate them. Guidance can be common, while supervision should remain close to the firms and markets where agents operate.

The principle to uphold through all of this is well established. Firms may organise how they act through employees, systems or algorithms, but they cannot outsource responsibility for what is done in their name. Autonomy can be delegated to machines. Accountability cannot.



GILLIAN HAMILTON

Cloud Regulatory Response
Lead – Google Cloud

The agentic horizon: balancing innovation and safety in European finance

The European financial sector stands at a critical juncture, where long-term productivity and global competitiveness must be balanced with systemic safety. Unlike traditional generative AI, which processes inputs to generate static text, agentic AI systems act as autonomous digital entities capable of planning and utilizing external tools, and executing complex workflows. In finance, these "digital workers" can revolutionize operations.

In practice, financial institutions are already deploying agentic systems across three highly impactful domains:

Institutions like HSBC are utilizing the Gemini Enterprise platform to deliver automated, personalized wealth advisory services at scale, augmenting human advisors.

Platforms are integrating Google's Agents Payment Protocol (AP2) standard. For example, utilizing a mandate-based security architecture to let AI agents securely execute transactional payments across platforms, ensuring agent actions align strictly with user intent.

Chief Risk Officers are deploying agentic AI to combat threats like Authorized Push Payment (APP) fraud—which bypasses traditional multi-factor authentication—by deploying real-time behavioral analytics and safety-by-design threat modeling.

Yet, this autonomy expands the risk landscape. Google's research identifies two primary risk categories: rogue actions (unintended behaviors driven by prompt injection or reasoning errors) and sensitive data disclosures. Mitigating these risks requires a hybrid posture combining deterministic controls with dynamic, reasoning-based defenses.

Safe deployment requires alignment on three core principles for secure agent design. First, agents must have well-defined human controllers to ensure accountability and prevent unapproved autonomous actions in critical financial scenarios. Second, agent powers must be limited, enforcing robust guardrails, the principle of least privilege through sandboxing and strict credential management so capabilities cannot be escalated. Third, agent actions and planning must be fully observable, utilizing robust, tamper-evident logging and standardized APIs to ensure monitoring, compliance and auditability.

For European financial entities, deploying AI agents is guided by the complementary frameworks of DORA and the EU AI Act. Under DORA, the focus shifts to systemic operational resilience, where the official designation of major cloud providers as critical ICT third-party service providers (CTPPs) underscores their role in supporting financial services with secure, resilient infrastructure. Within this model, financial entities maintain oversight of cloud-hosted workloads—including AI agents—in alignment with the shared responsibility model. Simultaneously, the EU AI Act's risk-based classification ensures high-risk applications are subject to strict data governance, while provenance and labeling tools preserve trust in the financial ecosystem.

As financial institutions transition toward agentic architectures, cloud service providers play a vital role in providing the secure and modern infrastructure required for safe deployment. To operationalize this, Google Cloud offers a secure-by-design foundation tailored for agentic workflows. Through the Gemini Enterprise Agent Platform, Google Cloud treats agents as digital workers requiring the same oversight as any workforce member. Every agent is assigned a unique

identity to enable granular permissions, while a centralized registry and built-in observability tools provide a clear audit trail. Google Cloud published the Secure AI Framework (SAIF 2.0), which offers structured guidance on agentic risks and controls to help organizations have a global risk approach. Google Cloud will not use customer data to train or fine-tune any AI/ML models without the customer's prior permission or instruction. For highly sensitive workloads, customers can leverage Google's sovereign solutions, including Distributed Cloud (GDC) Air-gapped which provides a fully sovereign, isolated run-time environment. Finally, Google Cloud's "shared fate" approach includes generative AI copyright indemnification and the elimination of data egress fees to prevent vendor lock-in.

The successful governance of agentic AI requires a collective commitment to outcomes-based standards.

The successful governance of agentic AI requires a collective commitment to outcomes-based standards, rejecting overly prescriptive mandates in favor of a harmonized, risk and principles-based framework. Google Cloud is fully committed to supporting this transition, offering our Secure AI Framework (SAIF 2.0) and sovereign cloud capabilities to ensure that European financial innovators can safely scale agentic systems. By co-designing these boundaries as trusted partners, we can successfully build a secure, competitive, and highly resilient digital ecosystem.



HEDVA BER

Global Chief Operations
Officer & Deputy Chief
Executive Officer – etoro

The competitiveness test for Europe

Artificial intelligence is no longer a forward-looking experiment in financial services. It is already embedded in core operations. The real debate is not whether AI will scale, but how to scale it responsibly without undermining competitiveness.

AI is no longer a layer etoro adds on top of its platform, it is the platform. From etoro's AI assistant Tori's proactive intelligence supporting millions of investors, to the systems running underneath the business, AI now touches nearly every decision etoro makes and every product etoro ships. What started as experimentation has become the operating model. It's why the mobile app itself was rebuilt AI-first from the ground up. It's why sub-accounts let investors organize every goal separately, why etoro Edge brings professional-grade tools to active traders, and why agent-powered portfolios let investors build or copy agents that trade around the clock, all kept compliant and in check automatically. And it's why an entire App Store and builders' ecosystem has opened up, inviting developers, quants and everyday investors to build the next generation of AI-powered tools on top of the platform..

These are not incremental productivity tweaks. They are structural shifts in how a retail-facing financial institution operates. AI improves efficiency, strengthens oversight and enhances customer experience simultaneously. As I often say

internally, "The biggest risk today is not using AI." The competitive environment is moving too quickly for hesitation.

Across the financial sector, AI is already widely used in anti-money laundering and fraud detection, credit scoring, portfolio management and automated reporting. Generative AI is expanding these applications further, enabling personalized client communication, enhanced analytics and increasingly sophisticated risk modelling. Looking ahead, agentic AI systems capable of autonomous decision sequences may further streamline operational workflows and compliance monitoring.

Will this lead to fundamental change or primarily incremental improvement? The answer is both. In the short term, AI delivers measurable efficiency gains and better customer service. The most profound transformation will not necessarily be visible to customers, but in how institutions manage risk, compliance and product development at scale.

Europe's regulatory context is central to this evolution. The EU is refining its approach through the Digital Omnibus on AI, clarifications around the AI Act. The direction is clear: strengthen legal certainty, improve supervisory coordination and ensure sufficient infrastructure to support AI deployment.

The Digital Omnibus proposals signal a shift toward implementation flexibility by linking high-risk obligations to the availability of harmonised standards. At the same time, the integration of AI into the revised Product Liability Directive reflects a desire for liability clarity without regulatory overlap.

**Europe's ability to
scale AI responsibly
will define its global
financial leadership.**

In this context, the EU's framework provides a strong basis for responsible AI deployment in retail financial services. The AI Act's risk-tiered approach is conceptually sound. The priority now should be effective implementation, supervisory guidance and coordination rather than additional layers of sector-specific legislation.

The answer to AI risk is not prohibition; it is proportionate, risk-based governance. At etoro, AI deployment follows four steps: defining governance policies, classifying each use case by risk level,

establishing a risk-based approval process and implementing tailored controls and monitoring.

When AI is used on the front end to help clients make better decisions, even if it is technically a tool rather than a regulated advisory service, firms that host the ecosystem cannot externalise responsibility. Embedding AI directly into platforms allows supervision, monitoring and market surveillance to apply across the full client journey, rather than leaving users to navigate disconnected environments with no linked governance.

Europe now faces a strategic choice. The debate should not be framed as innovation versus safety. With a clear, proportionate and coordinated framework, AI can enhance competitiveness, reinforce compliance and deliver better consumer outcomes. Implemented too slowly, however, AI adoption could become a business model risk for European firms competing globally.

Scaling AI responsibly is therefore not just a technical or regulatory exercise. It is a competitiveness strategy. The institutions that succeed will be those that combine ambition with governance and say yes to AI, responsibly.



ELENA ALFARO

Head of Global AI
Adoption – BBVA

Agentic banking at scale: The case for governed self- service agents

From answering to acting

The defining shift in agentic AI is not conversational but operational. AI is moving from informing and recommending to using tools and data, coordinating steps and taking action. For banks, this creates new opportunities for productivity, personalisation and process redesign. It also shifts risk from systems that advise to systems that act.

At BBVA, we do not see this as a fresh start. Over the past 2.5 years, we have built a mature foundation for AI adoption. More than 125,000 employees have access to our general-purpose AI tools, OpenAI ChatGPT and Google Gemini. Around three quarters use them intensively, including across our branch network. Employees have created and shared over 16,000 assistants used monthly.

The impact is tangible and measurable. We have completed more than 60 impact assessments, showing savings of 60-90% on specific tasks. An internal survey (+10,000 responses) suggests average savings of 2.6 hours per employee/week, with a satisfaction score of 4.5 out of 5. Yet technology alone does not transform organisations. Adoption requires access, training, communities, governance and

time for people to redesign how work gets done. Agentic AI must build on that foundation.

A portfolio of agents

There will not be one type of agent or development model. The challenge is combining diversity, speed and control. Self-service agents, created on enterprise low-code platforms, can handle workflows such as reviewing documents, preparing reports, reconciling information or moving data between applications. Custom agents will remain essential for customer-facing, highly differentiated or high-volume cases, or when deep integration and precise control over coordination, response time and cost are required. Custom agents may deliver greater impact per case, while self-service agents can scale faster.

Strategic importance alone should not determine the technology choice, and self-service is not only for non-technical employees. A technical team may also find low-code the best solution. The key questions are risk, required control, workflow complexity and total cost and value, not just token consumption.

The first wave of agents is already supporting employees and coordinating internal work: research, document review, reporting, controls and reconciliations. The next wave will reshape customer journeys through faster service, more consistent decisions, fewer hand-offs and greater personalisation. As impact on customers, financial outcomes or regulated processes grows, controls and human accountability must grow with it.

**Banks should neither
centralise every agent nor
decentralise every risk.**

The self-service multiplier

A major opportunity to reshape companies lies in self-service agents. They put creation in the hands of people who know the process and see the opportunities, allowing business teams to test ideas in days. They can unlock thousands of small opportunities that, together, create significant value but would never justify separate software projects.

Self-service cannot mean that everyone builds and publishes whatever they want. We envision a distributed but

structured model: business areas own a prioritised portfolio; trained builders create and maintain agents; and a central team provides approved platforms, reusable skills, secure connectors, evaluations, monitoring and governance. Low-code agents can later move to custom development when volume, risk or optimisation justifies it.

Governed decentralisation is possible

Agent governance should be continuous and proportionate to risk. It should follow what an agent can do, not which product or area it belongs to. The key dimensions are autonomy, data sensitivity, potential impact, scale of exposure and whether the agent can read, write, transact or trigger actions without human supervision.

Banks should start with low-risk cases: read-only access, reversible actions and human approval for material decisions. Before deployment, agents need scenario-based evaluations, safe connector testing and evidence that they behave correctly in edge cases, not only expected situations. Testing must continue when models or platforms change.

Once deployed, agents should have only the minimum access required, approved actions, policy checks before execution and human approval for high-risk activity. Every execution should be traceable across the agent, user, data, tools and resulting action. Monitoring must cover quality, changes in behaviour, anomalies, cost and incidents, with reconciliation, suspension and decommissioning mechanisms.

The sustainable advantage will not come from any single model or tool, but from a common control layer that connects agents safely to the banking ecosystem and provides permissions, reusable skills, evaluations and monitoring. To make the most of this era, banks should neither centralise every agent nor decentralise every risk. Business-led creation, common infrastructure and proportionate controls can turn agentic AI from a portfolio of projects into a capability embedded across the bank.



TOMAS KLENKE

Managing Director and
Head of AI Solutions & Data
Insights – Commerzbank AG

Agentic AI in banking: advancing innovation with trust and responsibility

Artificial intelligence is becoming a strategic capability for Europe's financial sector. In a period marked by geopolitical uncertainty, demographic change and increasing demands on productivity, banks need to adapt fast. At the same time, trust remains the foundation of banking. This is particularly relevant as the industry moves from generative AI towards agentic AI.

For Commerzbank, AI is a core element of our Momentum 2030 strategy. We use it to improve customer experience, risk management and operational excellence. To this end, we have established a strategic AI program which drives the implementation of prominent AI use cases across the bank together with the technological and data foundation to scale them efficiently. Until 2030, we plan to invest around €600 million in AI and expect our initiatives to contribute around €500 million per year from 2030 onwards.

Commerzbank already started to explore and implement AI use cases ten years ago. Today, the bank uses artificial

intelligence across the entire value chain. Our virtual assistant Ava supports customers around the clock by answering everyday banking questions and ensuring a smooth digital banking experience. Internally, AI automates standard tasks, increasing efficiency and effectiveness while relieving employees of repetitive administrative work. In advisory centers, digital assistants support employees in real time by summarizing conversations and proposing suitable solutions. Throughout the organization, AI is expected to free up and partly redeploy around 10% of capacity.

Agentic AI is a particularly promising next step. While generative AI primarily creates content, agentic AI can execute complex workflows or manage entire problem-solving processes within defined guardrails. Today, it is still an emerging technology. Current applications therefore require realism and discipline. This includes case-by-case decisions as to which type of AI represents the best solution from a cost-benefit perspective.

The transformative potential of agentic AI lies in moving beyond single-task automation. Banking processes often involve multiple handovers across systems, documents, control functions and client interactions. Agentic AI can help orchestrate these steps, reduce breaks in the process and provide employees with better prepared outputs.

Our AI agents will essentially cover three areas of application:

Agile Delivery Agents will support our technology teams in software development, for example through code generation and testing.

The transformative potential of agentic AI lies in moving beyond single-task automation.

Business Support Agents will assist our employees with day-to-day tasks such as documentation, meeting preparation and follow-up, as well as research.

Operations and Process Agents will orchestrate and execute complex end-to-end processes, from Know-Your-Customer checks to document reviews and contract preparations.

However, in a regulated and trust-sensitive industry, the key question is not only what agentic AI can do, but under what conditions it should operate. The EU AI Act provides the central regulatory

framework, with a risk-based approach. Commerzbank's AI governance framework ensures compliance with applicable regulatory requirements and is continuously developed further to reflect latest AI innovations and regulatory updates alike.

At Commerzbank, the effective deployment of agentic AI is therefore ensured by a disciplined governance approach across the entire AI lifecycle, differentiating the treatment of use cases by risk level. Use cases are assessed by risk and purpose, with stronger validation and oversight for higher risk classes. Testing goes beyond technical performance and cover data quality, model risks, fairness, security and failure modes. Monitoring is continuous, because agentic systems interact dynamically with processes and systems. Above all, human accountability remains explicit: AI can automate and accelerate processes, but responsibility must remain under human control. At Commerzbank, humans therefore always define the boundaries.

Policy makers and supervisors can support responsible scaling in three ways:

First, the EU AI Act should be implemented in a risk-based and proportionate manner.

Second, supervisory expectations should allow controlled experimentation before scaling.

Third, requirements across AI governance, data protection and financial-sector regulation should be aligned to avoid overlapping documentation and control burdens. Similarly, inconsistencies and overlaps in supervision should be avoided. This could be achieved by coordinating responsibilities and consolidating supervisory activities wherever possible.

Agentic AI can make banking more efficient and more personal. But it will only succeed if innovation and governance develop together. The future of banking will be shaped by a thoughtful interaction between human and artificial intelligence – with AI deployed responsibly and in a way that benefits both our organisation and our clients.



KELLY DEVINE

President –
Mastercard Europe

Trust is Europe's advantage in an AI-powered payment ecosystem

Most of us are already getting used to AI helping us do things faster: summarising information, cutting through a crowded inbox, planning a trip or finding precisely the product we had in mind. It is not a big leap to imagine AI helping us shop and pay. Consumers will increasingly be able to ask AI-powered agents to search, compare and buy on their behalf. Businesses will use AI to improve customer experiences, reduce friction and find new ways to grow. As this happens, payments will sit at the centre of a more intelligent digital economy.

That creates a significant opportunity for Europe. It also raises a simple question: how do we encourage innovation while protecting consumers and earning public confidence?

The answer is not to slow innovation down. It is to build the trust that allows innovation to scale. If we treat trust only as a compliance obligation, we will miss the bigger opportunity. For AI-powered commerce to work, three things must be clear: who is acting, what has been authorised, and who is accountable.

Payments have always been built on trust. For decades, Mastercard has helped develop the infrastructure, standards and

security that underpin digital commerce. In an AI-powered economy, that means being able to verify identity, capture consent and understand responsibilities at every step.

Digital identity is therefore more than a technical requirement. It is one of the foundations of consumer confidence. It can help confirm that the right person or business – or their authorised agent – is taking the right action, while making digital experiences simpler and reducing fraud.

At Mastercard, we are developing the capabilities that will be required in the future: connected, multi-rail experiences that are ready for AI. Payments will not be defined by one technology, one rail or one provider. Consumers and businesses will continue to choose between cards, account-to-account payments, digital wallets and emerging payment models. What matters is that these choices work together securely and seamlessly. In a diverse European payments landscape, interoperability is not optional; it is essential.

Agentic commerce makes this even more important. Consumers may soon instruct AI agents to act within clear limits: finding the best travel option, reordering household items or comparing products before making a purchase. But AI should not blur responsibility. A consumer should never be left wondering whether they made a choice, or whether an algorithm made it for them.

AI will transform commerce. Trust is the foundation for scaling that transformation.

Through solutions such as Mastercard Agent Pay, and through our work with banks, merchants and technology partners, we are helping to develop secure ways for AI-enabled transactions to take place. The principle is straightforward: AI can make commerce faster and easier, but people should know what they have authorized, who is acting for them and how they can stay in control.

This is not theoretical. Payments already depend on real-time trust at significant scale. In 2025 alone, we processed and secured 175 billion transactions, using advanced analytics and network intelligence to identify emerging threats. Over recent years, our fraud prevention solutions have helped prevent more than €9 billion in fraud across Europe.

Europe is well positioned to lead this transformation, not by choosing between innovation and protection, but by showing how the two can reinforce each other. The EU already has a strong foundation, from payments and data protection to cybersecurity, digital identity and the AI Act. As AI evolves, Europe should build on that foundation with rules that are technology-neutral, focused on outcomes and designed to avoid unnecessary fragmentation.

We continue to invest in that trust infrastructure. Since 2019, Mastercard has invested more than €10 billion globally in cybersecurity, anti-fraud and AI capabilities. The aim is not only to stop fraud, but to make legitimate commerce faster and easier.

Investment is necessary, but it is not enough. Public-private cooperation will be critical, because no single organization can solve these challenges alone. Fraud increasingly starts outside the traditional payments ecosystem: a fake advert, a fraudulent message or an impersonation call may all occur long before a payment is made. Protecting consumers will require faster intelligence sharing, trusted data-sharing and digital identity frameworks that work across borders and sectors. Cooperation needs to move from principle to practical execution.

Mastercard's role is to help turn that ambition into action: bringing trusted technology, secure infrastructure, open standards and strong partnerships together across Europe.

If we get this right, AI will not simply change the way people pay. It can help unlock growth, security and opportunity for consumers and businesses across Europe. But trust is built in the details: in the way identity is checked, consent is captured, risk is detected and responsibility is understood. That is where Europe can lead.