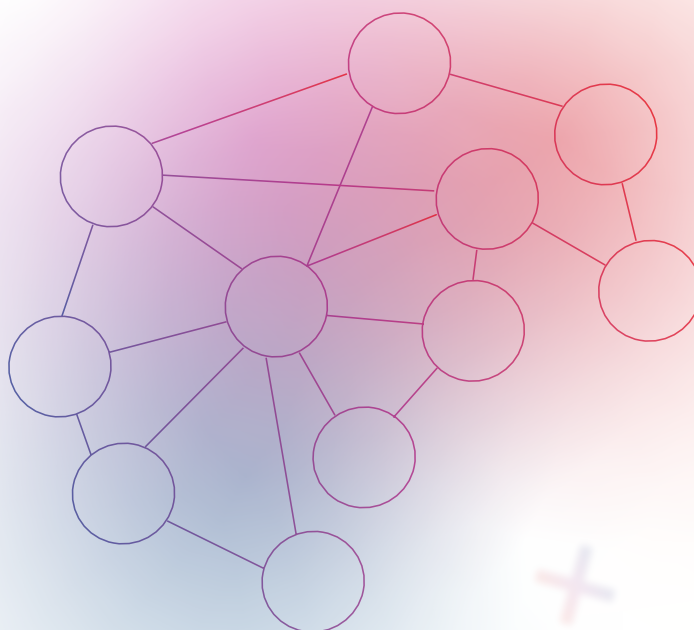


POLICY PAPER



Agentic AI: **Deployment, Adoption and Impacts**



June 2026



Generative AI has recently evolved from models that primarily produce text, images, or code into systems capable of autonomous decision-making and proactive behavior within human-defined constraints. These agentic AI systems, possess the capacity to perceive the environment, plan actions, execute them, and reflect on outcomes, often operating continuously with varying degrees of autonomy and adaptive behavior.

We argue that agentic AI should be proactively shaped to maximize beneficial outcomes while limiting societal, security, and environmental impacts. Rather than solely documenting the current landscape of agentic AI adoption, this paper aims at advancing specific policy recommendations designed to shape the trajectory of agentic AI development and deployment.

This document is a high-level excerpt of the analytic policy paper identifying key evolutions to anticipate and actions to take. This paper was developed through iteration with many actors (academics, industry- and government representatives), in order to develop a comprehensive and overarching panorama. In this regard, all elements regarding contributions and details will be in the analytical paper.

—————> What is Agentic AI?

Agentic AI is an emerging paradigm for which there is no broadly accepted definition yet, neither legally, nor technically. However, through usage and analysis, specific key characteristics can be identified, drawing out common architectural patterns: autonomy & execution; planning & reasoning; memory & statefulness; action and tool use; multi-agent coordination. **Most importantly, Agentic AI systems have a level of autonomy to achieve goals, based on orchestrated multi-agent collaboration and connection with their environment.**

—————> How is it used?

We consulted 30 companies across G7, covering sectors from energy and finance to software. 83.3% of respondents expressed a "very strong" intention to explore or invest in agentic approaches.

Agentic AI presents opportunities in streamlining, detailing, personalizing or accelerating processes. Agentic systems are finding early traction in domains where tasks are structured, outputs are verifiable, and failures are recoverable. In customer-facing functions, agents are deployed for intelligent service orchestration, autonomously handling interactions and escalating only ambiguous or high-stakes cases to human operators. In financial services, Agentic AI is being used for real-time risk monitoring and compliance automation. In software engineering, agentic systems are being integrated into Continuous Integration/Continuous Deployment pipelines as autonomous quality gates, capable of generating test cases, detecting regressions, and proposing corrective patches.

—————> Why should we take this paradigm shift into account?

Use-cases are still nascent, but the transformative potential is well identified, especially for internal processes. Some broad objectives in the deployment of Agentic AI can be set, anticipating risks and pitfalls, to make its deployment a success.



1. Traceability and explainability

Agentic AI systems possess long-term memory that shape outputs, which can make it difficult to reconstruct the past context that influenced a specific result. Furthermore, emergent behaviors can spread across agents, complicating the attribution and reproducibility of decisions. Silent model updates can invalidate explanations that were previously accurate, while evolving data sources further challenge reliability. These characteristics make thorough auditing difficult.

—————> **Strategies for successful deployment:**

- Comprehensive logging of raw prompts and system messages, memory reads and writes, tools and API calls with inputs and outputs, and intermediate reasoning traces.
- Versioning of all components, including models and decoding parameters (e.g. temperature, top-k/top-p, random seeds, and hardware configuration).
- Controlled randomness, supporting fixed random seeds where possible and allowing deterministic decoding modes (e.g. greedy).

Explainability methods for long-term memory systems, causal attribution in agentic systems and formal verification of autonomous reasoning represent present scientific gaps.



2. Security, safety and harm prevention, including controllability in multi-agent organization

Agentic AI operates as a loop, which can lead errors to cascade and biases to accumulate, thereby influencing decisions. Over time, this can lead to deviation from the intended objectives. Furthermore, Agentic AI presenting a broader surface open to its environment, it is vulnerable to adversarial inputs, such as attacks, malicious prompts, that may divert it from its original goals.

Controllability, as inspired by control theory, sets the objective to ensure humans can steer the system's trajectory through oversight, constraints, and intervention, which can be a challenge with systems that have little explainability. Indeed, through observability, this approach can help counter internal dynamics, such as the system adapting to feedback, which can cause its actions to gradually diverge from the intent, a challenge amplified in long decision chains.

→ Strategies for successful deployment that limit error amplification or misuse:

- Multi-sources validation with retroaction loops to detect and correct errors, hallucinations or model redundancy such as clear allocation of control authority, defining which actors are responsible for setting objectives, enforcing constraints, and supervising system operation.
- Repetition of objectives at several stages, differentiating entries between user- and system- instructions, control the access to tools and regular testing.
- Hold the collection of private data to a minimum and plan a strict conservation duration.
- Adapt the security measures to the Agentic Systems, which are more complex.
- Architectural control points, enabling inspection and intervention at key stages of the process.
- Constraint enforcement mechanisms, such as policy layers, sandboxed execution environments, and access controls for external tools.
- Regular testing of control mechanisms, ensuring that intervention procedures remain effective and that operators are prepared to respond to abnormal behavior.

Error propagation in agentic systems, formal methods for controllability and alignment maintenance and objective drift detection methods represent current scientific gaps..



3. Limit energy consumption

Energy Efficiency & Agentic Flow: Distributed, autonomous, and interconnected agentic systems (via APIs, cloud services, external tools) make deadlocks and uncontrolled action chains harder to detect and resolve. Their integration into broader infrastructure and intertwined multi-agent nature further obscures their specific energy footprint.

Optimization Under Constraints: Agentic AI's potential lies in surpassing human process efficiency by optimizing action trajectories and leveraging information processing for energy savings, especially under external limits such as energy scarcity, material costs or regulations. However, an important uncertainty of these efficiency gains lies in the potential rebound effects which might hinder these gains.

Energy impact of multi-agents' architecture choice and agents' conception: Multi-agent architecture (coordination and consensus mechanisms) influences the number of LLM calls, communication overhead and memory complexity, with an impact on energy use. Agents' model and architecture also matter: smaller, well tuned LLMs with information preprocessing can deliver equal or better task performance than large models while consuming less energy.

Reasoning Models' Energy Use: Energy consumption is less predictable, since it depends on model size as well as on the reasoning path complexity and the associated token usage. First benchmarks indicate that reasoning models consume more energy than traditional ones.

Black Box & Societal Awareness: Automation and opacity reduce public awareness of AI's energy impact.

→ Strategies for successful deployment:

- Pre-deployment energy/performance assessment.
- Selection of energy-efficient data centers.
- Robust loop detection and termination criteria.
- Calibrate model's size and reasoning depth of each entity to the expected outcome.

Energy-aware agentic architectures and lifecycle energy assessment measuring and attributing energy consumption across distributed agent pipelines warrant scientific attention.



4. Anticipate the impact on workers

Agentic AI has the potential to ease processes, but could represent a major shift in work relations and organization. Agentic AI reinforces the anthropomorphism of AI, which can cause over-reliance on those systems and harm human relations on the workplace

→ Strategies for successful deployment:

- Evaluate the human impacts of deployment ex-ante and train workers to the use of agentic AI.
- Anticipate the impacts on the organizational structure and on workers' well-being ex-ante.
- Evaluate ripple effects of organizational changes, such as impacts on tax incomes.

How are regulations covering Agentic AI?

Agentic AI can be used to “delegate” decisions or powers by professional and non-professional users (for instance to negotiate and sign contracts), which is a configuration that was rarely anticipated: regulatory approaches such as manipulation of human behavior, human oversight and assumption of human-machine interaction seem ill-suited to this evolution. However, precision and adaptation of the current framework appears as a better approach rather than additional regulations.

Data protection: The multi-agent and autonomous nature of these systems and their openness to the outside world can lead to non-consensual access and use. The number of actors accessing the data may expand without notice, as can the purposes for which the data is used. In addition, obsolete data might be retained and used at a later stage, potentially by external or new actors.

Competition law and consumer protection: Agentic AI can strongly impact market transparency if used to choose products on a market. If it used for pricing or logistics, it might generate involuntary collusion by levelling strategies. Given the barriers to entry of the Agentic AI market, concentrations or unfair advantages might emerge.

Liability: Given the multi-agent characteristic of Agentic AI, a difficulty to attribute responsibilities can emerge, including intellectual property. In addition, in international settings, it can create difficulties to identify the applicable regulations.

→ Strategies for successful deployment:

- Existing regulations (including soft law), such as the EU framework appear as sufficient, but might require a dedicated analysis, to identify potential gaps and how it can apply to Agentic Systems.
- Analyze existing regulations applying to AI (including soft law) in the wake of Agentic AI and potential gaps, looking at regulations applying specifically to AI, but also competition, data, etc. to add specifications, in order to cover Agentic AI.
- Include regulatory constraints in the very architecture of Agentic AI systems, in a logic of “code is law” to ensure compliance.

How can decision makers ensure a beneficial deployment of agentic AI? – Ordered by urgency

Priority 1 - Ensure sovereignty - The current Agentic AI landscape is dominated by a small number of providers. Decision makers can invest in domestic Agentic AI solutions across the entire stack (foundation models, orchestration platforms, data infrastructure, and execution environments), while diversifying suppliers to ensure autonomy.

Priority 2 - Establish a clear and operational definition of Agentic AI - A definition grounded in the technical specificities of Agentic AI systems is necessary to provide a regulatory anchor. It could be developed by a consensual authority, to be largely adopted.

Priority 3 - Foster a fertile ground for Agentic AI research - Given the rapid evolution of the field, it is essential to establish agile research funding mechanisms. Allocating short-term grants reflecting the inherent uncertainty of Agentic AI research, would enable timely generation of knowledge. Researchers must have access to state-of-the-art tools and platforms.

Priority 4 - Ensure legal security for users and developers - The current regulation of Agentic AI by regulatory frameworks, including competition law and data protection warrants precision and sometime adaptation to cover specificities of Agentic AI. Regarding the AI Act, a targeted review of Articles 5, 9, 14, 15, 50 and 55 would be particularly advisable to take into account the specificities of multi-agent interactions. Clauses addressing liability and specific safety risks appear necessary.

Priority 5 - Adopt an impact-centered governance approach - Developing governance frameworks requires an interdisciplinary standpoint from the outset, involving technical experts, but also lawyers, linguists, psychologists, sociologists, and environmental scientists.

Priority 6 - Define common ground to compare energy consumption - Develop an energy score, based on common benchmarks, to give insights to developers and users on the energy impact of the chosen models and multi-agents architecture. Such an approach should prevent retro-engineering practices revealing industrial secrets, but would require auditing to ensure the veracity of the scores.

Priority 7 - Address cybersecurity risks specific to Agentic AI - Agentic AI's open-loop interaction with its environment introduces novel cyberattack risks. The opacity of many commercial models further prevents the deep auditing required in regulated sectors. Dedicated analysis beyond the present work is required to anticipate the scale and diversity of these threats.

Priority 8 - Anticipate workforce transformation and cognitive impacts - The development of agentic AI systems calls for dedicated research into workforce transformation, skill erosion, and emerging phenomena such as “cognitive surrender.” Social sciences must play a central role in analyzing these dynamics and informing policies, including education, training, and labor market adaptation strategies.

Priority 9 - Strengthen public awareness and education on Agentic AI - Public education and AI literacy can ensure critical use for the broad public and more specifically for workers using Agentic systems. This includes clarifying how agentic systems differ from other IT tools and raising awareness of risks such as over-reliance, reduced human agency, and hidden resource consumption.

What is Inria's input?

An interdisciplinary task force at Inria's Program Agency has drafted a policy paper on this subject, to deepen the elements presented here. Through an iterative process, the initiative engages diverse stakeholders (including industry, academia, and policymakers) to address the challenges and deliver actionable, forward-looking insights for decision makers. They organized 5 workshops with over 80 participants in total from academia, the public and the private sector from all G7 countries, to collect feedback and a perspective from practitioners.

