

Why agentic AI projects fail—and what it takes to make them work

September 8, 2026

An analysis of the gap between experimentation and enterprise adoption

Bringing AI agents from experimentation to real-world use within a company is a challenge that goes far beyond choosing the right technology: processes, data quality and accessibility, integration with existing systems, governance, skills, and much more all come into play. Moving from an effective prototype to a reliable and sustainable solution therefore requires addressing both technological and organizational dimensions. **Massimiliano Luca**, a researcher at the Center for Augmented Intelligence, **and Alberto Purinan**, from FBK's Corporate Relationship Office, bring together the perspectives of research and industry engagement to analyze the barriers to agent adoption and the conditions needed to generate concrete value.

An **AI agent** is a system designed to pursue a goal through a sequence of activities. For example, an AI agent can gather information, assess the necessary steps, consult applications and databases, generate an output, or initiate actions within defined limits.

Consider a manufacturing company. When asked, "How many microchips do we have in stock?", a chatbot simply provides the latest data. An **agent**, on the other hand, can recognize that the number has fallen below the minimum threshold and, rather than simply providing an answer, take action: it checks the available budget in the company's management system, compares prices through suppliers' APIs, pre-populates the purchase order, and submits it to the person responsible for final approval. In this example, the agent operates within a well-defined framework of purpose-built safeguards and controls, handling specific and always verifiable activities. Its true value emerges above all in more complex business processes, which require orchestrating heterogeneous information sources, carrying out a variable sequence of steps, and adapting flexibly to cases that are never exactly the same.

Today, the market offers many ready-to-use solutions, most of them general-purpose, alongside numerous experiments at the *proof-of-concept* level, while the transition to large-scale production remains an open challenge. According to McKinsey's State of AI 2025 survey, 23% of surveyed organizations are scaling at least one AI agent system, while 39% are still in the experimental phase. The share of organizations reporting that they have successfully deployed agents at scale does not exceed 10%. Although these figures are from 2025 and the technology is

evolving rapidly, the reasons behind this gap between experimentation and adoption remain. Proofs of concept work because they often operate on selected data and linear workflows in protected environments. General-purpose commercial solutions, on the other hand, must contend with the realities of the enterprise: access permissions that vary across departments, legacy systems, and archives that are not always complete or well organized. As a result, these solutions often lack transparency, struggle to integrate with business data, and make costs difficult to predict. Overcoming these challenges requires technology and organizations to evolve together.

One of the main strategic mistakes, therefore, is to treat agentic AI as a product to be installed rather than as a **process to be built** together with the organization. No matter how advanced, technology cannot independently understand the exceptions, informal rules, and quality criteria that are unique to every company. Before introducing an agent, it is therefore essential to carry out preparatory work to define strategic objectives, redesign workflows, and establish which activities can be delegated, who should intervene in exceptional cases, and which indicators should be used to measure success. Without this analysis, there is a risk of automating inefficiencies without realizing it. In fact, an agent may work all too well on a process that should never have been replicated in the first place. The outputs of AI systems often appear polished and convincing, but this polished presentation can mask inaccurate content. Developing the ability to critically assess these outputs within operational functions requires training and change management, both of which are necessary and fundamental conditions for innovation to deliver the expected results.

A good Large Language Model alone does not make a good agent. The **technical challenges** arise primarily from the way the agent uses the model within the process.

A first problem, which also extends to the organizational dimension, is **cumulative error**. An agent can correctly perform several steps and still produce an inadequate result if it misinterprets information, uses the wrong source, or selects the wrong tool. Validations and checkpoints are therefore needed at critical stages.

A second issue is **traceability**: understanding why an agent produced a particular result. Sources, instructions, and tools used must be recorded in order to correct errors and meet audit and security requirements.

Memory must also be carefully designed, because company information cannot depend on the memory of past conversations. Instead, it must come from up-to-date, governed sources with clearly defined permissions.

Finally, there is the question of choosing the **most suitable model** for the task. Large general-purpose models offer broad capabilities but are often inefficient for repetitive or specialized tasks, where smaller models can provide faster responses at lower cost. The best solution often combines multiple approaches, including language models, document retrieval, databases, and deterministic rules. The guiding principle remains the same: use only the level of complexity necessary to achieve the required result.

Over the next three years, it is realistic to expect a shift from isolated tools to coordinated services, with defined roles and controlled interactions among multiple software components. Multi-agent systems, in which several specialized agents collaborate and coordinate to achieve a common goal, will be useful when there is a clear division of tasks: for example, one component for document retrieval, one for verifying rules, and one for preparing a proposal for validation. Without

this clear division of responsibilities, increasing the number of agents risks adding costs and complexity without improving the outcome.

People's work will increasingly focus on supervision, exception management, data quality, and process design. An agent may reduce the time spent collecting and preprocessing information, but decisions about what constitutes an acceptable outcome, which exceptions require different treatment, and where a process needs to be corrected will continue to depend on domain expertise and professional accountability that the system does not possess.

In a market that offers ready-to-use general-purpose solutions but still struggles to move experiments into production, the distinctive value of a research center like FBK is not to provide prepackaged tools, but to **advance the frontier of knowledge through specialized, tailor-made solutions**, in a process that enables technology and organizations to evolve together. Companies bring their concrete problems, data, and operational constraints. FBK brings scientific expertise, specialized architectures, and rigorous evaluation methods. The goal is to provide the evidence and expertise needed to bring efficient systems into production, integrate them with company data, and ensure full compliance with regulatory and data sovereignty requirements. It is this collaboration between research and industry that makes the difference between an agentic AI project that fails and one that truly works.

PERMALINK

<https://magazine.fbk.eu/en/news/why-agentic-ai-projects-fail-and-what-it-takes-to-make-them-work/>

TAGS

- #agenti
- #agentic ai
- #artificialintelligence
- #augmentedintelligence
- #aziende
- #llm
- #mobs
- #multi-agent

AUTHORS

- Massimiliano Luca
- Alberto Purinan