

MASSIMO TAMIATTI

FROM ALGORITHMIC INVASION TO THE **AGENTIFICATION** OF WORK

Governance, ethics and the future in the age of AI



What is it? This is a scientific dissemination website that studies the impact of new technologies on the world of work. Changes today are very fast and exponential. Professions are not static but very dynamic, so the ongoing evolutions can make them obsolete, transform them, or create entirely new ones. What skills will be necessary to face all this? Upskilling and reskilling, but not only.

Why? After 25 years of research, I have realized that to go beyond the boundaries, numbers are not enough. Traditional approaches are obsolete and tied to the old way of doing research. Today, we need to study the weak signals that are right before our eyes but that we see without really noticing. Stopping to consider these signals and imagining multiple possible futures will be the task of Frontiere. Traditional sources respond in real-time but only photograph the present. Today, however, there is a need to make predictions at least five years ahead, without neglecting the past and without forgetting that every topic addressed has a history to consider.

Who is it for? It's easy to make long-term predictions that cannot be verified. I am convinced that the research world needs an immediate confrontation with the world of education and orientation, and with young people who risk studying for professions that may not exist in a few years. They will be here to exchange opinions with the undecided, the disoriented, and those who intend to start over. Young people need to understand that there is a world above and a world below in the job market. A world of those who program algorithms (few and well-paid) and those who work for algorithms (many and poorly paid). We need to understand what to do to escape a job market that risks condemning young people to €1,000 per month salaries for life.

How is it? Each topic is treated with light publications and a special focus on the aspect of temporal depth. The publications will be accompanied by short videos that get to the point quickly. Frontieretecnologia.com positions itself as a space for information, guidance, and discussion precisely on these themes.



Linkedin: <https://it.linkedin.com/in/massimo-tamiatti-77605168>

Facebook: <https://www.facebook.com/profile.php?id=100006516969397>

Telegram: <https://t.me/frontieretecnologialavoro>

Instagram: <https://www.instagram.com/tamiattimassimo/>

PRESENTATION

After the explosion of generative AI in 2022-2023 a new technological frontier is emerging: autonomous AI agents. Unlike chatbots that respond to individual requests, these systems perceive the environment, plan complex strategies, and act in the real world without constant supervision. This essay explores three critical dimensions of the agentification of work: (1) the implications for governance and legal responsibility in contexts where decisions are delegated to probabilistic systems; (2) the radical transformation of required skills, from the logic of "doing" to the logic of "being"; (3) three plausible scenarios for the Italian labor market by 2030, between opportunities for human augmentation and risks of social polarization.

Table of Contents

From Assistants to Builders: What Changes	p. 6
The Role of Governance and Human Control	p.12
Where Are Agents Being Used?	p.17
What Are the Unresolved Problems?	p.21
Who Is Needed to Govern These Problems?	p.22
Rethinking Skills: From the Logic of Doing to the Logic of Being	p.25
The Return of the Philosophers: Why Humanities Skills Are the New Engineering?	p.32
The Crossroads of 2030: Three Scenarios for an "Agentified" Future	p.34
Conclusions	p.36

FROM ALGORITHMIC INVASION TO THE AGENTIFICATION OF THE LABOR MARKET

In recent years, much has been said about Artificial Intelligence, but today the great new topic is AI agents. They are being discussed everywhere: there is "hype," enthusiasm, fears, promises of revolutions in work and productivity. But before letting ourselves be carried away by enthusiasm, it is worth pausing and better understanding what these agents really are, what they can do today, and what their real limits are.

From Assistants to Builders: What Changes?

Until recently, when we thought of Artificial Intelligence in the workplace, we imagined assistants to talk to: ChatGPT that writes the email for you, a chatbot that answers frequently asked questions, a text generator that produces content on demand. All tools that wait for our input, process something in their "digital brain," and return an output: a text, an image, an idea.

But while ChatGPT waits for you to ask it to write an email, an AI agent can autonomously monitor your inbox, identify urgent messages, draft context-appropriate response drafts, consult company databases to verify information, and even schedule meetings, all without you having to give specific instructions for each step.

In short, we are moving from systems that "suggest" to systems that "act" autonomously. This is an incredible leap forward. According to

the McKinsey Global Survey 2025, 23% of organizations are already scaling agentic systems¹ and an additional 39% have begun experimenting with them.

AI agents¹ represent a substantial rupture from this model. They don't just respond: they act. They receive a complex objective, break it down into successive steps, mobilize various tools (APIs, databases, online services), make decisions on the fly, and arrive at a final operational result, all without asking for confirmation at every step.

Take platforms like Devin or Manus: the user says "make me a site with these features" and within 40-50 minutes receives a functioning site. Devin, in particular, acts as a "full stack AI developer"² capable of writing code, testing it, and deploying the application—roles that would normally require different people with distinct skills.

So we are no longer just talking about "generating a text" or "summarizing a document": we are talking about doing a complete job, which can include writing code, testing it, querying databases, orchestrating other digital entities, and negotiating between tools and systems.

The real promise in the business context is direct and provocative: in a company, you could have an agent covering multiple professional roles, at a lower cost than one or more human resources. From this

¹ "scaling agentic systems" refers to growing or adapting systems composed of intelligent agents efficiently, maintaining or improving performance, reliability, and control as complexity increases (more agents, more users, more data) or when operating conditions change.

² A full stack AI developer is a versatile professional who manages the entire development cycle of AI-based applications, covering both technical and operational aspects. This figure combines traditional software development skills with specific AI expertise, enabling end-to-end solutions without relying on specialized teams.

point on, the natural question is: but is it really that simple? Can we really delegate complex decisions to systems that, however intelligent they may be, remain software?

1. The Natural Evolution of AI: From Thought to Action

To understand how we got here, it is useful to quickly review the trajectory of Artificial Intelligence in recent years. From the explosion of ChatGPT at the end of 2022, generative AI has dominated the scene with large language models (LLMs),³ trained on enormous amounts of unstructured data, capable of answering questions, writing texts, translating, and coding. These models have demonstrated an extraordinary ability to "understand" natural language and adapt to scenarios on which they were not explicitly trained. But this was still a form of reactive AI: you receive a question and you respond. The algorithm has no purpose, no long-term memory, or decision-making autonomy. It is a tool that waits.

AI agents represent the natural evolution of this capability. They combine the linguistic and reasoning abilities of generative models with a capacity for autonomous action: perception of the environment, strategy development, action planning, and execution to pursue a specific objective without constant supervision. An agent perceives structured and unstructured data (what it finds on the web, in documents, in databases), develops strategies to achieve the objective, plans a sequence of actions, and executes them through

³ LLM: Large Language Model are advanced artificial intelligence systems capable of understanding, analyzing, and generating text in natural language. They are called "large" because they contain billions of parameters and are trained on enormous amounts of textual data, such as books, articles, and web content.

tools and APIs.⁴ Unlike traditional agentive systems, which required laborious programming based on rigid rules, modern agents built on generative foundation models have the potential to adapt to different scenarios: the same agent can plan a complex travel itinerary, manage logistics across multiple platforms, collaborate with other agents, and learn from feedback.

According to McKinsey and other research sources, this transition "from thought to action" represents the next frontier after the explosion of content generation. The economic value that can be glimpsed is enormous: McKinsey estimates that generative agents could bring between 2.6 and 4.4 trillion dollars of added value per year globally.

AI agents are therefore not a technology separate from language models: LLMs constitute the central nucleus (the "brain") of agents. The difference lies in the architecture that surrounds them: while a traditional LLM responds to individual requests, an LLM-based agent adds components to plan, memorize, and act in the real world.

How They Work: The Architecture of Decision-Making

A modern Artificial Intelligence agent can be described as a software system that makes decisions throughout an entire process, starting from an initial objective and arriving at an operational result. To do this, it combines various capabilities. First, it has a form of "environmental perception": it accesses structured and unstructured data, such as documents, databases, web information, or via APIs, which allow it to build an updated picture of the context in which it

⁴ API (acronym for Application Programming Interface) is a set of rules and protocols that allows different software applications to communicate with each other to exchange data, functionalities, and services.

must act. On this basis, language models intervene, interpreting the request, breaking down the task, planning the next steps, and modifying the plan based on feedback that arrives progressively from the environment or the user.

A key element is memory persistence: the agent maintains a state of the project and the conversation, remembering decisions made, constraints imposed by the user, attempts already made, and partial results, so as not to start from scratch each time. To this is added integration with external tools: the agent is not limited to producing text but is capable of calling APIs, querying databases, using software and web platforms to concretely execute the actions decided during planning. When these components are well orchestrated, the agent can proceed autonomously through multiple decision-making steps in sequence, without having to ask for human confirmation at every single action, while remaining under overall supervision.

This architecture becomes particularly powerful when organized according to an orchestra model. In this case, there is a central "orchestrator" that receives input from the user and coordinates a set of specialized agents, each focused on a specific task (e.g., data analysis, code writing, document research, quality control). The orchestrator breaks down the complex request into sub-tasks, assigns them to the most suitable agents, supervises their execution, and finally recomposes the results into a single, coherent output that the user sees as the final answer. From the user's point of view, all this appears very simple: one interacts with a single "front end," formulates a request, and receives a finished result, often in the form of a document, functioning code, or operational procedure ready for use. Behind the scenes, however, the dynamic is much more articulated: dozens of agents converse with each other, exchange intermediate data, refine results multiple times, and verify them

reciprocally, realizing in an automated way what, in a traditional organization, would be a genuine distributed workflow.

Agents in the Real World: From Browser to Enterprise Workflow

Browsers⁵ represent one of the most interesting fields of experimentation for AI agents. Until recently, developers integrated AI tools into browsers with mixed results. But recently, with the launch of OpenAI's ChatGPT Agent and Perplexity's Comet, the idea of a browser enhanced by an autonomous AI assistant has come back into vogue.

The two systems work differently. Comet is an independent browser: the user navigates normally on the web and then "summons" an AI assistant that can help them write an email or complete a routine activity. OpenAI built its tool inside ChatGPT itself: the user dialogues through the web interface and assigns tasks to the bot, which uses a virtual browser to execute them. In both cases, the agents can control cursors, enter text, click on links, and fill out online forms.

If this trend should consolidate, the web could transform into a ghost town where agents run rampant: while automatic software interacts 24/7 with services, databases and platforms, human beings venture

⁵ A browser is a program that allows access to Web resources (pages, images, videos) and navigation between Internet sites. A browser, also called a "navigator," is an application installed on computers or mobile devices that receives content from network servers and presents it in a form understandable to the user. It uses protocols like HTTP to request pages and a rendering engine to transform code (e.g., HTML, CSS, JavaScript) into what is seen on the screen.

ever more rarely into the network. The implications are fascinating and unsettling at the same time.⁶

In companies, the discussion is even more pragmatic. Agents are being integrated into real workflows to automate processes that until now required human coordination: from the onboarding⁷ of a new employee (the agent prepares documents, sends emails, configures accounts) to invoice management (reads accounting documents, classifies expenses, generates reports). The legal, tax and accounting sector is already experimenting with agents capable of analyzing contracts and regulatory documents. To avoid the risks of fraud, agents are already being used to identify suspicious patterns in financial data.

The Economic Promise and the Paradox

The promise is simple: more efficiency, lower costs, fewer human errors. An agent can work 24/7, doesn't get tired, doesn't forget, and can handle dozens of tasks simultaneously. At the organizational level, this means redesigning roles: fewer "executors of repetitive tasks," more "supervisors of agents" and "strategists who define objectives."

But here an interesting paradox emerges. On one hand, professionals who use AI in their work know that this change is coming. Many, especially in the legal, tax, and accounting sectors, are already

⁶ Data from 2024-2025 confirms the overtaking: AI bots now represent 51-52.3% of total web traffic, surpassing human activity for the first time in a decade. Of this automated traffic, 35% comes from crawlers for training LLMs (OpenAI, Anthropic, Google DeepMind), while 17% consists of operational bots that automate emails, calendars, and web scraping. Meta alone generates 52% of all AI bot traffic, more than Google (23%) and OpenAI (20%) combined.

⁷ Onboarding of a new employee is the structured process a company uses to welcome, integrate, and guide them in using a product or service.

familiarizing themselves with how GenAI can transform their daily work. On the other hand, many others are just beginning to understand what an AI-based future means for their profession and how much agents represent a further step into the unknown.

The reality we face is something deeply interconnected with our world, with its dynamics and its problems. It is not a science fiction vision: it is the automation of real processes that already today have concrete economic and social consequences.

The Role of Governance and Human Control

Here governance, regulation, and human control come into play. Europe has attempted to intervene with the AI Act, defining risk levels and transparency obligations. But truly integrating ethical principles within these models is complex, because ethics is not a technical parameter: it is an intangible asset that changes with culture, context, and reference values.

To move beyond the quicksand of theoretical debate, organizations need pragmatic tools to decide where and how to employ agents. Not all tasks are equal, and not all require the same level of human control.

An effective approach is to map each process onto a **Risk-Autonomy Matrix**, crossing the criticality of the activity with the degree of freedom granted to the agent. This matrix offers a golden rule for governance: as risk increases, the agent's autonomy must decrease proportionally. The most common mistake today is misalignment: companies applying "low-band" levels of autonomy to "high-band" processes (e.g., chatbots handling sensitive complaints without

supervision) or vice versa—those wasting human resources to control banal tasks that AI could handle alone.

Adopting this scheme means accepting that efficiency is not the only value. In the high band, the goal is not "to do it faster" but "to decide better." Here the agent does not replace the expert but enhances their analytical capacity, leaving the prerogative of judgment intact.

Beyond "Compliance": Three Principles of Responsible Design

Adopting a risk matrix is the first step, but it is not enough. To truly govern an agentic architecture, organizations must integrate three non-negotiable principles directly into process design, before the software is even turned on.

First, Auditability by Design. In a system where ten agents collaborate to produce a result, "who did what" quickly becomes a mystery. Traceability cannot be an afterthought. Every critical action of an agent—every query to a database, every generated text, every logical decision—must be recorded in an immutable, human-readable log. If a medical agent suggests the wrong therapy, we must be able to reconstruct what information (or hallucination) triggered that error. Without this accessible black box, there is no governance, only blind trust.

Second, Meaningful Human Control. The law and common sense require human supervision, but there is a risk it becomes a farce: the operator who clicks "Approve" a hundred times an hour without really reading. The principle of "meaningful supervision" requires designing workflows so that the human has the time, information, and real competence to be able to challenge the machine. If the pace imposed by the agent is too fast to allow reasoned doubt, then there is no

human control—only a human being used as a rubber stamp for legal responsibility.

Third, The Right to Override. No process managed by agents should be a dead end. There must always be a "red button," a clear and accessible procedure that allows the human operator to interrupt the automation, revoke a decision made by the agent, and take manual control of the situation. This right to override is not an emergency measure; it is a standard safety feature. A system that cannot be stopped or corrected mid-course by an authorized human is not an autonomous system: it is a system out of control.

Moreover, the competitive reality is ruthless. Regulation always arrives after the technology. AI has already become a strategic driver for companies and a factor that strongly influences financial markets. There is strong global competitive pressure, even though clear rules, subscribed to by all interested parties starting with the United States, have not yet been established. The result is that agents enter companies even though we know that their use still has enormous limits and entails possible risks.

For this reason, the fundamental element of any AI agent implementation must be this: **human supervision remains indispensable.** Not in the sense of "checking every single click"—that would be impossible and would neutralize the advantages of agents—but in the sense of:

- Having competent people who understand the domain (legal, tax, medical) and who know how to recognize when an agent's output is nonsensical or dangerous.

- Clearly defining the boundaries of the agent's autonomy: in which situations it can decide on its own and in which it must ask for human confirmation.
- Maintaining a trace (audit trail) of what the agent has done, so that in case of problems there is accountability and traceability.
- Being aware that agents are currently "premature" if imagined in complete autonomy. They lack self-awareness, a sense of limits, and a deep understanding of consequences.

The Puzzle of Responsibility: Whose Fault Is It When the Orchestra Plays Out of Tune?

Let's imagine a now plausible scenario: a company uses an "Orchestrator Agent" to manage the supply chain. This agent autonomously tasks a second agent (specialized in logistics) to book an urgent shipment and a third (specialized in payments) to settle the invoice. If the logistics agent makes a routing error causing goods to perish, or the payments agent sends funds to a fraudulent IBAN intercepted online, who pays?

In the old world of deterministic software, the answer was often linked to the "bug": if the code was poorly written, the fault lay with the vendor. In the world of AI agents, where behavior is probabilistic and non-deterministic, we enter a legal "no man's land."

The Liability Fragmentation. The main problem with agentic architectures is non-linearity. If Agent A (orchestrator) gives an ambiguous command to Agent B (executor) and B interprets it in the worst possible way causing damage, is it A's ambiguity or B's lack of prudence that is at fault? Emerging legal doctrines, in line with the

European AI Act, tend to shift the focus from the software producer (Provider) to the professional user (Deployer).

The Principle of "Culpa in eligendo" and "Culpa in vigilando." Initial legal analyses suggest that companies will be called to answer according to two ancient principles of law, adapted to the digital age:

- **Culpa in eligendo** (Fault in choosing): if a company decides to entrust a critical process (such as payments) to a system of autonomous agents, it assumes the business risk intrinsic to that choice. It cannot defend itself by saying "it was the software." Having chosen a probabilistic tool for a task that does not tolerate errors is, in itself, managerial negligence.
- **Culpa in vigilando** (Fault in supervision): here governance comes back into play. If the company has not set up adequate monitoring systems (Human-in-the-loop),⁸ it is responsible for not having stopped the agent in time. Responsibility does not lie with the agent that makes a mistake, but with the human who was not watching it.

The Provider vs. the User: The Case of "SaaS That Acts."⁹ A critical point concerns the Terms of Service (ToS). The major LLM providers (such as OpenAI, Google, Anthropic) specify clearly in their

⁸ Human-in-the-loop: It is an approach in which human beings actively intervene in artificial intelligence and machine learning processes to supervise, correct, or improve automated decisions. This method integrates human feedback into the training and operational cycle of AI models, ensuring greater accuracy, ethics, and reliability.

⁹ SaaS: Software as a Service. It is a software distribution model in which applications are provided by a provider via the Internet, accessible via browser without local installation. The provider manages infrastructure, updates and maintenance, while users typically pay on a subscription basis.

contracts that model output is provided "as is"—meaning that the results generated by the models are delivered in the raw state in which they are produced, without guarantees, modifications, verification, or further support from the provider, and therefore verification rests with the user. However, the new EU Directive on liability for defective products is extending the concept of "product" to software and AI as well. If an agent is sold as "autonomous" and is capable of performing a specific task and then fails spectacularly not due to user error but due to a training defect, the producer may no longer be able to hide behind liability waiver clauses.

Toward "Strict Liability" for AI? For high-risk sectors (healthcare, critical infrastructure), the doctrinal trend is moving toward strict liability, similar to that applied to hazardous activities (Article 2050 of the Italian Civil Code). Those who derive economic profit from the use of autonomous agents must answer for the damage these cause, regardless of specific fault or intent, simply as the cost of the risk introduced into the system.

In summary: the agent's autonomy is not a mitigating factor for the human, but an aggravating one. The more autonomy we give to the system, the more stringent our duty of guarantee toward third parties becomes.

Where Are Agents Being Used?

If the impact on sectors such as legal, tax, and accounting now appears to be a defined trajectory, it is by looking at three other crucial areas—healthcare, public administration, and marketing—that we grasp the real transformative scope of AI agents. The promise is always the same: more efficiency, process automation, and freeing up

qualified human resources. But it is precisely in the specificities of these worlds that the deepest implications for work and responsibility emerge.

Healthcare: The Agent as Clinical "Co-Pilot"

In the healthcare sector, agents are positioning themselves as powerful "co-pilots" for doctors and nursing staff. Their ideal field of action is orchestrating complex tasks with a strong procedural component:

- **Management of waiting lists:** an agent can cross-reference in real time the availability of different facilities, the clinical priority of patients (defined by a doctor), and the required diagnostic pathways, dynamically optimizing schedules.
- **Monitoring of chronic patients:** an agent connected to wearable devices can analyze the vital data of a diabetic or hypertensive patient, alerting the doctor only to anomalies that exceed a predefined risk threshold and require human intervention.

The advantage is evident: freeing up valuable time, reducing bureaucratic burden, and allowing healthcare personnel to focus on clinical care and the patient relationship. But it is precisely this advantage that defines, by contrast, the boundaries of non-delegable skills.

An agent can identify a statistical risk pattern, but it cannot take responsibility for:

- Formulating a complex diagnosis, where clinical intuition, experience, and understanding of what the patient leaves unsaid are fundamental—for example, distinguishing between

apparently similar symptoms that could be attributable to very different pathologies.

- Communicating a negative prognosis: a task that requires empathy, sensitivity, and an ability to adapt to the person's emotional context, which no algorithm can replicate.
- Making an ethical decision: how to allocate a scarce resource (an intensive care bed, an experimental drug) when guidelines are insufficient and human judgment on a "case-by-case" basis comes into play.

The most insidious risk here is not the algorithm's technical error, but its **blindness to context**. An agent trained to maximize efficiency might, in theory, suggest deprioritizing an elderly patient with comorbidities in favor of a younger one with greater chances of rapid recovery. It is the doctor, with their background of ethics and professional responsibility, who must impose a veto, affirming a principle of care that transcends mere statistical optimization. The human role thus shifts from data compilation to critical interpretation of that data, becoming the ultimate guarantor of the care process.

Public Administration: From Executor to Guarantor of Equity

In Public Administration, agents promise to dismantle long queues and bureaucracy for its own sake. A "virtual officer" can already today:

- **Guide the user in completing an application:** for example, for a building bonus, verifying in real time the correctness of cadastral data via APIs and flagging missing documents before submission.
- **Orchestrate an inter-office process:** such as the issuance of a commercial license, managing in parallel the checks by the

technical office, the health authority, and the chamber of commerce.

The role of the human official evolves: no longer a mere executor of procedures, but a validator of the process and a guarantor of equity. The risk, in fact, is replacing the "rubber wall" of the physical counter with the opacity of an "algorithmic filter" that applies rules rigidly but blindly.

Non-delegable human skills thus become:

- **Managing exceptions:** recognizing when an individual case, while not perfectly fitting the parameters, has the right to be treated. Think of an elderly citizen without SPID who needs to access an essential service.
- **Protecting against bias:** ensuring that the agent, perhaps trained on distorted historical data, does not discriminate against certain categories of citizens (by gender, ethnicity, income) through automated decisions that seem neutral but are not.
- **Interpreting rules in context:** applying the law with intelligence, not mechanically, because the same rule can have different effects depending on the social and economic context in which it is applied.

Marketing: The Agent as Algorithmic Creative Director

In marketing, the transformation is already underway. Agents are no longer limited to generating content, but orchestrate entire campaigns. An agent can analyze market trends, define the target of a new product line, generate creative content, plan advertising spending across multiple channels, and adjust the strategy in real time

based on conversion rates. Human creativity seems reduced to a supervisory role. But it is precisely here that the irreplaceable role of the strategist emerges. The agent is an optimizer par excellence, that is, given a metric (e.g. maximize clicks), it will always find the shortest path to achieve it. This can lead to a flattening of communication, where all brands converge toward the same "solutions that work," losing identity. The human strategist becomes the custodian of the long-term vision and the brand's values. Their task is not to choose the image with the highest CTR,¹⁰ but to decide whether that image is consistent with the brand's history and positioning. It is the person who vetoes a potentially viral campaign, but ethically questionable, sacrificing short-term gain to protect reputation. The agent wins the tactical battle; the human strategist wins the positioning war, ensuring that the brand does not become a slave to the autocracy of metrics.

What Are the Unresolved Problems?

But before concluding that AI agents will solve everything, caution is needed, because agents, however sophisticated, inherit all the problems of generative AI that we already know, and in some cases amplify them.

Hallucinations remain the first problem. Modern language models continue to produce incorrect answers, expressed in an extremely convincing manner. This is not a marginal defect: recent tests have shown that models celebrated for their "advanced reasoning" have hallucination rates as much as double those of

¹⁰ CTR stands for Click-Through Rate, i.e. the click rate calculated as (clicks / impressions) x 100, which measures the attractiveness of a visual element such as an image in ads or emails. A high CTR indicates that the image captures the attention of the target audience better.

previous models. Texts that are well-written but semantically incorrect. And if this was already risky for a chatbot writing an email, it becomes potentially dangerous when the agent acts on the basis of hallucinated information: it queries a database with invented data, sends requests based on "facts" that don't exist, makes decisions on false premises.

The black box problem worsens. We already know little about how generative models reach a conclusion. With agents, this opacity multiplies. If an agent is composed of ten specialized agents conversing with each other, coordinated by an orchestrator, who can really say what happened at the moment a wrong decision was made? The user sees input and output, but the internal process remains hidden. So the risk of "explainability washing" (pretending to be transparent by showing only what suits them) is far from remote.

Agents have no intrinsic ethics. This is perhaps the most important point. AI systems have no consciousness, self-awareness, sense of limits, or understanding of the ethical consequences of their actions. If ethics has not been designed and implemented by humans from the very beginning (what experts call "ethics by design"), the system cannot follow it. When an agent completely ignores the ethical impacts of a decision, it is not because it is "bad": it is because it was built without those constraints.

Who Is Needed to Govern These Problems?

If AI agents promise to transform processes and organizations, the almost inevitable consequence is the emergence of new professional figures. It is not enough to "use" agents: they must be designed, coordinated, evaluated, and embedded in ethical and legal

frameworks. Instead of limiting ourselves to the generic formula of "agent supervisors," we can begin to give a name and a more precise perimeter to some roles that, in the coming years, could become central.

Below are three emerging profiles that clearly mark the transition from purely executorial work to work of governing automation.

The "AI Agent Orchestrator"

If the "orchestra" architecture of agents truly spreads, someone will be needed to be the conductor of that orchestra. The AI Agent Orchestrator is the figure responsible for designing, coordinating, and monitoring the work of multiple specialized agents within an organization. Their task is not to write the code for each agent, but to translate the company's (or public entity's) strategic objectives into agentizable tasks, decide which agents to activate, with which permissions, based on which data; define the engagement rules between different agents (who does what, in what order, with what alternatives); oversee the overall performance of the system, identifying bottlenecks, redundancies, and conflicts between outputs.

This is a hybrid figure who must understand technology enough not to be hostage to technicians, but also know the application domain in depth (finance, healthcare, human resources, public administration). In a sense, they take the place of many "middle managers" who today coordinate human teams, but with a fundamental difference: their object of work is not people, but ecologies of agents. Yet the repercussions are deeply human, because how they orchestrate automation affects workloads, timelines, and the quality of services offered to people.

The "Automated Process Validator"

As agents enter decision-making processes—whether healthcare, administrative, or financial—a figure becomes indispensable whose explicit task is to put automation to the test. The Automated Process Validator is one who does not merely check whether an agent's output "works," but asks whether it is: correct on the merits; robust with respect to scenarios different from those of training; compatible with the norms, regulations, and rights of the people involved.

In practice, this role can take different forms:

- In healthcare, it can be the clinician who systematically evaluates how the "triage" agent's output compares to medical judgment, identifying patterns of recurring errors.
- In public administration, it can be the official who analyzes a sample of cases managed by the agent, verifying whether the algorithm treated similar situations fairly or whether systematic biases toward certain user groups emerge.
- In banking or insurance, it can be the analyst who checks whether an automated scoring system systematically excludes certain categories of applicants.

This is a profession that requires methodological skills (statistics, evaluation, audit), domain knowledge, and above all, independence of judgment: the validator must be able to say "no" to automation when it produces unacceptable effects, even if "on the numbers" it seems efficient. In a world where the European AI Act demands transparency, traceability, and risk management, the Automated Process Validator becomes a natural piece of new value chains.

The "Algorithm Ethicist"

Finally, there is a figure that already exists in part, but whose role the advent of agents makes even more indispensable: the Algorithm Ethicist. They do not merely draft ethical codes or guidelines to follow, but translate abstract ethical principles into concrete design and usage constraints. Their work operates on at least three levels:

- **Upstream**, in the design phase: defining what data can be used, which optimization objectives are acceptable and which are not (maximize profit? reduce average waiting time? guarantee priority to vulnerable subjects?).
- **In the middle**, during development: collaborating with data scientists, lawyers, and decision-makers to identify points where automation can conflict with fundamental rights (privacy, non-discrimination, access to essential services) and proposing safeguards.
- **Downstream**, in operational evaluation: monitoring the real impact of agents on people and social groups, collecting reports, problematic cases, unexpected effects, and updating policies accordingly.

In this perspective, the Algorithm Ethicist figure is anything but ornamental. It is someone who brings philosophical, legal, sociological, and historical skills into the organization and uses them to mount reasoned resistance to a purely efficiency-driven vision of AI. They are effectively a mediator between different worlds—between those who build systems, those who govern them, and those who are subject to them. It is no coincidence that this profile, in many contexts, arises from the meeting of technical-scientific competencies and a strong humanities background. In a context where AI agents promise to

automate more and more decisions, the question is not just "what can they do?" but "how far do we want to let them go?" The Algorithm Ethicist helps provide a non-improvised answer to this question, bringing into the innovation process the element that technology, alone, can never provide: a criterion that makes sense.

Rethinking Skills: From the Logic of Doing to the Logic of Being

If the new organizational roles—the AI Agent Orchestrator, the Automated Process Validator, the Algorithm Ethicist—describe what people will do in the next decade, it is equally important to understand how they will have to transform to inhabit these professional spaces. It is not enough to simply add "AI Literacy"¹¹ to the skillset: a more radical rethinking of what it means to acquire skills in the age of agents is needed—a rethinking that touches the very foundations of professional and human identity.

AI Agents represent an accelerated industrial revolution, with much faster implementation times than previous technological waves. The innovative rupture is inevitable, but those who adapt proactively, rather than passively suffering the change, can transform this challenge into a genuine opportunity for professional and economic growth. The question is not whether we will change, but how we will change and what meaning we will give to this change.

ILO (International Labour Organization) research highlights an alarming fact: significant gender differences in exposure to Generative

¹¹ AI Literacy: the ability to understand, use, and critically evaluate artificial intelligence technologies.

AI. In high-income countries, 8.5% of female employment is in jobs with high automation potential, compared to 3.9% of male employment. This pattern is not accidental, but reflects a greater concentration of women in administrative and support roles, particularly exposed to technology. However, the "augmentation" effect—the ability of agents to enhance human skills rather than replace them—shows more balanced trends, with significant opportunities for both genders in knowledge-intensive sectors. The real challenge, then, is not technical but political: ensuring equal access to training and professional retraining, so that the agentification of work does not become a new mechanism of inequality, but an opportunity for rebalancing and inclusion.

Immediate Skills and Long-Term Strategies: What to Do Now?

For professionals, the imperative is twofold and requires simultaneous action on two time levels.

In the short term (the next 12-18 months), it is essential to acquire fundamental skills:

- **Basic AI Literacy:** not about becoming technicians or machine learning experts, but understanding the potential and limits of Generative AI as conscious users, capable of recognizing both promises and risks.
- **Direct experimentation** with tools like ChatGPT, Copilot, and Claude for daily activities, transforming fear into operational familiarity.
- **Prompt Engineering:** the art of communicating effectively with AI systems, discovering that linguistic clarity and

communicative precision become valuable professional resources, not only in interaction with machines but also in human collaboration.

- **Investing in relational skills**—emotional intelligence, leadership, negotiation ability—becomes strategic precisely because they represent the unbreachable boundary of automation.

In the long term, three deeper positioning strategies emerge:

1. **Specializing in areas where AI is complementary rather than substitutive.** Not about avoiding sectors affected by automation—that would be illusory—but consciously identifying roles where technology amplifies human value rather than eroding it.
2. **Building hybrid skills:** combining technical expertise (such as reading AI output, evaluating its reliability, and identifying biases) with distinctive human capabilities like empathy, ethical judgment, and long-term strategic vision.
3. **Maintaining a continuous learning attitude and adaptability,** aware that the speed of technological transformation means no educational background remains valid for more than 3-5 years without significant updating.

For organizations, the challenge is structurally similar, but on a collective scale it requires cultural as well as technical transformation:

- **Gradual and strategic adoption** of agents begins with widespread training, not limited to frontline figures but extended to all employees, to create a shared cultural reference point around AI topics.

- **Circumscribed experimentation:** starting with pilot projects on non-critical processes, learning from failures, identifying success patterns, and then progressively extending the field of action.
- **Ethical management:** establishing clear guidelines for responsible AI use before the system is turned on, not as an afterthought when damage has already been done.
- **Proactive retraining:** not simply teaching employees how to use agents, but how to transform themselves to work with them, remaining central in decisions and quality control.

A change management approach that truly works involves workers in AI implementation decisions, promotes a culture of experimentation and innovation without punishing calculated failure, and above all focuses on enhancing human capabilities rather than mere automation and cost reduction. In this sense, the smartest organizations will see agents not as substitutes, but as allies that free up time and mental resources for higher value-added activities.

Prospects for Italy: A Window of Opportunity That Is Closing

Italy finds itself in a paradoxical and at the same time extraordinary position. On one hand, its industrial and creative DNA—Made in Italy, quality manufacturing supply chains, creative and artisanal districts—could be significantly enhanced by AI Agents, transforming the scale and speed constraints that have historically limited the global competitiveness of Italy's SME fabric. According to the Implement Consulting Group study for Google, widespread adoption of Generative AI could contribute to an 8% increase in Italian GDP over ten years, equivalent to 150-170 billion euros in annual added value. This is a

figure that reflects not mere economic growth, but a structural transformation of the country's productive capacity.

On the other hand, however, obstructive challenges remain that risk preventing Italy from seizing this opportunity:

- **The infrastructure and skills gap:** despite recent progress, the lack of quality digital infrastructure and the widespread shortage of basic AI skills remain significant bottlenecks, particularly in the regions of the Mezzogiorno (South).
- **Adoption among SMEs is still marginal:** only 7% of small and 15% of medium-sized enterprises have active AI projects, when precisely these categories could benefit most in terms of operational efficiency and access to global markets.
- **Educational and vocational training institutions** are not yet sufficiently aligned on the skills required by the agentic era, producing a mismatch between what schools teach and what the market demands.

The window of opportunity remains open, but it is visibly narrowing. Countries that invest massively now in training, digital infrastructure, and organized experimentation will have a decisive competitive advantage in 5-7 years. For Italy, this means that the investment choices we make in 2025-2026 will determine the country's competitive position in 2032-2033. The urgency, therefore, is not rhetorical: it is structural.

Beyond Skills: Being and Doing in the Age of AI

However, beneath these pragmatic and macroeconomic considerations, a deeper question emerges—one that transcends the mere accumulation of "technical" skills and touches the very meaning of human work. We are not simply adding a new subject to the school curriculum or the corporate upskilling program. We are facing an ontological transformation—that is, a transformation of our very nature and identity—of what it means to work, create value, contribute to a community, and be present in the world.

The distinction between **skills of being** and **skills of doing** becomes strategic precisely in this context. Technological transformation is not neutral from an existential point of view. As phenomenology teaches, every technology modifies our way of being in the world: it is not a tool we use while remaining unchanged, but a force that transforms us as we use it. Generative AI is not simply intelligent software, but a transformation of existence that requires a radical redefinition of skills in ontological terms.

From a practical point of view, a good educational guidance system should distinguish which skills to develop as a priority based on their ontological nature, knowing that different skills require different development strategies. Unfortunately, current guidance systems leave much to be desired, often limiting themselves to mechanical skill mapping or simple occupational projections. Yet, having correct guidance is fundamental to arriving at a professional strategy that allows professionals to consciously invest in "uniquely human" skills—irreducible to computational logic and algorithmic calculation—that are enriched and transformed through hybridization with AI.

From an anthropological point of view, finally, the being/doing distinction in the age of AI touches a fundamental existential

question: **what does it mean to be human when machines become "intelligent"?** The skills of being remind us that the human is not defined by what they do—by the tasks they perform, the products they generate, the results they achieve—but by how they exist, by their inalienable capacity to question the meaning of things, to open themselves to the otherness of others, to create meaning through narrative and dialogue. These are the skills that remain intact, indeed, that become even more precious, precisely because they are irreducible to automation, because they are not codifiable in an algorithm.

A New Pedagogical Paradigm: Education as Ontological Vocation

The being/doing distinction requires a radical rethinking of educational pathways—a rethinking that touches schools, universities, and organizations.

1. **Developing and protecting reflective capacity, emotional intelligence, and creativity** as constitutive dimensions of the person, not as marginal or supplementary "soft skills," but as the very core of authentic human formation.
2. **Mastering AI tools** not as substitutes for human capabilities, but as intelligent extensions of those capabilities that amplify their reach.
3. **Integrating being and doing** in a creative and responsible manner, knowing that the human added value in the next decade will reside in the ability to govern automation, give meaning to data, and take care of the human and social impact of technological choices.

As the Brazilian pedagogue Paulo Freire suggests, we are "beings who are becoming"—never completed, always in a constant process of formation and transformation. In the age of AI, continuous learning is not just a professional necessity, a skill to be periodically updated, but an ontological vocation: the way we realize our humanity in creative dialogue with AI, the way we remain alive, aware, and capable of growing.

Toward a Technological Humanism

The ontological analysis of skills in the age of Generative AI reveals that we are not witnessing a replacement of the human—the apocalyptic narrative of the robot taking your job—but a profound redefinition of it.

- **Skills of being** become even more precious precisely because they are irreducible to computational logic, because they reside in the realm of meaning, relationship, and responsibility.
- **Skills of doing** evolve in a collaborative and synergistic direction, where work is no longer a solitary exercise of an ability, but an intelligent orchestration between human intuition and the machine's computational capacity.
- **Hybrid skills** open new horizons of existential possibilities, allowing professionals to operate in spaces that were previously inaccessible due to lack of time, energy, or resources.

This perspective invites us to think of the future of work not as a human-machine competition—a battle where one wins and the other loses—but as the emergence of a new form of technological humanism, where technology is at the service of realizing the most authentic and profound human potential. In this evolutionary scenario,

investing in skills of being—in the ability to think critically, relate empathetically, create meaning—is not nostalgic conservatism or a romantic defense of an irretrievable past, but the most radical and necessary way to prepare a future in which technology amplifies what we hold most precious as a species: our capacity to be human in the world, to build communities, to give meaning to our living together.

The Return of the Philosophers: Why Humanities Skills Are the New Engineering?

For years we have heard a mantra repeated: "we need more STEM" (Science, Technology, Engineering, Mathematics). In a world where the digital infrastructure had to be built, it was true. But in the age of AI agents, the paradigm is turning around in an unexpected way. If tools like Manus¹² write code, test software, and manage databases better and faster than a junior developer, human added value no longer resides in the ability to speak the language of the machine (code), but in the ability to teach the machine to understand the human world. Humanities skills—philosophy, history, sociology, linguistics—cease to be a cultural ornament and become the "operating system" necessary to govern agents.

From Prompt Engineering to Context Engineering (Hermeneutics)

Interacting with an agentic architecture does not mean giving a command ("do this"), but defining a context of action. It requires the ability to anticipate ambiguities, clarify intentions, and interpret partial

¹² Manus is an AI platform capable of autonomously executing complex tasks, including writing code, testing software, and managing databases.

results. This is, in technical terms, hermeneutics: the art of interpretation. An engineer can optimize an agent's speed, but it takes a mind trained in linguistic and logical complexity to notice that the agent has "misunderstood" a cultural nuance in a marketing campaign or a diagnosis. Knowing how to ask the right question becomes more difficult—and more valuable—than knowing how to find the answer.

History Against Data Bias

AI agents are trained on the past (historical data). By definition, they tend to replicate the past. If historical data on hiring in a company is tainted by gender biases, the HR agent will tend to replicate them, mistaking them for "efficiency." This is where historical consciousness comes into play. Only those who have studied social evolutions can distinguish between a "statistical correlation" and a "historical cause." The historian knows that the past is not an infallible guide to the future, but an archive of choices to be contextualized. In a company, this means having people capable of telling the agent: "Ignore this pattern in the data because it is the result of a social context we no longer accept."

Applied Ethics: Beyond the "Trolley Problem"

Until yesterday, AI ethics was a theoretical exercise (the famous "trolley problem"). With autonomous agents, it becomes daily practice. A pricing agent must decide how aggressively to raise prices during a demand spike. Maximizing profit is the default instruction. But when does this become profiteering? The ability to navigate these gray areas is not learned on GitHub.¹³ It requires training in moral philosophy, capable of balancing utilitarianism (the greatest good for the greatest

¹³ GitHub is a web platform for version control and collaborative software development using Git, allowing developers to share, manage, and collaborate on code.

number) and deontology (there are rules that are never to be violated). Companies will need Chief Ethics Officers—not to look good, but to avoid reputational disasters caused by overly zealous agents.

In Summary

Paradoxically, the more sophisticated and "human" AI becomes in its interactions, the more we need humanists. Not to reject technology, but to provide it with that overview, that critical thinking, and that ethical depth that no algorithm, however advanced, can calculate.

The Crossroads of 2030: Three Scenarios for an "Agentified" Future

If agentification is the technological trend, its social impact is not deterministic. It depends on how we intersect the variables of regulation, access to technology, and corporate culture. Applying strategic foresight techniques, we can outline three possible horizons for the labor market at the end of the decade.

Scenario 1: The Symbiotic Augmentation (The Best-Case Scenario)

The Setting: Regulation keeps pace with technology. The EU AI Act is not a dead letter: it has been implemented with effective control mechanisms. Companies have invested in training and have not simply replaced employees but have redesigned roles, creating new professional figures (Orchestrators, Validators, Ethicists). The technology is advanced but transparent.

The World of Work: Humans and agents collaborate fluidly. Professionals delegate repetitive and low-value-added tasks to agents (data processing, scheduling, first-level reporting), freeing up time for strategic activities, creative thinking, and relationship management. The typical worker of 2030 spends 60% of their time on tasks that require human intelligence (interpretation of complex data, decision-making in conditions of uncertainty, management of human relationships) and 40% supervising and guiding their team of agents. The digital divide is reduced, and many routine jobs have been transformed rather than eliminated.

The Critical Issue: Keeping up requires continuous retraining. The professions most at risk are not the "manual" ones, but the "cognitive-routine" ones: clerks, administrative employees, first-level analysts. The winners are those who manage to hybridize their technical knowledge with advanced relational and critical thinking skills.

Scenario 2: The Opaque Network (The Fragile Automation)

The Setting: Regulation lags behind. Companies, driven by competitive pressure, deploy agents massively without adequate monitoring systems. The technologies are complex, poorly documented, and interconnected with each other through APIs that no one fully controls.

The World of Work: The system becomes too complex to be fully understood. Agents interact with each other, creating feedback loops that produce unexpected results. Errors are frequent, but no one can trace their origin. An agent might recommend a financial investment based on data that was generated by another agent that had hallucinated the source. The human supervisors, overwhelmed by the number of alerts and the speed of transactions, limit themselves to rubber-stamping outputs without truly understanding them. This is a

world of fast, incomprehensible transactions for humans. In the labor market, this creates a crisis of trust. Managers no longer trust reports because they don't know if they were generated by a cascade hallucination of ten different agents. "Digital Remediation" jobs emerge: teams of humans paid to deactivate runaway automations, clean datasets corrupted by synthetic content, and restore manual processes when digital systems collapse under the weight of their own complexity. It is a world that is efficient in appearance, but fragile and neurotic in substance.

Scenario 3: Digital Feudalism (The Concentration of Power)

The Setting: This is the critical scenario, where technology is mature but access is unequal. A few tech giants own the most powerful "Orchestrator Agents" (the "Super-Manager AIs"). Small businesses and independent professionals cannot afford these premium tools and remain cut off from the market, unable to compete in speed and costs.

The World of Work: Work polarizes: on one hand, an elite of "AI Feudal Lords" who govern automated empires with very few employees; on the other, a mass of "digital gig-workers" who perform tasks that cost agents too much in terms of computational energy or that require legal responsibility that machines cannot assume. In this scenario, the agent is not a collaborator but an algorithmic overseer that dictates pace and working conditions to increasingly interchangeable humans. Wages for non-specialized digital work plummet, while the economic returns from AI flow almost entirely to those who control the intellectual property of the most powerful agents. The middle class of knowledge workers—lawyers, accountants, consultants—is decimated, replaced by a few superstar professionals who use AI to do the work of hundreds.

The Weak Signal

Which of these three scenarios will prevail? It depends on the governance choices we make today. If we treat the agent as a magical "black box" (Scenario 2) or as exclusive property (Scenario 3), the future will be dystopian. If we design it as a transparent infrastructure at the service of human integration (Scenario 1), agentification could be the greatest liberation from repetitive work in recent history.

Conclusions

When we began this journey through the agentification of work, we started from the almost feverish enthusiasm of 2025—the year when a new tool promising revolutions in productivity, miraculous efficiency, and radical transformations was born every week. We expected to talk only about technology, algorithms, software architectures. Instead, we unexpectedly found ourselves confronting questions of moral philosophy, applied ethics, and human responsibility. This is not a coincidence, nor a deviation from the path. It is exactly what happens when a technology ceases to be a passive tool in our hands and becomes an autonomous actor in our economic and social ecosystem.

AI agents represent something profoundly different from everything we have known so far. They are not simply an incremental improvement on the software we used yesterday. They are a new species of digital entities that inhabit our world: they have the ability to perceive the surrounding environment, make complex decisions by breaking down objectives into sub-objectives, act autonomously without asking for confirmation at every step, and—perhaps more disturbingly—make mistakes in unexpected ways and amplify our unconscious biases on an industrial scale, transforming individual biases into systemic discrimination.

If we continue to treat these systems as simple utensils—as if they were more sophisticated digital hammers—we risk sliding inexorably toward the darker scenarios we have outlined. Think of the Opaque Network, where millions of agents interact in incomprehensible loops, making decisions that no human being can reconstruct or challenge, creating a kind of algorithmic bureaucracy infinitely more frustrating than the human one. Or Digital Feudalism, where a few tech giants control the most powerful orchestrator agents, while the majority of

workers and small businesses are relegated to the margins, reduced to performing residual tasks that are too costly for algorithms or that require a physical body and legal responsibility that machines cannot assume.

In both these dystopian scenarios, efficiency—that initial promise, that mantra repeated obsessively—actually becomes a subtle but pervasive form of tyranny. Automation does not free up time for creative and meaningful activities, but accelerates the pace until it becomes unsustainable. Responsibility is not clarified, but dissolves and fragments into a fog of code, incomprehensible logs, and contractual clauses written to exonerate providers from all obligations.

The Governance Challenge: Slowing Down to Accelerate Better

The real challenge facing those who lead organizations, those who make political decisions, those who design tomorrow's systems is not technological. It is not about learning to program agents, write better code, or train more powerful models. The challenge is learning to **govern** these systems—to establish clear boundaries, define red zones where automation cannot enter, build effective supervision and control mechanisms that are not mere bureaucratic formalities.

This requires, first of all, the courage to **slow down automation when the risk is too high**. Think of the healthcare sector: when an AI agent suggests a therapy, we cannot afford to trust blindly just because "the algorithm says so." We need a doctor who truly understands the patient's clinical context, who knows how to recognize when an apparently reasonable output is actually a statistical correlation devoid of causal meaning, who has the authority and competence to say "no, in this case the agent is wrong, and here is why..."

It also requires the clarity to **invest in those professional figures we discussed**—the Orchestrators, Validators, Algorithm Ethicists—who do not build the agents but govern them, test them, challenge them when necessary. These are not ornamental positions, created to give an impression of responsibility without substance. They are essential strategic roles that make the difference between an organization that uses AI to enhance human capabilities and an organization that is overwhelmed by automation without control.

This requires, above all, the intelligence to understand that **ethics is not a brake on innovation** or a bureaucratic obstacle to be overcome as quickly as possible. Ethics is the only safe navigation system when we venture into unexplored territory. Without clear ethical principles, codified from the design phase and then continuously verified during execution, we are simply navigating by sight on the open sea, hoping not to crash against rocks we cannot see.

Designing the Future, Not Suffering It

The final invitation—and this is perhaps the most important message of this entire essay—is **not to passively suffer agentification, but to actively design it**. The future of work will not be determined by a biological competition between humans and intelligent machines, as if we were in some kind of technological natural selection where the most efficient wins. It will instead be the result of a competition between different visions of automation: on one side, those who use machines to enhance the human, to free up time and mental energy for what truly matters—creativity, relationship, care, critical thinking; on the other, those who use them to replace the human, to reduce short-term costs and maximize profits without considering social consequences.

In this crucial game, technical skills are fundamental, because they build the engine and make automation possible. But it is the humanities skills—the ability to contextualize historically, reason philosophically, understand social dynamics, apply legal and ethical principles to concrete situations—that firmly hold the steering wheel in our hands. Without that direction, without that moral compass, the most powerful engine will simply take us in the wrong direction, faster.

Technological Humanism: A Possible Synthesis

What emerges from this analysis is neither apocalyptic pessimism nor naive optimism. It is instead the concrete possibility of a **technological humanism**—a future in which technology is placed at the service of realizing the most authentic and profound human potential, not in competition with it.

In this scenario, investing in skills of being, in the ability to think critically, relate empathetically, create meaning and give sense to experiences, and assume moral responsibility, is not an act of nostalgic conservatism nor a romantic attempt to defend an irretrievable past. It is instead the most radical and necessary way to build a future in which technology amplifies what we hold most precious as a species: our capacity to be fully human in the world, to build supportive communities, to give meaning and direction to our living together.

AI agents are offering us an extraordinary opportunity: to free ourselves from repetitive tasks, from bureaucracy for its own sake, from the cognitive fatigue of managing fragmented information. But for this promise to truly be realized, we must maintain control of the process. The steering wheel—the ability to decide where we want to go, what values we want to defend, what society we want to build—must remain firmly in our hands. In the hands of competent, responsible, and ethically aware human beings.

Only then can agentification be what it promises to be: not yet another wave of automation that replaces human work with more efficient machines, but the greatest collective liberation from repetitive work in recent history—an opportunity to rewrite the social contract between technology, work, and human dignity.

BIBLIOGRAPHY

From Assistants to Builders: What Changes

Alessandro Longo, "AI via l'economia degli AI Agent: l'automazione vale già miliardi" in "[unipd.it](https://thesis.unipd.it/retrieve/a1fc566d-b322-4ca0-8637-df66e7153978/Tesi%202072876%20Cavazzana%20Elena%202025.pdf)".
<https://thesis.unipd.it/retrieve/a1fc566d-b322-4ca0-8637-df66e7153978/Tesi%202072876%20Cavazzana%20Elena%202025.pdf>

McKinsey & Company. (2024). The state of AI in early 2024: Gen AI adoption spikes and starts to generate value. McKinsey Global Survey.
<https://www.mckinsey.com/capabilities/quantumblack/ourinsights/the-state-of-ai>

McKinsey Global Institute, "The Economic Potential of Generative AI: The Next Productivity Frontier" del 2023.
<https://www.mckinsey.com/capabilities/tech-and-ai/our-insights/the-economic-potential-of-generative-ai-the-next-productivity-frontier>

David G. Widder and Mar Hicks, "Watching the Generative AI Hype Bubble Deflate" in "[arxiv.org](https://arxiv.org/abs/2408.08778)" del 18 agosto 2024.
<https://arxiv.org/abs/2408.08778>

"L'AI Agentica vs l'AI Generativa: le differenze fondamentali" in "[astera.com](https://www.astera.com/it/type/blog/agentica-ai-vs-generative-ai/)" del 5 giugno 2025.
<https://www.astera.com/it/type/blog/agentica-ai-vs-generative-ai/>

Z. Xi, W. Chen, X. Guo, W. He, Y. Ding, B. Hong, M. Zhang, J. Wang, S. Jin, E. Zhou, "The rise and potential of large language model based agents: A survey," in "[arxiv.com](https://arxiv.org/abs/2309.07864)" del 14 settembre 2023.
<https://arxiv.org/abs/2309.07864>

<https://italianelfuturo.com/reports/mckinsey-2025-ai-lavoro-superagenzia/>

McKinsey, "Agents for growth: Turning AI promise into Impact" in "[mckinsey.com](https://www.mckinsey.com)" del 3 novembre 2025.
<https://www.mckinsey.com/capabilities/growth-marketing-and-sales/our-insights/agents-for-growth-turning-ai-promise-into-impact>

Kate Crawford, "L'impero del calcolo" in Wired n.113 del 24 giugno 2025.

The Role of Governance and Human Control

European Union Regulation (EU), Regolamento (EU) 2024/1689 in "[europa.eu](https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng)".
<https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>

Jessica Kelly and others, "Navigating the EU AI Act: A Methodological Approach to Compliance for Safety-critical Products" in "publicarest.fraunhofer.de" del dicembre 2024.
<https://publicarest.fraunhofer.de/server/api/core/bitstreams/2724ce80-442b-470e-84af-5069d66ea193/content>

Chaffer, T.J., von Goins II C., Cotlage D., Okusanya B., Goldston, J., "Decentralized Governance of Autonomous AI Agents" in "arxiv.org" del 2024.
<https://arxiv.org/abs/2412.17114>

Maslej N., Fattorini L., Brynjolfsson E., Etchemendy J., Ligett K., Lyons T., Manyika J., Ngo H., Niebles J.C., Parli V., Shoham Y., Wald R., Clark J., Perrault R., "Artificial Intelligence Index Report 2025" in

"[arxiv.org](https://arxiv.org/abs/2504.07139)" del 2025.
<https://arxiv.org/abs/2504.07139>

Santoni de Sio, F., Van Den Hoven J., "Meaningful Human Control over Autonomous Systems: A Philosophical Account" in "[frontiersin.org](https://www.frontiersin.org/journals/robotics-and-ai/articles/10.3389/frobt.2018.00015/full)" del 2018.
<https://www.frontiersin.org/journals/robotics-and-ai/articles/10.3389/frobt.2018.00015/full>

Crawford K., "L'impero del calcolo", "Wired Italia" n. 113 del 24 giugno 2025.

Where Are Agents Being Used?

Alberto Bozzo, "Agenti di intelligenza artificiale in ambito sanitario: applicazioni, automazione e tendenze future" in "[consultingpb.com](https://www.consultingpb.com/blog/diritto-rovescio/agentic-ai-in-sanita/)" del 3 giugno 2025.
<https://www.consultingpb.com/blog/diritto-rovescio/agentic-ai-in-sanita/>

Preeti Anchan, "I migliori strumenti di IA per gli operatori sanitari della comunità" in "[clickup.com](https://clickup.com/it/blog/541916/strumenti-di-ia-per-operatori-sanitari-di-comunita)" del 15 novembre 2025.
[https://clickup.com/it/blog/541916/strumenti-di-ia-per-operatori-sanitari-di-comunità](https://clickup.com/it/blog/541916/strumenti-di-ia-per-operatori-sanitari-di-comunita)

Kees Hertogh, "Agentic AI in action: Healthcare innovation at Microsoft Ignite 2025" in "[microsoft.com](https://www.microsoft.com/en-us/industry/blog/healthcare/2025/11/18/agentic-ai-in-action-healthcare-innovation-at-microsoft-ignite-2025/)" del 18 novembre 2025.
<https://www.microsoft.com/en-us/industry/blog/healthcare/2025/11/18/agentic-ai-in-action-healthcare-innovation-at-microsoft-ignite-2025/>

Pedram Hosseini, "AI in Healthcare: Key Takeaways from Menlo Ventures 2025 Consumer AI Report" in "[linkedin.com](https://www.linkedin.com/pulse/ai-healthcare-key-takeaways-from-menlo-ventures-2025-report-hosseini-ujinc)" del 3 luglio 2025.

<https://www.linkedin.com/pulse/ai-healthcare-key-takeaways-from-menlo-ventures-2025-report-hosseini-ujinc>

Agid, "Piano triennale per l'intelligenza artificiale nella PA" del 2025.

[https://www.agid.gov.it/sites/agid/files/2025-01/Piano Triennale per I informatica nella PA 2024-2026 Aggiornamento 2025.pdf](https://www.agid.gov.it/sites/agid/files/2025-01/Piano_Triennale_per_I_informatica_nella_PA_2024-2026_Aggiornamento_2025.pdf)

Fabrizio Davide, Pietro Torre, "Agenti AI per la PA del futuro: ecco il valore aggiunto" in "[agendadigitale.eu](https://www.agendadigitale.eu/cittadinanza-digitale/agenti-ai-per-la-pa-del-futuro-ecco-il-valore-aggiunto/)" del 12 novembre 2024.

<https://www.agendadigitale.eu/cittadinanza-digitale/agenti-ai-per-la-pa-del-futuro-ecco-il-valore-aggiunto/>

Matthew Finio, Amanda Downie, "Agenti di intelligenza artificiale nel marketing" in "[ibm.com](https://www.ibm.com/it-it/think/topics/ai-agents-in-marketing)" del 2025.

<https://www.ibm.com/it-it/think/topics/ai-agents-in-marketing>

What Are the Unresolved Problems?

Maurizio Carmignani, "Le Allucinazioni dell'AI aumentano: ecco il perché" in "[agendadigitale.eu](https://www.agendadigitale.eu/cultura-digitale/le-allucinazioni-dellia-aumentano-i-nuovi-modelli-sbagliano-piu-dei-vecchi/)" del 9 giugno 2025.

<https://www.agendadigitale.eu/cultura-digitale/le-allucinazioni-dellia-aumentano-i-nuovi-modelli-sbagliano-piu-dei-vecchi/>

Andrea Signorelli, "Le Allucinazioni dell'Intelligenza Artificiale" in "[zanichelli.it](https://ia.zanichelli.it/scoprire-intelligenza-artificiale/le-allucinazioni-dell-intelligenza-artificiale/)" del 2025.

<https://ia.zanichelli.it/scoprire-intelligenza-artificiale/le-allucinazioni-dell-intelligenza-artificiale/>

Celm Dilmegani e Aleyna Daldal, "AI Hallucination: Compare top LLMs like GPT-5.2" in "[aiultiple.it](https://research.aimultiple.com/ai-hallucination/)" del 15 Dicembre 2025.

Luca Barbanotti, "Oltre la scatola nera: come l'intelligenza artificiale agentiva sta ridefinendo la spiegabilità" in "[sas.com](https://www.sas.com/it_it/news/leading-art-innovation/innovation-sparks/tecnologie/agentic-ai-oltre-la-black-box-cambia-nostro-rapporto-con-spiegabilita.html)" del 2025.

"L'impatto dell'intelligenza artificiale sul processo decisionale" in "[digital4pro.com](https://www.digital4pro.com/2025/03/18/limpatto-dellintelligenza-artificiale-sul-processo-decisionale/)".

Thseen Zia, "La trappola degli agenti di intelligenza artificiale: le modalità di errore nascoste dei sistemi autonomi per cui nessuno si sta preparando." in "[unite.ai](https://www.unite.ai/it/when-ai-turns-rogue-exploring-the-phenomenon-of-agentic-misalignment/)" del 13 dicembre 2025.

Who Is Needed to Govern These Problems?

"Artificial Intelligence Index Report" 2025 in "[stanford.eu](https://hai.stanford.edu/assets/files/hai_ai_index_report_2025.pdf)".

Santoni De Sio, F., Van Den Hoven J., "Meaningful Human Control over Autonomous Systems" in "[frontiersin.org](https://www.frontiersin.org/journals/robotics-and-ai/articles/10.3389/frobt.2018.00015/full)" del 2018.

Philip Drammeh, "Multi-Agent LLM Orchestration Achieves Deterministic Response Quality" in "[arxiv.org](https://arxiv.org/abs/2511.15755)" del 19 novembre 2025.
<https://arxiv.org/abs/2511.15755>

Natalia Campana, "What Does An AI Auditor Do?" in "[freelancemap.com](https://www.freelancemap.com/blog/what-does-big-data-specialist-do/)" dell'11 settembre 2025.
<https://www.freelancemap.com/blog/what-does-big-data-specialist-do/>

"Who's Responsible When AI Gets It Wrong?" in "[bronson.ca](https://bronson.ca/algorithmic-accountability/)" del 7 ottobre 2025.
<https://bronson.ca/algorithmic-accountability/>

"The Role of the AI Officer: Guide to Responsible AI Leadership" in "[aiguardianapp.com](https://www.aiguardianapp.com/ai-officer-responsibilities#:~:text=Similar%20to%20a%20Chief%20Data,initiatives%20are%20developed%20and%20deployed)" del 2025.
<https://www.aiguardianapp.com/ai-officer-responsibilities#:~:text=Similar%20to%20a%20Chief%20Data,initiatives%20are%20developed%20and%20deployed>

Bradley Hook, "What is an AI Ethicist and How Do I Become One?" in "[bradleyhook.com](https://bradleyhook.com/what-will-a-world-teeming-with-ai-agents-be-like/)"
<https://bradleyhook.com/what-will-a-world-teeming-with-ai-agents-be-like/>

Rethinking Skills: From the Logic of Doing to the Logic of Being

ILO (International Labour Organization) (2023-2024). "Generative AI and Jobs: A global analysis of potential effects on job quantity and quality" in "[ilo.org](https://www.ilo.org/publications/generative-ai-and-jobs-global-analysis-potential-effects-job-quantity-and)".
<https://www.ilo.org/publications/generative-ai-and-jobs-global-analysis-potential-effects-job-quantity-and>

The Return of the Philosophers: Why Humanities Skills Are the New Engineering

Olojede H.T, "Arts and Humanities in the Age of Artificial Intelligence", in "philarchive.org" del 2025.
<https://philarchive.org/archive/POLIPO>

Yana Samuel, "Cultivation of human centered artificial intelligence: culturally adaptive thinking in education (CATE) for AI", in "frontiersin.org" del 1° dicembre 2023.
<https://www.frontiersin.org/journals/artificial-intelligence/articles/10.3389/frai.2023.1198180/full>

Lina Khair Rahme, Rasha Halat, "Fostering Ethical and Critical Minds in the AI Era: A Practical Approach" in "harvard.edu" del 6 novembre 2023.
<https://mepli.gse.harvard.edu/our-fellows-at-work/fostering-ethical-and-critical-minds-in-the-ai-era-a-practical-approach/>

The Crossroads of 2030: Three Scenarios for an "Agentified" Future

UK Government Office for Science, "AI 2030 Scenarios Report" in "gov.uk" del gennaio 2024. This report uses five critical uncertainties to build future scenarios on AI development, offering a robust methodological framework for scenario analysis.
https://assets.publishing.service.gov.uk/media/6808fc002a86d6dfb2b52772/AI_2030_Scenarios_Report.pdf

Deloitte, "Autonomous Workforce Planning: How Agentic AI Transforms Workforce Strategy" in "deloitte.com". The report explores how agentic AI transforms workforce planning from static models to

dynamic and adaptive systems.
<https://www.deloitte.com/us/en/insights/topics/talent/future-of-workforce-planning/autonomous-workforce-planning.html>

Felten, E.W., Raj M., Seamans R., "Forecasting the Impact of Artificial Intelligence on the Labor Market" in "[ssrn.com](https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4710299)". Proposes innovative methodologies to measure and forecast AI's impact on the labor market, using dynamic data from real-time online job postings.
https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4710299



Massimo Tamiatti

I have always been involved in the study and dissemination of topics related to new technologies, new jobs, and new skills. I do this in a completely new and different way compared to traditional research approaches because, in addition to consulting traditional sources, I consult sources from major strategic consulting firms, which do not have a static but dynamic vision of reality, as they focus on the future. I systematically use the web to search for the most up-to-date content, perhaps through webinars with privileged witnesses on the subject of the investigation or through some important influencer specialized in the topic being addressed. Obviously, today, I use Generative Artificial Intelligence tools. I do not neglect the present; I try to interpret the signals that we have before our eyes, which we are often used to looking at and not seeing, and which suggest what could happen. I do not neglect the past, that is, history, but in reality, many disciplines are mixed: economics, statistics, sociology, psychology, anthropology, and above all, I use imagination anchored to the facts that happen every day in the world. I consider numbers very important, but they are not everything; at the center, there must always be a particular attention to the privileged witnesses of the things that happen and their personal stories. Frontiere does not have the gift of foresight, but our team can make hypotheses, and among these hypotheses, I have realized that it is very likely to encounter the future. Spreading the culture of anticipation helps people understand what is happening in the world, especially in the world of work, helps them prepare, elaborating strategies, and organizing themselves to make very important conscious and correct choices for the effect they will have on their lives.

Institutional Profile

I was Chief Research Officer at the Piemonte Lavoro Agency, where I coordinated research groups focused on labor policies and vocational training until June 30, 2023. Since July 1, I have been part of the staff of the "Riccardo Revelli Laboratory" as an honorary member. I am an expert in public employment services and currently teach in the Human Resources Master's program at the University of Turin. I am a member of the Scientific Committee of ISMEL and serve on the Scientific Committee of Polis policy, an Academy of Advanced Training. Since 2022, I have been the founder and director of the online scientific dissemination site frontieretecnologialavoro.com and a member of the Association of Italian Futurists founded by Roberto Poli. I specialize in the use of forecasting techniques applied to local labor markets, collaborating with the start-up Skopia, which focuses on futures studies and anticipation to help organizations develop forward-looking strategies in support of complex decision-making under high uncertainty. I hold a degree in Literature and Philosophy with theses in Contemporary History and Political Science, both completed at the University of Turin, and I am the author of numerous publications related to the labor market, particularly on new labor market scenarios, emerging professions, and new skills.

The following collaborated with me on this project: Gabriele Lamberti, Benjamin Dafku, Ivan and Federico Giacobino e Fabio Prettico.