

When Artificial Intelligence Enters the Realm of Responsibility

Why I Propose Governance of Decision-Making

By Salvatore Martino — PENSAI

Artificial intelligence is rapidly changing the relationship between people, organizations, and decisions.

For years, I have asked myself a primarily technological question: **how intelligent can a machine become?**

Today, I believe this question is no longer sufficient.

When an answer produced by an AI system enters a real-world process and contributes to determining a decision, the issue changes.

I therefore ask myself:

who controls what AI produces, according to which rules, and with what responsibilities?

This question is the starting point of my reflection on decision governance.

A New Regulatory Development

The **Italian Legislative Decree of 9 September 2026, No. 160**, published in the Official Journal on 15 September 2026, represents another step in the development of Italian legislation concerning artificial intelligence.

The decree enters into force on **30 September 2026**.

I identify an especially significant element for my reflection: security, human oversight, and responsibility are acquiring an increasingly concrete dimension in the relationship between AI systems and human activities.

I do not interpret this decree as automatic evidence of PENSAI's validity.

Rather, I consider it part of the context in which I am developing the project.

The legislator is addressing an issue that I consider fundamental:

how can artificial intelligence be used while maintaining control and responsibility?

The Problem Is Not Only the Intelligence of the Machine

When I analyze an AI system, I am not interested only in knowing how powerful it is.

I also want to understand what happens before and after its output.

Let us imagine an organization using an AI system.

The system receives data.

It produces an analysis.

It proposes several alternatives.

It indicates a solution.

At this point, I could consider the process complete.

But according to the methodology I am developing with PENSAI, several questions remain:

What was the initial problem?

What was the objective?

What alternatives were available?

Which alternatives were excluded?

According to which criteria?

Which constraints applied?

What level of risk was acceptable?

How uncertain was the result?

Who could authorize the action?

When should the process have stopped?

For me, these questions concern the **decision-making process**, rather than simply the AI model.

A Convincing Answer Is Not Necessarily a Verifiable Answer

Generative AI makes it extremely easy to obtain sophisticated answers.

But I do not consider a well-formulated answer automatically correct, complete, or suitable for a specific context.

When an output is transferred directly into an operational process, the risk increases.

For this reason, I propose separating at least five stages:

generation → evaluation → decision → action → verification.

This is precisely the space in which I place PENSAI.

PENSAI: My Concrete Proposal

PENSAI is not an artificial intelligence model.

It is not a chatbot.

It was not created to replace the human decision-maker.

I do not present it as a certification of compliance with the AI Act.

I define it as a **decision-governance framework**.

With PENSAI, I propose structuring the process through which a situation is transformed into a decision.

The structure I have developed begins with:

**STATE → OBJECTIVE → OPTIONS → CONSTRAINTS → EVALUATION → DECISION
→ ACTION → VERIFICATION**

My objective is not simply to produce more information.

I want to structure the way in which information is used to make decisions.

What PENSAI Offers Today

On this point, I want to be particularly clear.

I do not want to present as an already available product something that still belongs to future development.

Today, PENSAI is primarily a methodological framework for decision governance, with a formalized and implemented deterministic decision-making core.

Using the PENSAI methodology, I can structure a decision-making situation starting from the problem and objective, identifying alternatives, constraints, evaluation criteria, risks, and the conditions required to reach a decision.

I can also define the resulting actions and provide for subsequent verification.

The value I currently attribute to PENSAI is therefore primarily **methodological**.

The software tools intended to make the framework increasingly operational represent a development phase.

This distinction is essential to me.

I prefer to describe accurately what I have today rather than attribute capabilities to the project that it does not yet possess.

The Single Path

One of the central concepts I have developed within PENSAI is the **Single Path** — *Via Unica*.

By this term, I do not mean that there is always one objectively correct solution.

I mean that once the problem has been defined, objectives and constraints established, alternatives considered, and explicit criteria applied, the process should be able to reach a reconstructable operational choice.

First, I consider the possibilities.

Then I filter them.

Next, I evaluate them.

Finally, I identify the operational path compatible with the defined conditions.

For me, the Single Path is therefore not a shortcut.

It is the outcome of a process.

I Can Also Stop the Process

Another characteristic that I consider fundamental is the possibility of **not making a decision**.

If essential information is missing, a critical constraint is violated, or the risk exceeds acceptable conditions, PENSAI can reach:

STOP.

I consider this principle important because I do not want to build a methodology that is forced to produce an answer every time.

I prefer a system capable of stating:

“There are not sufficient conditions to proceed.”

For me, this is also a decision.

Suspension does not necessarily represent failure.

It can represent a form of control.

Determinism and Transparency

The PENSAI decision-making core uses a deterministic logic.

This allows me to distinguish the decision-making process from the probabilistic behavior typical of generative systems.

I do not consider determinism a guarantee of correctness.

If I enter incorrect data or use inappropriate criteria, I can obtain an incorrect decision even through a perfectly deterministic process.

The advantage I seek is different:

to make the decision-making path reconstructable.

I want to be able to understand which conditions, criteria, and evaluations led to the outcome.

How I Govern Uncertainty

I do not consider it realistic to build decisions on the assumption that all information is certain.

Uncertainty exists.

For this reason, within the PENSAI methodology, I try to make uncertainty explicit rather than hide it behind apparent precision.

I am particularly interested in identifying the variables that can genuinely change the outcome.

In other words, I do not try to pretend that I know what I do not know.

I try to establish **how much what I do not know can influence the decision.**

PENSAI Together with Artificial Intelligence

I do not believe PENSAI should replace AI.

I believe it can work together with it.

AI can help me generate information, scenarios, analyses, and alternatives.

PENSAI can provide a methodological structure for organizing and evaluating them and subjecting them to explicit conditions before a decision is made.

The separation I propose is therefore:

AI → generates possibilities

PENSAI → structures and governs the process

HUMAN BEING → decides and retains responsibility

POST-CHECK → verifies the outcome

This is the position I am developing through PENSAI.

I Do Not Want to Confuse Oversight with Control

In my view, simply placing a person at the end of a process does not necessarily mean that genuine human oversight exists.

If a person receives an already-packaged result and can only approve or reject it without understanding the path that produced it, control risks becoming merely formal.

For this reason, I consider it necessary for the decision-maker to be able to:

- understand the process;
- verify the conditions;
- challenge the result;
- modify the decision;
- suspend the process;
- prevent the action when sufficient conditions are absent.

This is the space in which I want to place PENSAI governance.

PENSAI Does Not Certify Compliance

I also want to be clear about what **I do not offer**.

PENSAI does not certify an AI system.

It does not certify compliance with the AI Act.

It does not replace a legal assessment.

It does not replace an audit.

It does not automatically guarantee that an organization complies with applicable legislation.

I propose something more specific:

a methodology for structuring and governing the decision-making process in which AI is used.

This distinction is fundamental.

I do not want to sell PENSAI as a shortcut to compliance.

I want to develop it as a methodological tool that can contribute to making decision-making processes more controllable.

From Framework to Software

I see three distinct levels ahead of me.

Today

I have a **methodological framework for decision governance**, with a formalized and implemented deterministic decision-making core.

Development

I am moving the methodology toward software tools capable of making it easier to use, repeat, and integrate.

Evolution

I want to develop PENSAI as a possible **decision-governance layer** that can also be applied to AI systems and agentic systems.

However, I do not consider these three levels equivalent.

The first exists.

The second is under development.

The third represents the project's intended direction.

PENSAI and AI Agents

It is especially with the evolution of AI agents that I believe this issue will become even more interesting.

A system that only produces an answer is one thing.

A system capable of planning, using tools, and carrying out sequences of actions is another.

As AI moves from generation to action, I consider the following questions increasingly important:

Who authorizes the action?

Under which conditions?

What limits are applied?

When must the system stop?

What is recorded?

How can I subsequently verify what happened?

From this need, I am developing the concept of **PENSAI AGENT-GOV**.

The logic I propose is:

Decision Gate → Action Gate → Execution Trace → Post-Check

In this model, I do not want to prevent an agent from being operational.

I want to introduce a layer that governs the transition from a proposed decision to an authorized action.

A Possible Future Infrastructure

My vision can be represented as follows:

DATA / INPUT

↓

AI SYSTEM

↓

PENSAI GOVERNANCE LAYER

↓

EVALUATION / FILTER / RISK / CRITERIA

↓

DECISION GATE

↓

HUMAN DECISION OR AUTHORIZED ACTION

↓

POST-CHECK / TRACE

This is not a description of a fully available platform.

It is the architecture toward which I intend to take the project.

Why I Consider This a Potential Market

If artificial intelligence continues to move from simple content generation toward operational processes and agentic systems, I believe the need will also grow to distinguish between:

what AI can suggest

and

what an organization can authorize.

This is the space in which I see a potential market area for PENSAI.

I envision applications in:

- AI Decision Governance;
- Decision Intelligence;
- multicriteria evaluation;
- uncertainty management;

- process supervision;
- decision traceability;
- post-decision verification;
- governance of agentic systems.

I do not want to compete with the major generative models.

I want to work at a different level.

My Thesis

My thesis is simple.

Artificial intelligence can produce a possibility.

But a possibility is not yet a decision.

A decision is not yet an action.

And an action should not be separated from the responsibility of the person who authorizes it.

For this reason, I propose PENSAI as a methodological layer between the ability of AI to generate possibilities and the human responsibility to decide.

Conclusion

The new Italian regulatory framework makes an issue that I consider central even more evident: artificial intelligence cannot be assessed solely according to its capabilities.

I must also consider **how it is used, controlled, and incorporated into decision-making processes.**

This is the need from which I developed PENSAI.

Today, PENSAI is a decision-governance framework, with a methodological and deterministic core that has already been formalized and implemented.

I am working on its evolution toward more complete operational tools and, subsequently, toward a possible governance architecture applicable to agentic systems as well.

I do not present PENSAI as a new artificial intelligence.

I do not present it as a certification.

I do not present it as a guarantee of perfect decisions.

I propose it as **a method for governing the process through which an answer becomes a decision.**

My fundamental idea is this:

I do not want AI simply to make decisions faster.

I want it to be possible, whenever AI contributes to a decision, to understand **why that decision was made, under which conditions, within which limits, and with what possibility of verification.**

Technology produces possibilities.

I propose governing the transition from possibility to decision.

This is PENSAI.

Disclaimer

I wrote this essay with the support of artificial intelligence tools.

I used AI as support for research, information organization, analysis, and drafting. I verified the information used on the basis of the sources indicated and invite readers to consult those sources directly.

The assessments, interpretations, and proposal concerning PENSAI represent my own work and project development. When I describe future developments, I present them as my hypotheses and not as achievements already accomplished.

I present PENSAI as a project and decision-governance framework under development. I do not present it as a certification of regulatory compliance, nor as a substitute for legal, technical, or professional advice.

The hypotheses concerning the market, technological development, and possible future applications of PENSAI are not forecasts or guarantees of results.

I am Salvatore Martino, and PENSAI is the project I am developing to address a specific question: how can we govern decision-making in the age of artificial intelligence?

Salvatore Martino — PENSAI