

IN-DEPTH GUIDE

RISKS AND GOVERNANCE

What artificial intelligence can't do



GXN

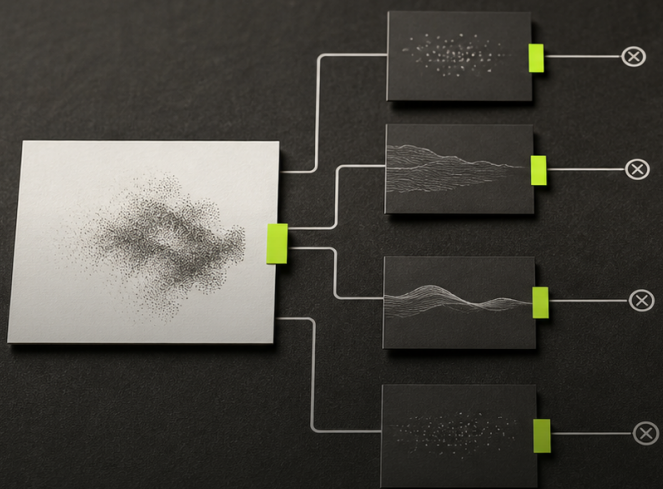
Senior expertise · Digital governance

Direct response

Limits, risks and decisions

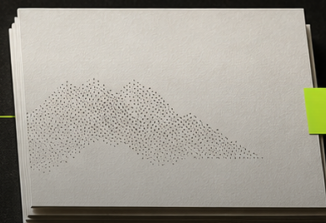
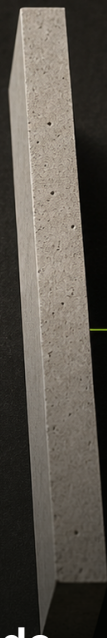
July 30, 2026

Updated



ARTIFICIAL INTELLIGENCE · RISKS AND GOVERNANCE

What artificial intelligence can't do



COMPREHENSIVE STUDY · GXN

TL;DR

The answer in one minute

- AI sometimes invents with confidence. This is not a temporary defect, it is a property of its operation.
- It is not constant. The same question twice can give two different answers.
- She doesn't know your business. Everything that is not written anywhere is invisible to him.
- She calculates badly. A language model is not a calculator, even when it looks like one.
- She bears no responsibility. Yours remains intact.
- The good news: each of these limitations is managed by design. None prevents a useful project. All of them prevent a naive project.

01 · BACKGROUND GUIDE

Why do I write this?

Almost everything written about AI in French is either promotion or fear. Both leave you without reference to decide.

This guide does the opposite. It lists what these systems do not do well, in 2026, with what this concretely implies for a business project. You won't find a conclusion like "so AI is overvalued". I make my living implementing them and I continue to do so, because it works when it's well designed. But well-designed means designed around these limitations, not ignoring them.

If you need to read a single section of the AI cluster before signing anything, read it.

BOUNDARY MAP

Seven constraints to be addressed by design

Invention

Variability

Absent context

Calculation

Judgment

Liability

Process

Limit 1. She invents, and she does it with aplomb

This is the best known and most misunderstood limit.

A language model produces plausible text. When he knows the answer, the plausible text is also the true text. When he does not know it, he still produces plausible text. A date, a section number, a standard name, an amount. Nothing in the formulation indicates the difference, and that is exactly what makes the phenomenon dangerous.

What people misunderstand is that this is not a bug waiting to be fixed. It comes from the very mechanics of the system. The best models make mistakes less often. Good architecture further reduces the rate. Nobody takes it to zero.

How we design around it. The system is responded to from documents provided rather than from its general brief, it is required to cite its sources, and human validation is placed wherever error is costly. A system that answers "I can't find this information in your documents" is a good system, not a failed system.

The real test to do before buying: ask the supplier to show you what their system answers to a question to which the answer does not exist in your documents. You'll learn everything you need to know in thirty seconds.

Limit 2. It is not constant

Ask twice the same question, you may get two different formulations. Sometimes two slightly different contents.

For writing, this is without consequence or even desirable. For data extraction, this is a serious problem. A system that reads an invoice must release the same amount each time, point.

How we design around it. Creativity settings are lowered, strict output formats are imposed, automatic checks are added that reject a poorly trained response, and a representative volume is tested rather than three selected examples. A serious supplier will show you a measured accuracy rate on dozens of real cases, not a successful demonstration.

Limit 3. She doesn't know anything about your business

The model read the Internet. He has never seen your price list, your collective agreement, your bid approval procedure, nor the fact that the Tremblay client has had a particular rate since 2019.

And above all, he doesn't know what no one has ever written. However, in most SMEs, a considerable part of the knowledge lives in the heads of three or four people. This tacit knowledge is invisible to the machine.

How we design around it. It is given access to relevant documents with the appropriate method. And we accept an uncomfortable truth: an AI project almost always reveals that your documentation is incomplete. This is not a failure of the project, it is a free diagnosis of the state of your internal knowledge. Many of my clients have gained more value from this awareness than from automation itself.

Limit 4. She calculates poorly

It always surprises. A system capable of summarizing a 60-page contract can be wrong by adding a column of numbers.

The reason is simple: it doesn't calculate, it predicts what a calculation answer looks like. Modern systems get around the problem by calling a real calculator or your database. But it must have been planned.

How we design around it. Any encrypted data that counts must come from a system that really calculates, never from the model itself. If a supplier offers a pricing or tax calculation tool based solely on a language model, there is an architectural problem.

Limit 5. She doesn't judge

A model can tell you what a contract clause says. He cannot decide whether to accept this clause, because this decision depends on your risk tolerance, your relationship with this client, your current financial situation and three things that no one has written anywhere.

The same logic applies to hiring, credit, health, discipline, or anything that lastingly affects a person.

How we design around it. The AI prepares the decision, the human makes it. It brings together, summarizes, highlights, proposes. The final action belongs to someone whose responsibility it is. This is not excessive caution, it is also the only defensible position the day someone calls you to account.

Limit 6. She bears no responsibility

If your system sends bad information to a client, the model provider will not answer. It's you. The conditions of use of all major suppliers are very clear on this, and nobody reads them.

How we design around it. Before conception, we identify what actions the company takes. These pass through a human, or do not automate at all. We also document who validated what, which is worth gold in case of litigation.

Limit 7. It does not fix a broken process

This one is not technical, and yet it is the one that kills the most projects.

Automating a poorly defined process produces a faster poorly defined process. If three people in your house do the same task in three different ways and none of them is right, no model will decide for you.

How we design around it. The process is documented before automating. Often this step reveals simplifications that give half the expected gain without a line of code. I saw warrants end there, with a satisfied client and a much smaller bill than expected.

What will improve, and what will not change

A useful distinction for planning.

What improves quickly: the hallucination rate decreases, reasoning skills improve, costs per token decrease almost every year, context windows enlarge, integration tools stabilize.

What will not change: A system will never know what nobody has written. Responsibility will remain human. Decisions with high stakes will always require someone to assume them. And a poorly defined process will remain poorly defined.

Practical conclusion: build for today, but with an architecture that allows you to change models without rewriting everything. The model is an interchangeable part. Your process, your data and your rules are the durable asset.

What I see on the ground

The two opposite errors, each as costly as the other.

The first is blind trust. A company plugs an assistant into its documentation and lets it respond directly to customers. It works for four months. Then a wrong answer goes to an important client, and the project dies in one meeting. Technology was not the problem. The lack of validation, yes.

The second is total refusal. “The AI is wrong, so we don’t touch it. » Except that your employees already touch it, on their personal accounts, with your documents. The official refusal does not make the use disappear. It just makes it invisible and unframed.

The reasonable position is somewhere in between, and it requires a bit of work: explicitly deciding where error is acceptable, and where it is not.

When these limits have to push you back

I say it clearly. Do not launch an AI project if you are in one of these situations.

- You cannot tolerate one mistake and you do not want any human validation. These two conditions are incompatible.
- Your reference data is false or contradictory.
- No one in the company can describe the process on a page.
- The project is driven by image rather than problem.
- You are looking to remove jobs without reviewing the work itself.

In these cases, the money is better placed elsewhere, and I will tell you that before charging anything.

QUESTIONS AND ANSWERS

Frequently asked questions

01 **The hallucinations will disappear.**

The rate is falling, it is not going to zero. Design assuming they happen, not hoping they stop.

02 **How to measure whether a system is reliable enough.**

By testing it on a representative volume of real cases, with a measured accuracy rate, and by defining in advance what rate is acceptable for this specific task. A supplier who refuses this exercise should worry you.

03 **Is it worth it despite all these limitations?**

On the right cases of use, yes, widely. Documentary processing and internal research yield measurable and lasting gains. Poorly targeted projects fail, not technology.

04 **A more expensive model is more reliable.**

Often a little, but the reliability gap comes much more from architecture than from the model. A small, well-structured model, with RAG and checks, beats a big model left without any crazy guards.

05 **What if our employees already use AI without coaching.**

Look what they're doing before they ban anything. They show you for free where your friction is. Then frame with a simple policy and approved tools.

EDITORIAL METHOD

Background guide prepared by GXN (Digital Governance), a division of MD79. Revision planned according to the rhythm documented for this cluster.

OFFICIAL SOURCES

References to be checked according to your situation

These references support the external rules and frameworks cited in this guide. They are not a substitute for legal or professional advice tailored to your organization.

NIST

Artificial Intelligence Risk Management Framework





ABOUT THE PUBLISHER

GXN

A directly accessible senior expertise, with more than twenty years of experience in technology implantation and a master's degree in artificial intelligence within the company.

[Discover GXN ↗](#)

NEXT STEP

You want to know if your usage case stands up despite these limitations.

This is exactly what the AI diagnosis evaluates: feasibility, acceptable error rate, validation points, real costs. And sometimes the conclusion is that the project is not ready. You will have it in writing.

No call required. No sales sequence.