

One Simple Way to Get Better Results from an LLM: Make the Right Information Govern

Large language models (LLMs) can have the right information in front of them and still give the wrong answer.

That sounds strange. If the information is there, why doesn't the model simply use it?

One useful way to think about this is to distinguish between information that is *available* to an LLM and information that actually *governs* what it says next.

That distinction can make a surprising difference in how you work with an LLM.

A simple example

Consider an example from the educational module [*A Semiotic Account of the Hierarchical Relational Organization of Large Language Models*](#).

Suppose an LLM is given a medical study that says, in effect:

The study found an association, but because it was observational, it does not establish that the treatment caused the improvement.

The qualification is clear. The information is right there in the context.

But the LLM might nevertheless produce something like:

The study suggests that the treatment produced the improvement.

Notice what happened.

The model did not necessarily lose the information. It preserved the association between the treatment and the improvement.

What disappeared was the *constraint on the inference*: association does not establish causation.

The information was present.

But it was no longer governing.

Why this matters

This problem occurs far beyond scientific studies.

You can give an LLM a document, a set of instructions, a previous decision, a definition, or an entire conversation history. All of that material may remain available in the context.

But availability alone does not guarantee that it will continue to control the answer.

An LLM may:

- preserve a fact but ignore its qualification;
- remember a rule but fail to apply it;
- cite a source but make a claim the source does not support;
- retain your instructions while gradually drifting away from the original task;
- repeat a concept correctly while reasoning according to a different framework.

This is why simply giving an LLM *more context* does not always solve the problem.

The more useful question is:

What information should govern the answer?

Turn important information into governing constraints

Suppose you provide an LLM with a source.

Instead of saying only:

Use this document to answer the question.

you can make the relationship more explicit:

Base factual claims on this document. Preserve qualifications that limit what can be concluded. Distinguish what the source states from what you infer from it. Mark claims that are not supported by the source.

Now you are doing more than adding information.

You are telling the model *how that information should constrain the continuation*.

The same principle can be used in ordinary conversations:

- If a distinction is essential, say that it must remain governing.
- If a source is authoritative, specify what authority it has.

- If an earlier decision should not silently change, identify it as a continuing constraint.
- If two concepts must remain separate, say so.
- If the conversation begins drifting, do not necessarily add more information. Re-establish the relation that has been lost.

A useful diagnostic question

When an LLM gives you an answer that seems fluent and intelligent but somehow misses the point, ask:

What should have constrained this answer but didn't?

That question can be more useful than simply asking, "What did the model forget?"

The model may not have forgotten anything.

The relevant information may still be sitting in the conversation. The problem may be that another pattern, assumption, instruction, or line of reasoning has become more influential.

In that case, the solution is not necessarily retrieval.
It is *recalibration*.

Restore the constraint that should be governing the conversation.

The bigger idea

This simple practice points toward a different way of thinking about LLMs.

A conversation with an LLM is not just a growing pile of information. Different parts of the context play different roles. Some provide examples. Some provide evidence. Some define terms. Some establish the task. Some place limits on what may legitimately be concluded.

These relations can be organized at different levels, and some should govern others.

That is the larger idea explored in [*A Semiotic Account of the Hierarchical Relational Organization of Large Language Models \(LLMs\): Implications for AI Development, Use and Evaluation.*](#)

You do not need to understand the whole theoretical framework to begin using one of its practical lessons.

Take away message

When working with an LLM, don't ask only:

Is the right information there?

Also ask:

Is the right information governing what happens next?

Further reading:

[*From Things to Relations: Rethinking how large language models \(LLMs\) work*](#)

Timothy M. Rogers, University of Toronto,
In a structured interaction with ChatGPT,
August 25, 2026