



What every professional needs to know about how AI works

From hallucinations to context windows: the technology behind AI made understandable

A whitepaper by Laurens Vandebergh, newTera
Business Development Manager AI and Project Manager
2026

Executive Summary

You do not need to be an engineer to work effectively with AI. But you do need to understand how the technology is built, what its fundamental properties are, and where its limits lie. Those who do not know this overestimate AI at the wrong moments and underestimate it at the right ones.

This whitepaper offers a thorough but accessible overview of how large language models (LLMs) work: from the three technological breakthroughs that made them possible, to the five properties every professional needs to know in order to use AI responsibly and effectively.

The goal is not technical completeness. The goal is a sharp and practical mental model that helps you make better decisions about when and how to use AI, what to expect from it, and when human judgment remains irreplaceable.

AI Fluency does not start with tools. It starts with understanding. Those who know how AI works, work with it in a fundamentally different way.

The problem: AI as a black box

Many professionals use AI tools every day without a clear picture of what is happening under the hood. That is understandable: the technology is complex and the tools are designed to hide that complexity. But that lack of understanding comes at a cost.

Those who do not understand how an LLM works also do not understand why it sometimes performs brilliantly and sometimes fails spectacularly. Why it presents incorrect facts with great confidence. Why it forgets earlier instructions. Why the same request produces a different answer twice.

These seemingly random behaviors are not bugs. They are direct consequences of how the technology works. Those who understand this can account for it. Those who do not will be caught off guard.

The greatest risk with AI is not that it fails. It is that you do not know when it fails.

How does a large language model work?

A Large Language Model (LLM) is an AI system that generates text by predicting which word or phrase is most likely to follow what is already there. That sounds simple. The implications are far-reaching.

LLMs are not programmed with rules or knowledge bases. They learn by recognizing patterns in enormous amounts of text: billions of documents, books, articles, and web pages. What comes out is not understanding in the human sense of the word. It is an extraordinarily refined statistical model of how language works and how ideas relate to each other.

Three technological breakthroughs that made LLMs possible

The current generation of AI did not appear out of nowhere. Three developments converged at the right moment:

- **Algorithms:** The transformer architecture (2017) made it possible to draw connections between words that are far apart in a text. This was the key to coherent, contextual language processing.
- **Data:** The digitization of enormous amounts of human knowledge gave models the training material to learn complex patterns in language and reasoning.

- **Computing power:** The explosive growth of GPU (Graphics Processing Unit) clusters made it possible to train models with hundreds of billions of parameters.

None of these three elements was sufficient on its own. The breakthrough came from their convergence. The result is systems with emergent capabilities: skills that were not explicitly programmed, but arose naturally from the scale of training. Reasoning, translating, writing code, drawing analogies: no one explicitly taught LLMs to do any of this.

Five concepts every professional needs to know

The technical workings of LLMs have direct practical consequences. The five concepts below are not academic curiosities. They determine, every day, what AI can and cannot do for you.

CONCEPT	Hallucination
WHAT IT IS	An LLM generates text based on statistical probability, not factual verification. When a model has insufficient training data about a specific topic, it fills the gap with plausible-sounding but factually incorrect information. We call this hallucination.
IN PRACTICE	Always treat factual claims in AI output as hypotheses that require verification, especially with specific names, dates, statistics, and citations. The more specific the claim, the greater the risk of hallucination.

CONCEPT	Context window
WHAT IT IS	An LLM processes information within a fixed window: the maximum amount of text it can 'see' at once. Information outside that window does not exist for the model. In long conversations or documents, early information disappears from the window and can no longer be used.
IN PRACTICE	In long conversations or complex documents: periodically repeat crucial information or summarize the beginning of a conversation. Never assume the model 'remembers' what was said earlier in the session.

CONCEPT	Non-determinism and temperature
WHAT IT IS	Unlike traditional software, an LLM does not always give the same answer to the same question. The model makes probabilistic choices at each step about the next token (word fragment). 'Temperature' is a setting that determines how much randomness goes into those choices: high for creativity, low for precision.
IN PRACTICE	For creative tasks, variability is an advantage. For factual or critical tasks, consistency is preferable. Some platforms offer temperature settings; in others this is hidden. Be aware of the effect.

CONCEPT	Knowledge cutoff
WHAT IT IS	LLMs are trained on data up to a specific date. After that date, the model knows nothing about new developments, unless those are provided via external tools such as search functions. A model trained on data through early 2024 knows nothing about events or publications after that date.
IN PRACTICE	Always check the training cutoff date of the model you are using. For current information, market data, or recent legislation: use models with web search capability or verify independently.

CONCEPT	Pre-training and fine-tuning
WHAT IT IS	LLMs go through two learning phases. Pre-training is the broad phase: the model learns language and knowledge from enormous datasets. Fine-tuning is the targeted phase: the model is steered toward specific behavior, such as helpfulness, safety, or domain-specific focus. Fine-tuning largely determines how a model behaves in practice.
IN PRACTICE	Not every model is suited to every task. A general model performs differently from a domain-specific model. For critical applications, it is worth evaluating which model best fits your specific use case.

What AI is not: the human benchmark

One of the most persistent misconceptions about AI is the assumption that a system that masters human language also thinks in a human way. It does not. LLMs have no intentions, no consciousness, and no understanding in the phenomenological sense. They are exceptionally good at mimicking the outcomes of understanding. That is fundamentally different.

This distinction is not academic. It has direct practical implications for how you deploy AI, what you expect from it, and where you deliberately apply human judgment as a counterbalance.

AI IS STRONG AT	PEOPLE ARE STRONG AT
• Speed and scale • Pattern recognition across large datasets • Consistently reproducing structures • Drawing broad connections across domains • Executing tasks without fatigue	• Judging with incomplete information • Ethical reasoning and contextual understanding • Creative leaps beyond known frameworks • Bearing responsibility • Building trust relationships

The most effective uses of AI are not a replacement for human work, but an amplification of it. AI handles scale and speed. People bring judgment, context, and responsibility. Organizations that are clear on this distinction build lasting advantage.

What does this mean for your organization?

A technical understanding of AI translates directly into better organizational choices. Three implications deserve particular attention:

Model selection is a strategic decision

Not every AI system is the same. Models differ in training date, specialization, context window size, safety measures, and data processing policies. Organizations that choose tools at random run risks around quality, privacy, and compliance. A deliberate model choice, based on the specific use case, is a competitive advantage.

Human oversight is not an optional extra

Given the properties of LLMs, particularly hallucination and non-determinism, human oversight of AI output is not a precaution but an architectural principle. Workflows in which AI output leaves the organization without review are inherently vulnerable. The question is not whether errors will occur, but when.

Technological literacy as an organizational competency

Organizations that invest in broad technological literacy, not just in IT but across all employees who work with AI, make better decisions. They recognize the limits of the technology, ask the right questions, and know when to fall back on human judgment. That is not a nice-to-have. It is a risk management measure.

The AI landscape is changing fast

The speed at which LLMs are developing is unprecedented. Limitations that apply today may be resolved tomorrow. Context windows are growing. Hallucinations are decreasing. New architectures are emerging. Those who understand the fundamental workings can interpret these developments and anticipate them. Those who only know the tools are always playing catch-up.

How do you translate this into practice?

Technical understanding only has value if it leads to different choices and behavior. Three concrete steps to get started today:

- **Step 1** — Know your tools: for every AI system your organization uses, find out what the training cutoff date is, how large the context window is, and what data processing policy applies. These are basic facts every user should know.
- **Step 2** — Build in verification: make fact-checking a standard step in AI-supported work, especially for external communications. Not out of distrust, but as a professional standard.
- **Step 3** — Invest in understanding: ensure that employees who work intensively with AI know and can apply the five core concepts from this whitepaper. Technological literacy is a skill, not an innate trait.

An AI Readiness Assessment maps out how strongly your organization scores on technological understanding and AI Fluency as a whole, and which targeted interventions will have the most impact.

Conclusion

AI is neither magic nor a threat. It is a powerful technology with specific properties, strengths, and fundamental limitations. Those who know these properties can deploy the technology where it adds value and exercise caution where human judgment is irreplaceable.

Technological literacy is the quiet engine behind effective AI adoption. The organization that wins is not the one with the most tools. It is the one that enables its people to understand, evaluate, and steer those tools. That is precisely what AI Fluency is about.

Want to know how AI-literate your organization is?

newITera's AI Readiness Assessment maps out where your people, processes, and governance stand across all dimensions of AI Fluency.

Contact us at laurens.vandebergh@newitera.nl or visit www.newitera.nl

About newITera & Laurens Vandebergh

newITera is a Dutch IT consultancy based in Beuningen, Gelderland. Founded in 2011, initially as a specialized SAP firm, newITera has grown into a broad digital transformation partner. Today, newITera combines deep SAP expertise with services in data analytics, AI, cloud strategy, and low-code platforms. With a team of more than a hundred consultants, newITera serves leading

enterprises and multinationals across their full digital transformation journey. What sets newITera apart is its combination of early adoption and deep specialization, as an independent partner that values human connection just as highly as technical excellence.

Laurens Vandebergh is Business Development Manager AI and Project Manager at newITera. With a broad background in technology and organizational change, he guides organizations in defining their AI strategy, conducting an AI Readiness Assessment, and realizing concrete use cases.

Contact: laurens.vandebergh@newitera.nl | www.newitera.nl