

Language, Simulation, Imagination

Large Language Models as Social Technology

Francesco Zorgno

Working Paper

*based on research conducted within the Degree Programme in Philosophy and Digital Transformation,
University of Udine*

Abstract

Large language models (LLMs) have reopened fundamental questions in the philosophy of language and in the ethics of artificial intelligence. Conversational systems such as ChatGPT, Claude, Gemini and comparable models now produce fluent, context-sensitive discourse without possessing consciousness, intentionality or semantic understanding in the human sense. Their relevance cannot therefore be captured either by treating them as artificial minds or by dismissing them as merely deceptive linguistic devices.

This working paper interprets LLMs as social technologies of communication. Drawing on Daniel Dor's account of language as a technology for the instruction of imagination, it argues that the philosophical significance of LLMs lies less in whether they 'really understand' and more in the way they transform the conditions under which meaning is generated, attributed and validated. The paper situates this claim between innatist and functionalist accounts of language, emphasizing both the limits of anthropomorphism and the practical effectiveness of artificial linguistic performance.

The paper also distinguishes between pure deception, in which users mistakenly attribute human-like understanding to the system, and conscious deception, in which the fiction of dialogue is knowingly accepted as a pragmatic interface. The conclusion proposes a relational account of intelligence: LLMs are not intelligent subjects, but human-LLM interactions may participate in intelligent processes when they generate meaningful, adaptive and critically controlled outcomes.

Keywords: large language models; philosophy of language; artificial intelligence; Daniel Dor; imagination; social technology; deception; authorship; relational intelligence; ChatGPT.

Note on the present version

This working paper is based on the research document *Linguaggio, simulazione, immaginazione. I large language model come tecnologia sociale*, developed at the University of Udine, within the Degree Programme in Philosophy and Digital Transformation.

The present version does not reproduce the research in full. It offers a synthetic and reorganized English formulation of its central argument, intended for academic circulation and discussion.

1. Introduction

Large language models have become one of the most visible and socially consequential expressions of contemporary artificial intelligence. In only a few years, conversational systems such as ChatGPT, Claude, Gemini, Grok, Llama, Mistral and DeepSeek have entered practices of writing, learning, translation, programming, research and decision-making. Their diffusion is not merely the result of technical performance. It depends on a much older and more familiar medium: natural language.

The philosophical difficulty begins precisely here. These systems address users in a form that resembles ordinary dialogue. They answer questions, reformulate arguments, simulate uncertainty, provide explanations and adapt to the tone of the interlocutor. Yet they do so without consciousness, intention or human-like access to meaning. The result is a technological phenomenon that appears simple at the level of use but is conceptually disruptive at the level of interpretation.

The questions raised by LLMs are therefore not entirely new. They reactivate classical problems concerning the nature of linguistic understanding, the relation between language and thought, and the distinction between competence and performance. What is new is the scale and social immediacy of the problem. A system that generates persuasive language without being a responsible speaker forces us to reconsider what exactly is being produced when language is produced.

The central claim advanced here is that LLMs should not be interpreted primarily as artificial minds. Nor should they be reduced to mere frauds, tricks or linguistic illusions. They are more accurately understood as social technologies of communication: systems that operate by producing linguistic instructions capable of activating the imagination of human interlocutors.

This approach makes it possible to avoid two symmetrical errors. The first is uncritical enthusiasm, which mistakes fluency for understanding and conversational smoothness for intelligence. The second is radical dismissal, which concludes that because LLMs do not understand as humans do, their linguistic performance has no philosophical significance. The more interesting question is not whether machines think as humans think, but what happens to human linguistic practices when machines become capable of producing language at scale.

The argument proceeds in seven steps. First, it clarifies the technical and conceptual status of LLMs as systems of linguistic performance. Second, it situates the debate between innatist and functionalist approaches to language. Third, it introduces Daniel Dor's account of language as the instruction of imagination. Fourth, it interprets LLMs as a technological implementation of this function. Fifth, it examines simulation and deception. Sixth, it discusses authorship and epistemic responsibility. Finally, it proposes a relational understanding of intelligence suited to the age of conversational AI.

2. Large Language Models and Linguistic Performance

From a technical standpoint, an LLM is a computational system based on neural networks and trained on very large textual corpora to generate coherent linguistic sequences. Its apparent generative capacity is rooted in prediction: given a sequence of tokens, the model calculates statistically plausible continuations and produces text step by step. The result can resemble ordinary discourse, but the underlying operation remains a form of probabilistic modelling.

The term large refers both to the scale of the model's parameters and to the magnitude of the data on which it has been trained. Contemporary LLMs are then refined through procedures such as supervised fine-tuning and reinforcement learning from human feedback, which align their outputs with instructions, social expectations, safety constraints and conversational norms (Christiano et al., 2017).

The decisive technological shift behind contemporary LLMs is the Transformer architecture introduced by Vaswani et al. (2017). Through self-attention, Transformer-based systems can process long-range relations within a text and dynamically assign relevance to different elements of the linguistic context. This architecture made it possible to move beyond earlier systems based on fixed rules or limited statistical sequences and to generate outputs that are more coherent, flexible and context-sensitive.

However, technical success should not be confused with human understanding. LLMs do not possess embodied experience, communicative intention or first-person awareness. They do not know what they are saying in the way a human speaker knows, or at least experiences, what he or she is saying. They generate language by modelling statistical and structural regularities in human-produced texts.

This distinction is crucial. If the question is whether LLMs understand in the human sense, the answer should remain negative. If the question is whether they can produce outputs that generate understanding in human users, the answer is more complex. LLMs can summarize difficult texts, translate between languages, propose analogies, identify argumentative structures, generate hypotheses and support reasoning processes. Their effectiveness lies not in subjective comprehension, but in the effects produced by their outputs within a human interpretive environment.

Chomsky's classical distinction between competence and performance can be usefully reformulated in this context (Chomsky, 1965). In human language, competence refers to the implicit knowledge that allows a speaker to produce and understand sentences; performance refers to the actual use of language in concrete situations. In LLMs, competence is not an innate grammar but a statistical and structural mapping of linguistic regularities acquired during training. Performance, by contrast, emerges in the interaction with the user, where the model adapts its output to the prompt, the context and the feedback it receives.

LLMs therefore force us to examine linguistic performance without presupposing human-like understanding. They show that communicative effectiveness and subjective comprehension are not identical. A machine may fail to understand in the human sense and nevertheless produce language that is socially, cognitively and practically effective.

3. Innatism, Functionalism and the Status of Artificial Language

The debate on LLMs can be situated within a broader philosophical conflict between innatist and functionalist conceptions of language. The innatist tradition, associated above all with Noam Chomsky, treats language as a specifically human biological faculty grounded in an internal generative structure. From this perspective, LLMs may appear as powerful but superficial systems: they manipulate linguistic forms without accessing meaning, truth, explanation or genuine understanding.

This critique remains indispensable. It resists the tendency to anthropomorphize AI systems and warns against the confusion of fluency with intelligence. The fact that a model can generate a persuasive answer does not mean that it understands the answer, believes it or can assume responsibility for it. The machine has no point of view, no lived experience and no intentional relation to the world.

At the same time, the innatist critique risks evaluating LLMs according to an inappropriate standard: their resemblance to human cognition. A printing press does not speak, a telescope does not see and a calculator does not understand mathematics. Yet each of these technologies transforms human practices by extending, reorganizing or externalizing certain functions. LLMs may need to be approached in the same way: not as artificial persons, but as linguistic technologies that alter the conditions of communication.

Functionalist approaches focus less on internal consciousness and more on what systems can do. If a model can summarize, translate, classify, explain, draft, simulate dialogue, write code or support research, then its practical role cannot be dismissed merely because it lacks inner experience. From this standpoint, the relevant question is not whether the model possesses human understanding, but whether it performs useful operations within a given system.

A purely functionalist account, however, also has limits. It may underestimate the ethical and epistemic risks created by systems that act linguistically without understanding. A tool that produces language can influence beliefs, decisions, emotions and social relations. If its outputs are mistaken for the speech of a responsible subject, the result may be misplaced trust, epistemic dependence, manipulation or intellectual laziness.

A more balanced interpretation is therefore needed. LLMs should be assessed neither as failed human minds nor as neutral instruments. They are technological mediators of linguistic interaction. Their value depends on reliability, transparency, controllability and the quality of the human practices they enable.

4. Language as the Instruction of Imagination

Daniel Dor's theory of language offers a productive framework for understanding this problem. Dor (2015) describes language not primarily as an internal cognitive faculty, but as a socially constructed communication technology. Its function is to instruct the imagination of interlocutors. Speakers provide linguistic cues that allow others to construct experiences, images, relations and meanings that are not directly present.

This view shifts attention from the mind of the speaker to the effect of language on the interlocutor. Communication does not require the direct transfer of mental contents from one mind to another. It requires a socially shared system of instructions through which one subject guides the imaginative activity of another. Meaning emerges from this guided reconstruction, not simply from the private intention of the speaker or from the material text alone.

A simple example clarifies the point. When a speaker says, 'Imagine a red house at the end of a narrow street', the words do not contain the house. They invite the listener to construct a scene. The meaning emerges through imaginative cooperation. Language works not by transporting ready-made meanings, but by coordinating the imaginative activity of different individuals.

This approach is especially useful in the case of LLMs. A model that explains a scientific concept through a metaphor, drafts a clause, writes a poem or summarizes a philosophical argument does not possess the corresponding experience or understanding. Yet its output can activate images, relations, concepts and inferences in the human reader. The meaning does not reside inside the machine. It emerges in the interaction between the linguistic output and the human interpreter.

In this sense, LLMs can be understood as technologies for the artificial production of instructions for imagination. They do not think as humans think, but they can help humans think. They do not understand in the human sense, but they can produce linguistic structures that support understanding in humans. They are not autonomous minds; they are instruments that reorganize the social distribution of linguistic production.

This also explains why LLMs are both powerful and dangerous. They operate at the level of language, and language is not merely a tool for describing reality. It shapes attention, constructs scenarios, frames problems, generates expectations and coordinates social action. When LLMs produce language at scale, they do not simply automate writing. They amplify the social technology through which human beings imagine, reason and act together.

5. LLMs as Social Technology

If language itself can be understood as a social technology, LLMs may be interpreted as one of its most advanced artificial implementations. They do not create language from outside human culture. They are trained on human language, shaped by human texts, refined through human feedback and used within human institutions. Their outputs are therefore deeply dependent on social, historical and cultural conditions.

This point prevents two symmetrical errors. The first is to treat LLMs as autonomous intelligences detached from human society. The second is to treat them as mere machines with no social meaning. In reality, they are socio-technical systems. Their functioning depends on data, infrastructure, corporate design, user practices, legal norms, economic incentives and cultural expectations (Bommasani et al., 2021).

The conversational interface is not a secondary feature of these systems; it is the condition of their social success. Users do not interact with an LLM primarily by writing code or configuring formal rules. They speak to it. They ask questions, give instructions, express preferences, correct errors, request reformulations and receive answers in natural language. The resulting experience resembles dialogue, even though the system is not a dialogical subject.

This interface creates a new mode of human-machine interaction. Earlier digital tools required users to adapt themselves to the logic of the machine: commands, menus, forms or programming languages. LLMs, by contrast, adapt the machine to the ordinary form of human linguistic interaction. The user does not need to know how the system is programmed. The user needs to know how to ask, how to frame a problem, how to evaluate the answer and how to revise the output.

The relevant competence therefore changes. It is no longer only the ability to produce a text from scratch, but the ability to guide, evaluate and transform machine-generated text. The human role moves from direct production to critical orchestration. The question is no longer simply whether one can write

a given text, but whether one can formulate the right request, detect the limits of the answer, verify its claims and assume responsibility for the final result.

For this reason, LLMs should be evaluated less by their supposed resemblance to human minds than by the quality of the socio-technical relations they produce. Do they increase understanding or reduce it? Do they support autonomy or dependency? Do they strengthen critical judgment or replace it with passive acceptance? Do they democratize access to knowledge or amplify existing inequalities? These are not merely technical questions. They are philosophical, ethical and political questions.

6. Simulation and the Ethics of Deception

The most delicate ethical issue raised by LLMs is deception. Because conversational AI systems imitate dialogue, they can lead users to attribute intentionality, understanding, authority or empathy to a system that possesses none of these in the human sense. This risk recalls the Eliza effect first associated with Joseph Weizenbaum's early chatbot (Weizenbaum, 1976). Even when interacting with a relatively simple pattern-matching program, users tended to perceive understanding and psychological depth.

Contemporary LLMs intensify this problem dramatically. Their fluency, adaptability and apparent contextual awareness make the illusion stronger. The machine does not merely answer; it seems to listen, reason, remember, explain and care. Bender and Gebru (2021) famously warned against the danger of treating language models as if their fluency entailed understanding or epistemic reliability.

This creates the risk of pure deception. Pure deception occurs when the user believes, implicitly or explicitly, that the system is a subject capable of understanding and epistemic responsibility. The user may treat the model as an authority, a companion, a teacher, a therapist or a co-author without recognizing the limits of the system. In such cases, simulation becomes misleading because the user does not understand the nature of the interaction.

Not all forms of simulation, however, are ethically equivalent. A distinction can be drawn between pure deception and conscious deception. Conscious deception occurs when the user knowingly accepts the fiction of dialogue in order to obtain a practical result. The user does not believe that the machine is human-like; rather, the user treats the simulation as an interface, a useful convention or a pragmatic game.

The distinction is important because human societies constantly rely on symbolic mediations, conventions, roles and fictions. A theatre spectator does not believe that the actor is truly Hamlet, but accepts the fiction in order to participate in the experience. A reader does not believe that a novel is literally happening, but allows the text to instruct imagination. In a similar way, the user of an LLM may accept the fiction of dialogue without being deceived by it.

The ethical problem is therefore not simulation as such, but ungoverned simulation that produces misplaced trust. A mature relationship with LLMs should not aim to eliminate the conversational fiction entirely. That would be unrealistic, since the interface itself is built on the imitation of dialogue. The more realistic task is to transform pure deception into conscious deception.

This transformation requires at least three conditions: transparency regarding the nature of the system; verification of outputs against reliable sources, especially in high-stakes contexts; and critical education in the use of AI-generated language. Under these conditions, conscious deception can

become a productive mode of interaction. The user knows that the dialogue is simulated, but uses the simulation to think, write, test ideas, explore alternatives and reorganize knowledge. The fiction becomes not a trap, but a tool.

7. Authorship and Epistemic Responsibility

The rise of LLMs also forces a reconsideration of authorship. When a text is generated through interaction between a user and a model, who is the author? The traditional model of authorship presupposes an identifiable subject who expresses a point of view, assumes responsibility and can be held accountable. LLM-generated text destabilizes this model.

The machine produces language without having a point of view. The user may guide the output without writing every sentence. The final text may therefore result from a distributed process involving training data, model architecture, interface design, corporate policies, user prompts and human revision. Authorship becomes less a matter of isolated originality and more a matter of responsibility for selection, validation, interpretation and use.

This does not mean that the human author disappears. On the contrary, the human role becomes more important precisely because the machine cannot assume responsibility. The user must decide what to ask, what to accept, what to reject, what to verify, what to rewrite and what to publish. The ethical centre of authorship moves from production alone to critical control.

The same issue changes the debate on plagiarism. The central problem is not always the copying of a specific human author. In many cases, the more relevant problem is the false attribution of competence: a person presents as his or her own a text, argument or analysis that he or she has not critically produced, understood, verified or re-elaborated.

In the age of LLMs, intellectual originality cannot mean the absence of external influence, since all thought develops within networks of language, culture and previous texts. Nor can it simply mean the manual production of every sentence. It should rather be understood as transparent and responsible re-elaboration: the capacity to appropriate, transform, verify and integrate linguistic material within a personally assumed argumentative project.

This point is particularly important in academic contexts. The use of LLMs does not automatically eliminate authorship, but it modifies its conditions. A student, researcher or writer who uses an LLM responsibly remains accountable for the final text. The key questions become: has the author understood the argument? Have the sources been checked? Has the contribution of the tool been critically evaluated? Is the final text intellectually owned by the person who signs it? Authorship in the age of LLMs is therefore not abolished. It becomes more demanding.

8. Toward a Relational Definition of Intelligence

LLMs invite a reconsideration of the concept of intelligence itself. If intelligence is defined only as an intrinsic property of a conscious individual, then LLMs are not intelligent in the human sense. They do not possess awareness, intentionality, embodiment or subjective experience. They do not know the world as human beings know it.

If intelligence is understood relationally, however, as the capacity to produce adaptive and meaningful effects within a system of interaction, then LLMs may participate in intelligent processes without being intelligent subjects. An LLM may not understand, but the interaction between a human user and an LLM can generate understanding. The intelligence is not located simply inside the model. It emerges from the relation between model, user, context, language and verification practices.

This relational view avoids both anthropomorphism and reductionism. It avoids anthropomorphism because it does not attribute consciousness to machines. It avoids reductionism because it recognizes that LLMs produce real cognitive and social effects. The model is not a mind, but it can become part of an intelligent process.

The perspective also helps preserve what remains distinctively human. The human being is no longer unique merely because he or she can produce text. Machines can now produce text at extraordinary speed and scale. What remains human is responsibility: the capacity to evaluate meaning, situate language within ethical and social contexts and answer for the consequences of communication.

The question is therefore not whether machines can become human. The more urgent question is whether humans can remain responsible within the linguistic environment they have created.

9. Conclusion

Large language models do not simply add a new tool to the existing landscape of digital technologies. They transform our relationship with language itself. By producing fluent and adaptive discourse without human understanding, they separate linguistic performance from consciousness and force us to rethink meaning, authorship, deception and intelligence.

The most productive interpretation is not to see LLMs as artificial minds, nor as mere frauds, but as social technologies for the instruction of imagination. Their outputs work when they activate human interpretive processes. Their risks arise when this activation is mistaken for the presence of a conscious speaker or an authoritative knower.

The ethical task is therefore not to reject LLMs, but to govern their use. Conscious deception, when transparent and critically managed, can become a mature mode of interaction with conversational AI. Pure deception, by contrast, must be reduced through design, regulation and education.

In the end, LLMs reveal something important not only about artificial intelligence, but about language. Language has always been a technology for coordinating imaginations. Artificial intelligence does not abolish this function; it amplifies it. The challenge is to ensure that this amplification strengthens human knowledge rather than weakening critical judgment.

What is at stake is not whether machines can think as humans think. What is at stake is whether human beings can continue to think critically, responsibly and imaginatively in dialogue with machines that speak.

Suggested Citation

Zorgno, F. (2026). *Language, Simulation, Imagination: Large Language Models as Social Technology*. Working paper based on the research document *Linguaggio, simulazione, immaginazione. I large language model come tecnologia sociale*, University of Udine.

References

- Bender, E. M., & Gebru, T. (2021). On the dangers of stochastic parrots: Can language models be too big? Proceedings of FAccT '21.
- Bommasani, R., et al. (2021). On the Opportunities and Risks of Foundation Models. Stanford Center for Research on Foundation Models.
- Chomsky, N. (1957). *Syntactic Structures*. The Hague: Mouton.
- Chomsky, N. (1965). *Aspects of the Theory of Syntax*. Cambridge, MA: MIT Press.
- Christiano, P., Leike, J., Brown, T. B., Martic, M., Legg, S., & Amodei, D. (2017). Deep reinforcement learning from human preferences. *Advances in Neural Information Processing Systems*.
- Dor, D. (2015). *The Instruction of Imagination: Language as a Social Communication Technology*. Oxford: Oxford University Press.
- Floridi, L. (2023). *The Ethics of Artificial Intelligence: Principles, Challenges, and Opportunities*. Oxford: Oxford University Press.
- Floridi, L. (2023). AI as agency without intelligence: On ChatGPT, large language models, and other generative models. *Philosophy & Technology*.
- Manzotti, R., & Rossi, S. (2023). *Io & IA. Mente, cervello & GPT*. Soveria Mannelli: Rubbettino.
- Natale, S. (2021). *Deceitful Media: Artificial Intelligence and Social Life after the Turing Test*. Oxford: Oxford University Press.
- Raschka, S. (2025). *Sviluppare Large Language Model. Costruire da zero LLM su misura*. Milano: Apogeo.
- Vaswani, A., Shazeer, N., Parmar, N., et al. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*.
- Weizenbaum, J. (1976). *Computer Power and Human Reason*. San Francisco: W. H. Freeman.