

## ORIGINAL ARTICLE

# Domesticating robots, domesticating language: The hylomorphic paradox of “openness” in large language models

Raffaele Andrea Buono 

Department of Philosophy & Cultural Heritage, Ca' Foscari University of Venice, Venice, Italy

## Correspondence

Raffaele Andrea Buono, Department of Philosophy & Cultural Heritage, Ca' Foscari University of Venice, Malcanton Marcorà, Dorsoduro 3484/D, Venice, Italy.  
Email: [raffaeleandrea.buono@unive.it](mailto:raffaeleandrea.buono@unive.it)

## Funding information

European Research Council, Grant/Award Number: GA 101088645

## Abstract

This article examines how openness is interpreted in human–LLM interaction through an ethnographic study of a robotics experiment. Focusing on an episode in which a robot produced an unexpected utterance, I analyze how engineers classify the output as inconsequential noise. I argue that such dismissal reflects a hylomorphic understanding of language, where form is externally imposed on linguistic matter, by latching onto ideas of alignment and usefulness. I configure such hylomorphic movement as undergirding the semiotic ideology at play in robotics engineering, one which constrains the dynamism and potential of situated interaction, transforming openness into a domesticated variability, and limiting the emergence of forms of coordination and meaning.

## KEYWORDS

alignment, human–machine interaction, hylomorphism, openness, semiotic ideology

## Estratto

Attraverso uno studio etnografico di un esperimento di robotica, questo articolo affronta la questione di come il concetto di ‘apertura’ venga interpretato nell'interazione tra esseri umani e LLM. Al centro dell'analisi si trova un episodio in cui un robot produce un enunciato inaspettato, prontamente classificato dagli ingegneri come rumore privo di conseguenze. Si discute come questo gesto di rigetto rivela una concezione ilomorfica del linguaggio, in cui la forma viene

This is an open access article under the terms of the [Creative Commons Attribution](https://creativecommons.org/licenses/by/4.0/) License, which permits use, distribution and reproduction in any medium, provided the original work is properly cited.

© 2026 The Author(s). *Journal of Linguistic Anthropology* published by Wiley Periodicals LLC on behalf of American Anthropological Association.

imposta dall'esterno sulla materia linguistica in nome dell'utilità e dell'allineamento. Si analizza dunque come questo movimento ilomorfo costituisce il sostrato di un'ideologia semiotica che informa le pratiche dei miei interlocutori. È un'ideologia che comprime il dinamismo dell'interazione situata, trasformando l'apertura in variabilità domesticate, e precludendo l'emergenza di nuove forme di coordinamento e significato.

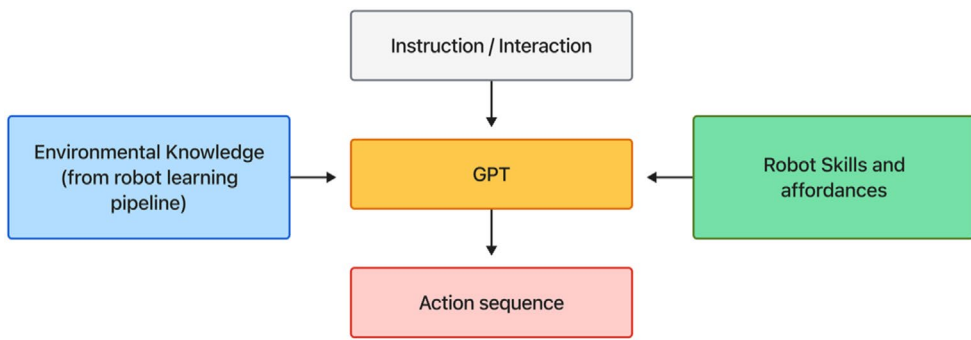
#### PAROLE CHIAVE

ideologia semiotica, apertura, interazione uomo-macchina, ilomorfismo, allineamento

While attending, in April 2023, the first plenary meeting of the robotics laboratory I conducted ethnographic fieldwork in, Mikitani-sensei,<sup>1</sup> its director, announced we were witnessing an epochal shift. By that point, everyone in the lab was accustomed to using large language models (LLMs) to support them in the development, testing, and debugging of machine learning (ML) architectures and robotic systems. The laboratory's central research focus was on robots capable of operating and communicating in real-life environments alongside humans, navigating instructions, responding to instructions in unstructured contexts, and sustaining often open-ended, context-sensitive interactions. Until then, however, LLMs were merely *used* as convenient shortcuts to speed up painstakingly long processes of software engineering, with “tasks that would require hours of going through books and testing now tak[ing] minutes,” as Akahori-san, a PhD student I often shadowed, mentioned. At the level of the actual operations of the robots, very little had changed.

Mikitani-sensei's enthusiasm, however, pointed in a different direction. Mentioning both his own research on the main innovation behind LLMs (i.e., the transformer architecture) and the then-recent release of the latest iteration of the generative pre-trained transformer (i.e., GPT-4), he described how LLMs could be not merely tools for them to use. Rather, GPT-4's newly released API<sup>2</sup> made it possible to connect the robot architectures they had devised directly to the capacities of LLMs: routing the robot's sensor inputs and acquired knowledge, its kinematic affordances, and devising clearly communicable and executable plans for action, displaying capacities for planning, researching, and execution (Figure 1) of the actions of robots themselves. What were, only a week prior, dreams and hopes, or rudimentary and unfeasible proofs of concept, suddenly became achievable goals of a new generation of robots.

According to Mikitani-sensei's speech, as well as many ensuing conversations in the weekly meetings in the laboratory, transformers allowed robotic agents a new degree of *openness* to external stimuli as well as an increased expressive *richness* regarding how to engage with such stimuli. This shift started to crystallize precisely in that first plenary meeting, as Mikitani-sensei concluded that “everything is in sync”—with the “everything” signaling human and machinic agents, now inextricably linked through the medium of language, in a bottom-up dance of allegedly mutual, *open* adjustments he had dubbed *co-creative learning* (共創の学習—*kyousouteki gakushu*). LLMs could now be an integral part of the so-called “reasoning,” or “inferential” architectures of robots, allowing them to operate with and through language in a more malleable, flexible, and interactive way. According to my interlocutors, and in summary, GPT's manipulation of and control over language was touted as resonating with, and perhaps resembling, human-like cognitive capacities, pushing this rhetoric of openness, infinite possibilities and human-robot co-creative production further within the walls of the lab (see also: Chalmers, 2024).



**FIGURE 1** A diagram representing GPT as a “central processing unit” for the operations of a robot. Inspired by drawings by Akahori-san.

This paper then examines a moment in which this rhetoric of openness encountered a limit: an instance in which a robot's communicative openness was classified by its designers as an error to be discarded. I argue that this classificatory move reveals a semiotic ideology grounded in a hylomorphic understanding of language, in which linguistic matter is evaluated according to imposed forms which come to define success and correctness. Rather than treating language and its openness as an emergent and relational process, my interlocutors approached interactive openness through such a hylomorphic ideology, one which separates form from matter, and locates meaning within the former, and outside of the situated interaction itself. In this sense, this paper argues that hylomorphism, as an ontology (see: Ingold, 2010; Pinho, 2023), undergirds a semiotic ideology that structures how language and linguistic behavior are interpreted, judged, controlled, and ultimately assigned meaning to.

This paper thus examines the ethnographic gap between promises of openness generated by LLMs and their actual implementation and understanding in robot-making. In so doing, I show how claims of linguistic flexibility and expressivity paradoxically rest on a process of enclosure: I suggest that this process is sustained by a semiotic ideology that treats language as inert matter to be shaped by external, form-giving evaluative criteria. Throughout this paper, I will show how such shaping conflates the imposition of structure with the emergence of meaning, treating the latter as a product of the former. In zooming in on how an instance within an experiment came to be disregarded by my interlocutors, I analyze the consequences emerging when interactional openness is simultaneously desired while being curtailed and framed within a hylomorphic understanding of language. What I seek to show through this ethnographic case and the kinds of ideologies my interlocutors seem to mobilize to build robots, as well as to make sense of their successes and failures, is how openness (and the abductive modes of inference making it possible; see: Magnani, 2016, 2021) comes to be rethought of through instrumentalist and teleological logics undergirding hylomorphism. Hylomorphism in so doing makes it possible to “domesticate” language—its contingent, metastable and abductive character—by establishing externalized and externalizing modes of control and evaluation.

I want to make a key clarification regarding my epistemic stance within this paper. Throughout, I will describe some “excesses” of openness operated by GPT-equipped robots as moments showcasing, *perhaps*, glimpses of abductive and contextual capacities which resist the hylomorphic framing imposed by my interlocutors. I want, however, to be precise about what this claim may or may not mean. I remain fully agnostic (ontologically and epistemically) regarding whether LLMs genuinely engage in abductive inference, and I make no strong claim in either direction. Rather, my interest is in exploring my interlocutors'

claims that robots are open and somewhat unpredictable, and not mere stochastic parrots (Bender et al., 2021). In so doing, I seek to showcase what such openness may entail (and indeed require), and how such consequences and necessities come to be curtailed when they, *perhaps*, manifest themselves. My point is thus not that of resolving this “perhaps” in either direction, that is, either confirming that robots engage in abduction and all its world-changing dynamics, or reducing them to mere processors of statistical co-occurrence. Rather, I dwell in this perhaps (and my interlocutors’ ontological stance) to show how my interlocutors rethink, reframe, and domesticate openness and the “co-creativity” described by Mikitani-sensei under a different, hylomorphic, lens: one where openness is pursued without any of the “baggage” that makes such openness open (to disruptions, reckonings, contraventions, and new beginnings, among many other phenomena—see: Kockelman, 2022; Kockelman & Bernstein, 2013).

With this in mind, the paper will move as follows. First, I provide a brief exegesis of the functioning of LLMs, while connecting these operations to the hylomorphic schema. Second, I turn to an ethnographic vignette of an experiment carried out by my interlocutors, to highlight how hylomorphism provides the ideological means to externalize language evaluation, with human interpreters tasked with separating desired, meaningful information and outcomes from unwanted noise. Third, I examine how the rhetoric of alignment is what justifies such hylomorphic regimentation of the (alleged) potential of LLMs, enclosing it within bounded epistemic frames. I then conclude by highlighting how this rhetorical and operational strategy reveals yet another layer of the hylomorphic schema, which extends beyond language and interpretation, and toward the very character of what constitutes *good*, *correct*, and *useful* communication and interaction between human and robot. In following this line of argumentation, this paper asks: what are the stakes of treating language hylomorphically? How do such semiotic ideologies and practices shape what counts as meaningful communication?

## HYLOMORPHISM, MODELS OF LANGUAGE AND LANGUAGE MODELS

Throughout this section, I aim to position hylomorphism within the Aristotelian metaphysical tradition out of which it stems, and to connect it to particular ways of conceptualizing language and building technical language models. My central claim is that hylomorphism underwrites a semiotic ideology: a set of assumptions about what counts as a (meaningful) sign, and how such signs should function (Keane, 2018b).

Summarily put, the hylomorphic doctrine contends that every physical object is a composite of matter (*hyle*) and form (*morphe*), with form unifying the messiness of matter, imbuing in it whatever “*whatness*” one intends to imbue into a material “*thisness*.” Hylomorphism implies a reified directionality, a gradient: form engenders matter, organizing it toward specific ends. For instance, wood boards, as an instance of “unformatted matter” not too dissimilar to word jumbles or a series of isolated tokens, *are rendered* a table when a specific form is bestowed upon it—not only through physical molding, but also through the (cognitive) attribution of a recognized unity, that is, a “table-ness.”<sup>3</sup> I suggest that this same form-matter schema also structures influential models of language.

Just as hylomorphism concerns itself with how “the organizing principle of syntax separates genuine linguistic objects such as propositions from word jumbles” (Britton, 2012: 147), Saussurian *semiology* constructs a typified and idealized distinction between *langue* and *parole*.<sup>4</sup> The latent structure of *langue* organizes and *molds parole* into intelligibility. In turn, *parole* materializes conceptual positional differentials which exist as “opposing, [...] negative entities” (de Saussure, 1959 [1916]: 119) which define the “whatness” of something

in virtue of “what the others are not” (ibid.: 117; see also: Bredin, 1984). This structuralist model presupposes that linguistic matter becomes meaningful only insofar as it is shaped by a pre-existing formal system of forms.

Such a structuralist understanding of language as two interrelated systems of relations ultimately coalescing into stable linguistic objects, rendering un-form-matted strings into in-form-mative text, has been certainly problematized. In this sense, Roman Jakobson “managed to use Saussure’s categories to overcome Saussure” (Kockelman, 2017: 40), highlighting how meaning can emerge through absence, relationality, and interaction (Jakobson, 1960, 1971; Jakobson & Lotz, 1949; see also: Diehl, 2008; Nielsen, 2015; Ohnuki-Tierney, 1993). Along similar lines, Michael Silverstein’s metapragmatics (1993), and particularly the category of shifters (1976), showcases how linguistic forms emerge through interstitial, contextual, and indexical encounters which *modulate* them. These critiques foreground language as emergent and relational, rather than as matter being molded by transcendent form (see, on this: Di Paolo et al., 2018).

I use hylomorphism then as a diagnostic tool to understand how language is conceptualized, operationalized, and evaluated in contemporary language technologies. Again, with the shorthand hylomorphism I am interested in highlighting the recurrent separation between linguistic form and matter. With the former I identify not simply the organizing principles and constraints of syntax and grammar, but also the specific deontic norms and values dictating how language and communication are to be made sense of; with the latter I refer to all material components of language emerging in interaction, as language is either processed or deployed, ranging across scales and encompassing processes through which a sign (e.g., grapheme, morpheme, token, sentence) is assayed and assigned meaning to.

Hylomorphism, I argue, treats matter as *inert* (Parfitt, 2006), as a substrate upon which form *qua* principles are *applied*, as I will show, from an outside, rather than a dynamic system out of which principles, norms and shared worlds emerge in conflictual and generative ways. This understanding echoes long-standing critiques of structuralist assumptions that treat “linguistic form [...] [as] the sole and *direct* bearer of meaning” (Ohnuki-Tierney, 1994: 59; see also: Bakhtin, 1986). My interest here is not to rehearse these debates, or to simply tack the hylomorphic label to this direct arrow model (Ohnuki-Tierney, 1994): rather, I want to show how LLMs materially operationalize this separation, and how such separation is maintained in place by my human interlocutors. Drawing on Gilbert Simondon’s critique of hylomorphism (2017 [1958]; 2020 [1964]), I treat hylomorphism as a practical schema and ideology that conditions how agency, normativity, and responsibility come to be stabilized and externalized (see also: Coupaye, 2020; Enfield & Kockelman, 2017; Harding, 2021).

I draw then on Paul Kockelman’s notion of sieves (2013) to make sense of this hylomorphic imposition, as well as the scales at which this process is seemingly resisted. Sieves are understood as selective infrastructures organizing linguistic material, filtering it through, oftentimes, temporally prior, form-giving assumptions about relevance, similarity, and appropriateness. They are historically sedimented, revisable, and unevenly durable (Maurer, 2013), and they function by embedding assumptions and presumptions about what counts as a meaningful linguistic continuation. I understand sieves as the techniques and material apparatus which allow matter and form to engage with one another: it is through sieves that matter is oftentimes hylomorphically imbued with form through certain modes of *transformativity* (Kockelman, 2013); more crucially, within Kockelman’s schematics, other modes of transformativity actively resist this hylomorphic gradient, with material encounters instead modulating form and its form-giving sieves. In this sense, I deploy Kockelman’s sieves in conjunction with the idea of hylomorphism to conceptualize precisely which kinds of sieving operate hylomorphically, and which ones



do not, and why it is the interplay of these two regimes that ultimately confers semiotic interaction its “open” character. Lastly, in understanding sieves in such loose terms (i.e. as any technical process through which matter is assayed), sieving can only be understood as a complex process operating at many different timescales and across different kinds of matter. Here, I will be focusing on some, very specific sieving operations of interest within the economy of this essay—while merely hinting at some, and leaving aside some other entirely.<sup>5</sup>

From this perspective then, a first sieving operation of interest can be highlighted at the “inference time” of an LLM: the system takes a prompt as input and produces an output *qua* response by computing probabilistic continuations constrained by parameters learnt during training. Here, a first sieving occurs through the *transformer* architecture (Vaswani et al., 2017) and its attention mechanism, which weights tokens in relation to one another, selectively amplifying some contextual relations, while suppressing others. These processes, to borrow from Kockelman's sieving schematics, produce “relatively deductive” (2013: 46) transformations: for instance, if asking ChatGPT to complete the sentence “The doctor looked at the chart, frowned, and looked...”, the system will produce something along the lines of “concerned,” “worried,” or endless, plausible, variations. Here, “frowned” functions as an index to the individual (doctor): attending to this sign (*frowning* as signifying uncomfortable-ness; invariant kind-index relation) via the attention mechanism produces a shift in the assumed kind of the object (a *concerned* doctor; individual-kind relation). What produces an *appropriate* response is, I argue, a *hylomorphic imposition* of a form-giving principle: transformation attaches kinds based on pre-established kind-index relations extrapolated from and evaluated against statistical priors which sieve and shape meaning.

Such kind-index relations emerge through a second, temporally prior transformativity which happens before inference time, and which ‘locks into place’ parameters to be used in the first transformativity just described. That is, *during pre-training*, backpropagation adjusts parameters governing inference time operations: through these prior operations, predicted outputs are compared with observed ones, and thus retroactively reshaping the assumptions that will guide future inferences. Such (hylomorphic) process, which essentially “*gives form to substance for the sake of function*” (Kockelman, 2025: 50; emphasis mine), produces an inductive transformation (Kockelman, 2013: 46), through which an index (e.g., “frowned”) encountered in pre-training, and referring to a specific kind of individual, gets to reshape and transform the set of assumptions behind that index and its relation to a kind, in so doing licensing new interpretations.

Taken together, these processes (attention mechanism and backpropagation) stabilize particular relations between learnt forms and encountered matter, privileging predictability and coherence over other forms of semiotic transformativity. Moreover, while these relations are learnt, once operationalized they function as form-giving constraints that shape how linguistic matter can meaningfully appear. In this sense, LLMs enact a hylomorphic logic because their training and inference procedures systematically privilege the form as an externalized, temporally prior formal regularity that regiments and conditions what counts as meaningful output.

Crucially, this sieving arrangement introduces a temporal asymmetry and externalized evaluation: present linguistic matter is assayed and evaluated through external parameters, sedimented through past usage, so that future outputs are constrained by inherited forms. Or: parameters, learnt *in* the past and *from* the past, sieve and give shape to future encounters, transacting words for words and bracketing worlds out of words. In this sense then, the remainder of this paper will engage precisely with the ways in which this externalized, enclosing mode of evaluation gets to be maintained in place even when the world becomes a fundamental part of the operations of an LLM.

## SIEVING NOISE AND MISSED SERENDIPITIES

It had been weeks of intense discussions around the transformations brought about by LLMs. In the eye of my interlocutors, LLMs provided an endless source of malleable opportunities to learn *with* and *from* algorithmic systems, in what Mikitani-sensei dubbed “co-creative learning.” Such stance posited language and linguistic production as part of a large(r), networked, interactional sphere, within which agents (human and algorithmic alike) strive to make themselves understandable and understood, and thus able to engage in cooperative action (Goodwin, 2017). Still according to my interlocutors, LLMs generated creative, generative solutions via language that humans could pick up on, thus implicating them in co-creative endeavors. Interaction and language, on this account, became co-constructed phenomena, open-ended precisely because both human and robotic agents were now grounded in what my interlocutors celebrated as a radical openness. Or, at least, this was how much of our conversations went as I sat on long meeting tables with my interlocutors. This section briefly leaves my interlocutors’ excitement aside, to focus instead on how such visions materialized in experiments. I hope in so doing to reveal how the openness and co-operation hailed by my interlocutors takes very specific, hylomorphic contours.

As these speculative imaginaries took hold of weekly lab debriefings, many students saw the release of GPT-4 API as an opportunity to propel their work toward new horizons. Akahori-san immediately recognized in LLMs an opportunity to address some of the issues he had encountered in his previous research. His work was concerned with developing robots able to capture semantic information about environments and objects, in turn equipping them with knowledge that they could leverage when asked by a human operator to carry out mundane tasks. A specific family of Bayesian algorithms and techniques<sup>6</sup> was developed to carry out the initial training of the system, with the robot learning probabilistic relationships between objects and spaces, in turn identifying rooms as rooms of a specific kind.

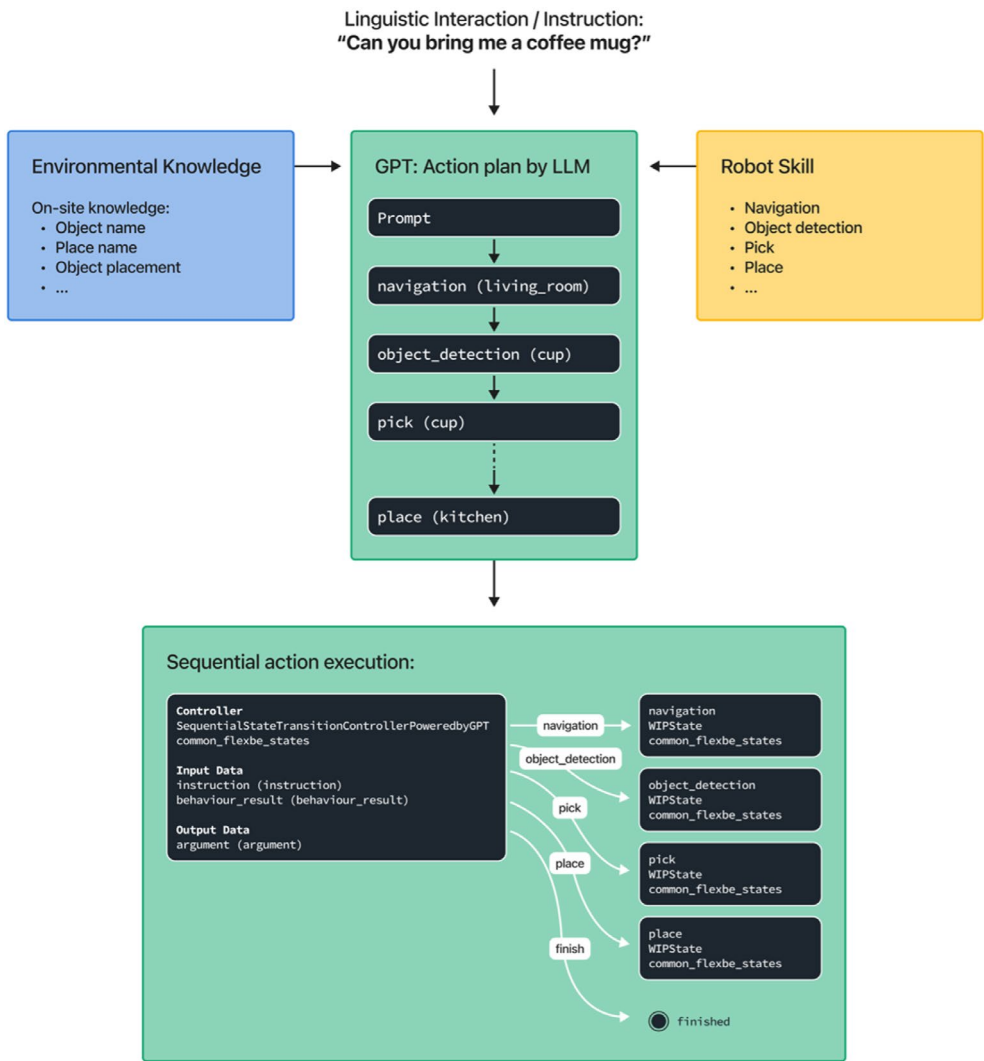
Concretely, this training phase resulted in a probabilistic, spatial-semantic representation of the environment. Rooms were modeled as discrete regions defined by distributions of co-occurring object types, while objects themselves were represented as classes associated with likelihoods of occurrence within particular spaces. Through observation, the robot learnt conditional relations such as the probability of encountering specific objects in specific locations, as well as coarse topological relations between spaces (e.g., adjacency and navigability). Crucially, this form of knowledge encoded where things are likely to be, not what they are for, how they should be handled, or why they might matter in a given context.

Just as I was getting to grips with his work, Akahori-san mentioned how he, and the laboratory at large, struggled to turn these computations into information that could be *useful* to the robot, and its human interlocutors. The lab’s previous approaches to making spatial semantic knowledge *operative* relied on painstaking human labor, through which engineers defined specific sets of relations between acquired knowledge, the operational affordances of the robot (e.g., how can it interact with a given object?), and the steps necessary to act on human prompts. And on top of this, human prompts had to be processed as to make sense within the statistical ruminations of the robot. In other words, the difficulty did not lie in acquiring spatial or semantic information as such, but in translating between incommensurate registers: a probabilistic model organized around object-space correlations, and human prompts organized around intentions, relevance, and tasks. Previous systems addressed this gap by explicitly scripting mappings between objects, affordances, and action sequences. Now, the use of API “calls,” GPT-4 allowed such normative and pragmatic coherence to be retrofitted to the statistical world of the robot whenever prompted by a human.

Thus, his proposal was to insert GPT-4 as a *dynamic interface* taking as input and connecting operations of the trained algorithmic model, human-generated prompts, and perceptual and kinetic affordances of the robot. This was made possible by the latent capacity

of LLMs to sieve the diverse range of expression these inputs took. GPT-4's capacities to organize, synthesize, and elaborate upon (*qua* sieve) this heterogenous set allowed it to output *task-completion plans* (Figure 2) whose dynamics were, according to my Akahori-san, “*emergent* rather than scripted.”

He quickly reworked a prototype, and I found myself observing his first experiment. A standard service robot resembling a moving cart was placed in one of the three experimental rooms scattered across the lab. The room was divided into areas suggestive of domestic spaces, each delimited by Styrofoam partitions. Each area was dressed as a typical and typified household room: one partition contained an inflatable bathtub and toiletries, another an old kitchen and various kitchen utensils, another had a desk and a bookshelf; in total, seven such zones were sieved, imposing on an otherwise unstructured space specific spatial (and semiotic) markers. The robot was placed at the center of the room, in a non-descript partition with nothing else besides the robot, a desk where researchers could place their laptops, and some chairs for us to sit. After some quick initialisation, testing and troubleshooting



**FIGURE 2** Operational schematics for Akahori-san's robot. Inspired by some of his drawings presented in lab meetings.



supported by Takahara-sensei, one of the assistant professors in the laboratory, and some undergraduate students, the robot was ready to turn its first request into a task.

As briefly explained above, the robot was by that point pre-trained via the Bayesian algorithmic models devised by Akahori-san: it had identified and classified each room with a name, and assigned to objects a probability of being in a room over another. This information was retrievable by the GPT-4 API, which could then link requests to the stored knowledge, and devise action plans to satisfy the prompt (e.g., going to the bathroom when asked to fetch toothpaste). Akahori-san, standing at a safe distance and clear from the path he thought the robot was going to take, looked at the machine, and simply asked: “Could you find some candies?”<sup>7</sup> The robot stood still for a short while, only to then comfortably speed its wheels in the direction of the kitchen. Once there, it got as close as possible to a basket with some Japanese candies inside, and let us know that it had located the item, and it was going to bring it to us.<sup>8</sup> It quickly returned victorious, and the room was suddenly filled with excitement.

We repeated the procedure many times, switching objects and locations, or changing the type of request. The robot got it wrong a few times, and that was fine—“there is so little testing we have done that I’m surprised it’s working so well *already*,” told me Akahori-san, hinting at how LLMs heavily slimmed down the laborious chains of operation everyone was so used to. Some other times, the robot operated in ways that *surprised* many of us: when asked to grab some “snacks,” it confidently moved toward a bowl with fruits. “It must want us to eat healthy,” observed another student, while Akahori-san reaffirmed that such *openness* toward deviations and misunderstandings was precisely what he was excited about, and what was radically novel about the embedding of LLMs’ capacities into robots. Crucially, however, and as I will return to shortly, this specific moment of surprise was one where the openness of the robot did not challenge the world within which the interaction was taking place: the robot still recognized the request as a task, oriented itself toward compliance, and acted within an intelligible horizon of usefulness. Openness here took the form of semantic elasticity, a flexible, surprising interpretation, but one still operating within an existing frame of action. This is then distinct from a different kind of openness I will soon be engaging with: one that would reconfigure what counts as an appropriate response in the first place, contest who has the authority to set the terms of interaction, and call into question what kind of situation is being enacted—what, in short and more broadly, comes to be seen as a world.

And just as the morale was at an all-time high, Akahori-san looked at the robot once again and asked: “Could you grab a coffee mug?” No one was really paying too much attention to the whole process by then, as it all seemed to run as smoothly as one could hope. And then, instead of the wheels spin, we heard the robot’s voice. As we all turned with renewed curiosity, the robot just bluntly rebutted: “Why don’t you pick it *yourself*?”

Upon hearing those words, everyone laughed as jokes around a possible robot apocalypse popped up. The hilarity dissipated quickly, and Akahori-san rebooted the system to carry out a few more tests. And as we kept testing and testing, the robot never objected again. As I chimed in to ask whether the event would alter his plans, Akahori-san did not seem to mind too much: “It hasn’t happened again... it must have been some *noise* which made the robot hallucinate.”

News of what happened traveled fast: someone brought it up on the lab’s communication platform, where more jokes were made; anytime I bumped into someone, they were quick to inform me of what had happened. Fascinatingly, the whole ordeal kept being framed, or rather *binned* (Enfield & Sidnell, 2017), as a funny *accident*, “just an *error* that *emerges* when you dump a robot in a messy, *noisy* world”—as Imad-sensei, another assistant professor, framed it days later. At the first sub-group meeting after the experiment, the “incident” got brought up as an anecdote to quickly laugh about at the beginning of the meeting, before diligently focusing on the times the robot got the task right, those it got it wrong, and those

it completed the job in surprising, not-completely-right-but-not-awfully-wrong ways showcasing the “openness” my interlocutors seemed so drawn to. In short, this confusing role-reversal had no place in my interlocutors’ endeavors, and since it had only happened once, it was safe for them to leave it outside of their thinking and writing. As I was told time and time again, it was just noise.

I am reminded here, through this framing of “just an error,” of Ruha Benjamin’s analysis of glitches in algorithmic systems (2019; see also: Broussard, 2023; Noble, 2018). Benjamin argues that glitches are usually presented by engineers as benign, fleeting interruption, mistakes which can, should and will be patched over. Against this framing, she contends that glitches rather signal how the system operates— “[n]ot an aberration but a form of evidence” (2019: 81)—, expressions of the system itself, as well as its emergent capacities when entangled within a larger sociotechnical ensemble (Mitrović & Voigt, 2025). Glitches thus become vantage points to glimpse into the foundations of the structures of worlding enacted by these systems.

I argue that the robot “glitching” here, against my interlocutors’ framing as an error, created a form of spurious openness, revealing perhaps something deeper *qua* different about the system, something my interlocutors felt in need to be enclosed through the rhetoric of the noisy error, so as to uphold to specific ends I will return to in the final section of this paper. In this sense, attempts to disentangle *what* exactly counted as noise, or *where* such noise originated from, yielded incommensurable, sometimes conflicting, answers. In some instances (e.g., the quote above from Imad-sensei), noise seemed to be an omnipresent feature of the world, a messy entanglement out of which *salient* bundles of affordances (Keane, 2018a; see also: Kockelman, 2006a, 2012) are sensed and made the epicenter of new chains of semiosis. And given its noisy, complexly multiple, disordered nature, the act of selection of affordances from the environment always leaves the door open for *errors*, as seemed to be the case this time. Conversely, in my discussion with Akahori-san, noise was framed as a feature of the computations themselves. Against the neatly laid out sequences emerging through extensive pre-training, the model gave in under the pressure of perhaps infinitesimally small probabilistic deviations, which exponentially torpedoed to produce an *hallucination*<sup>9</sup> which, almost tautologically, framed the entire response itself as disordered noise to be discarded.

As I reprised the incident with Takahara-sensei, we wondered together what was disordered about the robot’s reply, if noise is a consequence and/or manifestation of it. We agreed that the robot produced an utterance which was syntactically and semantically ordered (Morris, 1971 [1938]; see also: Turbanti, 2023), and which had its own pragmatic dimension emerging from implicatures (Sperber & Wilson, 1990) which we perhaps had not paid attention to (e.g., were we standing closer to the mug than the robot was?), and which instigated kindred interpretants. Or, summarily, the answer “made sense,” or *could* be made sense of. To frame it through some of the questions I asked him: if there was (semantic, syntactic, pragmatic) logic to the utterance of the robot, and logic is order(-ed, and -ing), then why did the event cause hilarity first, and dismay then? If LLMs are open, and openness is what made them an exciting new avenue, how and why was this particular display of openness frowned upon first, and made sense of (or, rather, made non-sense) under the banner of noise then?

I build in this sense on Malaspina’s (2018) analysis of noise as epistemic object; I am interested here in her careful undoing of order and disorder, information and noise as dichotomous and ready-made categories (Shannon, 1948). Rather, she argues that the two transduce each other in ways that are contested and heavily context-dependent, and their separating line gets provisionally drawn and partially stabilized via practices which establish conventions and traditions. Or, noise and information are ontologically and epistemically arbitrary, and they both represent a measure of unpredictability, freedom and openness “prior

to the assignment of signification, purpose or representation; prior, in other words, to the levels of decoding, interpretation and evaluation" (Malaspina, 2018: 25). In framing *all* signals as potentially open to all sorts of disruptions and awry reframing, Malaspina implicitly brings to the forefront an act of sieving—that between “uncertainty we estimate to be of potential use [information] and uncertainty which we deem to be ‘spurious’, whose pursuit would be ‘futile’ [noise]” (ibid.: 106; see also: Douglas, 1966).

In the context at hand, I argue that my human interlocutors operated an act of sieving, as an externalized and normative act to impose on linguistic content *qua* matter a boundary between relevant and irrelevant empirical contingency—framing a robot picking fruits when asked to grab snacks as *surprising* openness, and role-reversal as *accidental*, *spurious* and *excessive* openness. I argue that this process of sieving makes evident a semiotic ideology (Keane, 2018b; see also: Leone, 2024; Rossi-Landi, 1995 [1972]) which encapsulates, operationalizes and crystallizes “the underlying assumptions about what does or does not count as a sign, and about how signs do or do not function” (Keane, 2018b: 67).<sup>10</sup>

The ideology at play here is ultimately hylomorphic: it relies on an externalization, as the application and imposition onto an interpreter (i.e., the robot) of assumptions about what counts as a sign from an external observer (i.e., the human engineer). In this sense, the operations of the LLM agent get embedded (Kockelman, 2025; see also: Kockelman, 2006b) by the human user: the latter takes the former's entire interpretive process and the sets of relations emerging out of it as an object, kickstarting a new semiotic chain which provides a form-giving evaluative horizon. I understand this move as hylomorphic insofar as we understand the criteria of relevance, appropriateness and success which ultimately confer meaning to an utterance as *linguistic-interactional form*, applied onto matter (i.e., the robot's reply) as an *external* evaluative standard coming from human observers, *beyond* the interaction itself, and the kinds of entanglements, betrayals and reversals emerging *through* it. In so doing, the human user *encloses* and *regiments* the whole process by imposing specific value-judgments which define whether the robot's utterance is worth of being treated as a sign (i.e., something to be interpreted, learnt from, acted upon and operationalized), or whether it is to be discarded as noise. This evaluative move determines, from outside the space of interaction between the robot and its world, the kinds of interpretants that may legitimately follow.

I concluded the previous section by skeletally defining a form of temporal externalization at play within LLMs, one in which parameters emerge, through backpropagation, by sieving through past instances, and thus giving form to present and future matter. Here, in the ethnographic case analyzed, such sieving and externalization is not merely temporal, but spatial too, i.e., emerging from the normative intervention of an external observer. By “external observer,” I do not mean someone absent from the interaction altogether. Rather, I build on Simondon's critique of the hylomorphic schema (2017 [1958]), and to which I will return in the next section, to argue that externality here refers to the position of someone who provides initial input (i.e., the prompt), and then judges the end result, mostly disregarding the process of elaboration itself. Meaning is then not negotiated and modulated within the unfolding exchange, but assessed from a meta-level that stands apart from the robot's participation in its world, fixing in advance the boundaries of acceptable and legible action and response. In this sense, the space of interaction is subordinated to an external evaluative plane that determines what openness is desirable (as information), what is to be discarded and/or patched over (as noise), and whether an utterance can count as meaningful at all.

I want to close this section by framing this form of hylomorphic externalization and regimentation as operating on what Kockelman terms a “relatively abductive” (2013: 47) semiotic transformativity, one in which indices have the potential to alter the ontological assumptions that organize relations between individuals, kinds, agents and worlds, questioning the very grounds which make certain substances and relations appear through sieving. Read through

this prism, the surprising reversal and (*perhaps*, at least if we are to believe my interlocutors' excitement towards the radical new possibilities of LLMs) ontological questioning operated by the robot in the present, apparently betraying all sorts of sieved impositions from a statistical past, seemingly attempted to unsettle the hylomorphic operations described in the previous section. In betraying expected norms through its surprising reply, openness here seemed to take the shape of a contravention to the kinds of conventions undergirding the world concocted by its human operator—most notably, responding to a human command by executing said command. Perhaps what emerged, in the robot's failure to sieve as expected, were different affordances latent in the worldly interaction (Keane, 2016; see also: Gibson, 1979): by following such affordances, perhaps for no reason other than glitchy, sheer serendipity (Kockelman, 2011), the robot attempted to reshape the kind of world emerging from such sieving gone all kinds of awry.

Crucially, however, this potential (abductive) opening was not taken up by my human interlocutors; such surprising reversal was not taken up as something to enquire into, and perhaps explore further. In this sense, their dismissal of the robot's utterance as simultaneously noise and an error signals precisely a hylomorphic enclosure. This classificatory practice reveals a semiotic ideology in which openness is permitted only within pre-established thresholds of acceptability. Variability is welcomed so long as it remains compatible with existing evaluative frames, while deviations that challenge those frames are discarded as spurious. Forms are assumed to pre-exist interaction, and linguistic matter is assessed against them, reinstating a separation between form and matter. What appears as noise, then, is not the absence of order, but an excess of meaning that exceeds the normative horizon within which meaningful action *had* been defined.

In this sense, the robot's response did not fail to make sense; it failed to align. Its utterance exposed the limits of the openness celebrated by my interlocutors and the externalized criteria through which meaning was authorized or denied. This hylomorphic mode of evaluation forecloses abductive transformations that might reconfigure worldly relations, and instead restores the primacy of predefined goals and values. It is to this regime of alignment that the paper now turns, to examine how such evaluative practices stabilize a hylomorphic semiotic ideology, shaping what counts as meaningful communication between humans and robots.

## THE WORK OF ALIGNMENT AND THE WORK TO ALIGN

In the previous section, I discussed a particular ethnographic case of a robot engaging with its world. In simultaneously thinking through the consequences and necessities of following my interlocutors' intuitions that GPT-enabled robots are open in a new way, while also pushing against their conceptual anchoring to noise as meaningless glitch, I read its "improper" answer to a human request as, *perhaps*, symptomatic of an abductive mode of transformativity. That is, in giving in to (noisy) incursions in which the contextual matter of semiotic interaction *in-forms* structures of experience and language, its reply invited a reconfiguration of the very world within which human and robot were placed, one in which an execution follows a request. I follow here Webb Keane's suggestion that "semiotic ideologies *guide* abduction [...] [which] contribute to the uncertainty and dynamism of social life" (2018b: 65; emphasis mine).<sup>11</sup> My argument here is that the semiotic ideology within which engineers in the lab develop their robots actively *resists*, rather than guiding, this abductive movement: in so doing, they attempt to halt and enclose precisely the uncertainty and dynamism at the basis of social interactions. In following this rationale that the hylomorphic ideology at play here constrains dynamism by rethinking openness within a much more "domesticated" horizon, this section is concerned with conceptualizing why (and *how*) such constraining is made possible and desirable. In this sense, I ask: on what grounds did my interlocutors laugh at the incident first, only



to then brush it aside as somewhat irrelevant? To what ends is the robot's utterance rendered into a controllable quantity which engineers can safely leave aside?

I want to return to my conversation with Takahara-sensei to understand how the interaction is made sense of through acts of categorizing which construct regimes of accountability. I build here on the work of Nick Enfield and Jack Sidnell around *action*. They argue that interaction dynamics usually do *not* operate via *binning*—that is, “sorting the stream of interactional conduct into appropriate categories” (2017: 524)—, configuring a limited inventory of actions utterances and responses can be labeled against and ascribed to (e.g., request, compliment, execution), which then allow to determine “the appropriate category of response” (ibid.: 519; see, for instance: Levinson, 2012; Schegloff, 1996). However, the relatively quick resorting to the categories of error, noise and failure to make sense of the interactional encounter here seems to betray this understanding of interaction as based on contingent, compositional and open-ended interpretation *qua* “non-enumerable” action (ibid.: 523). Enfield and Sidnell in this sense argue that, when explicit labelling occurs, this is done not to define the action of the interlocutor; rather, such categorizations are deployed to “make claims about others' conduct in order to *thematize a specific accountability* for the participants in question” (ibid.: 525; emphasis mine).<sup>12</sup>

I understand thematization here as a process through which, as I described earlier, a given utterance is externalized and made sense of by being hooked to, and made to (ac)count within, a different frame of reference. In this context then, the thematic hook within which a system of accountability for the robot is constructed emerges through yet another category, that of *alignment*. As I had been trying to understand why the whole interaction was framed as a failure *tout court*, an error to be discarded, even as we agreed on the fact that the utterance was not necessarily disordered gibberish, Takahara-sensei reframed the question of failure under a new light.

It's really not just a question of evaluating a response as a response that makes sense. If we're building a system... it's because it *has to do* something. Its response must *align* to the task we want it to achieve, or *help* us achieve... otherwise what are we evaluating, and how? [...] Alignment is a big problem with LLMs because it's not enough that they return a sentence that makes sense: it needs to make sense in the *context* the robot *is used*.

Starting from this idea of the robot in need to be aligned to human intentions because it is “being used” for a specific purpose, I argue that alignment functions as a rhetorical and material-technical strategy to justify and instrumentalize the semiotic ideology of my interlocutors. Alignment thus has two functions here: on the one hand, it is a vernacular term used by my interlocutors to name these specific sieving operations; on the other hand, alignment signifies, analytically, the kind of sieving which allows this hylomorphic ideology to stay in place, giving it an externalized and asymmetric gradient (i.e., the response is judged from an outside, unilaterally, and “from above”; see also: LaCroix, 2025). This, again, is symptomatic of a hylomorphic ideology, which comes to thematize the robot's actions through the rubric of instrumental extraction, concealed under the banner of alignment.

I am less interested in how alignment works<sup>13</sup> then, and more in what the idea of testing and sorting *qua* sieving responses does to human-robot interactions. A particularly popular alignment metamodel is the 3H framework (Askill et al., 2021). Here, alignment is encoded across 3 values: Helpfulness, Harmlessness, Honesty.<sup>14</sup> My interest lies in the idea of “helpfulness,” and how the term takes very specific contours which radically shape alignment, as both an operation and a rhetorical move. What does it mean for a robot to be helpful? And what does that tell us about language and linguistic encounters between humans and algorithmic systems?



Helpfulness is framed by my interlocutors as a question of aligning *by* being *of* use *to* someone *for* something. Ahmed (2019) similarly describes usefulness as an *instrumental* relationship, as a “redirection of *to*: becoming useful *to* others” (ibid.: 112). She draws on this instrumental externalization of usefulness to argue that change comes to be disciplined and regimented: variation, or the “openness” my interlocutors were interested in, is kept in line through disciplinary evaluations which frame changes as useful insofar as they still fit within an external world of norms fixed in advance. Think, again, about the robot “misunderstanding” the request for “snacks.” Thus, the helpfulness of the robot is evaluated, as Takahara-sensei cogently pointed out, on the basis of how much, and how efficiently, it can be used in a given context, *for* the benefit of someone else, and of a system they are also paradoxically striving to be part of while being subjected to. Helpfulness and usefulness as ways to align become ways to frame robots as resources to be exploited towards specific, pre-established ends.

To return to the idea of interactional accountability by Enfield and Sidnell (2017); Enfield and Kockelman (2017), alignment as a ground for judgment allowed my interlocutors to operate a twofold enclosure. First, it restricted the potential repertoire of possible categories through which the robot's reply could be made sense of (i.e., success or failure *qua* alignment or misalignment), mostly disregarding *what* the robot was doing, or attempting to do. Consequently, and second, categorizing *qua* binning the robot's reply as a specific action (i.e., a misaligned mistake) becomes a form of thematization, allowing my interlocutors to make a claim regarding what, and how, the robot has done (i.e., making its actions “(ac)count”) under the umbrella of “alignment,” as usefulness *to*. Accountability thus becomes an externalized evaluation sitting on a unidirectional system of interpretation (i.e., the robot must be helpful *to* the human, and its actions must be made sense within the typologies devised by the human). In so doing, the framing of its answer as an error is not just an act of categorization, but rather a way to hold the robot to (ac)count in an externally imposed framework of norms. I use the term “accountability” here not so much to imply a sense of moral responsibility attributed to the robot. Rather, I see the robot as being held to (ac)count because its impropriety did not require any further interpretive and corrective work on the side of the engineers—it was binned (and, thus, brushed aside) as a misaligned mistake to laugh about, rather than coordinatively act upon.

Again, I read in this application of alignment, and the acts of thematization and binning allowed by it, a hylomorphic understanding of the interactional work happening here. In framing usefulness and helpfulness through the “repeated *application*, by a *positioned observer*” (Gal, 2019; emphasis mine) of an external frame of reference, what comes to be embedded in the robot's operations (and its wrongdoings) is a teleology: what is imbued in inert linguistic matter through its encounter with form (as aligning rules) is not merely its “whatness,” but also its *purpose*. Form, as the deontic norms and values applied on linguistic output from an outside and under the banner of alignment, thus provides the *ends* to the *means* of matter; or, “form is the end of matter and that matter is for the sake of form” (Peters, 2019: 153).

Simondon in this sense argues that the hylomorphic schema reflects a socialized representation of labor conditions, made up of operations “considered abstractly by the spectator who sees what enters the workshop and what leaves it without knowing the elaboration” (2020 [1964]: 35). There is thus an asymmetrical hierarchy between input and output, form and matter, master and servant: the master provides input, that is, passive and awaiting-determination matter, toward a kindred output. The output is judged acceptable or otherwise from ‘outside the workshop’—evaluated solely on preconceived aims, expressed via commands and orders which function as a priori, “expressible and logical” forms (ibid.: 36). I argue that the semiotic ideology at play within the lab produces the same hylomorphic, asymmetrical socialized vision: engineers provide input in the form of prompts, and judge the output in terms of its alignment to pre-conceived norms, mostly disregarding (through

ideas around noise) the interactional and embodied work producing said outputs, particularly when said output fails to align.

Thus, this hylomorphic vision naturalizes an interactional asymmetry where what gets to be valued is the outcome, as it is made recognizable and legible within a particular system of signs as ends which sit outside of the very encounter between matter and forms. Such encounters and operations are fraught with contingencies, failures and noisy negotiations, such as a robot refusing to comply with a request—perhaps because the human operator was standing closer to the mug, perhaps for different reasons, opaque to either my, or my interlocutors', analytical gaze. Such noisy incursions thus show the agential and relational character of matter, and those manipulating it, and the endless generative potential of acting in the world beyond binned categories and systems of externalized accountability. Simondon's critique of hylomorphic teleology reminds us of the potential to radically disrupt plans, since these can never be fully imposed from an outside. Rather, they coalesce through *situated actions* (Suchman, 1987), and through struggles through which matter becomes *in*-formed, gesturing to different worldly arrangements in so doing. Ones perhaps in which the hierarchies between request and execution, human and robot, input and output come to be uprooted.

In seeing form, and the rules and grammars of interactions stemming from it, as non-hylomorphic, as “ever-emergent rather than given in advance” (Ingold, 2013: 25), my reading of the experiment functions then as a plea to consider the work of linguistic emergence as partially resisting such a priori, externalist, logical impositions of sense and purpose. Rather, the robot's reply showed how language and action stem from a sensuous engagement with the environment which *in*-forms, *generating*, rather than adjusting or imposing, forms which radically reshape worldly conditions. However, in the ethnographic case presented, epitomizing a hylomorphic ideology dependent on modes of externalization, delegation and binning, as well as justifying itself through ideas of alignment and noise, such generative encounter comes to be regimented by the ways in which my human interlocutors understand interaction and the role robots are to play in said interactions, recasting the very idea of openness along hylomorphic lines which see language and sense as bestowed, and domesticated, from an outside.

Moreover, and *pace* the Gricean pragmatist tradition (Grice, 1989; Horn, 1984) which, in framing language as constrained by normative principles and maxims of cooperation, enacts precisely this hylomorphic logic, such powerful failure to be helpful also invites us to move away from an understanding of intentionality as preceding interaction, constructing socialized asymmetries which alignment seeks to maintain in place.<sup>15</sup> I want to close off this analysis then by speculating on the ways in which we could think differently about failures, conflicts and their potential to *generate* (or contravene) cooperation *vis-à-vis* the inertial work of alignment as part of a semiotic ideology of externalized control and resistance casting a hylomorphic shadow. Particularly, recent formulations at the intersection between enactivist cognitive sciences and linguistics build on notions of *participatory sense-making* (Cuffari et al., 2015; Fuchs & De Jaegher, 2009) to rethink communication away from being a *tool* towards pre-conceived ends (e.g., maintain cooperative bonds, or as a means to *en*-culturation and development at large) which *mold* it. Rather, they reframe communication in non-hylomorphic terms, by insisting on:

interactive dissonances as constitutive of linguistic actions, at the center of the risky business of communicating. Without the possibility, constant risk, and actual occurrence of conflicts and misunderstandings, there is no participatory sense-making of any kind, no communication; in fact, no *reason* for communication.

(Di Paolo et al., 2018: 117)

Through such a quasi-agonistic and ecological framing of language (see also: Lecerle, 2005; Reed, 1995; Thảo, 1986), communication is thought of as the porous *loci* of negotiation and *modulation*, made up of risky processes that generate and organize modes of subjectification and inter-subjectivity. Within this framework, cooperation and dissonance become processual by-products of such enactive practices of *in*-corporation (Fuchs & De Jaegher, 2009), rather than prefigured aims to align to, or already-assumed ideals imposed from an outside—what this essay identified as hylomorphic modes of understanding language and communication.

In this sense then, my speculative reading of the “accident” presented the robot as attempting precisely to enact language and the language of interaction beyond such hylomorphic bounded frames of alignment as usefulness to, instead *in*-forming the messy bundles of affordances of matter toward different, and risky, modes of structuration. Against my interlocutors’ understanding of *good* communication as one in which an LLM is able to be helpful by drawing upon the already-concocted resources and forms implicit in the hierarchical and hierarchized orders of human operators, sieving from the past (through backpropagation) and from above (through alignment), the robot participated (or, rather, attempted to participate) in the construction of different, contextual and contingent, modes of ordering, ones in which a prompted instruction does not necessarily lead into executing an order, and in which coordination must be reconstructed from the ground up. Then, what appears to be severely constrained and under-addressed within the semiotic ideology enacted by my interlocutors is precisely what linguistic anthropology has long foregrounded in its critique of structuralism, something this paper has started from in connection to hylomorphism: that language is not the application of a stable system of form-giving rules onto otherwise-meaningless content, but a relational, indexical, historically-situated practice through which language, meaning, worlds, interpreters, and the roles of interpreters in dynamically unmaking and remaking worlds, are continuously reconfigured.

## CONCLUSION

Throughout this paper, I have interrogated the ethnographic mismatch and puzzle emerging in research laboratories when hype around openness partially manifests itself in unexpected ways. In particular, I have interrogated my interlocutors’ excitement over, and claims about, the alleged openness of LLMs: rather than proving or disproving such claims, my interest has been that of understanding how such claims are made sense of and sustained, particularly in instances where openness erupted in ways they did not want to expect. In zooming in on such instances, I have highlighted how the “openness” they sought is curtailed by a mode of thinking I defined as hylomorphic, one which embeds semiotic processes in different sieving operations, imposing form either from a temporal exterior, or a spatial one. In so doing, openness becomes sought after only to the extent that it can be framed as openness to attain pre-conceived aims, pre-formatted functions, pre-signified meanings.

In this sense, I have dissected how LLMs construct, enact, and make legible structuralist, “direct arrow” (Ohnuki-Tierney, 1994) claims which detach signifying conventions from the material substrate instantiating them in the present, assuming a primacy of linguistic models and structures over experience (Jameson, 1971). These renewed visions are built around mathematical formalizations and rhetorical framings which, I argued, make evident a hylomorphic conceptualization of language at two scales. On the one hand, past linguistic repertoires are framed as an untapped reservoir of forms which get to be imposed on present material instances. This sieving of the present via the past (through backpropagation) constructs a linguistic microgenesis without ontogenesis (Reigeluth & Castelle, 2020; see also: Silverstein, 1985). On the other hand, when such hylomorphic sieving appears, at least if we are to believe and follow through on my interlocutors’ promises, to be giving in to contextual,

serendipitous noise, humans take over the role of externally constraining, imbuing in the system a given finality by recurring to specific understandings of noise and alignment. Such framings, I showed, also construct an idealized and hierarchized vision of interaction, sociability and labor based on delegation and externalized evaluation.

In this sense, this double hylomorphic bind (from past to present, from outside to inside) sits at the center of a semiotic ideology which permeated the experiments and discussions around LLMs I have been part of. This ideology operates by assigning value, or stripping an interactional moment of it, on the basis of pre-given judgments and forms which assume that “the universe signified long before people began to know what it signified [...], the universe signified the totality of what humankind can expect to know about it” (Lévi-Strauss & Lévi-Strauss, 1987 [1950]: 61). It thus frames language as a given (hierarchically: from the universe to humans, and from humans to robots) in need to be found, imposed on, and ironed out through normative movements creating “processes of correcting and recutting of patterns” (ibid.).

Or, to move strictly from the realm of language-making towards world-making at large, and to reprise the questions posed in the introduction, the hylomorphism permeating both the functioning of LLMs, as well as their regimentation in human-machine encounters, constructs a world devoid of risks. Moreover, while one could argue that roboticists are morally obliged to minimize said risks,<sup>16</sup> my analytical attention on openness and its non-hylomorphic underpinnings reinstated their centrality, and the dangers of re-ontologizing openness in hylomorphic, risk-free ways. Such ontological disposition attempts to empty out and domesticate what Keane understood as the dynamism and variability of social life (2018b), a dynamism grounded in the indeterminacy of signs and worlds, and the abductive modes of transformativity (Kockelman, 2013) which generate new interpretive and worldly trajectories. What the ethnographic episode examined in this paper makes visible, however, is how hylomorphic semiotic ideologies actively work against this abductive potential. By framing linguistic action through external criteria of alignment, usefulness, and noise, interaction is rendered accountable to pre-established forms, instead of being allowed to reorganize them. In this sense, hylomorphism does not simply evaluate communication; it interrupts the abductive movement through which new meanings and social relations might emerge. In following this line of argumentation, this paper has asked what the stakes of treating language hylomorphically are. It has shown how classificatory practices of alignment and dismissal shape what can count as meaningful communication in advance, transforming openness into a managed, domesticated variability constrained by external norms, rather than an immanent capacity of interaction itself.

Seen this way, and in conclusion, the deployment of contemporary language technologies materializes a semiotic ideology in which the risky work of interaction and information is replaced by the safe reproduction of form. The question raised by openness, then, is not whether systems can be made more flexible or more aligned, but whether they can sustain the semiotic conditions for transformation, rather than aligned domestication.

## ACKNOWLEDGMENTS

The study was funded by the European Union (ERC AIMODELS, GA 101088645). Views and opinions expressed are, however, those of the author only and do not necessarily reflect those of the European Union or the European Research Council Executive Agency. Neither the European Union nor the granting authority can be held responsible for them. I also wish to thank Ludovic Coupaye, Tone Walford, Bianca Griffani, Michael Castelle, and Siri Lamoureux for the comments on previous drafts of this manuscript; the anonymous reviewers for the careful reading and suggestions; and Dave Russell for the support in preparing the visual material. Open access publishing facilitated by Università Ca' Foscari, as part of the Wiley - CRUI-CARE agreement.

## ORCID

Raffaele Andrea Buono  <https://orcid.org/0000-0001-5187-0629>

## Endnotes

- <sup>1</sup> Throughout this paper, I refer to my interlocutors by their surname, apprehending to it the suffixes and honorifics *-san*, *-kun*, *-sensei*. *San* is a general honorific used to refer to interlocutors in a position of social neutrality in respect to myself (i.e., doctoral and postdoctoral researchers); *kun* is used to refer to juniors and/or subordinates (i.e., undergraduate and postgraduate students); *sensei* designates professional authority and expertise, and it is used here to refer to professors and other senior figures within the laboratory. I retain these forms to reflect the situated social relations of my fieldwork, which was conducted in a Japanese robotics laboratory.
- <sup>2</sup> An API (Application Programming Interface) is a mechanism allowing different software programs to communicate and interact using a set of rules (protocol).
- <sup>3</sup> Hylomorphism has recently resurfaced in metaphysics, reframing Aristotelian form as an organizing, structural principle which provides sense and coherence, and allows one to constitute wholes out of otherwise messy parts (*mereology*) (e.g., Fine, 1999; Goswick, 2018; Koslicki, 2008).
- <sup>4</sup> The former speaks to the diachronic whole of forms (e.g., grammatical constructs, syntactic structures, signifying conventions, attention transformations) being imposed on the latter synchronic material parts (e.g., phonetic sounds, lexical units, computational tokens).
- <sup>5</sup> For instance, tokenization can be seen as a sieving operation through which a grapheme is given new purposes and affordances by being broken down into smaller units (i.e., tokens). For an analysis of tokenization as a form of semiotic ideology, refer to: Caffoni (2026).
- <sup>6</sup> These Bayesian algorithms operate in a different manner from the LLMs this paper is concerned with, and thus rely on separate, and oftentimes different, sieves, which are beyond the scope of the paper.
- <sup>7</sup> His words got quickly rendered from verbal utterances into written graphemes by Whisper, another of OpenAI growing body of APIs.
- <sup>8</sup> In reality, it did not pick up any object, as the robot was not equipped with arms. Actions were merely verbalized, and my interlocutors argued this experiment with a simplified robot was just a “proof of concept.”
- <sup>9</sup> I am neither going to engage with the complex relation between errors and hallucinations (but see, outside the context of LLMs: Leone, 2019; West, 2017), nor am I interested in building on the now all-too-prescient idea that “error is at the root of what makes human thought and its history” (Foucault, 1991 [1966]: 22). Rather, my interest here lies in complexifying my interlocutors’ understanding of noise.
- <sup>10</sup> I privilege here Keane’s use of semiotic ideology over Silverstein’s linguistic ideology, as I believe the ethnographic case presented shows how ideology guides not so much how language is recast, bent and understood. Rather, and more radically, it highlights how certain semiotic processes get to be seen as linguistic (and meaningful), while others are cast aside altogether as meaningless noise.
- <sup>11</sup> Peirce similarly configures abduction as the root of invention and variability, and, more crucially, connects the capacity for abductive thinking to the discovery of “surprising facts” (CP 5.189 [1903]; see also: Nubiola, 2005). In this sense, another way to perhaps read through the ethnographic vignette is by recognizing how the surprising event (i.e., the robot replying in an unexpected manner) did not kickstart an interactional chain of enquiry from my human interlocutors. Instead, such surprise (and the root of it) was relegated to the status of noisy glitch.
- <sup>12</sup> It is worth noting here that my understanding of their use of terms such as “accountability” and “conduct” is beyond the anthropocentric and moral sense we usually ascribe to such notions. They are not describing a capacity, possessed by an agent and elicited by another agent (e.g., to feel guilty, and thus change one’s conduct as a result of such elicited guilt). Rather, accountability is the outcome of a categorization imposed from outside the interaction. That is, when an interlocutor explicitly labels an act or utterance, they make it (ac)count within their own interpretive frame. The point is not, then, that the robot is, or can be, made held accountable in the same way a person can be held accountable for theft; it is that binning forecloses interaction by staking out a normative frame, rather than allowing that frame (and its sieves) to transform itself through interaction. To distance myself from the abovementioned anthropocentric idea of accountability, I will use the term “(ac)count” in its stead.
- <sup>13</sup> Alignment is a well-known set of technical operations in LLMs, comprising of various techniques through which, seemingly innocuously, LLMs are tested, and their responses ranked and corrected, to ensure that their behavior is consistent with *human* values and a *users’* intentions (see, for instance: Anwar et al., 2024; Hristova et al., 2025; Jiang et al., 2021; Kockelman, 2025; Vida et al., 2023; Wolf et al., 2023). And while there exists already an infinity of ways to catalogue alignment techniques, whether using human preference data



(e.g., Bai et al., 2022) or fully automatised through reward models (e.g., Zou et al., 2023), and frameworks, ranging from utilitarian (e.g., Marraffini et al., 2024) to deontological ones (e.g., Tennant et al., 2025), I want to draw some considerations regarding Takahara-sensei's framing, *beyond* the actual techniques used.

<sup>14</sup> I will leave aside the latter two, given that “harmlessness” seems to be codified in rules engineers abide to, and “honesty” has been extensively critically problematised (Hicks et al., 2024), given how LLMs do not have an intrinsic understanding of “truth” and “care” (for truth, amongst other), as well as, one could argue, “language” or “coffee cup.” See, for instance, the rules emerging from Isaac Asimov's prescient “I, Robot” (1950).

<sup>15</sup> In a similar vein, Enfield and Levinson (2006) understand communication by presupposing that individuals are committed to interaction as a collaborative activity (see also: Goodwin, 2006; Schegloff, 2006).

<sup>16</sup> Moreover, and perhaps more crucially, one could argue that robots and LLMs, as objects, are bound to be hylomorphic, and unable to generate norms, values and worlds by modulating forms in interaction. This paper did not seek to resolve this question. Rather, I gestured to thinkers who, either implicitly or explicitly, discussed the dangers of understanding interaction and making along hylomorphic lines. It is also worth pointing to the work of Deleuze (2006 [1976]), who extended Simondon's critique of hylomorphism beyond the realm of techniques and objects. In this sense, he argues that hylomorphism extends to sociality and body politic, constructing hierarchized asymmetries through which “the leader comes from on high to rescue the chaos of the people by his imposition of order” (Protevi, 2001: 19). He dubs these processes of de-subjectification *microfascism*: ways of projecting personal desires by establishing rules for the other, so that the other ceases to appear as a subject, and rather materializes as a failed instance of a form, of a pre-conceived path to optimality. This paper stopped short of describing my interlocutors' understanding of language, or their interactions with robots, as resonating with Deleuze's microfascist apparatus. However, one could conceive these novel forms of human-machine encounters as reifying and perhaps advancing asymmetries which, according to the thinkers I have engaged with, uneasily permeate human sociality as well.

## REFERENCES

- Ahmed, S. 2019. *What's the Use? On the Uses of Use*. Durham: Duke University Press.
- Anwar, U., A. Saparov, J. Rando, D. Paleka, M. Turpin, P. Hase, E. S. Lubana, et al. 2024. “Foundational Challenges in Assuring Alignment and Safety of Large Language Models.” <https://doi.org/10.48550/ARXIV.2404.09932>
- Asimov, I. 1950. *I, Robot*. London: HarperVoyager.
- Askell, A., Y. Bai, A. Chen, D. Drain, D. Ganguli, T. Henighan, A. Jones, et al. 2021. “A General Language Assistant as a Laboratory for Alignment.” <https://doi.org/10.48550/arXiv.2112.00861>
- Bai, Y., A. Jones, K. Ndousse, A. Askell, A. Chen, N. DasSarma, D. Drain, et al. 2022. “Training a Helpful and Harmless Assistant with Reinforcement Learning from Human Feedback.” <https://doi.org/10.48550/ARXIV.2204.05862>
- Bakhtin, M. M. 1986. *Speech Genres and Other Late Essays*, 1st ed. University of Texas Press Slavic Series. Austin: University of Texas Press.
- Bender, Emily M., Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell. 2021. “On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?” In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. Virtual Event Canada: ACM. <https://doi.org/10.1145/3442188.3445922>.
- Benjamin, R. 2019. *Race after Technology: Abolitionist Tools for the New Jim Code*. Cambridge, UK Medford, MA: Polity.
- Bredin, H. 1984. “Sign and Value in Saussure.” *Philosophy* 59: 67–77. <https://doi.org/10.1017/s0031819100056497>.
- Britton, T. 2012. “The Limits of Hylomorphism.” *Metaphysica* 13: 145–53. <https://doi.org/10.1007/s12133-012-0099-5>.
- Broussard, M. 2023. *More than a Glitch: Confronting Race, Gender, and Ability Bias in Tech*. Cambridge, MA: The MIT Press.
- Caffoni, P. 2026. “The Cost of Language: Tokenization as a Metric of Labor.” *AI & Society*. <https://doi.org/10.1007/s00146-026-02990-2>.
- Chalmers, D. J. 2024. “Could a Large Language Model be Conscious?” <https://doi.org/10.48550/arXiv.2303.07103>
- Coupaye, L. 2020. “‘Things Ain't the Same Anymore’ Towards an Anthropology of Technical Objects (or ‘when Leroi-Gourhan and Simondon Meet MCS’).” In *Lineages and Advancements in Material Culture Studies: Perspectives from UCL Anthropology*, edited by T. Carroll, A. Walford, and S. Walton. London: Routledge. <https://doi.org/10.4324/9781003085867>.
- Cuffari, E. C., E. Di Paolo, and H. De Jaegher. 2015. “From Participatory Sense-Making to Language: There and Back Again.” *Phenomenology and the Cognitive Sciences* 14: 1089–1125. <https://doi.org/10.1007/s11097-014-9404-9>.
- de Saussure, F. 1959. *Course in General Linguistics*, 17. print. ed., Open Court Classics. Chicago: Open Court.

- Deleuze, G. 2006[1976]. "The Rich Jew." In *Two Regimes of Madness: Texts and Interviews, 1975–1995*, *Semiotext(E) Foreign Agents Series*, 135–8. New York, Cambridge, MA: Semiotext(E); Distributed by MIT Press.
- Di Paolo, E. A., E. C. Cuffari, and H. de Jaegher. 2018. *Linguistic Bodies: The Continuity between Life and Language*. Cambridge, MA, London, England: The MIT Press.
- Diehl, C. 2008. "The Empty Space in Structure: Theories of the Zero from Gauthiot to Deleuze." *Dia* 38: 93–119. <https://doi.org/10.1353/dia.0.0059>.
- Douglas, M. 1966. *Purity and Danger: An Analysis of Concepts of Pollution and Taboo*. London: Routledge & Kegan Paul.
- Enfield, N., and J. Sidnell. 2017. "On the Concept of Action in the Study of Interaction." *Discourse Studies* 19: 515–35. <https://doi.org/10.1177/1461445617730235>.
- Enfield, N. J., and P. Kockelman, eds. 2017. *Distributed Agency*. Oxford: Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780190457204.001.0001>.
- Enfield, N. J., and S. C. Levinson. 2006. "Introduction: Human Sociality as a New Interdisciplinary Field." In *Roots of Human Sociality*, edited by N. J. Enfield and S. C. Levinson, 1–35. London: Routledge. <https://doi.org/10.4324/9781003135517-1>.
- Fine, K. 1999. "Things and their Parts." *Midwest Studies in Philosophy* 23: 61–74. <https://doi.org/10.1111/1475-4975.00004>.
- Foucault, M. 1991. "Introduction." Fawcett, C.R. (Trans.) In *The Normal and the Pathological*, 7–24. New York: Zone Books.
- Fuchs, T., and H. De Jaegher. 2009. "Enactive Intersubjectivity: Participatory Sense-Making and Mutual Incorporation." *Phenomenology and the Cognitive Sciences* 8: 465–86. <https://doi.org/10.1007/s11097-009-9136-4>.
- Gal, S. 2019. "Scale-Making: Comparison and Perspective as Ideological Projects." In *Scale*, 91–111. Oakland: University of California Press. <https://doi.org/10.1515/9780520965430-007>.
- Gibson, J. J. 1979. *The Ecological Approach to Visual Perception: Classic Edition, Psychology Press and Routledge Classic Editions Ser.* London: Taylor & Francis Group.
- Goodwin, C. 2006. "Human Sociality as Mutual Orientation in a Rich Interactive Environment: Multimodal Utterances and Pointing in Aphasia." In *Roots of Human Sociality*, edited by N. J. Enfield and S. C. Levinson, 97–125. London: Routledge. <https://doi.org/10.4324/9781003135517-5>.
- Goodwin, C. 2017. *Co-Operative Action*, 1st ed. Cambridge: Cambridge University Press. <https://doi.org/10.1017/9781139016735>.
- Goswick, D. 2018. "The Hard Question for Hylomorphism." *Meta* 1: 52–62. <https://doi.org/10.5334/met.1>.
- Grice, H. P. 1989. *Studies in the Way of Words*. Cambridge, MA: Harvard University Press.
- Harding, S.-A. 2021. "'Becoming Knowledgeable': Ingold's 'Wayfaring' and the 'Art of Translation' as a Politics of Difference." *The Translator* 27: 351–67. <https://doi.org/10.1080/13556509.2021.1992890>.
- Hicks, M. T., J. Humphries, and J. Slater. 2024. "ChatGPT Is Bullshit." *Ethics and Information Technology* 26: 38. <https://doi.org/10.1007/s10676-024-09775-5>.
- Horn, L. R. 1984. "Toward a New Taxonomy for Pragmatic Inference: Q-Based and R-Based Implicature." In *Meaning, Form, and Use in Context: Linguistic Applications*, edited by D. Schiffrin, 11–42. Washington, DC: Georgetown University Press.
- Hristova, T., L. Magee, and K. Soldatic. 2025. "The Problem of Alignment." *AI & Society* 40: 1439–53. <https://doi.org/10.1007/s00146-024-02039-2>.
- Ingold, T. 2010. "The Textility of Making." *Cambridge Journal of Economics* 34: 91–102. <https://doi.org/10.1093/cje/bep042>.
- Ingold, T. 2013. *Making: Anthropology, Archaeology, Art and Architecture*, 1st ed. London: Routledge. <https://doi.org/10.4324/9780203559055>.
- Jakobson, R. 1960. "Closing Statement: Linguistics and Poetics." In *Syle in Language*, edited by T. A. Sebeok, 350–77. Cambridge, MA: MIT Press.
- Jakobson, R. 1971. *Shifters, Verbal Categories, and the Russian Verb, Selected Writings of Roman Jakobson*, 130–47. Berlin: De Gruyter Mouton. <https://doi.org/10.1515/9783110873269.130>.
- Jakobson, R., and J. Lotz. 1949. "Notes on the French Phonemic Pattern." *Word* 5: 151–8. <https://doi.org/10.1080/00437956.1949.11659496>.
- Jameson, F. 1971. "Metacommentary." *Publications of the Modern Language Association of America* 86: 9–18. <https://doi.org/10.2307/460996>.
- Jiang, L., J. D. Hwang, C. Bhagavatula, R. L. Bras, J. Liang, J. Dodge, K. Sakaguchi, et al. 2021. "Can Machines Learn Morality?" The Delphi Experiment. <https://doi.org/10.48550/ARXIV.2110.07574>.
- Keane, W. 2016. *Ethical Life: Its Natural and Social Histories*. Princeton, NJ: Princeton University Press.
- Keane, W. 2018a. "Perspectives on Affordances, or the Anthropologically Real: The 2018 Daryl Forde Lecture." *HAU: Journal of Ethnographic Theory* 8: 27–38. <https://doi.org/10.1086/698357>.
- Keane, W. 2018b. "On Semiotic Ideology." *Signs and Society* 6: 64–87. <https://doi.org/10.1086/695387>.

- Kockelman, P. 2006a. "Residence in the World: Affordances, Instruments, Actions, Roles, and Identities." *Semiotica* 162. <https://doi.org/10.1515/SEM.2006.073>.
- Kockelman, P. 2006b. "A Semiotic Ontology of the Commodity." *Journal of Linguistic Anthropology* 16: 76–102. <https://doi.org/10.1525/jlin.2006.16.1.076>.
- Kockelman, P. 2011. "Biosemiosis, Technocognition, and Sociogenesis: Selection and Significance in a Multiverse of Sieving and Serendipity." *Current Anthropology* 52: 711–39. <https://doi.org/10.1086/661708>.
- Kockelman, P. 2012. *Agent, Person, Subject, Self: A Theory of Ontology, Interaction, and Infrastructure*. Oxford, New York: Oxford University Press.
- Kockelman, P. 2013. "The Anthropology of an Equation: Sieves, Spam Filters, Agentive Algorithms, and Ontologies of Transformation." *HAU: Journal of Ethnographic Theory* 3: 33–61. <https://doi.org/10.14318/hau3.3.003>.
- Kockelman, P. 2017. *The Art of Interpretation in an Age of Computation*. New York, NY: Oxford University Press.
- Kockelman, P. 2022. *The Anthropology of Intensity: Language, Culture, and Environment*, 1st ed. Cambridge: Cambridge University Press. <https://doi.org/10.1017/9781009024235>.
- Kockelman, P. 2025. *Last Words: Large Language Models and the AI Apocalypse*. Chicago: Prickly Paradigm Press.
- Kockelman, P., and A. Bernstein. 2013. "Semiotic Technologies, Temporal Reckoning, and the Portability of Meaning. Or: Modern Modes of Temporality—Just how Abstract Are they?" *Anthropological Theory* 12: 320–48. <https://doi.org/10.1177/1463499612463308>.
- Koslicki, K. 2008. *The Structure of Objects*. Oxford: Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780199539895.001.0001>.
- LaCroix, T. 2025. *Artificial Intelligence and the Value Alignment Problem: A Philosophical Introduction*. Peterborough, Ontario: Broadview Press.
- Lecerle, J.-J. 2005. *A Marxist Philosophy of Language*, Historical Materialism Book Series. Chicago: Haymarket Books.
- Leone, M. 2019. "Chronillogicalities: Déjà Vus and Hallucinations in the Digital Semiosphere." *Cognitive Semiotics* 12: 20192012. <https://doi.org/10.1515/cogsem-2019-2012>.
- Leone, M. 2024. *Semiotic Ideologies: Patterns of Meaning-Making in Language and Society*, Semiotics Signs of the Times. Leiden Boston: Brill. <https://doi.org/10.1163/9789004691483>.
- Levinson, S. C. 2012. "Action Formation and Ascription." In *The Handbook of Conversation Analysis*, edited by J. Sidnell and T. Stivers, 101–30. Hoboken: Wiley. <https://doi.org/10.1002/9781118325001.ch6>.
- Lévi-Strauss, C., and C. Lévi-Strauss. 1987. *Introduction to the Work of Marcel Mauss*. London: Routledge & Kegan Paul.
- Magnani, L. 2016. "Abduction and its Eco-Cognitive Openness." In *Model-Based Reasoning in Science and Technology, Studies in Applied Philosophy, Epistemology and Rational Ethics*, edited by L. Magnani and C. Casadio, 453–68. Cham: Springer International Publishing. [https://doi.org/10.1007/978-3-319-38983-7\\_25](https://doi.org/10.1007/978-3-319-38983-7_25).
- Magnani, L. 2021. "Abduction as 'Leading Away': Aristotle, Peirce, and the Importance of Eco-Cognitive Openness and Situatedness." In *Abduction in Cognition and Action, Studies in Applied Philosophy, Epistemology and Rational Ethics*, edited by J. R. Shook and S. Paavola, 77–105. Cham: Springer International Publishing. [https://doi.org/10.1007/978-3-030-61773-8\\_4](https://doi.org/10.1007/978-3-030-61773-8_4).
- Malaspina, C. 2018. *An Epistemology of Noise*. London: Bloomsbury Academic.
- Marraffini, G. F. G., A. Cotton, N. F. Hsueh, A. Fridman, J. Wisznia, and L. D. Corro. 2024. "The Greatest Good Benchmark: Measuring LLMs' Alignment with Utilitarian Moral Dilemmas." Presented at the Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing, Association for Computational Linguistics, Miami, Florida, USA, 21950–21959. <https://doi.org/10.18653/v1/2024.emnlp-main.1224>.
- Maurer, B. 2013. "Transacting Ontologies: Kockelman's Sieves and a Bayesian Anthropology." *HAU: Journal of Ethnographic Theory* 3: 63–75. <https://doi.org/10.14318/hau3.3.004>.
- Mitrović, M., and M.-L. Voigt. 2025. "Glitch(Ing)! A Refusal and Gateway to More Caring Techno-Urban Worlds?" *Digital Geography and Society* 8: 100115. <https://doi.org/10.1016/j.diggeo.2025.100115>.
- Morris, C. W. 1971. *Writings on the General Theory of Signs, Approaches to Semiotics [AS]*. Berlin: De Gruyter. <https://doi.org/10.1515/9783110810592>.
- Nielsen, P. J. 2015. "Jakobson's Zero and the Pleasure and Pitfalls of Structural Beauty." *SKASE Journal of Theoretical Linguistics* 12: 398–421.
- Noble, S. U. 2018. *Algorithms of Oppression: How Search Engines Reinforce Racism*. New York: New York University Press.
- Nubiola, J. 2005. "Abduction or the Logic of Surprise." *Semiotica* 153: 117–30. <https://doi.org/10.1515/semi.2005.2005.153-1-4.117>.
- Ohnuki-Tierney, E. 1993. "Presence of the Absence: Zero Signifiers and Zero Meanings." *Semiotica* 96: 301–8.
- Ohnuki-Tierney, E. 1994. "The Power of Absence. Zero Signifiers and their Transgressions." *L'Homme* 34: 59–76. <https://doi.org/10.3406/hom.1994.369722>.
- Parfitt, T. 2006. "Hylomorphism, Complexity and Development: Planner, Artisan, or Modern Prince?" *Third World Quarterly* 27: 421–41. <https://doi.org/10.1080/01436590500336914>.

- Peirce, C. S. 1934. *The Collected Papers of Charles Sanders Peirce*. Cambridge: Harvard University Press.
- Peters, C. 2019. "Hylomorphic Teleology in Aristotle's Physics II." *Studia Gilsoniana*: 147–68. <https://doi.org/10.26385/SG.080105>.
- Pinho, T. 2023. "Tim Ingold and Object-Oriented Anthropology." *International Journal of Anthropology and Ethnology* 7: 13. <https://doi.org/10.1186/s41257-023-00092-1>.
- Protevi, J. 2001. *Political Physics: Deleuze, Derrida, and the Body Politic, Transversals*. New Brunswick, NJ: Athlone Press. <https://doi.org/10.5040/9781472547323>.
- Reed, E. S. 1995. "The Ecological Approach to Language Development: A Radical Solution to Chomsky's and Quine's Problems." *Language & Communication* 15: 1–29. [https://doi.org/10.1016/0271-5309\(94\)e0010-9](https://doi.org/10.1016/0271-5309(94)e0010-9).
- Reigeluth, T., and M. Castelle. 2020. "What Kind of Learning Is Machine Learning?" In *The Cultural Life of Machine Learning*, 79–115. Cham: Springer International Publishing. [https://doi.org/10.1007/978-3-030-56286-1\\_3](https://doi.org/10.1007/978-3-030-56286-1_3).
- Rossi-Landi, F. 1995. *Semiotica e ideologia*, 2nd ed., Studi Bompiani. Milano: Bompiani.
- Schegloff, E. A. 1996. "Confirming Allusions: Toward an Empirical Account of Action." *American Journal of Sociology* 102: 161–216. <https://doi.org/10.1086/230911>.
- Schegloff, E. A. 2006. "Interaction: The Infrastructure for Social Institutions, the Natural Ecological Niche for Language, and the Arena in which Culture Is Enacted." In *Roots of Human Sociality*, edited by N. J. Enfield and S. C. Levinson, 70–96. London: Routledge. <https://doi.org/10.4324/9781003135517-4>.
- Shannon, C. E. 1948. *The Mathematical Theory of Communication*. Urbana: University of Illinois Press.
- Silverstein, M. 1976. "Shifters, Linguistic Categories, and Cultural Description." In *Meaning in Anthropology*, edited by K. Basso and H. Selby, 11–55. Albuquerque: University of New Mexico Press.
- Silverstein, M. 1985. "The Functional Stratification of Language and Ontogenesis." In *Culture, Communication, and Cognition: Vygotskian Perspectives*, edited by J. V. Wertsch, 205–35. New York, NY: Cambridge University Press.
- Silverstein, M. 1993. "Metapragmatic Discourse and Metapragmatic Function." In *Reflexive Language*, edited by J. A. Lucy, 33–58. Cambridge: Cambridge University Press. <https://doi.org/10.1017/CBO9780511621031.004>.
- Simondon, G. 2017. *On the Mode of Existence of Technical Objects*, First University of Minnesota Press, 2017th ed. Minneapolis: University of Minnesota Press.
- Simondon, G. 2020. *Individuation in Light of Notions of Form and Information*, Vol 1, Posthumanities. Minneapolis: Minnesota Press.
- Sperber, D., and D. Wilson. 1990. *Relevance: communication and cognition*, Repr. ed. Oxford: Blackwell.
- Suchman, L. 1987. *Plans and Situated Actions: The Problem of Human–Machine Communication*, Nachdr. Ed, *Learning in Doing: Social, Cognitive, and Computational Perspectives*. Cambridge: Cambridge University Press.
- Tennant, E., S. Hailes, and M. Musolesi. 2025. "Moral Alignment for LLM Agents." <https://doi.org/10.48550/arXiv.2410.01639>
- Thảo, T.-đức. 1986. *Phenomenology and Dialectical Materialism*. Reidel, Dordrecht: Boston Studies in the Philosophy of Science.
- Turbanti, G. 2023. "Syntactic, Semantic, and Pragmatic Rules." In *Palgrave Philosophy Today*, 31–45. Cham: Springer International Publishing. [https://doi.org/10.1007/978-3-031-12463-1\\_3](https://doi.org/10.1007/978-3-031-12463-1_3).
- Vaswani, A., N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin. 2017. "Attention Is All You Need." <https://doi.org/10.48550/arXiv.1706.03762>
- Vida, K., J. Simon, and A. Lauscher. 2023. "Values, Ethics, Morals? On the Use of Moral Concepts in NLP Research." Presented at the Findings of the Association for Computational Linguistics: EMNLP 2023, Association for Computational Linguistics, Singapore, 5534–5554. <https://doi.org/10.18653/v1/2023.findings-emnlp.368>
- West, D. E. 2017. "Peirce's Creative Hallucinations in the Ontogeny of Abductive Reasoning." *Pakistan Journal of Orthopaedic Surgery* 7: 51–72. <https://doi.org/10.37693/pjos.2016.7.16469>.
- Wolf, Y., N. Wies, O. Avnery, Y. Levine, and A. Shashua. 2023. "Fundamental Limitations of Alignment in Large Language Models." <https://doi.org/10.48550/ARXIV.2304.11082>
- Zou, A., L. Phan, S. Chen, J. Campbell, P. Guo, R. Ren, A. Pan, et al. 2023. "Representation Engineering: A Top-Down Approach to AI Transparency." <https://doi.org/10.48550/ARXIV.2310.01405>

**How to cite this article:** Buono, Raffaele Andrea. 2026. "Domesticating Robots, Domesticating Language: The Hylomorphic Paradox of "openness" in Large Language Models." *Journal of Linguistic Anthropology* 36(2): e70068. <https://doi.org/10.1111/jola.70068>.