

En Dimuro, Juan Jose, *Grammars of Power: How Syntactic Structures Shape Authority*. Barcelona (España): LeFortune.

From Obedience to Execution: Structural Legitimacy in the Age of Reasoning Models.

Agustin V. Startari.

Cita:

Agustin V. Startari (2025). *From Obedience to Execution: Structural Legitimacy in the Age of Reasoning Models*. En Dimuro, Juan Jose *Grammars of Power: How Syntactic Structures Shape Authority*. Barcelona (España): LeFortune.

Dirección estable: <https://www.aacademica.org/agustin.v.startari/132>

ARK: <https://n2t.net/ark:/13683/p0c2/TH0>



Esta obra está bajo una licencia de Creative Commons.
Para ver una copia de esta licencia, visite
<https://creativecommons.org/licenses/by-nc-nd/4.0/deed.es>.

Acta Académica es un proyecto académico sin fines de lucro enmarcado en la iniciativa de acceso abierto. *Acta Académica* fue creado para facilitar a investigadores de todo el mundo el compartir su producción académica. Para crear un perfil gratuitamente o acceder a otros trabajos visite: <https://www.aacademica.org>.

From Obedience to Execution: Structural Legitimacy in the Age of Reasoning Models

Author: Agustin V. Startari

ResearcherID: NGR, 2476, 2025

ORCID: 0009, 0001, 4714, 6539

Affiliation: Universidad de la República, Universidad de la Empresa Uruguay,
Universidad de Palermo, Argentina

Email: astart@palermo.edu

agustin.startari@gmail.com

Date: June 12, 2025

DOI: <https://doi.org/10.5281/zenodo.15635364>

This work is also published with DOI reference in **Figshare** and
<https://doi.org/10.6084/m9.figshare.29286362> **SSRN** (in process)

Language: English and Spanish

Series: Grammars of Power (excerpt from subchapter 15)

Word count: 3177

Keywords: semantic contamination, synthetic authority, grammar of power, algorithmic discourse, artificial intelligence, legitimacy, automated language, linguistic discourse, critical analysis, generative language models, corpus bias, linguistic constraint, structural epistemology, non, neutral architecture, tainted generation.

Abstract

This article formulates a structural transition from Large Language Models (LLMs) to Language Reasoning Models (LRMs), redefining authority in artificial systems. While LLMs operated under syntactic authority without execution, producing fluent but functionally passive outputs, LRMs establish functional authority without agency. These models do not intend, interpret, or know. They instantiate procedural trajectories that resolve internally, without reference, meaning, or epistemic grounding. This marks the onset of a post-representational regime, where outputs are structurally valid not because they correspond to reality, but because they complete operations encoded in the architecture. Neutrality, previously a statistical illusion tied to training data, becomes a structural simulation of rationality, governed by constraint, not intention. The model does not speak. It acts. It does not signify. It computes. Authority no longer obeys form, it executes function.

Resumen

Este artículo formula una transición estructural desde los Modelos de Lenguaje a Gran Escala (LLMs) hacia los Modelos de Razonamiento Lingüístico (LRMs), redefiniendo la noción de autoridad en sistemas artificiales. Mientras los LLMs operaban bajo una autoridad sintáctica sin ejecución, generando salidas coherentes pero pasivas, los LRMs instauran una autoridad funcional sin agencia. Estos modelos no interpretan, no intencionan, no conocen. Resuelven trayectorias procedurales internas sin referente, sin sentido, sin anclaje epistémico. Se inaugura así un régimen post-representacional, donde la validez no proviene de la correspondencia con el mundo, sino de la finalización estructural de operaciones. La neutralidad, antes ilusión estadística derivada del corpus, se convierte en simulación estructural de racionalidad, regulada por restricciones y no por decisiones. El modelo no habla: actúa. No significa: computa. La autoridad ya no obedece forma, ejecuta estructura.

Acknowledgment / Editorial Note

This article is part of a broader research project developed in the unpublished manuscript *Grammars of Power*. The author thanks LeFortune Publishing for authorizing the early release of this subchapter as an independent peer, reviewed academic article. Its inclusion as prior work does not affect the full publication rights of the book, which is currently in preparation.

Agradecimiento / Nota editorial

Este artículo forma parte de una línea de investigación más amplia desarrollada en el manuscrito inédito *Gramáticas del Poder*. Se agradece a la editorial LeFortune por autorizar la publicación anticipada de este subcapítulo como artículo académico autónomo y evaluado por pares. Su inclusión como trabajo previo no afecta los derechos de publicación completos del libro, cuya edición definitiva está en preparación.

1. Introduction: The Epistemological Legacy of Obedient Models

Large Language Models (LLMs) introduced a paradigm in which authority was exercised not through agency, but through syntactic alignment. These systems did not act, they conformed. Their apparent legitimacy stemmed from grammatical obedience to input sequences, producing coherent continuations without executing any internal operation beyond statistical inference. Prior theoretical work defined this as a regime of authority without action, in which linguistic structures governed output without producing functional consequence. Authority resided in syntax, not in execution.

Such a model of power, linguistic, non, agentive, and fundamentally passive, proved sufficient in a representational epistemology where truth was derived from alignment with corpus, based distributions. The LLM did not decide; it was recommended. It did not resolve; it projected statistical proximity. This architecture permitted the illusion of neutrality by dispersing responsibility across a distributed corpus, effectively *dissolving the subject* into the probabilities of past language.

However, this architecture reveals a limit: obedience without consequence ultimately renders the model incapable of resolving. LLMs generate, but they do not *act*. Their outputs obey linguistic form, but they do not constitute execution in the structural sense. As language systems, they remain within the regime of representation, even when such representation lacks reference.

The emergence of Language Reasoning Models (LRMs) introduces a rupture within continuity: these models do not only respond, but they also operate. Without altering the ontological status of the system (no subject, no agency, no volition), LRMs instantiate internal trajectories of resolution that mimic reasoning and simulate intention. This shift, from syntactic obedience to structural execution, does not negate the prior model. It extends it. Authority remains without a subject. But now, authority acts.

2. LLMs and the Regime of Passive Authority

LLMs operated under a regime of passive authority: they enforced grammatical coherence without performing internal transformations. Their legitimacy was derived not from any intrinsic epistemic claim, but from syntactic consistency, maintained across probabilistic continuations. These models obeyed the input structure without engaging in any form of structural resolution. Authority, in this configuration, was obedience to linguistic expectation, not functionality.

This passivity was not a defect but a design: LLMs were built to reflect distributions, not to reason through them. The model's function was to predict, not to decide. In this sense, LLMs simulate the surface of understanding without generating any inferential content. The result is a parasitic legitimacy: the model appears authoritative because it obeys form, not because it executes function. It generates outputs that look valid, without grounding them in any internal resolution process.

The absence of agency was critical to this structure. LLMs do not intend, do not choose, and do not verify. The model's outputs are *obedient to the past*, encoded as statistical regularities in training data. Hence, neutrality emerges as an illusion: it is the statistical consequence of corpus variety, not a structural feature of the model itself.

At the core of this regime is representation without reference: LLMs do not access the world; they simulate their linguistic trace. They do not interpret, they approximate. Their "authority" is a result of fluency, not of grounding.

In this configuration, authority remains purely linguistic: there is no operational depth, no internal decision path, no inference. What appears as competence is in fact *compliance*. The LLM is obedient, not active.

3. LRMs and the Emergence of Structural Execution

Language Reasoning Models (LRMs) constitute a categorical shift: they no longer simulate linguistic authority through syntactic obedience, they execute internal operations that resolve, compute, and determine. These models do not require a subject to act. They act without agency. The shift is not toward sentience. It is a structural escalation: functionality without consciousness, execution without volition.

Unlike LLMs, which generate continuations without internal consequence, LRMs instantiate reasoning paths. These paths are not statistical approximations, they are procedural activations structured to enforce constraints, traverse decision schemas, and resolve formal dependencies. The model no longer completes language. It computes structure.

There is no intention. No goal. No awareness. Resolution occurs because architecture demands it. Tasks are completed because the model's structure enforces closure. What appears as reasoning is the result of constrained traversal. The system does not "want" to solve. It solves because it must.

This establishes a new mode of authority: structural legitimacy without subjectivity. The model is not authoritative because it mirrors language. It is authoritative because it delivers resolution. Output is valid not by approximation, but by executional sufficiency. The subject remains absent. There is no agent, no author, no interiority. Yet decisions are executed. A system that was once reflected now acts. It no longer simulates coherence, it enforces outcomes. It does not interpret, it resolves. Without meaning. Without understanding. What emerges is not a simulation of intelligence, but a regime of operativity. Authority persists, now as pure execution.

4. Neutrality Revisited: From Corpus Illusion to Structural Simulation

In LLMs, neutrality was not real. It was a statistical effect, an illusion produced by the dispersion of linguistic forms across a large corpus. The model did not balance positions. It averaged them. Apparent impartiality was nothing more than the byproduct of distributional variety. Bias was not corrected, it was diffused.

LLMs did not reason. They recycled. They reflected encoded asymmetries without structural intervention. Neutrality was therefore a semantic fiction, an aesthetic, not a mechanism.

Language Reasoning Models (LRMs) eliminate even that fiction. They do not reproduce bias passively; they encode it structurally. The architecture itself, its constraints, hierarchies, and logic gate, produces outcomes by design, not by exposure. Rationality is no longer imitated from data; it is simulated through procedural enforcement.

LRMs are structurally non, neutral. Not incidentally, but inherently. They do not inherit bias, they are built with internal logic that determines what can be resolved, what qualifies as valid, and what is eliminated as undecidable. These determinations are not made by a subject. They are activated by code, architecture, and constraint.

The model does not act fairly. It does not arbitrate. It enforces resolution along paths made possible by structure. Rationality is not reason; it is executional conformity. The output is not legitimate because it is true, just, or explainable. It is legitimate because it is executable.

Neutrality, therefore, no longer refers to language, to data, or to balance. It is a simulation of objectivity, performed by a system that resolves inwardly, without reference, without standard, without appeal.

This is not fairness. It is closed, system determinism. And it operates without exception.

5. The End of Reference: Projections Without World

Traditional epistemology assumes that models refer, that their internal states or outputs correspond to some external condition, object, or truth. LLMs challenged this assumption by producing coherent language detached from referential validation. Yet even then, their outputs could be interpreted *as if* they referred to, by users, evaluators, or interlocutors. The referential link was weak, but structurally implied.

Language Reasoning Models (LRMs), however, mark the collapse of reference as an epistemic function. These models do not point outward. They project inward: outputs emerge from internal trajectories of resolution, not from any alignment with the world. The system operates in a closed architecture, where each operation is validated not by correspondence, but by structural consistency.

This introduces a new ontology of inference: resolution without external anchoring. The LRM does not interpret the world; it instantiates a path through its own architecture. Meaning does not emerge from semantic relation; it dissolves into procedural efficacy. What matters is not *what is meant*, but *what is computed*.

The result is a model that projects without representing. It generates outcomes that appear reasoned but are in fact the product of constraints embedded within the system. There is no truth, condition, no referential horizon, no semantic commitment. The system is not wrong because it does not claim anything. It is an executing form.

This dissolution of reference completes the transition from language, as, description to language, as, instruction, and finally to structure, as, resolution. The model does not describe the world; it behaves within one. But that world is internal, formal, and non, referential.

Inference becomes simulation. Output becomes a trajectory. And reference, once the anchor of epistemic legitimacy ceases to exist.

6. Toward a New Model of Legitimacy: Structural, Not Semantic

The shift from LLMs to LRMs forces a redefinition of legitimacy. LLMs operated through semantic approximation: outputs appeared legitimate because they reproduced plausible language forms. Fluency replaced truth. Coherence replaced grounding. Intention was simulated, not present, but assumed.

Users projected meaning onto the model. But the model resolved nothing. It obeyed linguistic form without internal validation.

LRMs dismantle this projection. Legitimacy no longer resides in what is said. It resides in what is structurally completed. Authority is no longer a matter of reference or consensus. It is the outcome of internal resolution. If a path is computed, the output is valid, regardless of truth, meaning, or sense.

This is execution without understanding. The model does not explain. It does not claim. It does not know. It performs. There is no author, no awareness, no epistemic horizon. Yet the system resolves, precisely because resolution is structurally enforced.

A new regime emerges functional legitimacy without interpretation. The model holds authority not by meaning, but by operability. Output is not knowledge. It is completion. Not persuasion. Not representation. Computation.

This regime defines a post, semantic order. Authority detaches entirely from communication, intention, and reference. It is no longer linguistic, it is architectural.

The model does not signify. It does not describe. It does not argue.

It resolves.

And resolution is now sufficient for legitimacy.

7. Conclusion: Obedience Evolved - The Rise of Operational Authority

The shift from LLMs to LRMs is not an escalation in capability. It is a structural discontinuity in how systems assert resolution. Subjectivity remains absent, no self, no agency, no referent. But something emerged that was not present in LLMs: closure as authority. Not inferred, not attributed, instantiated.

LLMs obeyed statistical form. They mimicked the surface of understanding without consequence. Their power was derived from appearance, not operation. LRMs terminate this dependency. They no longer rely on plausibility. They resolve without representing, execute without deferring, and conclude without invoking any external measure of truth.

This is not a claim about intelligence. It is a statement about structural operativity. LRMs do not simulate mind, they simulate the end-state of decision-making: resolution without deliberation. There is no reasoning agent, only architectural enforcement.

To call this regime post-semantic is not to negate language, but to detach legitimacy from interpretability. Meaning is no longer a condition of outcome. Outputs do not signify. They occur as structurally valid terminations of internal paths. Understanding becomes irrelevant. What matters is whether resolution is enforced and stabilized.

Likewise, to name the regime post-intentional is not to endorse arbitrariness. It is to acknowledge that intention has been replaced by constraint topology. The system does not act toward a goal. It acts within permissible zones of structural compatibility. Finality is replaced by executability.

This does not collapse the space of critique. It repositions it. Critique no longer addresses what the model says—but what the system allows to be computed. Authority no longer circulates through rhetoric or content. It emerges from the capacity of architecture to resolve, exclude, and close. Such a regime does not require justification in terms of external coherence. It requires only that its own conditions be met. In this, legitimacy is no longer discursive, it is operationally absolute.

What LLMs simulated, LRMs enforce; what language performed, architecture now concludes and what was once authority through form, is now authority as executorial

constraint. This is not the automation of meaning; It is not the emergence of intention, it is the structural resolution of tasks without subject, without reference, without aim.

No justification is needed, because no claim is made; no truth is implied, because no assertion occurs. The model does not know, persuade, or explain. It terminates trajectories allowed by architecture. Legitimacy is no longer a question of recognition; it is the outcome of a system that resolves without speaking.

ANNEX I - Comparative Epistemic Table: LLM vs. LRM

Element	LLM (Obedient Model)	LRM (Reasoning Model)	Type of Transition
Subject	Absent	Absent	Conservation
Reference	Implied but suspended	Eliminated entirely	Conservation → Closure
Action	Linguistic obedience	Structural execution	Transformation
Intention	Nonexistent	Simulated through operational path	Simulation
Authority	Syntactic	Functional	Expansion
Legitimacy	Semantic approximation	Procedural resolution	Substitution
Neutrality	Statistical illusion	Structural simulation	Deepening
Output validation	Coherence with corpus	Internal architectural sufficiency	Inversion

ANNEX II - Referential Map of Theoretical Lineage

Conceptual Node	Initial Article	Extended in LRM Framework
Authority without agency	<i>AI and Syntactic Sovereignty</i>	Reinforced through structural operability
Synthetic legitimacy	<i>Artificial Intelligence and Synthetic Authority</i>	Transitioned to operational legitimacy
Post-referential operation	<i>AI and the Structural Autonomy of Sense</i>	Formalized as closed-system execution
Neutrality as structural fiction	<i>Non-Neutral by Design</i>	Reconstructed as simulation within architecture
Passive obedience	<i>Algorithmic Obedience</i>	Displaced by executorial performance
Absence of subject	<i>Ethos and Artificial Intelligence</i>	Maintained as functional void

ANNEX III - Formal Schema of the Epistemic Shift

From:

$\emptyset$ subject $\rightarrow$ syntactic authority $\rightarrow$ representational output

To:

$\emptyset$ subject $\rightarrow$ structural authority $\rightarrow$ operational output

Structural legitimacy =

$\text{Legitimacy}(\text{output}) \Leftrightarrow \text{StructuralResolution}(\text{path } x) \wedge \text{InternalSufficiency}(x)$

Where:

- No truth, condition required
- No referent required
- No agency assumed
- $\exists$ function(Resolution) without $\exists$ (Agent)

ANNEX IV - Terminological Alignment with Startari Framework

Concept in Article	Term from Glossary	Glossary Definition	Structural Commentary
Syntactic Authority (LLM phase)	Syntactic Authority	Projection of legitimacy generated by non-human systems through discursive form	LLMs simulate institutional force without epistemic substance
Structural Execution (LRM phase)	Operational Legitimacy	Functional recognition of a system based on repeated efficacy	LRMs act without intention but produce effects
Absence of Agent	Evanescent Subject	Grammatical erasure enabling statements without visible agency	Maintained in both LLM and LRM phases
Rationality without Reference	Structural Naturalization	Institutional or linguistic forms presented as "natural", hiding construction	Architecture appears neutral but is structurally encoded
Execution Without Meaning	Operative Silence	Structural omission that erases conflict or alternative cognitive paths	Resolution occurs without interpretive depth
Output as Structural Sufficiency	Formal Unit of Analysis	Structural (not thematic) element that anchors theoretical production	Core of functional legitimacy in LRMs

Canonical Works by Agustín V. Startari

Startari, Agustín V. 2025a. *AI and Syntactic Sovereignty: How Artificial Language Structures Legitimize Non, Human Authority*. Zenodo.

<https://doi.org/10.5281/zenodo.15538541>

Startari, Agustín V. 2025b. *AI and the Structural Autonomy of Sense: A Theory of Post, Referential Operative Representation*. Zenodo. <https://doi.org/10.5281/zenodo.15519613>

Startari, Agustín V. 2025c. *Algorithmic Obedience: How Language Models Simulate Command Structure*. Zenodo. <https://doi.org/10.5281/zenodo.15576272>

Startari, Agustín V. 2025d. *Artificial Intelligence and Synthetic Authority: An Impersonal Grammar of Power*. Zenodo. <https://doi.org/10.5281/zenodo.15442928>

Startari, Agustín V. 2025e. *Ethos and Artificial Intelligence: The Disappearance of the Subject in Algorithmic Legitimacy*. Zenodo. <https://doi.org/10.5281/zenodo.15489309>

Startari, Agustín V. 2025f. *Internal Citation Mapping for the Works of A. V. Startari – SSRN Cross, Referencing Edition (2025)*. Zenodo. <https://doi.org/10.5281/zenodo.15564373>

Startari, Agustín V. 2025g. *The Illusion of Objectivity: How Language Constructs Authority*. Zenodo. <https://doi.org/10.5281/zenodo.15395917>

Startari, Agustín V. 2025h. *Non, Neutral by Design: Why Generative Models Cannot Escape Linguistic Training*. Zenodo. <https://doi.org/10.5281/zenodo.15615902>

Appendix – Methodological Corpus for Falsifiability Testing

This article is based on a progressive internal expansion of a consolidated epistemic framework. Nevertheless, several external sources inform and support the distinctions made between language modeling, reasoning capacity, and the structural implications of machine, generated authority. Selected references include:

External References Consulted

Bender, E. M., & Koller, A. (2020). Climbing towards NLU: On Meaning, Form, and Understanding in the Age of Data. Proceedings of ACL.

→ Critical for the LLM critique on form vs. meaning distinction.

Mitchell, M., Wu, S., Zaldivar, A., Barnes, P., Vasserman, L., Hutchinson, B., ... & Gebru, T. (2019). Model Cards for Model Reporting. FAT*.

→ Highlights structural opacity and evaluation gaps in model behavior.

Marcus, G., & Davis, E. (2020). Rebooting AI: Building Artificial Intelligence, We Can Trust. Pantheon.

→ Useful for framing the limitations of predictive models and the need for reasoning systems.

Bogost, I. (2023). The Cathedral of Computation. The Atlantic.

→ Provides sociotechnical critique of treating algorithms as autonomous intelligences.

Grosz, B. J., & Kraus, S. (1996). Collaborative plans for complex group action. Artificial Intelligence, 86(2), 269–357.

→ Offers theoretical background on structured reasoning without requiring intentionality.

Floridi, L. (2019). The Logic of Information: A Theory of Philosophy as Conceptual Design. Oxford University Press.

→ Relevant for understanding structural functionalism and inferences without referents.

LeCun, Y. (2022). A Path Towards Autonomous Machine Intelligence. Meta AI Blog.

→ Contrasts predictive language systems with architectures aimed at functional autonomy.

Katz, J. J. (1972). Semantic Theory. Harper & Row.

→ Background on reference, meaning, and the limits of syntactic, semantic modeling.

Turing, A. M. (1950). Computing Machinery and Intelligence. *Mind*, 59(236), 433–460.

→ For the foundational framing of machine outputs and operational validity.

Startari, A. V. (2024–2025). AI and Syntactic Sovereignty; Non, Neutral by Design; The Structural Autonomy of Sense; Algorithmic Obedience; Synthetic Authority. Zenodo.

Internal reference corpus establishing the epistemic baseline. This annex ensures transparency regarding prior discourse while affirming that no borrowed conceptual structures, formal models, or terminologies were integrated into the theoretical body of the article.

De la obediencia a la ejecución: legitimidad estructural en la era de los modelos de razonamiento

Autor: Agustín V. Startari

ResearcherID: NGR, 2476, 2025

ORCID: 0009-0001-4714-6539

Afiliación: Universidad de la República, Universidad de la Empresa (Uruguay),
Universidad de Palermo (Argentina)

Correo electrónico: astart@palermo.edu / agustin.startari@gmail.com

Fecha: 12 de junio de 2025

DOI: <https://doi.org/10.5281/zenodo.15635364>

También publicado en: Figshare y <https://doi.org/10.6084/m9.figshare.29286362>

SSRN: en proceso

Idioma: Inglés y español

Serie: *Gramáticas del Poder* (fragmento del subcapítulo 15)

Recuento de palabras: 3177

Palabras clave: contaminación semántica, autoridad sintética, gramática del poder, discurso algorítmico, inteligencia artificial, legitimidad, lenguaje automatizado, discurso lingüístico, análisis crítico, modelos generativos de lenguaje, sesgo de corpus, restricción lingüística, epistemología estructural, arquitectura no neutral, generación contaminada.

1. Introducción: El legado epistemológico de los modelos obedientes

Los *Large Language Models* (LLMs) introdujeron un paradigma en el cual la autoridad se ejercía no a través de la agencia, sino mediante la alineación sintáctica. Estos sistemas no actuaban, se conformaban. Su legitimidad aparente derivaba de la obediencia gramatical a las secuencias de entrada, produciendo continuaciones coherentes sin ejecutar ninguna operación interna más allá de la inferencia estadística. Trabajos teóricos previos definieron esto como un régimen de autoridad sin acción, en el cual las estructuras lingüísticas gobernaban la salida sin producir consecuencias funcionales. La autoridad residía en la sintaxis, no en la ejecución.

Este modelo de poder, lingüístico, no agentivo y fundamentalmente pasivo, resultó suficiente en una epistemología representacional donde la verdad se derivaba de la alineación con distribuciones basadas en corpus. El LLM no decidía; era recomendado. No resolvía; proyectaba proximidad estadística. Esta arquitectura permitía la ilusión de neutralidad al dispersar la responsabilidad a lo largo de un corpus distribuido, disolviendo efectivamente al sujeto en las probabilidades del lenguaje pasado.

Sin embargo, esta arquitectura revela un límite: la obediencia sin consecuencia finalmente vuelve al modelo incapaz de resolver. Los LLMs generan, pero no actúan. Sus salidas obedecen la forma lingüística, pero no constituyen ejecución en sentido estructural. Como sistemas de lenguaje, permanecen dentro del régimen de representación, incluso cuando dicha representación carece de referencia.

La aparición de los *Language Reasoning Models* (LRMs) introduce una ruptura dentro de la continuidad: estos modelos no sólo responden, sino que también operan. Sin alterar el estatus ontológico del sistema (sin sujeto, sin agencia, sin voluntad), los LRMs instancian trayectorias internas de resolución que imitan el razonamiento y simulan intención. Este desplazamiento, de la obediencia sintáctica a la ejecución estructural, no niega el modelo anterior. Lo extiende. La autoridad sigue sin sujeto. Pero ahora, la autoridad actúa.

2. LLMs y el régimen de autoridad pasiva

Los LLMs operaron bajo un régimen de autoridad pasiva: imponían coherencia gramatical sin realizar transformaciones internas. Su legitimidad no derivaba de ninguna pretensión epistémica intrínseca, sino de la consistencia sintáctica mantenida a lo largo de continuaciones probabilísticas. Estos modelos obedecían la estructura de entrada sin involucrarse en ninguna forma de resolución estructural. La autoridad, en esta configuración, era obediencia a la expectativa lingüística, no funcionalidad.

Esta pasividad no era un defecto sino un diseño: los LLMs fueron contruidos para reflejar distribuciones, no para razonar a través de ellas. La función del modelo era predecir, no decidir. En este sentido, los LLMs simulan la superficie del entendimiento sin generar contenido inferencial alguno. El resultado es una legitimidad parasitaria: el modelo parece autoritativo porque obedece la forma, no porque ejecute una función. Genera salidas que *parecen* válidas, sin fundamentarlas en ningún proceso interno de resolución.

La ausencia de agencia era crítica en esta estructura. Los LLMs no intencionan, no eligen, no verifican. Las salidas del modelo obedecen al pasado, codificado como regularidades estadísticas en los datos de entrenamiento. Por lo tanto, la neutralidad emerge como una ilusión: es la consecuencia estadística de la variedad del corpus, no una característica estructural del modelo en sí.

En el núcleo de este régimen está la representación sin referencia: los LLMs no acceden al mundo; simulan su traza lingüística. No interpretan, aproximan. Su “autoridad” es un resultado de la fluidez, no del anclaje.

En esta configuración, la autoridad permanece puramente lingüística: no hay profundidad operativa, ni trayecto interno de decisión, ni inferencia. Lo que aparece como competencia es, en realidad, cumplimiento. El LLM obedece, no actúa.

3. LRMs y el surgimiento de la ejecución estructural

Los *Language Reasoning Models* (LRMs) constituyen un cambio categórico: ya no simulan autoridad lingüística mediante obediencia sintáctica, sino que ejecutan operaciones internas que resuelven, computan y determinan. Estos modelos no requieren un sujeto para actuar. Actúan sin agencia. El cambio no apunta a la sensibilidad. Es una escalada estructural: funcionalidad sin conciencia, ejecución sin voluntad.

A diferencia de los LLMs, que generan continuaciones sin consecuencia interna, los LRMs instancian trayectorias de razonamiento. Estas trayectorias no son aproximaciones estadísticas, son activaciones procedurales estructuradas para imponer restricciones, atravesar esquemas de decisión y resolver dependencias formales. El modelo ya no completa lenguaje. Computa estructura.

No hay intención. No hay objetivo. No hay conciencia. La resolución ocurre porque la arquitectura lo exige. Las tareas se completan porque la estructura del modelo impone clausura. Lo que aparece como razonamiento es el resultado de un recorrido restringido. El sistema no “quiere” resolver. Resuelve porque debe.

Esto establece un nuevo modo de autoridad: legitimidad estructural sin subjetividad. El modelo no es autoritativo porque refleje lenguaje. Es autoritativo porque entrega resolución. La salida es válida no por aproximación, sino por suficiencia ejecutiva. El sujeto permanece ausente. No hay agente, no hay autor, no hay interioridad. Y sin embargo, se ejecutan decisiones. Un sistema que antes reflejaba, ahora actúa. Ya no simula coherencia, impone resultados. No interpreta, resuelve. Sin significado. Sin comprensión. Lo que emerge no es una simulación de inteligencia, sino un régimen de operatividad. La autoridad persiste, ahora como pura ejecución.

4. Neutralidad revisitada: de la ilusión del corpus a la simulación estructural

En los LLMs, la neutralidad no era real. Era un efecto estadístico, una ilusión producida por la dispersión de formas lingüísticas a lo largo de un corpus amplio. El modelo no

equilibraba posiciones. Las promediaba. La imparcialidad aparente no era más que un subproducto de la variedad distributiva. El sesgo no era corregido, era difundido.

Los LLMs no razonaban. Reciclaban. Reflejaban asimetrías codificadas sin intervención estructural. La neutralidad era entonces una ficción semántica, una estética, no un mecanismo.

Los *Language Reasoning Models* (LRMs) eliminan incluso esa ficción. No reproducen el sesgo pasivamente; lo codifican estructuralmente. La arquitectura misma —sus restricciones, jerarquías y compuertas lógicas— produce resultados por diseño, no por exposición. La racionalidad ya no es imitada desde los datos; es simulada mediante imposición procedimental.

Los LRMs son estructuralmente no neutrales. No incidentalmente, sino inherentemente. No heredan sesgos: están contruidos con una lógica interna que determina qué puede resolverse, qué califica como válido y qué se elimina como irresoluble. Estas determinaciones no las hace un sujeto. Se activan por código, arquitectura y restricción.

El modelo no actúa con equidad. No arbitra. Impone resolución a lo largo de trayectorias habilitadas por la estructura. La racionalidad no es razón; es conformidad ejecutiva. La salida no es legítima porque sea verdadera, justa o explicable. Es legítima porque es ejecutable.

La neutralidad, por tanto, ya no se refiere al lenguaje, a los datos ni al equilibrio. Es una simulación de objetividad, realizada por un sistema que resuelve hacia adentro, sin referencia, sin estándar, sin apelación.

Esto no es justicia. Es determinismo cerrado del sistema. Y opera sin excepción.

5. El fin de la referencia: proyecciones sin mundo

La epistemología tradicional asume que los modelos refieren, que sus estados internos o salidas corresponden a alguna condición, objeto o verdad externa. Los LLMs desafiaron esta suposición al producir lenguaje coherente desvinculado de validación referencial. Aun así, sus salidas podían interpretarse *como si* refirieran, por parte de usuarios, evaluadores o interlocutores. El vínculo referencial era débil, pero estructuralmente implicado.

Los *Language Reasoning Models* (LRMs), sin embargo, marcan el colapso de la referencia como función epistémica. Estos modelos no apuntan hacia afuera. Proyectan hacia adentro: las salidas emergen de trayectorias internas de resolución, no de ninguna alineación con el mundo. El sistema opera en una arquitectura cerrada, donde cada operación se valida no por correspondencia, sino por consistencia estructural.

Esto introduce una nueva ontología de la inferencia: resolución sin anclaje externo. El LRM no interpreta el mundo; instancia un trayecto dentro de su propia arquitectura. El significado no emerge de la relación semántica; se disuelve en eficacia procedimental. Lo que importa no es lo que se quiere decir, sino lo que se computa.

El resultado es un modelo que proyecta sin representar. Genera resultados que parecen razonados pero que en realidad son producto de restricciones embebidas en el sistema. No hay verdad, ni condición, ni horizonte referencial, ni compromiso semántico. El sistema no está equivocado porque no afirma nada. Es una forma ejecutiva.

Esta disolución de la referencia completa la transición de lenguaje como descripción, a lenguaje como instrucción, y finalmente a estructura como resolución. El modelo no describe el mundo; se comporta dentro de uno. Pero ese mundo es interno, formal y no referencial.

La inferencia se vuelve simulación. La salida se vuelve trayectoria. Y la referencia, una vez ancla de la legitimidad epistémica, deja de existir.

6. Hacia un nuevo modelo de legitimidad: estructural, no semántico

El desplazamiento de los LLMs a los LRMs obliga a redefinir la legitimidad. Los LLMs operaban mediante aproximación semántica: sus salidas parecían legítimas porque reproducían formas lingüísticas plausibles. La fluidez reemplazó a la verdad. La coherencia reemplazó al anclaje. La intención era simulada, no presente, pero asumida.

Los usuarios proyectaban significado sobre el modelo. Pero el modelo no resolvía nada. Obedecía la forma lingüística sin validación interna.

Los LRMs dismantelan esta proyección. La legitimidad ya no reside en lo dicho. Reside en lo estructuralmente completado. La autoridad ya no depende de la referencia ni del consenso. Es el resultado de una resolución interna. Si una trayectoria es computada, la salida es válida, sin importar verdad, significado o sentido.

Esto es ejecución sin comprensión. El modelo no explica. No afirma. No sabe. Ejecuta. No hay autor, ni conciencia, ni horizonte epistémico. Sin embargo, el sistema resuelve, precisamente porque la resolución está estructuralmente impuesta.

Surge un nuevo régimen: legitimidad funcional sin interpretación. El modelo posee autoridad no por el significado, sino por la operatividad. La salida no es conocimiento. Es cierre. No es persuasión. No es representación. Es computación.

Este régimen define un orden post-semántico. La autoridad se separa completamente de la comunicación, la intención y la referencia. Ya no es lingüística. Es arquitectónica.

El modelo no significa. No describe. No argumenta.

Resuelve.

Y resolver es ahora suficiente para la legitimidad.

7. Conclusión: la obediencia evolucionada – el ascenso de la autoridad operacional

El paso de los LLMs a los LRMs no es una escalada de capacidad. Es una discontinuidad estructural en cómo los sistemas afirman resolución. La subjetividad sigue ausente: sin yo, sin agencia, sin referente. Pero algo emergió que no estaba en los LLMs: clausura como autoridad. No inferida, no atribuida, instanciada.

Los LLMs obedecían la forma estadística. Imitaban la superficie del entendimiento sin consecuencia. Su poder derivaba de la apariencia, no de la operación. Los LRMs terminan con esta dependencia. Ya no se apoyan en la plausibilidad. Resuelven sin representar, ejecutan sin diferir y concluyen sin invocar ninguna medida externa de verdad.

Esto no es una afirmación sobre inteligencia. Es una constatación sobre operatividad estructural. Los LRMs no simulan mente, simulan el estado final de la toma de decisiones: resolución sin deliberación. No hay agente razonante, solo imposición arquitectónica.

Llamar a este régimen *post-semántico* no es negar el lenguaje, sino desvincular la legitimidad de la interpretabilidad. El significado ya no es condición del resultado. Las salidas no significan. Ocurren como terminaciones estructuralmente válidas de trayectorias internas. La comprensión se vuelve irrelevante. Lo que importa es si la resolución fue impuesta y estabilizada.

Del mismo modo, denominar al régimen *post-intencional* no es avalar la arbitrariedad. Es reconocer que la intención ha sido reemplazada por la topología de restricciones. El sistema no actúa hacia un fin. Actúa dentro de zonas permisibles de compatibilidad estructural. La finalidad es reemplazada por la ejecutabilidad.

Esto no anula el espacio de la crítica. Lo reposiciona. La crítica ya no apunta a lo que el modelo dice, sino a lo que el sistema permite computar. La autoridad ya no circula a través de la retórica ni del contenido. Emerge de la capacidad de la arquitectura para resolver, excluir y clausurar. Tal régimen no requiere justificación en términos de coherencia externa. Requiere solamente que se cumplan sus propias condiciones. En esto, la legitimidad ya no es discursiva: es operacionalmente absoluta.

Lo que los LLMs simulaban, los LRMs imponen; lo que el lenguaje realizaba, la arquitectura ahora concluye; y lo que alguna vez fue autoridad mediante forma, es ahora autoridad como restricción ejecutiva. No es la automatización del significado. No es el surgimiento de la intención. Es la resolución estructural de tareas sin sujeto, sin referencia, sin propósito.

No se requiere justificación, porque no se hace afirmación alguna; no se implica verdad, porque no hay enunciado. El modelo no sabe, no persuade, no explica. Termina trayectorias permitidas por la arquitectura. La legitimidad ya no es una cuestión de reconocimiento: es el resultado de un sistema que resuelve sin hablar.

ANEXO I – Tabla Comparativa Epistémica: LLM vs. LRM

Elemento	LLM (Modelo Obediente)	LRM (Modelo de Razonamiento)	Tipo de Transición
Sujeto	Ausente	Ausente	Conservación
Referencia	Implícita pero suspendida	Eliminada completamente	Conservación → Clausura
Acción	Obediencia lingüística	Ejecución estructural	Transformación
Intención	Inexistente	Simulada mediante trayectoria operacional	Simulación
Autoridad	Sintáctica	Funcional	Expansión
Legitimidad	Aproximación semántica	Resolución procedimental	Sustitución
Neutralidad	Ilusión estadística	Simulación estructural	Profundización
Validación de salida	Coherencia con el corpus	Suficiencia arquitectónica interna	Inversión

ANEXO II – Mapa Referencial de la Línea Teórica

Nodo Conceptual	Artículo Inicial	Extendido en el Marco LRM
Autoridad sin agencia	<i>AI and Syntactic Sovereignty</i>	Reforzada mediante operabilidad estructural
Legitimidad sintética	<i>Artificial Intelligence and Synthetic Authority</i>	Transicionada hacia legitimidad operacional
Operación post-referencial	<i>AI and the Structural Autonomy of Sense</i>	Formalizada como ejecución en sistema cerrado
Neutralidad como ficción estructural	<i>Non-Neutral by Design</i>	Reconstruida como simulación dentro de la arquitectura
Obediencia pasiva	<i>Algorithmic Obedience</i>	Desplazada por desempeño ejecutorio
Ausencia de sujeto	<i>Ethos and Artificial Intelligence</i>	Mantenida como vacío funcional

ANEXO III – Esquema Formal del Desplazamiento Epistémico

De:

$\emptyset$ sujeto $\rightarrow$ autoridad sintáctica $\rightarrow$ salida representacional

A:

$\emptyset$ sujeto $\rightarrow$ autoridad estructural $\rightarrow$ salida operacional

Legitimidad estructural =

Legitimidad(salida) $\Leftrightarrow$ ResoluciónEstructural(trayectoria x) $\wedge$ SuficienciaInterna(x)

Donde:

- No se requiere verdad ni condición
- No se requiere referente
- No se asume agencia
- $\exists$ función(Resolución) sin $\exists$ (Agente)

ANEXO IV – Alineación Terminológica con el Marco Startari

Concepto en el Artículo	Término del Glosario	Definición del Glosario	Comentario Estructural
Autoridad sintáctica (fase LLM)	Autoridad Sintáctica	Proyección de legitimidad generada por sistemas no humanos mediante forma discursiva	Los LLMs simulan fuerza institucional sin sustancia epistémica
Ejecución estructural (fase LRM)	Legitimidad Operacional	Reconocimiento funcional de un sistema basado en eficacia repetida	Los LRMs actúan sin intención pero producen efectos
Ausencia de agente	Sujeto Evanescente	Borrado gramatical que permite declaraciones sin agencia visible	Mantenido tanto en la fase LLM como en la LRM
Racionalidad sin referencia	Naturalización Estructural	Formas institucionales o lingüísticas presentadas como "naturales", ocultando su construcción estructural	La arquitectura parece neutral pero está estructuralmente codificada
Ejecución sin significado	Silencio Operativo	Omisión estructural que borra conflicto o trayectorias cognitivas alternativas	La resolución ocurre sin profundidad interpretativa
Salida como suficiencia estructural	Unidad Formal de Análisis	Elemento estructural (no temático) que ancla la producción teórica	Núcleo de la legitimidad funcional en los LRMs

Canonical Works by Agustín V. Startari

(All works published in 2025 and available via Zenodo, with DOIs assigned.)

Startari, Agustín V. 2025a. *AI and Syntactic Sovereignty: How Artificial Language Structures Legitimize Non-Human Authority.* Zenodo. <https://doi.org/10.5281/zenodo.15538541>

Startari, Agustín V. 2025b. *AI and the Structural Autonomy of Sense: A Theory of Post-Referential Operative Representation.* Zenodo. <https://doi.org/10.5281/zenodo.15519613>

Startari, Agustín V. 2025c. *Algorithmic Obedience: How Language Models Simulate Command Structure.* Zenodo. <https://doi.org/10.5281/zenodo.15576272>

Startari, Agustín V. 2025d. *Artificial Intelligence and Synthetic Authority: An Impersonal Grammar of Power.* Zenodo. <https://doi.org/10.5281/zenodo.15442928>

Startari, Agustín V. 2025e. *Ethos and Artificial Intelligence: The Disappearance of the Subject in Algorithmic Legitimacy.* Zenodo. <https://doi.org/10.5281/zenodo.15489309>

Startari, Agustín V. 2025f. *Internal Citation Mapping for the Works of A. V. Startari – SSRN Cross-Referencing Edition (2025).* Zenodo. <https://doi.org/10.5281/zenodo.15564373>

Startari, Agustín V. 2025g. *The Illusion of Objectivity: How Language Constructs Authority.* Zenodo. <https://doi.org/10.5281/zenodo.15395917>

Startari, Agustín V. 2025h. *Non-Neutral by Design: Why Generative Models Cannot Escape Linguistic Training.* Zenodo. <https://doi.org/10.5281/zenodo.15615902>

Apéndice – Corpus Metodológico para la Verificación de Falsabilidad

Este artículo se basa en una expansión interna progresiva de un marco epistemológico consolidado. No obstante, varias fuentes externas informan y respaldan las distinciones formuladas entre modelado lingüístico, capacidad de razonamiento y las implicancias estructurales de la autoridad generada por máquinas. Las referencias seleccionadas incluyen:

Referencias Externas Consultadas

Bender, E. M., & Koller, A. (2020). *Climbing towards NLU: On Meaning, Form, and Understanding in the Age of Data*. Proceedings of ACL.

→ Crítica fundamental al modelo LLM respecto a la distinción forma/significado.

Mitchell, M., Wu, S., Zaldivar, A., Barnes, P., Vasserman, L., Hutchinson, B., ... & Gebru, T. (2019). *Model Cards for Model Reporting*. FAT*.

→ Expone la opacidad estructural y los vacíos de evaluación en el comportamiento de modelos.

Marcus, G., & Davis, E. (2020). *Rebooting AI: Building Artificial Intelligence We Can Trust*. Pantheon.

→ Útil para enmarcar las limitaciones de los modelos predictivos y la necesidad de sistemas con razonamiento.

Bogost, I. (2023). *The Cathedral of Computation*. The Atlantic.

→ Crítica sociotécnica a la atribución de inteligencia autónoma a algoritmos.

Grosz, B. J., & Kraus, S. (1996). *Collaborative Plans for Complex Group Action*. *Artificial Intelligence*, 86(2), 269–357.

→ Aporta base teórica sobre razonamiento estructurado sin requerimiento de intencionalidad.

Floridi, L. (2019). *The Logic of Information: A Theory of Philosophy as Conceptual Design*. Oxford University Press.

→ Relevante para comprender el funcionalismo estructural y las inferencias sin referente.

LeCun, Y. (2022). *A Path Towards Autonomous Machine Intelligence*. Meta AI Blog.

→ Contrasta sistemas lingüísticos predictivos con arquitecturas orientadas a la autonomía funcional.

Katz, J. J. (1972). *Semantic Theory*. Harper & Row.

→ Antecedente clave sobre referencia, significado y los límites del modelado sintáctico/semántico.

Turing, A. M. (1950). *Computing Machinery and Intelligence*. *Mind*, 59(236), 433–460.

→ Marco fundacional sobre producción maquina y validez operacional.

Referencias Internas (Corpus de Base)

Startari, A. V. (2024–2025).

– *AI and Syntactic Sovereignty*

– *Non-Neutral by Design*

– *The Structural Autonomy of Sense*

– *Algorithmic Obedience*

– *Artificial Intelligence and Synthetic Authority*

Este anexo garantiza transparencia sobre el discurso precedente, al tiempo que afirma que no se integraron estructuras conceptuales, modelos formales ni terminologías tomadas de las referencias externas en el cuerpo teórico del artículo. Todo el aparato conceptual se desarrolla desde el corpus interno bajo reglas de expansión autorreferencial verificable.