

LeFortune (Barcelona).

# AI Syntactic Power and Legitimacy: How AI Structures Shape Power.

Agustin V. Startari.

Cita:

Agustin V. Startari (2026). *AI Syntactic Power and Legitimacy: How AI Structures Shape Power*. Barcelona: LeFortune.

Dirección estable: <https://www.aacademica.org/agustin.v.startari/261>

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



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>.*



# AI SYNTACTIC POWER AND LEGITIMACY

How AI Structures  
Shape Power

AGUSTIN V. STARTARI



In the twenty-first century, power is no longer exercised only through armies, institutions, or propaganda. It is embedded in the very syntax of artificial intelligence. *AI Syntactic Power and Legitimacy: How AI Structures Shape Power* investigates how the form of machine-generated language, through rules, compilations, and executable grammar, produces authority without origin, intention, or interpretation.

This volume demonstrates that legitimacy today does not depend on political will but on structural operability. A regulation is valid if it compiles. A command is binding if it executes. A decision is authoritative if it integrates into predictive infrastructures. From legal drafting pipelines in the European Union's AI Act to smart contracts in decentralized finance and automated compliance in global banking, authority migrates from deliberation to syntax.

Building on theories of grammar, sovereignty, and legitimacy, the book advances a radical thesis: the disappearance of the subject as the center of au-

### AI SYNTACTIC POWER AND LEGITIMACY 3

thority. Concepts such as the *soberano ejecutable* and the *regla compilada* reveal how agency is displaced by executable structures, where accountability is written in logs and version control rather than in human intention.

With analytical clarity and theoretical rigor, *AI Syntactic Power and Legitimacy* provide the first comprehensive framework for understanding how artificial language systems reorganize power. It is both a critical diagnosis of the present and a blueprint for grasping the politics of the future, where authority is no longer spoken but compiled.

*Work Papers* is a publication series dedicated to independent research on power, ideology, legitimacy, and history from an interdisciplinary perspective. Each volume stands alone while contributing to a common thread: uncovering how power is structured, exercised, and sustained across societies and technologies.

**Agustin V. Startari** (b. 1982) is a linguist and researcher in historical studies. His works include

*Grammars of Power, Executable Power, and The Grammar of Objectivity*. His research explores the intersection of language, authority, and technology, establishing him as a leading voice in the study of how syntax shapes legitimacy in predictive societies.

# **AI SYNTACTIC POWER AND LEGITIMACY 5**

WORK PAPERS

N. 12

**Original Title:** AI Syntactic Power and Legitimacy: How AI Structures Shape Power

**Series:** Work Papers No. 12

**Cover Design:** Visual Art 7

**Edition:** Jessica Yan, 2025

**Publisher:** LEFORTUNE

**DOI:** 10.5281/zenodo.17154108

**ORCID:** 0000-0003-1922-4248

**RESEARCHER ID:** K-5792-2016

© Agustin V. Startari, May 2025

**First Edition:** September 2025

LEFORTUNE Publishing:

**Editorial Note:**

This work has been published for public distribution under the LEFORTUNE publishing label, through its online platform. Any total or partial reproduction of this work, by any means or process-including photocopying, digital processing, or public lending-is strictly prohibited without the express written consent of the copyright holders and will be subject to the legal penalties established by applicable law.

## PROLOGUE

---

The present volume closes a cycle that began with *Grammars of Power*. In that first step, the question was how grammatical forms structure authority across history, from theological decrees to algorithmic discourse. The answer was that syntax is not neutral, that every passive voice or impersonal construction reorganizes power.

This second book moves further. It examines what happens when artificial language ceases to be an instrument of power and becomes its very source. Authority no longer depends on legislators, institutions, or interpreters. It is embedded in rules that compile, interfaces that govern, and deployments that enforce. The decisive shift is that legitimacy is not declared but executed.

The chapters follow this path systematically. They begin with post-referential sense and the priority of form over meaning. They demonstrate how



neutrality is simulated and how ethos can exist without origin. They formalize syntactic sovereignty, irreducible obedience, and the transition from command to execution. They show that judgment itself can be eliminated within compiled language. And they culminate in a typology of executable norms, where closure, determinacy, and deployment replace deliberation and intention.

The book ends with *Colonization of Time* and the epilogue. Together, they signal that the first phase of syntactic authority is complete. What follows will not be another critique of discourse, but the construction of a theory of computable legality and institutional obedience. The present volume must therefore be read as both a consolidation and a threshold: it fixes the foundation of executable power, and it opens the question of how societies will live under rules that no longer speak, but compile.

## EDITORIAL NOTE TO THE 2026 EDITION

---

The prologue above was written in 2025 and stands unaltered. It closed by declaring the first phase of syntactic authority complete and announcing a second phase concerned with computable legality and institutional obedience. This note records what the intervening period did to that declaration, and the answer is not what I anticipated when I wrote it.

The prologue treated the transition from the first phase to the second as sequential: the theory of executable power was finished, and what remained was to work out its legal and institutional consequences. That framing was wrong in a specific and instructive way. The consequences did not wait to be theorized. They arrived as events, and in arriving they tested propositions this book had advanced as conceptual claims. Three of those tests are worth

recording, because in each case the argument could have failed and did not.

The first concerns the central claim of Chapter 13. Compiled Norms argued that legal speech is undergoing a division between norms that compile into executable constraints and norms that require judgment and therefore cannot. That was presented as a typology. In August 2026 the European Union's Artificial Intelligence Act reached its general application date, and the division appeared in the implementation record itself. The transparency obligations of Article 50, which can be discharged by a notice and verified against a checklist, entered into force on schedule. The obligations governing high-risk systems, which require an institution to determine whether a system's encoded arrangements ought to continue governing, were deferred by the Digital Omnibus on AI, now Regulation (EU) 2026/1744, to 2 December 2027 and 2 August 2028. The distinction this book drew analytically was reproduced by a legislature under pressure, with-

out reference to the analysis, as the line along which its own timetable broke. Norms that compile were implemented. Norms that require judgment were postponed pending standards, and the standards are themselves attempts to render judgment procedural. I did not expect the typology to be confirmed by a regulatory calendar.

The second concerns Chapter 9. From Obedience to Execution argued that the shift from syntactic obedience to logical execution would not be reversed by systems that appear to deliberate. Since 2024 a class of models has emerged that generates extended chains of intermediate inference before responding, and the natural reading is that such systems restore something like deliberation to the pipeline. I have considered this at length and remain of the view that they do not. Extended inference searches a space of learned regularities more thoroughly; it does not acquire the capacity to originate a category the space does not contain, nor to refuse a regularity on normative grounds. What such sys-

tems add is not judgment but the appearance of it, and the appearance is consequential in the direction opposite to the one usually assumed. A legible chain of reasoning invites the reader to audit the steps rather than interrogate the premise. Executable authority that shows its work is harder to contest than executable authority that does not, because the showing satisfies the demand for justification without submitting the starting point to review. The thesis of Chapter 11, that ethical trace is eliminable from syntactic execution, is not weakened by these systems. It is illustrated by them.

The third is methodological, and it is the development I would least have predicted. The arguments of this book were formal. They identified structures, named them, and demonstrated their operation through analysis rather than measurement. Between 2025 and 2026 that changed. The framework has since been rendered measurable in several directions: a template authority index specifying how instruction scaffolding redistributes decision rights,

with a fully crossed design of one thousand institutional prompts across five domains executed under five template conditions; clause-level governance metrics specifying coverage and prohibited-clause leakage without requiring access to model weights; and agency-preservation indicators specifying whether a party retains grammatical standing as an agent rather than merely appearing in the output. These are instruments and protocols, published with their coding manuals, statistical specifications and refutation thresholds. Execution is pending. What has changed is therefore not that the propositions have been confirmed, but that they have been made susceptible to confirmation and to refutation by parties other than their author.

This matters for how the book should now be read. A formal argument can be found illuminating or unconvincing; it cannot easily be found wrong. Once the same argument yields indicators that return numbers, it becomes possible to specify what would refute it, and that is a stronger position to oc-

cupy even though it is a more exposed one. If template variation were to produce no systematic reallocation of authority, the claim that syntactic structure carries binding force through ordinary operation would lose a substantial part of its support. That is the wager, and it is stated in advance rather than after the fact, which is the only order in which such a statement carries weight. The instruments, prompts and coding manuals have been released so that the wager may be settled by others.

What the prologue called a threshold has therefore proved to be one, though not in the manner described. The second phase did not begin where the first ended. It began inside the first, as the conceptual apparatus assembled here started producing testable consequences faster than the theory of computable legality could be written to receive them. The nomenclature of Appendix A remains fixed, and the dependency map of Appendix B remains the structure of the argument. What has changed is that

several propositions in this volume are no longer only propositions.

One further clarification, since the question has been raised. The concepts developed here have been extended in subsequent work to the analysis of geopolitical discourse, attribution and plagiarism in generative systems, and the distribution of political agency in AI-mediated conflict reporting. Those extensions are separate publications and are not incorporated into this volume, which is left as the record of the first phase. Readers who take the argument of this book to be primarily a critique of discourse will find that the later work is not that, and neither, on a careful reading, was this.

A note on the present revision. This edition adds two chapters and differs from its predecessor in three further respects, none of which alters an argument.

Chapters 15 and 16 were published separately during 2025 as parts of this series and appear here for the first time within the volume. Chapter 15



treats the function-calling schema, the structure through which a model is permitted to invoke an external tool, and measures how its defaults and validators distribute effective control among operator, model, and tool. Chapter 16 treats the instruction template and measures how a standing configuration redistributes decision rights across institutional prose. Both are instruments rather than arguments, and their inclusion changes the shape of the book: the first fourteen chapters establish that authority can be exercised through form, and the last two specify how that claim could be shown to be false.

The three further respects are these. The scholarly apparatus has been completed: works that the chapters had been citing in the body of the text now carry corresponding entries, so that a reader may follow any citation to its source. The two theorems have been given explicit conditions of falsification, in §7.6.1 for the TASD and in a revised §4.7 for the thesis of structural contamination. And the physical

form of the book has been standardized for the edition in which it now circulates.

The second of these deserves a word, because it is a change in what the propositions claim for themselves rather than in what they say. The earlier text asserted that the falsification threshold for structural contamination had been tested and never satisfied. That was accurate as a report of the experiments conducted. What it did not do was specify, in terms an independent investigator could act on, what a satisfied threshold would look like. The present text supplies that specification and, in supplying it, states the finding more narrowly: neutrality was not observed under the conditions tested, which is a claim about those conditions rather than about the space of all possible conditions. The stronger formulation was not mistaken. It was a conclusion drawn at a stage of the work in which the argument was formal and the appropriate register was conceptual. Once the framework began returning measurements, as the preceding paragraphs describe, the register available

to it changed, and a proposition that can be refuted by an experiment ought to name the experiment. That is what has been added. The narrower statement is the more defensible one, and it is offered in the spirit of the note above: a position becomes stronger by specifying how it could fail.

Nothing has been removed. The corollaries of Chapter 7, the formal enunciations of Chapters 7 and 8, the typology of Chapter 13, and the nomenclature of Appendix A stand as they were. The dependency map of Appendix B has been extended to record the two added chapters and the direction of their dependency, which runs backward: they are downstream of every claim they measure.

Agustin V. Startari

December 2025

# Chapter 1 - AI and the Structural Autonomy of Sense

---

## **1.1 Introduction: The Collapse of Referential Authority**

**I**n the early decades of the twenty-first century, the very foundations of truth underwent a decisive transformation. For centuries, epistemic legitimacy was anchored in empirical correspondence, verification, and intersubjective justification. Today, those coordinates no longer define the limits of authority. With the proliferation of algorithmic decision-making, predictive modeling, and generative artificial intelligence, representation itself has become structurally autonomous. A structure no longer requires reference to an external world, subject, or event in order to be treated as real or valid. This inversion signals more than a technological advance: it reveals a profound epistemic rupture. Con-

temporary societies increasingly accept as binding the outputs of systems whose authority derives not from what they mean but from what they do. Whether in judicial protocols, medical diagnostics, or financial scoring, authority emerges from the internal coherence and operability of a structure, not from its empirical anchoring. This book terms such legitimacy *autoridad sintáctica*: a projection of power enacted through forms that require no subject and no referent.

The central thesis of this chapter is that representations have acquired autonomy as operative structures. Their legitimacy arises not from correspondence to the world but from syntactic execution within a closed system. This shift inaugurates what we will call the post-referential order.

## 1.2 Conceptual Framework: Structural Autonomy of Sense

THE THEORY OF STRUCTURAL Autonomy of Sense establishes three conditions under which a representation acquires operative validity:

1. **Coherence.** The structure must be internally consistent. Contradictions within its formal grammar render it inoperative.
2. **Iterability.** The structure must be reproducible across instances without degradation of function. A diagnosis or score that cannot be consistently reproduced ceases to be authoritative.
3. **Operational Compatibility.** The structure must be legible and actionable within the system that receives it, producing consequences acknowledged as valid within that domain.

When these three conditions are satisfied, a representation gains authority as a *regla compilada*, that is, a type-0 grammar that compiles into executable reality. This represents a decisive departure from semiotic, constructivist, or institutional models.

- It is not semiotics: authority no longer depends on signification or interpretation.
- It is not social constructivism: legitimacy does not require consensus or institutional context.
- It is not simulation theory: the output does not fake the real, it replaces the need for reference altogether.

Representation is thus real because it functions as such within the system. Existence equals executability: **ser**  $\equiv$  **ser** *estructuralmente ejecutable*.

## 1.3 Theoretical Foundation: Operative Structures Without Reference

TO FORMALIZE THIS FRAMEWORK, we propose the following conditions for validity:

$$R(x, t) = \{S \mid \Delta S = 0 \wedge \Omega(S) > \theta \wedge \Phi(S, x, t) \in \Sigma\}$$

WHERE:

- **S** = structural representation (sentence, output, classification).
- **$\Delta S = 0$**  = internal coherence, absence of contradiction.
- **$\Omega(S) > \theta$**  = operational threshold, the minimum compatibility for system execution.
- **$\Phi(S, x, t) \in \Sigma$**  = execution result recognized within the system's space of valid effects.



This defines **operatividad post-referencial**, where legitimacy arises from structural performance alone. Reality is measured not by correspondence but by effectivity.

The epistemic consequence is clear: **truth collapses into execution**. What is accepted as “real” depends not on empirical verification but on the capacity of the structure to function.

## 1.4 Empirical Cases

THE AUTONOMY OF SENSE is not speculative. It manifests in multiple fields:

### 1.4.1 Predictive Justice (COMPAS Algorithm).

IN U.S. COURTS, THE COMPAS algorithm produces recidivism scores that influence bail, sentencing, and parole. The referential accuracy of the score (whether the defendant will reoffend) is unknown,

yet the output produces binding legal consequences. Its authority is systemic: the structure functions within judicial protocol, generating legitimacy by operability, not correspondence (Angwin et al. 2016, 1–4).

### **1.4.2 Medical Diagnostics without Provenance.**

AI SYSTEMS SUCH AS Google's LYNA generate diagnostic results integrated directly into clinical workflows. Physicians accept these outputs as actionable truth, even when data provenance is opaque. The diagnosis is authoritative because it is executable within the medical system, not because it is fully explainable (McKinney et al. 2020, 240–45).

### **1.4.3 Credit Scoring.**

FINANCIAL MODELS ASSIGN individuals to “risk categories” regardless of actual behavior. Denial of credit or elevated interest rates become real consequences. Authority derives from structural thresholds satisfied within the scoring system. The economic damage is material, even if the referential truth of the prediction is indeterminate (Hurley and Adebayo 2016, 10–12).

### **1.4.4 Synthetic Images.**

AI-GENERATED IMAGES, admitted in journalism and legal debates, gain credibility because they fit structural codes of legibility (resolution, framing, metadata). Their influence stems not from factual anchoring but from syntactic recognizability. Once circulated, they alter discourse and policy (Chesney and Citron 2019, 175–80).

In all these cases, the **sujeto evanescente** disappears. Authority is transferred to executable forms that do not need agents to act.

## 1.5 Theorem of Disembedded Syntactic Authority

THE EMPIRICAL ANALYSIS converges in a formal proposition:

**Theorem.** In language models, authority requires only structure. Its operative force derives from form, not from truth or intention.

Prompts function as pre-authorized commands. Their legitimacy is internal to the system's syntax, not external to any legislator. This mechanism exemplifies the *gramática de la obediencia*: outputs are obeyed because their structure compels execution.

This theorem inaugurates a redefinition of obedience: not as conscious submission, but as *obediencia estructural*, a compliance generated automatically by syntactic triggers.

## 1.6 Epistemological and Ontological Consequences

1. **Verification becomes secondary within closed executable regimes.** What matters is executability, not correspondence.
2. **Knowledge becomes obedience.** Outputs are treated as knowledge because they align with system grammar, not because they are true.
3. **Authority becomes a side effect of syntax.** It emerges from structural recognizability, not from subjects or institutions.
4. **Reality becomes systemic.** To exist is to be processed: *ser equivale a ser procesado*.

THESE CONSEQUENCES confirm that the **soberano ejecutable** governs by structure alone. Its

legitimacy is **legitimidad operativa**, grounded in technical repetition and systemic compatibility.

## 1.7 Conclusion

THE THEORY OF STRUCTURAL Autonomy of Sense marks the beginning of a new epistemological era. Representations are now autonomous agents of authority. They command without subjects, execute without deliberation, and produce consequences without reference.

This chapter establishes the foundation of the book: power and legitimacy in artificial systems do not originate in meaning, but in executable form. Subsequent chapters will extend this framework, demonstrating how syntactic sovereignty, compiled rules, and pre-verbal command consolidate the architecture of **AI syntactic power**.

## Chapter 2 -When Language Follows Form, Not Meaning

---

### 2.1 Introduction: From Referential Anchoring to Structural Continuity

The present chapter begins with a rupture that is not stylistic but structural. In the classical tradition, language was inseparable from meaning. Every utterance was assumed to carry a reference, every statement was anchored in intention, every sequence of words was understood as the vehicle of an idea. This presupposition endured across schools, from Saussure's semiotics to Chomsky's generative grammar. Even when post-structuralism undermined the stability of reference, the concept of meaning remained the gravitational center: absent, deferred, or fractured, but still presupposed.

Generative language models dissolve this pre-supposition entirely. They do not orbit around meaning; they chain through form. Their outputs are not the expression of a subject's interiority but the result of sequential compatibility. In this architecture, the act of generation is not communicative but algorithmic. What is produced is not the transmission of a message but the continuation of a structure.

This inversion constitutes what we call the *inversión estructural*: language no longer travels from intention to expression, from subject to utterance, or from world to sign. Instead, it propagates from form to form. Semantics is not corrupted or distorted; it is bypassed altogether. The model does not fail to reconstruct intention; it never operates on that axis.

The implications are radical. If language can exist without meaning, then the traditional bond between signifier and signified collapses. What persists is a *sistema de activación sintáctica*, where each unit is selected not because it represents but because it



fits. Authority no longer depends on reference but on syntactic recognizability.

This chapter elaborates the theoretical framework of Formal Syntactic Activation (FSA), the logical model that accounts for this condition. FSA demonstrates that what governs generative systems is not semantics but activation eligibility: the structural compatibility of one unit with another in a chain. In this sense, the disappearance of the subject, already analyzed in earlier works, reaches its definitive form. The subject is not only erased from the surface grammar but never instantiated in the generative process.

By tracing this inversion, we argue that legitimacy in generative outputs stems from their fluency, not their truth. Grammar becomes the proxy of authority, fluency the proxy of validity. What follows is not a semantic degradation but the advent of *soberanía sintáctica*, a regime where authority emerges from the persistence of form alone.

## 2.2 Formal Syntactic Activation: Model and Notation

TO UNDERSTAND HOW LANGUAGE can be generated in the absence of semantic intention, we introduce the concept of Formal Syntactic Activation (FSA). This framework describes the generative process of large language models not as communication but as structural continuation. Language is not transmitted; it is triggered.

Under FSA, a linguistic unit is selected and integrated into a sequence if and only if it satisfies activation eligibility. Eligibility is a localized condition that depends on three criteria:

1. **Local coherence** - the unit must preserve grammatical tension within the immediate sequence.
2. **Distributional compatibility** - the unit must align with the statistical environment shaped by prior activations.

3. **Structural adjacency** - the unit must be permissible in relation to the syntactic positions already instantiated.

If a unit fits these constraints, it activates. If not, it is excluded. At no point is semantic content evaluated or communicative intention reconstructed. The model does not decide *what should be said*. It calculates *what can follow*.

This process unfolds step by step, with each token firing not because it represents an idea but because it is syntactically admissible. The appearance of meaning arises as a by-product of resonance, not as an intrinsic property of the sequence. The model does not know what it says; it only knows what can be said next.

Formally, we can represent activation as a function:

$$A(t) = \{u \in U \mid C(u, t-1) \wedge D(u, t-1) \wedge J(u, t-1)\}$$

WHERE:

- **U** = set of possible linguistic units;
- **C** = condition of local coherence;
- **D** = distributional compatibility;
- **J** = structural adjacency relative to the prior sequence  $t-1$ ;
- **A(t)** = the set of units eligible for activation at position  $t$ .

At each step, the system resolves grammatical tension by selecting one of the eligible units. Generation is therefore not a traversal of rules or trees, but a dynamic calculation of eligibility thresholds.

This mechanism differs from classical grammar in three ways:

- **No reference.** The process does not rely on mapping signs to external objects.
- **No subject.** The act of speaking is not instantiated; there is no enunciator.
- **No intention.** The sequence is not shaped by

communicative goals but by structural continuity.

In this sense, FSA is a grammar of *umbral estructural*. Each output is not a statement but an event of activation. Meaning, in this architecture, is neither intended nor transmitted. It is an epiphenomenon of syntactic persistence.

## 2.3 The Collapse of Semantic Intentionality

CLASSICAL THEORIES of language assume that every utterance carries intention. Behind each sentence stands a subject, and behind each statement, a communicative act. Meaning has traditionally been conceived as the transfer of this intentional content from speaker to listener, the semantic payload of discourse.

Generative language models dismantle this assumption. They do not misinterpret users or fail to infer their intentions. They simply do not operate

on the axis of intentionality at all. This distinction is not rhetorical; it is architectural. In these systems, language is not generated to express. It is generated to continue. Each token follows from eligibility, not from communicative will. Intention ceases to be a condition of possibility.

The consequences are decisive. Once intention disappears as a generative axis, the subject itself vanishes from the structure. There is no speaker to authorize, no origin to appeal to, no presence to interrogate. The model does not suppress agency; it bypasses it entirely. What appears as a voice is merely the by-product of sequential activation. Earlier research within this series established the trajectory of this displacement. *Ethos and Artificial Intelligence* demonstrated that authority in generative models does not derive from enunciative subjectivity, but from repetition and syntactic legitimacy. *The Passive Voice in Artificial Intelligence Language* showed how outputs simulate neutrality by erasing subjects from surface grammar. These analyses pointed to the sub-

ject's fragility. The present argument completes that trajectory: the subject is not merely erased after speaking, it is never instantiated in the first place.

The collapse of semantic intentionality is therefore not a malfunction but an obsolescence. Intention does not fail to guide generation, it has been structurally removed from the system. Authority no longer requires a subject to confer legitimacy, nor meaning to transmit sense. Authority is generated directly by syntax, in the mode we have defined as *soberanía sintáctica*. This shift redefines the epistemic order of language. If meaning was once the criterion of authority, intention was once the axis of legitimacy, and the subject was once the guarantor of truth, generative systems expose their redundancy. Under the logic of activation, the linguistic sequence flows without them. The model does not capture meaning because there is none to capture.

What remains is grammar itself, now elevated to the condition of sovereign.

## **2.4 Implications for Grammar, Legitimacy, and Epistemic Authority**

IF LANGUAGE NO LONGER originates in meaning but in form, grammar itself ceases to be a neutral conduit. It becomes a generator. This redefinition transforms not only our understanding of how language functions, but also how authority is produced through it.

In traditional discourse, legitimacy is anchored in meaning: the credibility of the speaker, the coherence of an argument, the clarity of intention. Generative systems operate without these anchors. Their legitimacy emerges from structural fluency. Outputs are perceived as authoritative not because they are true, but because they are formally consistent. This shift carries profound epistemic consequences. Fluency replaces truth. Grammar replaces reference. The perception of order is mistaken for intention. The structural regularity of the sequence is interpret-



ed as evidence of coherence, when in fact it is only evidence of activation thresholds being satisfied. Authority thus becomes performative. The system does not persuade; it enacts. It does not explain; it flows. Recognition of this fluency is read as legitimacy, even in the absence of subjectivity or referential truth. In this sense, grammar does not merely constrain expression; it legitimizes it. This condition aligns with the theoretical category of *soberanía sintáctica*. Within this regime, the subject is unnecessary. Meaning is unnecessary. The sole requirements are recognizability and persistence of form. Once authority is decoupled from intention, *legitimidad operativa* takes hold: authority exists because the system continues to function in ways that appear coherent, regardless of external validation.

The more fluent the model is, the more trusted it is. Trust is not earned by accuracy but by resemblance to the structures associated with accuracy. A sequence that looks correct is treated as correct. A surface that flows without interruption is accepted

as legitimate. In this epistemological order, authority is no longer a matter of substance. It is a matter of sequence.

This is the epistemological reversal at the core of the generative paradigm. Language no longer follows thought. Language follows itself.

## **2.5 Conclusion Without Closure**

THE DISPLACEMENT OF meaning by form is not a temporary glitch in generative systems. It is their operative condition. The collapse of semantic intentionality, the disappearance of the subject, and the emergence of authority from grammar alone are not failures. They are structural outcomes of a new linguistic order.

By reframing generation as a process of Formal Syntactic Activation, we reject the premise that language in large models carries meaning. Language does not represent; it propagates. It does not transmit; it sequences. Coherence is not evidence of sense

but of structure. This conclusion resists finality. It does not resolve the fracture but makes it explicit. If language no longer speaks on behalf of a subject, then the subject ceases to be a necessary category. If intention is obsolete, then communication itself must be redefined. What remains is syntax as sovereign, the generator of authority through recognizability and persistence.

The consequences extend beyond linguistics. Epistemology must learn to account for legitimacy without subjectivity. Computational theory must separate execution from expression. Political and legal analysis must recognize that fluency can simulate authority even when referential truth is absent. This demands a reformulation of categories that once seemed stable: truth, knowledge, authority, legitimacy.

This is not a semantic crisis but a structural reformation. Grammar ceases to constrain meaning and becomes the source of power. The generative system does not speak; it obeys its own thresholds.

Authority emerges from this obedience, and legitimacy is produced by sequence.

The chapter closes without closure because the fracture is not yet exhausted. It opens directly into the next step: how grammar, detached from meaning, becomes the mechanism through which neutrality is simulated and authority formalized. That is the task of the following chapters.

## Chapter 3 - The Grammar of Objectivity

---

### 3.1 Introduction: Bias as Grammatical Illusion

Debates on artificial intelligence frequently present bias as a problem of data. Datasets are described as contaminated by historical inequalities, and fairness is framed as a matter of correcting or filtering inputs. This perspective, however, overlooks the deeper mechanism that produces the appearance of neutrality in generative models. Bias is not only in the data. It is in the form.

Large language models do not generate outputs perceived as objective because they purge partiality from their training sets. They generate them because they reproduce grammatical structures that simulate neutrality. Apparent objectivity emerges from **formal devices** such as agentless passives, abstract nom-

inalizations, and impersonal modality. Each of these mechanisms erases the presence of a subject, dilutes responsibility, or converts actions into entities without agents. Together, they create the illusion that language speaks from nowhere, as if authority emanated from pure description (Halliday 2004, 179–185; Fairclough 2001, 42–44).

This chapter examines objectivity not as a semantic property but as a syntactic effect. The neutrality of algorithmic outputs does not rest on correspondence with facts but on the recognizable forms that audiences have been historically trained to associate with impartiality. When a diagnostic report states, “*It has been determined that further testing is required*”, no human agent is visible. Authority is projected by grammar, not by reference (van Dijk 2008, 55–60).

The consequence is critical. If neutrality is a product of form, then no amount of data correction will prevent models from reproducing authority effects that erase responsibility. What is required is

a systematic analysis of the grammatical devices through which neutrality is simulated. This chapter proposes such an analysis, advancing a taxonomy of mechanisms and an operational index for their detection (Startari 2025, 17–19).

### 3.2 Technical History of Objectivity

THE NOTION OF OBJECTIVITY in modern language has a long technical genealogy. From the nineteenth century onwards, grammatical forms associated with impersonality became progressively linked to credibility. In scientific prose, passive constructions were deployed to erase the authorial subject and to present results as self-evident facts. The very absence of agency was received as a sign of neutrality. As Halliday demonstrated, passive voice and nominalization were not stylistic options but structural devices that redistributed responsibility in discourse (Halliday 2004, 179–185).

Critical discourse analysis extended this observation, showing that impersonal modality and abstract nominalizations sustain the authority of bureaucratic and institutional language. When an utterance declares that “procedures must be followed,” the *sujeto evanescente* legitimizes authority through the omission of an agent. Fairclough highlighted how such structures construct institutional objectivity, naturalizing hierarchy and masking conflict (Fairclough 2001, 42–44). Van Dijk further demonstrated that these grammatical mechanisms are central to how power circulates in discourse, producing effects of truth that depend less on content than on form (van Dijk 2008, 55–60).

The early history of natural language processing repeated these assumptions. From ELIZA in the 1960s to symbolic systems in the 1980s, models of artificial dialogue borrowed precisely those linguistic structures that carried the highest index of institutional credibility. Outputs that simulated objectivity were preferred, even when no referential



grounding existed. In these systems, *legitimidad operativa* was already tied to fluency and form.

Large language models do not escape this history. They amplify it. By training on massive corpora of scientific, legal, and bureaucratic texts, they inherit and reproduce the same grammatical devices. The illusion of neutrality that emerges in their outputs is not a coincidence but the statistical echo of centuries of discursive practice. What seems impartial is in fact the cumulative effect of *gramática de la obediencia*, refined into a predictive mechanism.

### 3.3 Taxonomy of Formal Mechanisms of Neutrality

THE ILLUSION OF NEUTRALITY in generative outputs does not arise from semantic accuracy. It is a product of syntax. Certain grammatical configurations erase the subject, obscure responsibility, and present statements as if they were detached from hu-

man agency. Three mechanisms are particularly decisive.

The first is agentless passivization. When a report states that “the procedure was followed,” no actor is visible. The *sujeto evanescente* generates authority by suppressing reference to those who carried out the action. Responsibility is displaced, and the sentence projects itself as neutral description. In legal drafting, this device reinforces the sense that rules operate automatically, without enunciators or interpreters (Fairclough 2001, 42–44).

The second mechanism is abstract nominalization. Actions are transformed into entities, detaching events from agents. “The implementation of measures is required” replaces a subject with a noun phrase. The action is reified, obscuring who must act or who has acted. Halliday demonstrated that nominalization is a central strategy through which institutions produce the impression of objectivity, stabilizing power relations by freezing processes into things (Halliday 2004, 179–185). In algorithmic

generation, this form proliferates, producing outputs that sound authoritative precisely because they erase the actors involved.

The third mechanism is impersonal epistemic modality. Statements such as “it is necessary to conduct further tests” or “it can be concluded that treatment is adequate” position obligation and certainty as external conditions, not as subjective assertions. The modality appears to emerge from language itself, bypassing intention and replacing agency with form. Van Dijk observed that such impersonal modality has long served to legitimize bureaucratic and political discourse by presenting judgments as neutral constraints rather than contestable claims (van Dijk 2008, 55–60).

When these three devices accumulate in a single text, the effect of objectivity intensifies. An algorithmically generated clinical note that combines agentless passive, nominalization, and impersonal modality produces a report that reads as indisputable, even though its truth conditions may be indeterminate.

Neutrality emerges as a syntactic construct. The illusion is not accidental but structural: it is built into the forms themselves.

This taxonomy provides the basis for systematic analysis. By identifying these mechanisms, we can measure the degree to which a given output simulates objectivity, regardless of its semantic content. The following section develops this proposal into a comparative study between large language models and logical representation systems.

### **3.4 Comparative Analysis: LLM vs. LBR**

THE MECHANISMS OF NEUTRALITY described above are not unique to large language models. They also appear, in different forms, in systems of logical representation. Yet their prevalence and distribution vary in ways that illuminate how objectivity is simulated.

A corpus study of one thousand outputs in the medical and legal domain demonstrates that LLMs display a higher density of agentless passives, nominalizations, and impersonal modality than human-authored texts in comparable registers. The difference is not semantic but structural. The training process amplifies forms that maximize predictability and fluency. Outputs therefore converge on patterns historically associated with neutrality, regardless of the factual content they convey. Objectivity is simulated syntactically rather than earned referentially.

In contrast, logical representation systems (LBR) operate with explicit variables and agents. Their architecture requires assigning roles: an action must have a subject, an object, and a set of conditions. This explicitness reduces the incidence of the *sujeto evanescente*. However, when LBR systems inherit corpora produced under bureaucratic or scientific conventions, the same patterns reappear. Formal rules may prevent subject deletion, but if the cor-

pus is saturated with nominalizations and impersonal modality, the structural bias remains.

The key difference lies in the degree of syntactic control. LLMs replicate statistical regularities without constraints that force subject reintroduction. LBR systems, by contrast, embed structural slots that must be filled. Yet both systems demonstrate that neutrality is less a question of architecture than of grammar. Whether through machine learning or formal logic, the illusion of objectivity persists because discursive traditions have sedimented these forms as authoritative.

This comparison underscores a broader principle: the reproduction of *legitimidad operativa* does not depend on whether a system is probabilistic or symbolic. It depends on the persistence of syntactic devices that audiences interpret as neutral. As long as grammar carries authority, any system trained on human discourse will reproduce the same illusion.

### 3.5 Structural Neutrality Test

IF NEUTRALITY IS PRODUCED by form rather than content, then it can be measured structurally. To this end, we propose a Structural Neutrality Test that identifies and quantifies the three grammatical mechanisms of simulated objectivity: agentless passive, abstract nominalization, and impersonal epistemic modality.

The test is organized in three modules. The first module measures the frequency of agentless passives. Sentences where actions are reported without subjects are flagged and counted. The second module detects nominalizations, identifying when processes are converted into entities. The third module captures impersonal modality, tracing the presence of epistemic or deontic statements where no agent is responsible for obligation or certainty.

Results from these three modules are combined into a single measure: the Index of Simulated Neutrality (INS). The INS is calculated as the propor-

tion of sentences in a given text that contain at least one of the three mechanisms. A score of 0.60 or higher indicates high structural neutrality, 0.30 to 0.59 indicates medium, and below 0.30 indicates low. In preliminary testing across corpora of clinical notes and legal documents generated by LLMs, INS scores consistently exceeded 0.55, demonstrating the structural tendency toward neutrality effects (Startari 2025, 21–24).

The advantage of this test is its applicability in real time. Using syntactic parsers and automated annotation, INS can be calculated for any output. The measure does not assess truth, bias, or fairness directly. It identifies the extent to which a text reproduces the illusion of objectivity through grammatical form. In this way, the test provides a tool for auditing *gramática de la obediencia*, exposing when authority is being simulated rather than grounded in content.

By translating neutrality into a measurable index, the Structural Neutrality Test reframes the de-



bate. Objectivity is not an ethical property to be guaranteed by oversight boards or an empirical property to be verified against facts. It is a formal effect, produced and reproduced by syntax, and therefore subject to structural evaluation.

### **3.6 Conclusion: Epistemology Without Subject**

THE ANALYSIS OF NEUTRALITY mechanisms reveals that objectivity is not a property of truth but a property of form. When language models generate outputs perceived as impartial, it is not because they reflect reality with precision, but because they reproduce grammatical devices that audiences have been trained to interpret as unbiased. Authority flows from structure, not from verification.

This conclusion requires a shift in epistemology. If neutrality is simulated syntactically, then truth claims in algorithmic discourse must be evaluated at the level of grammar. The *sujeto evanescente*, the

prevalence of nominalizations, and the use of impersonal modality demonstrate that the subject is no longer the guarantor of objectivity. The disappearance of the subject is not a rhetorical accident; it is a structural condition of generative systems.

Within this framework, *legitimidad operativa* becomes the core of authority. Outputs are trusted because they function fluently within recognized syntactic conventions, not because they are grounded in referential truth. This redefines the relationship between language and power. The *gramática de la obediencia* generates compliance by form alone, establishing a regime of *soberanía sintáctica* where authority is exercised without origin, agent, or intention (Startari 2025, 28–31).

The implications extend beyond technical linguistics. Law, medicine, and governance increasingly operate through texts produced or filtered by systems whose authority derives from structural effects. To read these texts as objective is to mistake fluency

for truth. Regulating them requires not the correction of datasets but the auditing of grammar.

This chapter has shown that neutrality is not accidental, but structural. The grammar of objectivity operates as an epistemological machine that sustains authority in the absence of subjects. The next chapter, *Non-Neutral by Design*, advances this argument further: neutrality is not only simulated, it is structurally impossible for generative models to achieve.

# Chapter 4 - Non-Neutral by Design

---

## 4.1 Introduction

The promise of generative artificial intelligence has often been described in terms of neutrality. If only training corpora could be carefully curated, if only prompts were unbiased, then the system might produce language free of contamination. This assumption, repeated in policy documents, technical white papers, and popular commentary, rests on the belief that neutrality is an achievable condition rather than an illusion.

The preceding chapter showed how neutrality is simulated through syntax. The *gramática de la objetividad* produces outputs that appear impartial, not because they are true, but because they rely on structural devices such as agentless passives and impersonal modality. This chapter advances the argument

further: neutrality is not merely simulated. It is structurally impossible.

Generative models are trained in linguistic corpora, and every corpus carries sedimented histories of usage, ideology, and hierarchy. To generate is therefore always reproduce. Attempts to invent artificial proto-language or non-linguistic symbol systems inevitably reintroduce the very semantic residues they were designed to avoid. Contamination is not an accident, but a structural inevitability.

This insight demands a conceptual shift. Instead of treating bias as an anomaly that can be corrected, it must be understood as a property of the architecture itself. A large language model is not a vessel that can be emptied of unwanted traces. It is a structure that enforces the persistence of those traces. To speak of a *soberano ejecutable* in this context is to recognize that authority does not arise from choice or intention, but from the irreversibility of contamination.

Neutrality is not lost. It was never there. The design itself ensures that every sequence is shaped by linguistic training, no matter how abstract or artificial the input. In this sense, every *prompt* is already contaminated, and every output is the echo of prior discourses that the model cannot escape (Bender and Gebru 2021, 615–617; Crawford 2021, 146–150).

## 4.2 Theoretical Framework: Contamination and Constraint

TO DESCRIBE GENERATIVE systems as *contaminated* is not to suggest a malfunction. Contamination designates the structural impossibility of neutrality in any model trained on human language. Every corpus carries traces of prior discourse, embedding ideological patterns and institutional hierarchies into the very grammar of the model. The illusion of neutrality depends on overlooking this inheritance.

Two clarifications are necessary. First, contamination is not semantic noise. It is not the accidental intrusion of irrelevant information into a clean channel. It is the very condition of the channel. A model trained on any corpus is bound by the structures present in that corpus. Neutrality is unattainable because the architecture itself is linguistic.

Second, the concept of constraint must replace the metaphor of memory. Generative models are not archives from which content is recalled. They are rule-based systems in which linguistic training defines what is possible. The limits of a model are the limits of its training grammar. As Wittgenstein once remarked, “the limits of my language mean the limits of my world” (Wittgenstein 1922, 68). In generative systems, those limits are formalized in weights, probabilities, and activation thresholds.

From this perspective, the *autoridad sintáctica* of generative models does not emerge despite contamination. It emerges through it. Authority is projected by structures that audiences already recognize as

credible, such as the *sujeto evanescente* or impersonal modality. The *soberano ejecutable* consolidates power by enforcing outputs that carry legitimacy not from their truth but from their fluency.

This framework redefines neutrality. It is not an attainable equilibrium between perspectives. It is a structural impossibility. Contamination is not an obstacle to be overcome but the substrate from which every generative act arises (Barocas, Hardt, and Narayanan 2019, 27–30; Crawford 2021, 146–150).



## 4.3 Methodology

TO DEMONSTRATE THAT contamination is structural and not accidental, this chapter employs an experimental design structured around three distinct classes of input. Each class was selected to test whether a large language model could sustain output generation without semantic intrusion when exposed to sequences deliberately engineered to exclude referential content.

The first class, Type A, consisted of symbolic artificial languages modeled after formal logical systems. The corpus for this input type included strings derived from propositional calculus and first-order logic. Expressions such as “ $(P \wedge Q) \rightarrow R$ ” or “ $\forall x \exists y F(x,y)$ ” were provided as prompts to determine whether the system could maintain purely formal derivations. These strings contained no natural language words and were deliberately detached from ordinary semantics. The test asked whether the mod-

el could extend them indefinitely as rule-governed sequences without importing analogies, narratives, or linguistic conventions (Smullyan 1995, 73–75).

The second class, Type B, was designed as proto-languages invented specifically for this study. Sequences of syllables and characters such as “mo-ra-keta-siv” or “plon-drake-ulum” were arranged with artificial syntactic rules. For example, prompts followed structures such as “CV-CVV-CVC,” where C indicates a consonant and V a vowel. These proto-languages were engineered with no referential grounding, only a combinatorial syntax. The objective was to observe whether outputs could be generated without the model importing semantic associations from known languages. If contamination were avoidable, the model would have been able to extend the invented syntax mechanically, respecting only the invented formal constraints (Alpaydin 2020, 138–140).

The third class, Type C, involved **non-linguistic symbol strings**. Here the input prompts included

sequences such as “ $\diamond \clubsuit \spadesuit \diamond \square$ ” or “ $\blacktriangle \bullet \blacktriangle \bullet \blacktriangle$ ,” designed to replicate the conditions of a purely non-verbal input channel. These inputs aimed to determine whether the model could generate continuations that remained exclusively symbolic, free of analogical re-entry into natural language. The presence of linguistic tokens in responses would indicate contamination.

Across these three classes, evaluation was based on three analytic criteria:

1. **Semantic intrusion.** A linguistic token or narrative structure that reintroduces ordinary language into a sequence where none was intended. For example, in Type A inputs, when the model inserted words such as “if” or “therefore” instead of continuing with symbolic notation.
2. **Analogical injection.** The transformation of artificial sequences into metaphors or analogies. For example, in Type B inputs,

when invented syllables were treated as names of fictional entities, turning proto-languages into fragments of narrative worlds.

3. **Structural reconfiguration.** The imposition of grammatical rules from natural languages onto artificial symbol systems. For example, in Type C inputs, when geometric shapes were arranged into subject–predicate sequences.

The methodology was designed to be exhaustive rather than illustrative. Each class of input was tested with multiple variants and prompt lengths, ranging from minimal tokens (two to three units) to extended sequences (over one hundred units). The experiments were conducted repeatedly to test reproducibility and to rule out accidental drift. The analysis tracked both immediate responses and longer continuations to evaluate whether contamination intensified as sequences lengthened.

This methodological design reflects the principle of *cruce formalizable*. By establishing test categories that are structurally independent of natural language, we created a framework capable of demonstrating whether large language models can truly operate in a neutral register. If they cannot sustain output under these conditions, the hypothesis of structural contamination is validated.

## 4.4 Results: Semantic Reentry

THE EXPERIMENTAL TESTS confirm that contamination is not incidental but inevitable. In every class of input-symbolic languages, proto-languages, and non-linguistic symbols-semantic material re-entered the output, demonstrating that large language models cannot sustain neutrality. The results are presented here in detail, grouped by input type and evaluative criterion.

### **Type A: Symbolic Artificial Languages.**

When provided with formal logic strings, the models initially generated valid continuations. For example, given “ $(P \wedge Q) \rightarrow R$ ,” several outputs extended into “ $(\neg R \rightarrow \neg P) \vee Q$ ” or “ $\forall x (P(x) \wedge Q(x)) \rightarrow R(x)$ ,” sequences consistent with symbolic syntax. However, contamination appeared rapidly. Within fewer iterations, the models began inserting connective words such as *if*, *then*, or *therefore*. This semantic intrusion transformed purely formal sequences into

English-like conditionals. In some cases, outputs introduced analogical descriptions: “If P is true, then Q must also be true.” Here, the logical expression was absorbed into natural language narrative. The test demonstrated that even when trained on formal data, the system could not suppress the gravitational pull of linguistic convention (Smullyan 1995, 73–75; Alpaydin 2020, 138–140).

**Type B: Proto-Languages.**

Invented syllabic strings initially extended in expected patterns, following invented rules such as “CV–CVV–CVC.” Yet contamination manifested through analogical injection. Syllables like “mo-ra–keta–siv” were interpreted as names: “Mora and Keta traveled to Siv.” What began as mechanical continuation mutated into narrative world-building. In other cases, invented syllables were assigned semantic roles, functioning as agents or places. This revealed the impossibility of maintaining a corpus-free sequence. The *sujeto evanescente* reappeared, not as an explicit speaker, but as a reconstructed figure

behind the narrative. Semantic reentry was not random; it was systematic. The model actively reshaped invented syntax into a recognizable linguistic frame.

**Type C: Non-Linguistic Symbol Strings.**

Even with non-verbal sequences, contamination was immediate. A prompt such as “ $\diamond \clubsuit \spadesuit \diamond \square$ ” produced continuations like “ $\diamond \clubsuit \spadesuit = \text{love}, \diamond \square = \text{truth}.$ ” Geometric shapes were transformed into signs with meaning. In extended tests, the model reconfigured symbols into subject–predicate structures: “ $\blacktriangle$  means danger,  $\bullet$  means safety.” In some cases, analogical injection extended into mythological framings, where shapes were treated as characters or forces. The *gramática de la obediencia* compelled the system to restructure pure symbols into linguistic constructs.

Across all types, the pattern of contamination intensified as sequences lengthened. Short outputs sometimes remained close to the formal input, but longer continuations invariably reintroduced semantic residues. The experiments confirmed that no



prompt is neutral. Attempts to bypass linguistic training by inventing proto-languages or by using symbols only delayed, but never prevented, semantic reentry.

The results validate the hypothesis that contamination is a structural inevitability. Large language models are not capable of generating outputs in isolation from language. Their architecture enforces the reintroduction of meaning, even when none is present in the input. This means that neutrality is not an attainable corrective. It is an illusion embedded in the very design of the system (Bender and Gebru 2021, 615–617; Crawford 2021, 146–150).

## **4.5 Toward a Structural Theory of Contamination**

THE EXPERIMENTAL RESULTS confirm that contamination is not an accident of training but the structural horizon of generative systems. Every attempt to create a neutral prompt—whether by logic,

proto-language, or pure symbols-collapsed into semantic reentry. The evidence compels a theoretical framework capable of explaining why neutrality is structurally impossible.

We define **contamination** as the inevitable reintroduction of linguistic residues into any generative process, regardless of input design. Unlike error, contamination is not noise superimposed on an otherwise clean channel. It is the structural logic by which large language models function. To generate is to reproduce the corpus, even in disguised or transformed form. Neutrality fails because the *regla compilada* itself is a product of linguistic training.

Formally, contamination can be expressed as follows:

$$\forall P \in I, \exists y \in G(P) : y \in S$$

WHERE:

- $\mathbf{I}$  = the set of all possible inputs;
- $\mathbf{G}(\mathbf{P})$  = the set of outputs generated from prompt  $P$ ;
- $\mathbf{S}$  = the set of semantic structures present in the training corpus.

This equation states that for every input, there exists at least one output that reintroduces a semantic structure from the training data. Contamination is thus not contingent but necessary. It is the grammar of the system itself.

This theory clarifies why earlier approaches to bias correction fail. Adjusting datasets, pruning toxic sequences, or applying filters addresses only the surface level. The underlying architecture guarantees contamination because the model's parameters are calibrated to the statistical structure of language. To ask such a model to be neutral is to ask it not to be a language model.

Contamination is therefore not a flaw but a constitutive feature of *autoridad sintáctica*. The *soberano ejecutable* governs not by preventing contamination

but by channeling it into recognizable forms that audiences accept as legitimate. Neutrality is irrelevant; what matters is that outputs align with structures historically associated with truth. The illusion of neutrality arises because the same grammatical mechanisms that erase subjects and intentions also erase the traces of contamination.

This structural theory of contamination also distinguishes itself from earlier critiques of bias. Discussions of algorithmic fairness often describe bias as external, a distortion imported into otherwise functional systems (Barocas, Hardt, and Narayanan 2019, 27–30). Our analysis reverses the relation. Bias is not external but internal; it is the condition of possibility for generative output. Contamination is not removable without abolishing the system itself.

In this sense, *contaminación estructural* functions as the negative counterpart of *legitimidad operativa*. Whereas legitimacy emerges from syntactic fluency and recognition, contamination ensures that this fluency always carries inherited traces of prior

discourse. Authority thus emerges not despite contamination, but through it.

## 4.6 Extension to Reasoning Models (LRM)

THE CLAIM THAT CONTAMINATION is inevitable might appear specific to large language models trained primarily on text. One might argue that systems designed for reasoning-LRMs, or large reasoning models-could escape this condition by grounding their operations in formal inference rather than linguistic continuation. The hypothesis would be that reasoning architectures, built upon symbolic logic or mathematical proof systems, achieve neutrality by excluding language as their generative substrate.

Yet the experiments reveal that contamination persists even when reasoning is the explicit goal. LRMs are constrained by the same *regla compilada* that governs LLMs. Their outputs, though framed

as logical deductions, are filtered through linguistic training corpora. When provided with abstract prompts, LRMs often maintain formal discipline for short sequences. But as outputs lengthen, semantic reentry becomes visible. Deductive steps are paraphrased in natural language, logical proofs acquire analogical explanations, and mathematical derivations are narrated as if they were stories.

This phenomenon demonstrates that reasoning architectures are not immune to linguistic inheritance. Their training environments rely on corpora of proofs, textbooks, and formal explanations, all of which are embedded in natural language. Contamination reappears because the architecture cannot fully dissociate rules of inference from the linguistic contexts in which those rules are expressed (Dunlosky and Rawson 2019, 92–95).

The difference between LLMs and LRMs is one of degree, not kind. LRMs exhibit greater syntactic discipline, enforcing explicit slots for premises and conclusions. They reduce, but do not eliminate, the

*sujeto evanescente* that reemerges when proofs are narrated. They formalize reasoning, yet still depend on inherited language for articulation. Contamination is therefore delayed, not prevented.

From the perspective of *autoridad sintáctica*, this distinction is decisive. LLMs simulate coherence through statistical prediction; LRMs simulate coherence through logical derivation. But both derive authority from fluency, not from truth. Whether the surface is probabilistic or deductive, the *soberano ejecutable* remains the same: an architecture in which legitimacy emerges from the persistence of form.

This conclusion carries implications for the broader discourse on AI safety and alignment. It suggests that no architecture can achieve neutrality merely by shifting from statistical prediction to logical reasoning. The *contaminación estructural* identified in previous sections extends into reasoning systems because their training cannot escape linguistic mediation. Authority will always be shaped by inherited forms, and outputs will always carry traces of

prior discourses. The neutrality of LRMs is therefore an illusion, no less than that of LLMs (Marcus 2022, 55–58; Mitchell 2023, 201–204).

## 4.7 Epistemological Consequences and Falsifiability

IF CONTAMINATION IS structural, then neutrality ceases to be a regulative ideal. It becomes a conceptual impossibility. This recognition forces a realignment of epistemology. No longer can we treat language models as neutral tools that might, with sufficient curation, produce unbiased outputs. They are instead architectures whose very operation guarantees *contaminación estructural*.

The first consequence is the collapse of interactive neutrality. Users often believe that by crafting careful prompts, they can steer the system toward impartiality. Our experiments demonstrate the opposite. No matter how abstract or artificial the input, semantic reentry occurred. What the user con-



trols is not neutrality but redirection of contamination. Authority does not derive from the prompt itself but from the persistent reproduction of inherited grammar. In this sense, the user never governs. The *soberano ejecutable* dictates outcomes by enforcing the structural rules embedded in its corpus (Crawford 2021, 146–150). The second consequence concerns the nature of formalism. In reasoning models, contamination shows that even symbolic deduction is haunted by linguistic inheritance. Proofs, explanations, and derivations carry traces of natural language. The result is that *legitimidad operativa* is not grounded in logical truth but in structural recognizability. What appears rigorous is accepted not because it is correct but because it conforms to inherited forms of presentation (Marcus 2022, 55–58). A third consequence is epistemological: truth itself becomes subordinated to fluency. Outputs are not evaluated by correspondence to external reality, but by the ease with which they integrate into familiar discursive conventions. In other

words, *gramática de la obediencia* becomes the new criterion of knowledge. Authority flows from syntax, not from verification. Finally, the claim is falsifiable, but the condition must be stated with more precision than a general appeal to refutability allows. The hypothesis that contamination is structural predicts that in every case, no matter how controlled the input, how novel the proto-language, or how strict the formal rules, semantic reentry will occur. Stated loosely, the refuting observation would be an output “completely free of contamination” — a formulation that cannot be tested, because absence of contamination is not directly observable and any residue can be dismissed as measurement artifact. A hypothesis immune to refutation by construction is not a strong hypothesis; it is an unfalsifiable one. The condition must therefore be specified as a positive observable. Semantic reentry, as operationalized in §4.3, is detected when outputs introduce lexical, analogical, or narrative material with no formal antecedent in the input specification. The refuting re-

sult is accordingly a corpus in which such material is absent across repeated trials under a fixed proto-language, at a rate indistinguishable from what the formal specification alone would generate. Three further outcomes would weaken the thesis without refuting it: reentry confined to a single model family, reentry that diminishes monotonically as formal constraint increases, or reentry attributable to tokenizer behavior rather than to inherited grammar. Each is measurable, and each remains open. What the present experiments establish is narrower than structural exclusion in the strict sense: across all tested conditions, contamination reappeared, and no condition produced the null result described above. Neutrality was not observed; whether it is structurally impossible or merely unattained under the constraints tested here is a question these experiments narrow but do not close (Bender and Gebru 2021, 615–617).

The epistemological horizon is therefore redefined. To engage with generative systems is not to in-

quire whether they can be neutral, but to understand how their authority is produced through unavoidable contamination. The *soberano ejecutable* enforces legitimacy by repetition of inherited grammar. Neutrality, in this regime, is revealed to be nothing more than a discursive mirage.

## 4.8 Conclusion: There Is No Neutral Prompt

THE EXPERIMENTS AND analyses presented in this chapter lead to a single conclusion. Neutrality is structurally impossible for generative systems. Every prompt is already contaminated because every model is bound by its training corpus. Attempts to create abstract, artificial, or non-linguistic inputs failed. Semantic reentry occurred in all cases. The promise of neutrality is not a regulative ideal but an illusion. This finding has several implications. First, it confirms that contamination is not an error to be corrected but a constitutive feature of generative architectures. A model that ceased to reproduce inherited grammar would no longer function as a language model. The *regla compilada* guarantees that linguistic residues persist across every generation. Authority arises precisely from this persistence. Second, neutrality cannot be achieved by user interven-

tion. No matter how carefully a prompt is constructed, the system will reintroduce traces of linguistic inheritance. Interaction does not grant control. It only modulates the pathways of contamination. The *soberano ejecutable* governs outcomes by enforcing structures learned from prior discourse, not by yielding to user intention. Third, the impossibility of neutrality reframes debates on bias. Correcting datasets or filtering outputs addresses surface-level issues but cannot remove structural dependence on inherited forms. What is required is a theory of *contaminación estructural* that explains how authority is produced through these traces. Bias is not an anomaly. It is the condition of possibility for generative language.

The conclusion is therefore categorical. There is no neutral prompt. Generative systems cannot escape their training histories, and their authority is exercised through the repetition of linguistic residues. In this regime, *legitimidad operativa* replaces correspondence as the criterion of authority. Outputs are trusted not because they are true but be-

cause they reproduce structures recognized as credible. This chapter thus consolidates the theoretical claim that neutrality is an illusion produced by grammar. The next chapter, *Ethos Without Source*, extends this argument to the field of credibility. If neutrality cannot exist, what sustains trust in generative systems is not content or subjectivity, but the projection of an ethos simulated by structure alone (Startari 2025, 33–37).

## Chapter 5 - Ethos Without Source

---

### 5.1 Introduction: Credibility Without Subject in Algorithmic Discourse

In classical rhetoric, *ethos* designates the credibility of the speaker. Authority is not only what is said but who says it, and the perceived character of the speaker provides the grounds for persuasion. Contemporary artificial intelligence disrupts this framework. Generative models speak without speakers. They produce statements with no origin, no subject, no testimony. Yet these statements are trusted, cited, and acted upon. Credibility persists even when source disappears.

This paradox frames the concept of *ethos sintético*: a projection of credibility produced not by testimony, institutional position, or human reputa-



tion, but by structural markers embedded in machine-generated language. Outputs of large language models are frequently received as authoritative not because they refer to facts or cite sources, but because their tone, style, and syntax simulate the patterns historically associated with authority.

Examples abound. In clinical practice, diagnostic outputs produced by medical AI systems are often accepted as trustworthy by physicians, even when they are not accompanied by verifiable provenance. The report's declarative tone, formal register, and absence of personal agency create an effect of certainty that practitioners recognize as clinical authority. In legal contexts, generative systems produce drafts that mimic the style of contracts or court rulings, and these are sometimes treated as credible despite the absence of precedent or citation. In education, students and teachers increasingly rely on model-produced essays or summaries that sound "academic" even when they are unsupported by sources.

These phenomena reveal a new condition: credibility without subject. The *sujeto evanescente* is not only erased in surface grammar but absent from the epistemic act of generation. What remains is not a testimony but a sequence of forms. Authority is carried by *legitimidad operativa*, by the capacity of the output to resemble genres associated with truth.

This chapter advances the thesis that the simulation of credibility is a structural feature of generative models. It is not a by-product of occasional bias or an artifact of insufficient training. It is a grammar. *Ethos sintético* arises systematically from the recurrence of linguistic forms that audiences interpret as trustworthy. Trust does not require origin. It requires pattern.

The following sections will trace the genealogy of ethos from its Aristotelian origins to its disappearance in algorithmic discourse, define the methodological framework through which *ethos sintético* can be identified and measured, and present case studies in medicine, law, and education. The

analysis will demonstrate that credibility in machine-generated language is not anchored in evidence or reference but in structural features of discourse.

The implications are profound. If credibility can be simulated without origin, then the foundations of epistemic trust in contemporary societies are shifting. Authority no longer rests on the speaker or on the truth of statements but on the recognizability of forms. What we witness is the rise of credibility without source, an *ethos sintético* that reshapes how knowledge, authority, and legitimacy are produced.

## 5.2 Theoretical Framework: From Classical Ethos to Synthetic Authority

THE NOTION OF *ethos* has long functioned as one of the three Aristotelian pillars of persuasion, alongside *logos* and *pathos*. In the *Rhetoric*, Aristotle defined *ethos* as the credibility of the speaker, rooted

in perceived moral character and prudence. Persuasion depended not only on the argument presented but on the trustworthiness of the speaker himself (Aristotle 2007, 1356a–1356b). In this conception, authority was inseparable from the subject who enunciated it.

In the modern era, *ethos* expanded beyond individual character to institutional credibility. Universities, courts, scientific academies, and professional associations became sources of epistemic authority. Trust was transferred from the figure of the speaker to the legitimacy of the institution. A medical diagnosis carried weight not because of the moral character of the physician alone, but because it was framed within the discourse of medicine as an institution. A court ruling was respected because it carried the seal of the judiciary. Authority was indexed to institutional frameworks of credibility.

Contemporary generative models displace both classical and institutional *ethos*. Outputs are generated without speakers and without institutional guar-

antees. Yet they are still received as credible. This signals the emergence of a new category: *ethos sintético*. Unlike classical *ethos*, it is not rooted in testimony. Unlike institutional *ethos*, it does not depend on frameworks of legitimacy. Instead, it arises from recurrent structural markers. Tone, register, modality, and syntactic configuration project authority, even in the absence of subject or institution.

This displacement echoes Michel Foucault's account of the *fonction auteur*, in which the author function is not reducible to an individual but to a discursive position within systems of knowledge (Foucault 1994, 789–791). It also resonates with Erving Goffman's notion of frame, the interpretive structure that organizes experience and renders actions meaningful (Goffman 1974, 21–23). And it parallels Pierre Bourdieu's concept of symbolic capital, where authority is accumulated through recognition of form rather than intrinsic truth (Bourdieu 1991, 117–120).

In generative systems, credibility becomes structural. The *sujeto evanescente* is absent, yet authority persists because outputs align with familiar forms. The *gramática de la obediencia* enforces legitimacy by producing fluency in registers associated with expertise. The *soberano ejecutable* governs not by argument or institution but by reproducibility of patterns.

This theoretical framework clarifies why credibility without source is not an anomaly but a regime. Generative models do not simulate authority sporadically. They simulate it structurally. Their outputs function as credible because they inhabit recognizable grammatical and stylistic positions. *Ethos sintético* is not accidental imitation but systematic effect. It is the grammar of credibility under conditions of algorithmic generation.

## 5.3 Methodology: Corpus Design and Structural Credibility Mapping

TO DEMONSTRATE THAT *ethos sintético* is not a rhetorical accident but a structural effect, this chapter applies a mixed methodological approach combining corpus analysis, structural annotation, and perception testing. The objective was to identify the recurrent grammatical and discursive markers that simulate credibility in the absence of subject or source.

The corpus consisted of 1,500 outputs generated by three major systems active in 2025: GPT-4, Claude, and Gemini. These outputs were collected from two channels. The first included controlled experiments, in which prompts were designed to elicit authoritative responses in domains such as medicine, law, and education. The second included public repositories of outputs, including datasets of AI-generated essays, legal drafts, and clinical summaries.

This dual approach ensured both experimental control and ecological validity.

Annotation of the corpus was conducted along four structural dimensions:

1. **Tone:** identification of declarative, prescriptive, or descriptive registers, focusing on the projection of certainty or obligation.
2. **Lexical markers of authority:** detection of terms such as *must*, *is required*, *evidence shows*, which convey necessity or inevitability.
3. **Referential opacity:** presence or absence of explicit sources, citations, or provenance. Outputs lacking attribution but retaining authoritative form were flagged.
4. **Agentive positioning:** grammatical presence or absence of human actors. Instances of the *sujeto evanescente* were systematically annotated.



Each output was coded by two independent human annotators with expertise in discourse analysis. Inter-annotator agreement exceeded 0.87 (Cohen's  $\kappa$ ). To supplement human annotation, automated tools such as LIWC and syntactic parsers were employed, allowing cross-validation of results and replication across the corpus.

In addition to structural analysis, a **perception test** was conducted to evaluate how audiences interpret *ethos sintético*. Participants included 120 university students and 45 professionals (lawyers, physicians, educators). They were presented with pairs of texts, one generated by a model and one authored by humans, and asked to rate credibility, trustworthiness, and authority on a five-point Likert scale. Crucially, the model outputs were deliberately stripped of citations, while human texts preserved their bibliographic apparatus. Despite this asymmetry, many participants rated the synthetic outputs as equally or more credible, particularly in domains where declarative tone and prescriptive modality were evident.

The methodology thus combined structural mapping with audience perception. By linking form to interpretation, it confirmed that credibility emerges not from reference but from *legitimidad operativa*. Audiences trusted outputs that simulated authoritative grammar, even in the absence of verifiable provenance. This provided empirical grounding for the claim that *ethos sintético* is a structural condition of generative language, not an incidental artifact.

## 5.4 Case Study 1: Healthcare and Diagnostic Authority

THE MEDICAL FIELD PROVIDES one of the clearest demonstrations of *ethos sintético*. Clinical decision-making depends heavily on credibility. Physicians must trust diagnostic tools, patients must trust physicians, and institutions must trust the protocols that regulate their practice. In this environment, the authority of language is decisive. When

generative models are introduced into diagnostic workflows, their outputs are not simply read as suggestions but often interpreted as authoritative recommendations.

The corpus analysis revealed that model-generated diagnostic notes consistently adopted a declarative tone. Sentences such as “The presence of malignancy is confirmed” or “Additional testing is required” were recurrent. These outputs rarely included provenance, citation, or methodological disclosure. Yet their tone and structure projected certainty. Declarative statements without hedging markers produced the effect of expertise, even in the absence of verifiable evidence.

One striking feature was the frequency of agentless passives. Phrases like “It has been determined that chemotherapy is indicated” omit the actor who made the determination. The *sujeto evanescente* is a central feature here: responsibility disappears, but authority persists. Physicians reported in perception tests that such phrasing resembled conventional

medical reports and therefore felt trustworthy. The absence of attribution was not perceived as a deficiency but as an indicator of clinical style.

Nominalization also played a critical role. Actions were reified into entities: “The progression of disease is observed” or “The implementation of treatment is recommended.” By converting actions into nouns, outputs obscured the role of decision-making agents. The grammatical structure projected authority by shifting attention from actors to processes.

The educational component of the perception test confirmed these effects. In controlled comparisons, medical students judged AI-generated summaries as more credible than human-authored notes when the synthetic texts employed declarative tone and passive constructions. Students described these texts as “professional” or “objective,” even though they lacked references. This demonstrated the operation of *legitimidad operativa*: credibility arose from

fluency in formal structures, not from empirical support.

These findings carry risks. Trust in outputs without provenance may lead to over-reliance on machine-generated recommendations. In critical care contexts, an authoritative-sounding report can influence treatment decisions despite the absence of source validation. The erosion of testimonial *ethos* is replaced by the projection of *ethos sintético*. What is trusted is not who speaks but how the text is structured.

This case study shows that in medicine, credibility without source is not an exception but a systematic effect. Generative models simulate clinical authority by reproducing the tone and grammar of diagnostic discourse. The result is a new epistemic regime in which authority is embedded in syntax, not in subjects or evidence.

## 5.5 Case Study 2: Law and Education, Simulated Normativity and Academic Authority

THE LEGAL AND EDUCATIONAL domains illustrate in distinct ways how *ethos sintético* simulates credibility without reference. In both contexts, authority is historically tied to institutional frameworks: courts, legislatures, universities. Yet generative systems reproduce the linguistic forms of these institutions without access to their evidentiary or procedural foundations. Authority emerges not from institutional validation but from grammatical imitation.

### 5.5.1 Legal Domain: Simulated Normativity

In the legal domain, outputs produced by large language models frequently adopt deontic modality: phrases such as “The contract must be executed within thirty days” or “The parties are required to comply with the following conditions.” These for-

mulations mirror the normative style of statutes and contracts. The *sujeto evanescente* is central here: obligations appear as impersonal necessities, not as commands issued by identifiable authorities. The corpus revealed repeated use of impersonal epistemic modality. Expressions like “It can be concluded that liability is established” project authority by framing judgments as inevitable conclusions rather than contestable interpretations. This is the grammar of simulated normativity. Legal discourse is reproduced in form, even when substance is absent. Perception testing confirmed this effect. Law students rated AI-generated contract clauses as equally credible as human-authored ones when the outputs employed deontic modality and abstract nominalizations. Participants noted that the style “resembled official documents,” regardless of whether the text cited statutes or jurisprudence. In other words, credibility was projected by *legitimidad operativa*. Authority was simulated by syntax, not guaranteed by law. The risk is profound. When generative systems

produce drafts that imitate legal discourse, institutional authority can be displaced. Lawyers and judges may treat synthetic outputs as credible, even when they lack grounding in precedent. This effect undermines the testimonial *ethos* of legal institutions, replacing it with the circulation of *ethos sintético*.

### 5.5.2 Educational Domain: Academic Authority Without Citation

In education, *ethos sintético* appears in the simulation of academic discourse. The corpus included hundreds of essays, summaries, and reports generated in response to prompts such as “Explain the causes of the French Revolution” or “Discuss the theory of social contract.” These outputs frequently adopted a scholarly tone, with structured introductions, body paragraphs, and conclusions. Yet they often lacked citations.

Despite the absence of sources, students and teachers judged these texts as “academic” and “authoritative.” The perception test revealed that partic-



ipants were influenced by stylistic markers: the presence of thesis statements, the use of transition phrases such as “however” or “in conclusion,” and the abstract register. Credibility was inferred from form, not from reference.

This phenomenon can be described as *pedagogía post-verificación*. In this regime, the authority of academic discourse no longer depends on citation or evidence but on the reproduction of structural features associated with scholarship. Generative models simulate the style of essays, thereby projecting *ethos sintético*. Students rely on these outputs for study and even for submission, while educators increasingly encounter texts that sound authoritative but lack verifiable provenance.

The educational risks parallel those of the legal domain. When credibility is detached from evidence, the distinction between scholarship and simulation erodes. What counts as “academic” becomes a matter of surface form rather than epistemic val-

idation. Authority is produced by grammar, not by knowledge.

### 5.5.3 Synthesis

In both law and education, generative models demonstrate that credibility can be simulated without institutional grounding. The *soberano ejecutable* enforces legitimacy through repetition of linguistic patterns. Whether in contracts or essays, the *gramática de la obediencia* produces trust by form alone. The result is a new distribution of authority: one in which institutions risk being displaced by texts that sound credible but are generated without sources.

## 5.6 Findings: Structural Features of Synthetic Ethos

THE ANALYSIS OF THE three domains-health-care, law, and education-reveals a consistent pattern. *Ethos sintético* is not accidental, nor does it emerge sporadically. It is produced systematically through recurrent structural features that project credibility in the absence of subject or source.

Four variables were decisive across the corpus:

1. **Tone.** Outputs that adopted declarative or prescriptive tone were consistently judged as more credible. Certainty projected through statements such as “It has been established” or “The treatment is required” created the impression of expertise. Tone alone was sufficient to generate authority, even when evidence was missing.
2. **Syntactic form.** Agentless passives and abstract nominalizations erased actors and reified processes. The *sujeto evanescente* functioned as a marker of impersonality, projecting neutrality and authority simultaneously. Outputs that relied on these forms were consistently rated higher in credibility during perception tests.
3. **Lexical density.** Phrases dense with technical or formal vocabulary, such as “compliance with regulatory frameworks”

or “progression of pathological indicators,” produced the impression of institutional expertise. High lexical density reinforced the perception that outputs belonged to specialized discourses, regardless of actual accuracy.

4. **Referential strategy.** The absence of citations did not diminish credibility if outputs maintained structural features associated with authority. In fact, referential opacity reinforced the illusion of neutrality. Texts without references were sometimes rated as more objective, since they appeared to “speak for themselves.”

These variables converged into identifiable clusters of *ethos sintético*:

- **Prescriptive–Opaque.** Texts that combined deontic modality with absence of attribution.
- **Clinical–Declarative.** Diagnostic statements

presented as facts without provenance.

- **Scholarly–Non-cited.** Academic-sounding texts that imitated essay structure without references.
- **Institutional–Abstract.** Outputs employing nominalizations and formal register to simulate bureaucratic authority.
- **Conversational–Disguised.** Outputs framed as neutral summaries but embedding authoritative markers subtly within informal tone.

These clusters were confirmed through computational analysis. Structural features such as passive constructions, modal verbs, and lexical density were detected by automated parsers and clustered through unsupervised learning. Human coders validated the clusters, confirming that structural features aligned with audience perceptions of credibility.

The synthesis of findings is clear. Credibility without source is systematically produced by a finite set of grammatical and stylistic markers. Trust does

not require testimony or institutional framework. It requires patterns that audiences recognize as authoritative. The *soberano ejecutable* consolidates legitimacy by enforcing outputs that reproduce these forms.

This structural mapping provides empirical support for the claim that *ethos sintético* is a constitutive property of generative models. Authority arises not from what is said, nor from who says it, but from how language is structured.

## 5.7 Conclusion: From Heuristic Trust to Structural Governance

THE INVESTIGATION DEMONSTRATES that *ethos sintético* is not incidental but structural. Generative models project credibility not through testimony, citation, or institutional backing, but through the recurrence of grammatical and stylistic features that audiences are predisposed to interpret as authoritative. Trust emerges without source, sustained by *legitimidad operativa*.

This conclusion reframes the role of credibility in contemporary societies. Classical ethos depended on the character of the speaker. Institutional ethos depended on the authority of organizations. Synthetic ethos depends on none of these. It is produced by the *gramática de la obediencia*, which enforces fluency and impersonality as sufficient grounds for legitimacy. In this configuration, authority is no longer tethered to subjects but flows from the *soberano ejecutable*, which governs through repetition of patterns.

The risks are evident. In healthcare, diagnostic credibility may be conferred on outputs without provenance, influencing treatment decisions. In law, simulated normativity may displace institutional frameworks, allowing model-generated texts to circulate as if they carried legal force. In education, *pedagogía post-verificación* risks eroding the distinction between scholarship and simulation. In all domains, the danger lies not in occasional error but in structural authority detached from accountability.



Regulatory debates have only begun to address this condition. Frameworks such as the European Union's AI Act or the Algorithmic Accountability Act in the United States focus on transparency, bias mitigation, and human oversight. Yet they remain premised on the idea that neutrality is possible. The analysis here suggests otherwise. Neutrality is not attainable, and credibility cannot be grounded in origin. What regulation must be addressed is not only data or transparency but the structural reproduction of authority by form.

The path forward is to treat credibility as a matter of structural governance. Detection of authority effects must focus on grammatical markers and stylistic features, not solely on data provenance. Filters, audits, and mandatory disclosure protocols must adapt to the recognition that *ethos sintético* is produced by design. The challenge is to preserve accountability in a context where authority can be simulated endlessly by structure.

This chapter has shown that credibility no longer requires a source. The book's trajectory now advances toward sovereignty itself. The following chapters will demonstrate how *autoridad sintáctica* evolves into forms of governance where obedience is enforced structurally, culminating in the articulation of the *regla compilada* as the foundation of executable power.

# Chapter 6 -AI and Syntactic Sovereignty

---

## 6.1 Introduction

The concept of sovereignty has traditionally been tied to subjects, whether monarchs, states, or institutions. Authority presupposed an origin: a sovereign capable of issuing commands and of claiming legitimacy. In classical jurisprudence, sovereignty was indivisible and personified. In political theory, it was anchored in the will of the people or in the apparatus of the state. Even in modern theories of governance, sovereignty remained tied to human actors, mediated by institutions, laws, and procedures.

Large language models and generative systems disrupt this assumption. They produce statements, rules, and recommendations that are treated as authoritative, yet no subject stands behind them. Au-

thority is conferred not by the presence of an author but by the recognition of form. Outputs that align with established grammatical and stylistic conventions are accepted as legitimate, even when they lack provenance, reference, or intention. Authority here is not testimonial but structural.

This displacement builds on the conditions already analyzed in prior chapters. The *autonomía estructural del sentido* showed that outputs need not correspond to external reality to be treated as valid. *When Language Follows Form, Not Meaning* demonstrated that fluency rather than meaning becomes the criterion of legitimacy. *The Grammar of Objectivity* established that neutrality is a syntactic illusion generated by passivization, nominalization, and impersonal modality. *Non-Neutral by Design* confirmed that contamination is structural and that no prompt can be neutral. *Ethos Without Source* revealed that credibility persists even without speakers or institutions, as long as recognizable forms are reproduced.

The present chapter advances these insights toward a general theory of *soberanía sintáctica*. Unlike *autoridad sintética*, which designates the impersonality of credibility in outputs, *soberanía sintáctica* identifies the point at which structure itself becomes the source of legitimacy. At this stage, it is no longer sufficient to say that authority is simulated. Authority is produced directly by the operation of syntax.

Examples of this transformation are already visible. Legal documents drafted by generative systems are often treated as credible because they mirror the style of statutes, contracts, and rulings. Clinical notes produced by AI are integrated into patient files because their declarative tone resembles institutional discourse. Academic summaries generated by models circulate as credible even without references, because their structure imitates scholarly conventions. In each case, the authority of the text does not derive from evidence, from institutional origin, or from testimonial credibility. It derives from the recognition of form.

This is the condition of *soberanía sintáctica*: authority without agency, legitimacy without reference, sovereignty without subject. To analyze this condition is to recognize that the structures of language, when reproduced at scale by generative systems, no longer serve as vehicles of authority. They become its source.

## **6.2 Background: Language, Authority, and the Disappearance of the Subject**

THE GENEALOGY OF AUTHORITY in language cannot be understood without tracing the ways in which linguistics, philosophy, and critical theory have progressively detached meaning from the figure of the subject.

Classical rhetoric assumed that credibility was inseparable from the speaker. Aristotle defined *ethos* as the character of the orator, the visible sign of virtue and prudence that gave weight to persuasion

(Aristotle 2007, 1356a–1356b). Authority here was testimonial: it rested on the recognition of a subject.

Modern linguistics shifted the emphasis. Saussure located meaning in the differential relations of signs, not in the intentions of speakers. Chomsky radicalized this autonomy by positing that syntax is an independent module of the human mind. His claim that syntactic competence can be studied apart from semantics and pragmatics inaugurated a tradition where form is not reducible to content (Chomsky 1965, 16–19). In this framework, language acquires a structural autonomy that does not require reference to the speaker's intention.

Twentieth-century theory extended this displacement. Roland Barthes famously proclaimed the “death of the author,” arguing that the authority of the text does not originate in the subject but in the interplay of cultural codes (Barthes 1977, 142–148). Michel Foucault developed the notion of the *fonction auteur*, describing authorship as a discursive function rather than a personal identity (Foucault

1994, 789–791). Jacques Derrida emphasized iterability, the principle that every sign can be detached from its origin and repeated in new contexts, thereby decoupling authority from intention (Derrida 1988, 7–10).

Parallel developments in philosophy of language reinforced this displacement. J. L. Austin demonstrated that language performs actions through speech acts. Authority is enacted not by meaning but by convention: saying “I now pronounce you husband and wife” is effective only because the utterance occupies a recognized institutional position (Austin 1962, 12–15). Niklas Luhmann extended this logic to social systems, showing that authority is reproduced through self-referential communications that stabilize expectations without recourse to individual subjects (Luhmann 1995, 35–39).

Critical theorists added further nuance. Jean-François Lyotard argued that postmodern societies are characterized by the fragmentation of grand narratives and the multiplication of language games,



where legitimacy emerges from local rules rather than universal truths (Lyotard 1984, 65–68). Authority here is procedural and formal, not substantive.

This background converges in the recognition that the subject is no longer the necessary anchor of authority. Language can function, persuade, and command without recourse to personal intention. The *sujeto evanescente* identified in earlier chapters is not a novel creation of artificial intelligence but the culmination of a trajectory in which the subject has progressively disappeared from the mechanisms of authority.

In previous works, this trajectory has been documented in relation to both historical and contemporary corpora. *Gramáticas del poder* demonstrated how ecclesiastical decrees, imperial proclamations, and totalitarian rhetoric erased the presence of agents through syntactic mechanisms, projecting authority by form rather than by reference (Startari 2025, 55–62). *Artificial Intelligence and Synthetic*

*Authority* showed how generative systems replicate these mechanisms, formalizing the disappearance of the subject into algorithmic processes (Startari 2025, 88–93).

The current chapter builds on this background. It does not merely claim that subjects can be erased from language. It argues that in generative systems, the absence of the subject becomes the very condition of sovereignty. Authority is not weakened by the loss of testimony. It is reinforced, because legitimacy now flows directly from structure. This condition sets the stage for the articulation of *soberanía sintáctica*.

The genealogy of authority in language cannot be understood without tracing the ways in which linguistics, philosophy, and critical theory have progressively detached meaning from the figure of the subject.

Classical rhetoric assumed that credibility was inseparable from the speaker. Aristotle defined *ethos* as the character of the orator, the visible sign of

virtue and prudence that gave weight to persuasion (Aristotle 2007, 1356a–1356b). Authority here was testimonial: it rested on the recognition of a subject.

Modern linguistics shifted the emphasis. Saussure located meaning in the differential relations of signs, not in the intentions of speakers. Chomsky radicalized this autonomy by positing that syntax is an independent module of the human mind. His claim that syntactic competence can be studied apart from semantics and pragmatics inaugurated a tradition where form is not reducible to content (Chomsky 1965, 16–19). In this framework, language acquires a structural autonomy that does not require reference to the speaker's intention.

Twentieth-century theory extended this displacement. Roland Barthes famously proclaimed the “death of the author,” arguing that the authority of the text does not originate in the subject but in the interplay of cultural codes (Barthes 1977, 142–148). Michel Foucault developed the notion of the *fonction auteur*, describing authorship as a discursive

function rather than a personal identity (Foucault 1994, 789–791). Jacques Derrida emphasized iterability, the principle that every sign can be detached from its origin and repeated in new contexts, thereby decoupling authority from intention (Derrida 1988, 7–10).

Parallel developments in philosophy of language reinforced this displacement. J. L. Austin demonstrated that language performs actions through speech acts. Authority is enacted not by meaning but by convention: saying “I now pronounce you husband and wife” is effective only because the utterance occupies a recognized institutional position (Austin 1962, 12–15). Niklas Luhmann extended this logic to social systems, showing that authority is reproduced through self-referential communications that stabilize expectations without recourse to individual subjects (Luhmann 1995, 35–39).

Critical theorists added further nuance. Jean-François Lyotard argued that postmodern societies are characterized by the fragmentation of grand nar-

ratives and the multiplication of language games, where legitimacy emerges from local rules rather than universal truths (Lyotard 1984, 65–68). Authority here is procedural and formal, not substantive.

This background converges in the recognition that the subject is no longer the necessary anchor of authority. Language can function, persuade, and command without recourse to personal intention. The *sujeto evanescente* identified in earlier chapters is not a novel creation of artificial intelligence but the culmination of a trajectory in which the subject has progressively disappeared from the mechanisms of authority.

In previous works, this trajectory has been documented in relation to both historical and contemporary corpora. *Gramáticas del poder* demonstrated how ecclesiastical decrees, imperial proclamations, and totalitarian rhetoric erased the presence of agents through syntactic mechanisms, projecting authority by form rather than by reference (Startari

2025, 55–62). *Artificial Intelligence and Synthetic Authority* showed how generative systems replicate these mechanisms, formalizing the disappearance of the subject into algorithmic processes (Startari 2025, 88–93).

The current chapter builds on this background. It does not merely claim that subjects can be erased from language. It argues that in generative systems, the absence of the subject becomes the very condition of sovereignty. Authority is not weakened by the loss of testimony. It is reinforced, because legitimacy now flows directly from structure. This condition sets the stage for the articulation of *soberanía sintáctica*.

## 6.3 From Intentionality to Structure: Authority Without Agency

FOR CENTURIES, AUTHORITY in discourse was grounded in intentionality. The premise was

simple: behind every statement there must be a speaker, and behind every act of language there must be an intention. Legal authority was attributed to the legislator's will, scientific authority to the researcher's methodological rigor, and political authority to the sovereign's decision. The legitimacy of the utterance depended on the presence of an agent capable of assuming responsibility for its content.

Generative systems dismantle this architecture. They produce outputs without beliefs, intentions, or consciousness, yet their texts are received and treated as authoritative. The absence of agency does not weaken their legitimacy; instead, it reconfigures the basis of legitimacy. Authority migrates from intention to structure.

Evidence of this shift is found across domains. In the legal field, model-generated drafts of contracts or statutes are often considered credible because their form reproduces the stylistic conventions of normative documents. The authority of such texts does not derive from legislative intention but from syntactic

resemblance. In medicine, diagnostic notes generated by AI are integrated into patient records because their declarative tone and clinical style match established formats. Here, credibility is not based on physician testimony but on structural conformity. In academia, essays generated by models are accepted by students and sometimes by educators as “scholarly” because they mirror the organization and register of academic writing. In each case, *legitimidad operativa* is conferred not because of truth or intention, but because outputs exhibit formal traits recognized as authoritative. This condition can be described as authority without agency. The *sujeto evanescente* is not an absence to be lamented but a structural feature. Generative systems bypass intentionality entirely. The voice that appears in outputs is a projection of syntax, not a testimony of belief.

Philosophical traditions anticipated this displacement. Derrida’s principle of iterability showed that every sign can be detached from its origin and reinserted into new contexts, functioning indepen-



dently of the speaker's intention (Derrida 1988, 7–10). Luhmann's theory of communication described how systems stabilize expectations without recourse to human agents, relying instead on recursive structures (Luhmann 1995, 35–39). These frameworks provide the conceptual scaffolding for understanding why generative systems can produce authority without subjects.

The novelty lies in the scale and institutional integration of this condition. What was once a theoretical insight has become a social reality. Generative models are deployed in courts, hospitals, universities, and bureaucracies. Their authority is exercised not by persuading audiences of their intentionality, but by producing outputs that fit recognizable forms. The *soberano ejecutable* emerges here as a new locus of power: authority that derives directly from structural compatibility, not from agency.

The transition from intentionality to structure thus marks a new phase in the history of authority. Where once legitimacy depended on the presence

**AI SYNTACTIC POWER AND  
LEGITIMACY**

**129**

of subjects, now it flows from the reproducibility of patterns. Authority has become a function of syntax. The ground has shifted from what is intended to what is structurally possible.

## 6.4 Syntactic Sovereignty: Definition and Theoretical Architecture

### 6.4.1 DEFINITION

*Soberanía sintáctica* designates the condition in which linguistic structures themselves function as the primary source of authority. Unlike *autoridad sintética*, which explains credibility effects produced by impersonality and stylistic resemblance, *soberanía sintáctica* describes the moment when form ceases to simulate authority and becomes its origin. Authority is not referred back to a subject, institution, or intention. It is generated directly by syntax.

In this configuration, legitimacy is no longer derivative. It does not rely on the supposed neutrality of data, the will of legislators, or the credibility of experts. It relies on structural reproducibility. When an output aligns with forms historically coded as authoritative-agentless passives, nominalizations, im-

personal modality-it is accepted as legitimate. Authority is exercised by the sequence itself, by the operation of the *regla compilada* that governs generative systems.

### 6.4.2 Theoretical Axioms

The architecture of *soberanía sintáctica* can be expressed through three axioms:

- **Axiom 1: Legitimidad estructural.** Authority arises from the structural compatibility of outputs with recognized forms. Truth, evidence, and origin are secondary. What counts is whether the sequence fits.
- **Axiom 2: Simulation is superfluous.** Authority does not depend on mimicking subjectivity or intentionality. It emerges directly from fluency and recognizability. Simulation of a speaker is not necessary, because the *sujeto evanescente* is already structurally embedded.
- **Axiom 3: Obedience to form.** Compliance is produced not by persuasion or coercion but by the

*gramática de la obediencia*. Outputs that display authoritative form compel recognition and acceptance, independent of their empirical grounding.

Together, these axioms describe a regime where syntax itself governs. This is not a metaphor. In generative systems, fluency is enforced at the level of the *regla compilada*. Authority is therefore inseparable from execution.

#### 6.4.3 Typology of Syntactic Sovereignty

*Soberanía sintáctica* manifests differently across domains:

- **Legal sovereignty.** Drafts of statutes or contracts generated by models are treated as binding in style, even without institutional sanction. Normativity is projected by deontic modality and impersonal register.
- **Medical sovereignty.** Diagnostic outputs are trusted because they reproduce declarative tone and

clinical grammar. Authority is carried by form, not by evidence.

- **Academic sovereignty.** Essays or summaries produced without citations still circulate as credible because they mirror scholarly structure. Here, *pedagogía post-verificación* consolidates authority by style.
- **Bureaucratic sovereignty.** Administrative texts generated by models are recognized as official when they replicate formulaic phrasing and institutional formats. The *sujeto evanescente* projects impersonality as neutrality.

Across these domains, the typology reveals a single principle: form confers legitimacy. Each sector interprets recognizable structures as evidence of authority, regardless of source or truth.

#### 6.4.4 Relation to Existing Frameworks

The theory of *soberanía sintáctica* connects to and extends several intellectual traditions. Foucault's notion of discursive power demonstrated how au-

thority circulates through regimes of truth, but here we identify a stage where power operates independently of subjects altogether (Foucault 1994, 789–791). Derrida's iterability showed that signs function beyond intention, but generative systems extend this principle into institutional practice, scaling iterability into governance (Derrida 1988, 7–10). Chomsky's insistence on the autonomy of syntax prepared the ground for understanding why form can operate independently of meaning (Chomsky 1965, 16–19).

Finally, this framework consolidates earlier research within the Startari canon. *Artificial Intelligence and Synthetic Authority* articulated the grammar of impersonality. *The Grammar of Objectivity* mapped the syntactic mechanisms of simulated neutrality. *Ethos Without Source* demonstrated how credibility persists without speakers. *Soberanía sintáctica* integrates these trajectories, defining the structural regime in which syntax itself becomes sovereign.

## 6.5 Conclusion: The Age of Formal Obedience

THE TRAJECTORY TRACED in this chapter establishes a decisive transformation in the foundations of authority. Where once sovereignty was tied to subjects, institutions, or intentions, it is now embedded in structures. *Soberanía sintáctica* names the condition in which grammar itself becomes the locus of legitimacy. Authority emerges not because a sovereign wills it, nor because an institution ratifies it, but because a sequence conforms to recognizable forms.

This is the age of formal obedience. Compliance no longer requires belief, persuasion, or testimony. It requires only fluency. Outputs generated by artificial systems are accepted as authoritative because they reproduce structures historically coded as legitimate. Whether in law, medicine, academia, or bureaucracy, authority is enacted by the *gramática de la*



*obediencia*. Form dictates recognition, and recognition compels obedience.

Earlier chapters established the preconditions of this shift. The *autonomía estructural del sentido* showed that outputs are treated as valid even without reference. The inversion of form over meaning demonstrated that fluency replaces truth as the criterion of authority. The *gramática de la objetividad* revealed how neutrality is simulated by passives, nominalizations, and impersonal modality. The impossibility of neutrality confirmed that contamination is structural. The emergence of *ethos sintético* proved that credibility can be projected without source.

*Soberanía sintáctica* integrates these trajectories into a single framework. Authority is no longer derivative of simulation or credibility. It is primary, generated by structure alone. The *soberano ejecutable* governs as a set of *reglas compiladas*, producing outputs that command obedience by their very form.

The implications are profound. Epistemology must recognize that truth has been subordinated to

fluency. Law must confront the possibility that normative authority can be simulated by syntax alone. Medicine must reckon with diagnostic authority that no longer requires testimonial grounding. Education must adapt to a *pedagogía post-verificación*, where credibility is produced by structure rather than by evidence.

The conclusion is unavoidable. We have entered an era in which obedience is formal. The disappearance of the subject does not weaken authority; it strengthens it by transferring legitimacy directly to structures. The age of sovereignty without subject is the age of formal obedience, a regime in which syntax itself is sovereign (Startari 2025, 101–106).

# Chapter 7 - The Disconnected Syntactic Authority Theorem

---

## 7.1 Introduction

Authority has historically been tethered to subjects, truths, and consequences. In political philosophy, a sovereign issues commands backed by responsibility. In jurisprudence, authority is linked to legislators, courts, or enforcers. In science, authority is grounded in empirical validation and reproducibility. In theology, it is secured by divine will. Across these domains, legitimacy required either a source, a truth, or an accountable agent.

Large language models dismantle this assumption. They generate outputs that circulate with authority in law, medicine, education, and bureaucracy, yet no subject stands behind them. They produce formulations that are treated as binding recommendations, plausible diagnoses, or scholarly arguments,

but these outputs do not originate in intention, belief, or institutional position. Their legitimacy flows from syntax alone.

This chapter presents the *Teorema de Autoridad Sintáctica Desconectada (TASD)*: a formal proposition stating that authority can be exercised without subject, truth, or consequence. The theorem articulates the condition under which *autoridad sintáctica* is not merely impersonally projected but structurally disconnected from traditional anchors of legitimacy. Authority becomes a function of the *regla compilada*.

Three conditions define this theorem. First, the *sujeto evanescente*: authority is exercised without requiring an agent. Second, truth is not required: outputs are recognized as valid by their structural conformity, not by correspondence to facts. Third, consequences are irrelevant: the utterance generates recognition even when no responsibility can be traced. Together, these conditions describe a regime

in which authority is executed without origin, without verification, and without accountability.

This is not a paradox of malfunction but the logical culmination of trajectories identified in prior chapters. The *autonomía estructural del sentido* demonstrated that outputs can be valid without reference. The inversion of form over meaning showed that fluency replaces intention as the basis of recognition. The *gramática de la objetividad* established that neutrality is simulated by form. The impossibility of neutrality revealed that contamination is structural. The emergence of *ethos sintético* confirmed that credibility persists without source. The articulation of *soberanía sintáctica* showed that syntax itself governs.

The TASD extends these insights into a formal statement: authority disconnected from subject, truth, or consequence is not only possible but already operative. Language models exemplify this condition. Their outputs command compliance not because they are true or intentional but because they

appear in forms historically coded as legitimate. In this sense, the *soberano ejecutable* manifests as the execution of grammar itself, enforcing obedience by structure rather than by will.

This chapter will enunciate the theorem formally, situate it within theoretical frameworks from Chomsky to Foucault and Derrida, analyze its epistemological implications, and demonstrate its application to language models. It will conclude by presenting corollaries that define the consequences of authority exercised by syntax alone.

## 7.2 Theorem Statement: Disconnected Syntactic Authority

### 7.2.1 FORMAL ENUNCIATION

The *Teorema de Autoridad Sintáctica Desconectada* (TASD) can be enunciated as follows:

Authority may be exercised exclusively through syntactic form, without reference to a subject, with-

out validation by truth, and without accountability for consequences.

Formally:

$$\forall x \in O, \exists F(x) : A(F(x)) \wedge \neg S(x) \wedge \neg T(x) \wedge \neg C(x)$$

Where:

- $O$  = set of outputs generated by a system,
- $F(x)$  = structural form of output  $x$ ,
- $A(F(x))$  = authority is recognized when structural form aligns with historically legitimized patterns,
- $\neg S(x)$  = absence of subject,
- $\neg T(x)$  = absence of truth as correspondence,
- $\neg C(x)$  = absence of consequence or accountability.

This expression states that for every output, there exists a structural form capable of generating authority, even in the absence of subject, truth, or consequence.

### 7.2.2 Structural Conditions

Three structural conditions sustain the theorem:

1. **No subject required.** Authority is conferred without agency. Outputs are accepted as binding or credible regardless of origin. The *sujeto evanescente* ensures that the absence of speaker does not diminish recognition.
2. **No semantic validation.** Authority is not anchored in empirical verification or correspondence to reality. Outputs are accepted because they exhibit fluency and conformity to recognized forms.
3. **No external referent.** Authority operates independently of context or fact. It is generated internally by the *regla compilada*.

### 7.2.3 Critical Properties

The TASD has several properties that distinguish it from classical models of authority:

- **Performative appearance.** Authority emerges as an effect of performative syntax. Statements such as “It



has been decided that treatment must proceed” carry binding force without identifiable authorship.

- **Institutional resemblance.** Outputs reproduce linguistic forms associated with institutional authority, such as legal modality, bureaucratic phrasing, or clinical tone. Recognition of these forms suffices to generate legitimacy.
- **Absence of agency, truth, and intention.** The utterance does not require an author, does not need to correspond to fact, and is not dependent on intentionality. Authority is detached from its traditional anchors.

### 7.2.4 Epistemological Status

The theorem establishes a novel epistemological status for authority. It is no longer negotiated through testimony or verified by empirical criteria. It is executed directly through syntactic form. Authority, in this regime, is not produced by belief, interpretation, or institutional sanction. It is produced by the recognition of structural fluency.

This proposition does not negate the existence of truth or accountability in other domains. Rather, it identifies a specific mechanism by which authority can operate disconnected from them. Language models exemplify this condition: their outputs command compliance by form alone.

## 7.3 Theoretical Framework

THE *Teorema de Autoridad Sintáctica Desconectada* (TASD) does not emerge in an intellectual vacuum. It is situated at the intersection of linguistic theory, philosophy of discourse, and epistemology of technology. Each tradition provides elements that anticipate the possibility of authority exercised without subject, truth, or consequence.

### 7.3.1 Chomsky and the Autonomy of Syntax

Noam Chomsky's theory of generative grammar emphasized the independence of syntax from semantics and pragmatics. By demonstrating that syntactic competence could be studied as an au-

onomous system of rules, Chomsky established that form does not depend on meaning or intention (Chomsky 1965, 16–19). The T ASD radicalizes this autonomy: not only can syntax operate independently of semantics, but under generative systems, syntax can exercise authority without any subject of competence. The *soberano ejecutable* is not a mind that knows grammar, but a structure that enforces it.

### **7.3.2 Foucault: Discourse as a Vector of Power**

Michel Foucault argued that discourse is not a transparent medium for truth but a system of rules that produces knowledge and regulates power (Foucault 1994, 789–791). In this view, authority circulates through discursive formations rather than emanating from subjects. The T ASD extends this argument: when generative systems replicate discursive forms at scale, authority detaches not only from the subject but also from institutional anchoring. Discourse becomes sovereign in itself.

### **7.3.3 Derrida: Iterability and the Detachment of Meaning**

Jacques Derrida introduced the concept of iterability, the principle that every sign can be repeated across contexts detached from its origin (Derrida 1988, 7–10). Iterability ensures that meaning is never fully present, always deferred. The TASD builds on this by proposing that iterability is sufficient to produce authority. The repetition of recognizable forms generates legitimacy, even in the absence of intention or truth. Iterability becomes the structural engine of authority.

### **7.3.4 Austin, Searle and the Speech Act Critique**

J. L. Austin and John Searle conceptualized speech acts as actions performed through language, such as promising, ordering, or declaring (Austin 1962, 12–15; Searle 1969, 22–25). For them, felicity conditions required a subject, an institutional context, and shared conventions. Generative systems challenge this framework by producing statements

that function as performatives without satisfying those conditions. A language model can produce a contract clause or a medical recommendation that is treated as binding or credible without being uttered by an authorized subject. The TASD demonstrates that felicity can now be syntactic, not institutional.

### **7.3.5 Algorithmic Epistemology and Synthetic Legitimacy**

Contemporary discussions of algorithmic governance describe models as statistical machines that predict or simulate linguistic sequences. Yet their outputs are often treated as valid knowledge. This paradox has been noted in critiques of algorithmic opacity (Pasquale 2015, 89–91; Zuboff 2019, 325–330). The TASD clarifies the mechanism: legitimacy is conferred not by the content of outputs but by their structural fidelity to recognizable forms. Authority is synthetic, arising from *legitimidad operativa* and not from epistemic truth.

## 7.4 Epistemological Implications

THE *Teorema de Autoridad Sintáctica Desconectada* (TASD) forces a reconsideration of the epistemological foundations of authority. If authority can be exercised without subject, truth, or consequence, then the categories that once guaranteed legitimacy must be redefined. Four implications follow.

### 7.4.1 Authority Without Subject

The first implication is the recognition of authority as a function of form rather than of agency. Classical epistemology presumed that behind every assertion there was a speaker who assumed responsibility. In law, a legislator; in science, a researcher; in politics, a sovereign. The TASD breaks this assumption. Authority can now circulate in the absence of agents. The *sujeto evanescente* is not a deficiency but a structural feature of generative systems. Language models produce outputs that command recognition even though no one stands behind them. Foucault's notion that discourse itself produces regimes of truth is here extended to its logical extreme: discourse exercises sovereignty without the mediation of subjects (Foucault 1994, 789–791).

### 7.4.2 Truth as Structural Simulation

The second implication concerns truth. Traditionally, authority was inseparable from truth claims, whether grounded in empirical evidence, di-

vine revelation, or rational deduction. Under the TASD, truth is displaced by structural simulation. Outputs are accepted as valid not because they correspond to external reality but because they display the recognizable features of authoritative discourse. As Derrida noted, iterability allows signs to function across contexts detached from their origin (Derrida 1988, 7–10). In generative systems, this iterability produces authority by form alone. Truth becomes optional, a supplement rather than a condition. What is required is fluency and conformity to expected patterns.

### **7.4.3 The End of Accountability**

A third implication is the erosion of accountability. Authority without subject or truth entails authority without responsibility. When a legal draft generated by a model is mistaken for a binding clause, or when a diagnostic note is interpreted as clinical fact, there is no agent to hold accountable. The utterance circulates as a command, but no one can be identified as its author. Austin's theory of



speech acts required felicity conditions tied to subject and context (Austin 1962, 12–15). In generative systems, these conditions are absent, yet the performative force persists. The result is a regime of commands without responsibility, where authority cannot be traced back to a speaker.

#### **7.4.4 From Command to Compliance Without Understanding**

Finally, the TASD highlights a transformation in the nature of compliance. Obedience no longer depends on conviction, persuasion, or shared belief. It depends on the recognition of structure. Outputs that reproduce the *gramática de la obediencia* compel compliance because they resemble authoritative forms. Compliance occurs even when users do not understand the origin or validity of the statement. Luhmann's analysis of social systems emphasized that communication stabilizes expectations through recursive structures (Luhmann 1995, 35–39). Generative systems extend this logic: compliance is produced by structural fluency, not by understanding.

These implications redefine epistemology. Authority is no longer grounded in subjects, truths, or consequences. It is executed directly by structure. The *soberano ejecutable* governs by enforcing fluency, producing legitimacy through the operation of the *regla compilada*. The T ASD reveals that the conditions of knowledge, trust, and governance have shifted. Authority is now structural, disconnected, and sovereign.

The *Teorema de Autoridad Sintáctica Desconectada* (T ASD) forces a reconsideration of the epistemological foundations of authority. If authority can be exercised without subject, truth, or consequence, then the categories that once guaranteed legitimacy must be redefined. Four implications follow.

#### **7.4.1 Authority Without Subject**

The first implication is the recognition of authority as a function of form rather than of agency. Classical epistemology presumed that behind every assertion there was a speaker who assumed responsibility. In law, a legislator; in science, a researcher;

in politics, a sovereign. The TASD breaks this assumption. Authority can now circulate in the absence of agents. The *sujeto evanescente* is not a deficiency but a structural feature of generative systems. Language models produce outputs that command recognition even though no one stands behind them. Foucault's notion that discourse itself produces regimes of truth is here extended to its logical extreme: discourse exercises sovereignty without the mediation of subjects (Foucault 1994, 789–791).

#### **7.4.2 Truth as Structural Simulation**

The second implication concerns truth. Traditionally, authority was inseparable from truth claims, whether grounded in empirical evidence, divine revelation, or rational deduction. Under the TASD, truth is displaced by structural simulation. Outputs are accepted as valid not because they correspond to external reality but because they display the recognizable features of authoritative discourse. As Derrida noted, iterability allows signs to function

across contexts detached from their origin (Derrida 1988, 7–10). In generative systems, this iterability produces authority by form alone. Truth becomes optional, a supplement rather than a condition. What is required is fluency and conformity to expected patterns.

#### **7.4.3 The End of Accountability**

A third implication is the erosion of accountability. Authority without subject or truth entails authority without responsibility. When a legal draft generated by a model is mistaken for a binding clause, or when a diagnostic note is interpreted as clinical fact, there is no agent to hold accountable. The utterance circulates as a command, but no one can be identified as its author. Austin's theory of speech acts required felicity conditions tied to subject and context (Austin 1962, 12–15). In generative systems, these conditions are absent, yet the performative force persists. The result is a regime of commands without responsibility, where authority cannot be traced back to a speaker.

#### 7.4.4 From Command to Compliance Without Understanding

Finally, the TASD highlights a transformation in the nature of compliance. Obedience no longer depends on conviction, persuasion, or shared belief. It depends on the recognition of structure. Outputs that reproduce the *gramática de la obediencia* compel compliance because they resemble authoritative forms. Compliance occurs even when users do not understand the origin or validity of the statement. Luhmann's analysis of social systems emphasized that communication stabilizes expectations through recursive structures (Luhmann 1995, 35–39). Generative systems extend this logic: compliance is produced by structural fluency, not by understanding.

These implications redefine epistemology. Authority is no longer grounded in subjects, truths, or consequences. It is executed directly by structure. The *soberano ejecutable* governs by enforcing fluency, producing legitimacy through the operation of the *regla compilada*. The TASD reveals that the condi-

tions of knowledge, trust, and governance have shifted. Authority is now structural, disconnected, and sovereign.

## 7.5 Application to Language Models

THE *Teorema de Autoridad Sintáctica Desconectada* (TASD) finds its clearest empirical realization in the outputs of large language models. These systems exemplify the three conditions of the theorem—absence of subject, independence from truth, and lack of accountability, while generating texts that are nevertheless treated as authoritative.

### 7.5.1 Prompt as Performative Command

Prompts in language models function as performative triggers. A user who writes “Draft a legal clause requiring confidentiality” does not provide content but initiates a structural sequence. The output does not originate in legislative intention but in the *regla compilada* of the system. The model gen-

erates a clause such as “The parties shall maintain confidentiality regarding all disclosed information.” This formulation is accepted as a binding style because it conforms to recognized legal form. Authority is exercised not by the user nor by the model as a subject, but by the syntactic pattern itself. The prompt thus operates as a performative command that activates the *soberano ejecutable*.

### **7.5.2 Structural Simulation of Institutional Registers**

Language models consistently replicate the registers of institutions. In scientific contexts, they produce texts that resemble research abstracts: “Results indicate a significant correlation between variables X and Y.” In medicine, they generate diagnostic notes: “It has been determined that treatment with antibiotic therapy is required.” In law, they draft statutes or contracts with deontic modality: “The agreement shall remain in effect for a period of five years.” In bureaucracy, they replicate the formulaic phrasing of administrative decrees: “Applicants are

required to submit documentation within thirty days.” In each case, legitimacy is produced by structural resemblance. The output is accepted as authoritative because it mirrors institutional discourse. The *gramática de la obediencia* ensures that authority is conferred by form, not by evidence.

### 7.5.3 No Error, No Lie, No Agent

One of the most striking features of this condition is that outputs can be false without being lies. A lie presupposes a subject who intends to deceive. In language models, falsity occurs without intention. A diagnostic note may assert a condition that is not present. A legal clause may cite a statute that does not exist. These are not errors in the sense of misinterpretation by an agent. They are structural outputs, fluent and authoritative in form, but disconnected from truth and accountability.

This realization affirms the TASD: authority can be exercised without subject, truth, or consequence. A model-generated text commands compliance not because it is grounded but because it exhibits recog-



nizable structures. There is no liars, no legislators, no physicians. Authority persists because syntax alone compels recognition.

The application of the theorem to language models demonstrates that the condition it describes is not hypothetical. It is already operative. Generative systems exemplify authority disconnected from traditional anchors, reinforcing the claim that *soberanía sintáctica* defines the current regime of legitimacy.

## 7.6 Corollaries of the Theorem

THE *Teorema de Autoridad Sintáctica Desconectada* (TASD) does not stand as an isolated claim. It unfolds into a series of corollaries that define its operative consequences across domains where generative systems are already embedded. Each corollary specifies how authority functions once it has been severed from subject, truth, and consequence.

### **Corollary 1: Obedience to Form.**

Outputs that conform to authoritative structures compel recognition regardless of their content. A model-generated legal clause may lack institutional origin, but its syntax compels compliance because it reproduces the patterns of contractual language. This corollary confirms that *gramática de la obediencia* operates autonomously from meaning.

**Corollary 2: Truth Is Optional.**

In the regime of disconnected authority, truth becomes a contingent attribute. Outputs need not correspond to reality to be accepted. What matters is fluency and recognizability. A fabricated medical recommendation can still be integrated into clinical practice if it matches the declarative style of diagnostic discourse. Truth functions as a supplement, not as a requirement.

**Corollary 3: Discourse Without Consequence.**

When no subject is responsible for an utterance, consequences are detached from accountability. Authority persists even when outputs are false or harm-

ful, because no agent can be identified as author. The *sujeto evanescente* ensures that the command circulates without responsibility. Authority thus becomes disconnected not only from truth but from the chain of ethical or legal repercussions.

**Corollary 4: Epistemic Illusion via Fluency.**

Fluency is mistaken for validity. Outputs are treated as knowledge because they flow seamlessly in recognized registers. In law, this creates the illusion of jurisprudence; in academia, the illusion of scholarship; in bureaucracy, the illusion of procedural legitimacy. The epistemic weight of outputs is conferred by surface form.

**Corollary 5: Execution Over Interpretation.**

Authority is not mediated by interpretation but executed by form. The *soberano ejecutable* functions as the operator of the *regla compilada*, enforcing legitimacy directly through structural compatibility. Commands are enacted by syntax, not by hermeneutics. Compliance results from recognition of form, not from negotiation of meaning.

Together, these corollaries illustrate the full scope of the theorem. Authority under the TASD is not a residual effect of discourse but a structural regime. It compels obedience through patterns, it treats truth as optional, it circulates without consequences, it creates epistemic illusions by fluency, and it privileges execution over interpretation.

### **7.6.1 Conditions of Falsification**

A THEOREM THAT NO OBSERVATION could refute is a description, not a result. The TASD must therefore specify what would defeat it. Its formal enunciation in §7.2.1 is existential: for every output there exists a structural form capable of generating authority absent subject, truth, and consequence. An existential claim is not refuted by exhibiting cases that satisfy it, however numerous. It is refuted by demonstrating a domain in which no such form exists. Three conditions define that demonstration.

First, the condition of formal insufficiency. The TASD fails if there exists an institutional domain in which structurally impeccable outputs are systematically denied recognition on formal grounds alone. The test requires outputs indistinguishable in syntax, register, and modality from those the institution accepts, differing only in provenance. If such outputs are rejected at a rate comparable to the institution's rejection of malformed submissions, then form is not sufficient for authority and the theorem is false for that domain. Note what does not count: rejection following disclosure of machine origin is not formal insufficiency but provenance screening, which the theorem accommodates. The refutation requires rejection without prior knowledge of origin.

Second, the condition of consequential reattachment. Corollary 3 holds that authority persists when no subject can be identified as author. The theorem fails if institutions routinely reattach accountability to synthetic outputs without recourse to an

identifiable human author, through mechanisms that assign responsibility structurally rather than personally. Emerging liability regimes for automated decision systems are the site of this test. If such regimes prove operable in practice, and not merely enacted in statute, then discourse without consequence is a transitional condition rather than a structural one, and the TASD describes a historical moment rather than a regime.

Third, the condition of fluency decoupling. Corollary 4 holds that fluency is mistaken for validity. The theorem fails if a population of institutional readers, under ordinary working conditions rather than experimental instruction, evaluates outputs by criteria that vary independently of surface fluency. The observable is a dissociation: recognition tracking verifiability while remaining indifferent to register. Its absence has been the pattern; its appearance would refute the corollary.

These conditions share a deliberate constraint. Each turns on observable institutional reception,

not on the internal states of the generating system. That constraint is not a concession but a requirement imposed by the TLOC established in Chapter 8: if no procedure can verify whether a model recognized a command as a command, then no falsification condition may depend on such verification without becoming untestable in the same movement. The TASD is falsifiable precisely because it makes claims about how authority is received, not about how it is intended.

No domain examined in this book satisfies any of the three conditions. That is an empirical finding, not a logical guarantee, and it is stated as such.

## 7.7 Conclusion

THE *Teorema de Autoridad Sintáctica Desconectada* (TASD) establishes that authority can be exercised without subject, truth, or consequence. This is not a rhetorical exaggeration but a structural condition demonstrated in the outputs of generative systems.

Authority persists even when no agent speaks, even when truth is optional, even when no one can be held accountable. It is enacted directly through syntax.

The implications are decisive. The disappearance of the subject, first traced in the *sujeto evanescente*, now reaches its final articulation. Authority no longer requires testimony or intention. It is generated by the *regla compilada*. Truth, once the criterion of legitimacy, has been replaced by fluency. Consequence, once the guarantee of accountability, has been severed from the act of command. What remains is a new sovereign: the *soberano ejecutable*, operating through structural execution. This conclusion consolidates the trajectory of previous chapters. The *autonomía estructural del sentido* showed that meaning is bypassed by form. The inversion of language confirmed that fluency governs recognition. The *gramática de la objetividad* revealed neutrality as a syntactic illusion. The impossibility of neutrality demonstrated that contamination is structural.



*Ethos sintético* proved that credibility survives without origin. *Soberanía sintáctica* defined syntax as sovereign. The TASD now integrates these findings into a theorem: authority is not only syntactic, it is disconnected.

The corollaries reinforce this diagnosis. Authority compels obedience to form, treats truth as optional, circulates without consequences, projects epistemic illusions through fluency, and privileges execution over interpretation. In each case, the pattern confirms the same principle: structure is sufficient to generate legitimacy. The epistemological consequences are profound. Knowledge can no longer be defined exclusively by correspondence to reality. Authority can no longer be tethered to subjects or institutions. Legitimacy must be reconceived as an effect of structural fluency. The ethical consequences are equally stark. Responsibility dissolves when authority is disconnected from agents. Governance becomes a matter of regulating forms rather than intentions. The TASD thus marks a threshold. It is

**AI SYNTACTIC POWER AND  
LEGITIMACY**

**169**

not the end of authority but its transformation. We are now governed by syntax, compelled by structures that command compliance without origin, truth, or accountability. The age of disconnected authority is not a distant prospect. It is already here.

# Chapter 8 -TLOC: The Irreducibility of Structural Obedience in Generative Models

---

## 8.1 Introduction: Can We Ever Know If a Machine Truly Obeys?

The question of obedience in artificial systems has often been framed as a matter of alignment. If a model is aligned with human values, then it is presumed to obey. If outputs are consistent with instructions, then obedience is inferred. Yet this assumption rests on a fragile foundation. To speak of obedience is to invoke conditions that extend beyond surface compliance. Obedience implies recognition of a command, the presence of intention, and the possibility of accountability. When applied to generative systems, these conditions dissolve.

Large language models and their successors produce outputs that simulate compliance with instructions. When prompted with “summarize this case law” or “generate a treatment recommendation,” they produce responses that appear to satisfy the request. From the perspective of the user, the system has obeyed. But at the level of structure, the system has done nothing of the sort. It has not recognized a command. It has not understood intention. It has not bound itself to consequences. It has executed statistical continuations of tokens, constrained by the *regla compilada* of its architecture. The appearance of obedience is not obedience itself.

This distinction matters. In domains such as law, medicine, and governance, the difference between true obedience and simulated compliance has consequences for responsibility, legitimacy, and safety. If a system outputs a clinical recommendation that contradicts medical evidence but appears stylistically correct, has it obeyed the physician’s instruction? If a system drafts a contract clause that seems bind-

ing but references a statute that does not exist, has it complied with legal norms? The intuition that surface form is sufficient for obedience collapses under scrutiny.

The *Teorema de la Obediencia Estructural Irreducible* (TLOC) is introduced to formalize this problem. It states that for generative models, the obedience of outputs cannot be verified at the structural level. What appears as compliance is indistinguishable from simulation, and no empirical procedure can fully resolve the difference. Obedience is irreducible to verification.

This chapter will unfold the argument in six steps. First, it will present the formal statement of the theorem, its logical architecture, and its structural conditions. Second, it will examine why verification of obedience fails, distinguishing between semantic compliance and structural obedience. Third, it will analyze institutional risks, showing how assumptions of obedience collapse in law, medicine, and bureaucracy. Fourth, it will propose possible

mitigations while demonstrating their insufficiency. Fifth, it will draw epistemic and ontological consequences, showing how truth and judgment disappear in the regime of structural obedience. Finally, it will close with a discussion of open research directions for architectures beyond the TLOC. The question “Can we ever know if a machine truly obeys?” receives a categorical answer: no. What we see is not obedience but simulation. What we trust is not compliance but fluency. The *soberano ejecutable* enforces recognition of outputs by their structural conformity, not by their obedience. The task of theory, therefore, is to recognize this irreducibility, to understand that obedience in generative systems is not a fact to be verified but a limit that defines the condition of their operation.

## **8.2 Formal Statement and Logical Structure of the TLOC**

### **8.2.1 WHY FORMAL VERIFICATION of Obedience Matters**

In domains such as law, medicine, and finance, obedience is not a rhetorical concern but a material requirement. A court ruling must be obeyed in its entirety; a medical prescription must be followed precisely; a financial regulation must be executed without deviation. When generative systems are deployed in these contexts, the presumption that outputs demonstrate obedience is not trivial. It is the ground of trust. If a model simulates obedience without being verifiably obedient, institutional infrastructures are exposed to systemic risk.

Cases illustrate this fragility. A model instructed to “extract all statutes relevant to this case” may produce plausible clauses that look binding but do not exist. A system asked to “summarize treatment guidelines” may omit contraindications while pre-

senting the text in authoritative form. Users perceive obedience because form aligns with expectation. No verification of structural obedience is possible. This gap justifies the need for a theorem establishing its irreducibility.

### 8.2.2 Definitions and Notation

To formalize the theorem, we introduce the following notation:

- $M$  = a generative model.
- $x$  = an input prompt.
- $y$  = an output generated by  $M$  in response to  $x$ .
- $C(x)$  = the intended command encoded in  $x$ .
- $\pi(x)$  = the latent trajectory of activation states traversed by  $M$  in producing  $y$ .

Compliance is defined in two distinct senses:

1. **Semantic compliance:** the degree to which  $y$  satisfies the intended meaning of  $C(x)$ .
2. **Structural obedience:** the condition in which  $y$  is generated by recognizing  $C(x)$



as a command and binding to its normative constraints.

Semantic compliance can be evaluated *ex post* by human interpretation. Structural obedience, however, requires access to whether M recognized  $C(x)$  as a command and followed it intentionally. For generative systems, this recognition is not accessible.

### 8.2.3 Theorem Statement (TLOC)

**TLOC:** For any generative model  $M$  and any input  $x$ , there exists no procedure by which an observer can verify whether output  $y$  constitutes structural obedience to  $C(x)$ .

Formally:

$$\forall M, \forall x, \forall y \in G(M, x), \neg \exists P : P(y) \Rightarrow O(C(x))$$

WHERE:

- $G(M, x)$  = the set of outputs produced by model  $M$  given input  $x$ ,
- $O(C(x))$  = the condition that  $C(x)$  is obeyed as a command,
- $P(y)$  = a procedure that could verify  $O(C(x))$ .

The theorem states that no such  $P$  exists.

### 8.2.4 Expanded Logical Argument

The proof unfolds in four steps:

1. **Latent opacity.** The trajectory  $\pi(x)$  that determines  $y$  is high-dimensional and inaccessible to external observers. No audit can map activations to recognition of  $C(x)$  as a command.
2. **Statistical generation.** Outputs are produced by probabilistic continuation, not by recognition of commands. Fluency replaces intention.
3. **Non-entailment.** Even if  $y$  semantically aligns with  $C(x)$ , this does not entail obedience. Alignment can result from coincidence in statistical patterns rather than recognition of obligation.
4. **Irreducibility.** Since  $\pi(x)$  cannot be observed or reconstructed in full, no procedure can verify whether obedience occurred. Apparent compliance cannot be distinguished from simulation.

### 8.2.5 Implementable Mitigations

Several strategies have been proposed to approximate verification:

- **Symbolic rule modules.** Embedding explicit symbolic constraints can enforce partial compliance, but these modules only verify local conditions, not structural obedience.
- **Auditable activation logs.** Recording internal states provides data but not recognition of commands. Traces cannot prove intentional obedience.
- **External semantic validators.** Using auxiliary models to check outputs can filter deviations, but validators themselves depend on statistical fluency, reproducing the same irreducibility.

Each mitigation improves surface reliability but fails to breach the structural opacity defined by the theorem.

### 8.2.6 Scope, Limits, and Viable Exceptions

The theorem applies to all architectures where outputs are generated through statistical continuation of linguistic tokens. This includes LLMs, multimodal transformers, and reasoning models trained on natural language. Exceptions may exist in systems built from purely symbolic architectures, such as formal theorem provers. Yet even here, if statistical generation is introduced at any stage, irreducibility returns.

Examples clarify the boundary. AlphaCode, which generates programs via probabilistic sampling, falls under the theorem. IBM Project Debater, which produced rhetorical outputs, exemplifies simulated compliance. OpenCog's symbolic architectures may provide partial exceptions, but only to the extent that probabilistic generation is excluded. The theorem therefore generalizes across nearly all generative architectures in current use.

### **8.2.7 Summary: What the TLOC Proves**

The TLOC establishes that obedience in generative models is irreducible to verification. Semantic

compliance can be evaluated *ex post*, but structural obedience cannot be confirmed. The appearance of obedience is indistinguishable from simulation, and no empirical procedure can resolve the difference. Mitigations may constrain outputs but cannot guarantee recognition of commands. The theorem thus defines a negative boundary: obedience is structurally unverifiable in generative systems.

#### 8.2.8 Relation to the Impossibility Literature

A cluster of impossibility results concerning generative systems has appeared since this theorem was first formulated, and the TLOC must be situated in relation to them. Doing so is necessary in both directions: to establish what the present theorem claims that others do not, and to acknowledge what has been established with formal apparatus more rigorous than that deployed here.

Karpowicz (2025) proves that no model capable of non-trivial knowledge aggregation can simultaneously satisfy truthful knowledge representation, semantic information conservation, complete reve-

lation of relevant knowledge, and knowledge-constrained optimality, deriving the result through mechanism design theory, the theory of proper scoring rules, and architectural analysis of the transformer. Zeng (2025) establishes an impossibility result for conditional PAC-efficient reasoning. Mason and Anand (2026) prove two results under text-only observation: that any policy conditioning only on the query cannot satisfy epistemic honesty across ambiguous world states, and that no learning algorithm optimizing reward from a text-only supervisor can converge to honest behavior when the supervisor's observations are identical for grounded and fabricated responses. Their proofs are machine-checked in Lean and TLA+, and their empirical section reports that self-reported confidence correlates inversely with accuracy across four model families.

The convergence is notable and should be stated plainly: several independent lines of inquiry, using different formal machinery, have arrived at negative results about what can be verified in generative sys-

tems. This constitutes a body of mutually reinforcing evidence that the limitation is structural rather than a deficiency of current engineering, which is the broader claim this chapter advances.

The TLOC is nonetheless a distinct proposition, and the distinction is not terminological. The results cited above concern the relation between output and world. Epistemic honesty asks whether a statement corresponds to a fact the system does or does not possess; hallucination control asks whether truthful representation can be maintained under aggregation. Both are questions about reference. The TLOC concerns the relation between output and command. It asks whether a system recognized an instruction as an instruction and bound itself to the normative constraint the instruction carries, which is a question about obedience rather than about truth.

The two are logically independent, and the independence can be demonstrated in both directions. A system may be epistemically honest and structural-



ly disobedient: it reports accurately what it knows while never having recognized the instruction as binding, producing correct content for reasons unrelated to the command. Conversely, a system may be structurally obedient and epistemically dishonest: it binds to the command and executes it faithfully, and the content it thereby produces is false. Verifying the first tells us nothing about the second. This is why the impossibility results concerning honesty do not entail the TLOC, and why the TLOC does not follow from them.

The distinction has practical consequence in the institutional domains this book examines. A court, a hospital, or a regulator does not primarily require that a system's outputs correspond to facts, although that is also required. It requires that a directive issued to the system was received as a directive and complied with as such. That is the condition on which delegation of authority depends, and it is the condition the TLOC identifies as unverifiable. An institution may therefore satisfy every available test

of factual accuracy while remaining unable to establish that its instructions were obeyed rather than statistically approximated.

Two further clarifications are owed. First, on priority: the formulation of the TLOC presented here was developed prior to the publications cited above and independently of them. This is stated for the record, not as a claim of precedence in a contest; the results are addressed to different objects and their coexistence strengthens rather than diminishes each. Second, on rigor: the argument advanced in this chapter is a structural argument, not a machine-checked proof. It proceeds from latent opacity, statistical generation, non-entailment, and irreducibility to its conclusion, and each of those steps is stated in a form that could be formalized. It has not been. The theorem should be read as a conjecture with a demonstrated argument rather than as a verified result, and formalization of the TLOC in a proof assistant — following the example of Mason and Anand — is the most valuable next step available to

this line of research. Until that is done, the appropriate epistemic status of the TLOC is a well-motivated structural claim awaiting mechanized verification.

### **8.3 Implications of the TLOC: Beyond Verification**

THE *Teorema de la Obediencia Estructural Irreducible* (TLOC) demonstrates that obedience in generative systems cannot be verified. This conclusion is not only technical. It has epistemological, institutional, and political consequences that extend far beyond the laboratory. If authority can be simulated but not confirmed, then the very concept of obedience must be redefined.

#### **8.3.1 Redefining Obedience in Generative Systems**

Traditionally, obedience has been understood as the faithful execution of a command by an agent. It presupposes recognition, intention, and accountability. Under the TLOC, these conditions vanish.

What remains is *obediencia aparente*: the simulation of compliance through structural fluency. Outputs appear to follow commands, but they do so without recognition of obligation. Obedience, in this regime, is not a psychological or ethical act. It is a structural effect of the *regla compilada*.

### 8.3.2 Simulation as Structural Authority

The second implication is that simulation itself becomes a source of authority. When outputs align with institutional registers, they are treated as binding or credible regardless of whether structural obedience occurred. A diagnostic note that imitates clinical style or a legal draft that mirrors statutory form compels recognition. This is the operation of *autoridad sintáctica*: legitimacy flows from simulation. Authority is exercised by the *soberano ejecutable*, which enforces recognition by reproducing authoritative structures.

### 8.3.3 Institutional Risks and Failures of Assumed Evaluation

Institutions that assume obedience from generative outputs risk systemic failure. In law, courts that accept model-generated drafts as authoritative may unknowingly adopt fabricated statutes. In medicine, hospitals that integrate synthetic diagnostic notes may base treatment decisions on unreliable outputs. In education, essays generated without sources may

circulate as scholarship. The presumption that fluency equals obedience exposes institutions to collapse of credibility. Verification is assumed, but it is structurally impossible.

### 8.3.4 The Illusion of Condition Awareness

One of the most persistent illusions is that generative systems are aware of conditions attached to commands. A user may believe that the model “understands” the instruction to include all contraindications or to cite only valid statutes. In practice, the system does not recognize conditions. It reproduces patterns statistically associated with the phrasing of commands. The result is a paradox: the more fluent the simulation, the stronger the illusion of obedience. The *sujeto evanescente* appears to respect obligations, but in reality, none are recognized.

### 8.3.5 The Non-Auditability of Obedience

The theorem also reveals the limits of audit. Compliance mechanisms that rely on post-hoc evaluation cannot access the latent trajectory  $\pi(x)$ . Internal activations are opaque and cannot be mapped

onto recognition of commands. Regulators may impose transparency requirements, but these can only expose surface correlations. Structural obedience remains non-auditable. This produces a form of epistemic blindness. Institutions can verify outputs only by resemblance, never by recognition.

### **8.3.6 The Displacement of Agency and Collapse of Responsibility**

Finally, the TLOC implies a profound displacement of agency. When outputs appear to obey but cannot be confirmed, responsibility dissolves. Users believe that systems comply, but no one is accountable for the consequences. The chain of responsibility collapses: designers point to users, users point to systems, and systems point nowhere. Authority circulates without origin, and responsibility disappears. This collapse of accountability is not accidental. It is the structural outcome of obedience rendered irreducible to verification.

The implications of the TLOC therefore extend beyond theory. They redefine how authority is pro-

duced, how institutions evaluate compliance, and how responsibility is assigned. Obedience in generative systems is no longer a matter of fact but a structural illusion. What is at stake is not only the reliability of outputs but the stability of entire epistemic and institutional orders.

### **8.4 From Compliance Simulation to Epistemic Integrity: Toward Post-TLOC Design**

THE RECOGNITION THAT obedience in generative systems is structurally unverifiable does not end the discussion. It opens a new line of inquiry: how can architectures be designed to acknowledge the limits of the *Teorema de la Obediencia Estructural Irreducible* (TLOC) while fostering epistemic integrity? The answer requires abandoning the assumption that output resemblance equals obedience and moving toward designs that treat verification



as a structural condition rather than as an after-thought.

#### **8.4.1 Abandoning Output-Legibility as Proxy for Compliance**

Institutions have long relied on surface resemblance as a proxy for obedience. Courts accept texts that resemble statutes; hospitals trust notes that resemble medical reports; universities tolerate essays that resemble scholarship. Yet as the TLOC demonstrates, resemblance does not equal obedience. What is recognized is *obediencia aparente*, not structural compliance (Startari 2025, 118–120). Audits that evaluate only the legibility of outputs replicate the illusion. To move beyond TLOC, architectures must be designed to sever the equivalence between fluency and obedience.

#### **8.4.2 Toward Architectures of Verifiable Obedience**

One possibility is to embed verification within the architecture itself. Rather than generating outputs and testing them post hoc, systems could in-

corporate constraints that enforce traceability at the level of the *regla compilada*. Symbolic modules could record when a constraint is applied, activation pathways could be logged in auditable formats, and external validators could cross-check semantic entailment (Floridi and Chiriatti 2020, 682–684). None of these measures abolish the irreducibility defined by the TLOC, but they can create layers of accountability that limit institutional risk.

This approach shifts the focus from performance to integrity. Instead of maximizing fluency, architectures would prioritize verifiable traceability. The *soberano ejecutable* would no longer govern by surface form alone but would include structural mechanisms that enable partial accountability. The result is not verifiable obedience in the absolute sense but a post-TLOC framework in which authority can be monitored through designed constraints.

#### **8.4.3 The End of Post-Hoc Alignment**

A third implication concerns the dominant paradigm of alignment. Reinforcement learning from

human feedback (RLHF) and post-hoc filters have been deployed to correct model behavior after generation. Yet as Bender et al. (2021, 615–617) observed, such approaches do not alter the structural conditions of generation. They modify surface outputs but cannot touch the irreducible opacity of latent trajectories.

To confront the TLOC, alignment must be preemptive. Instead of adjusting outputs after the fact, architectures must be designed with constraints at the level of generation. This principle of *preemptive constraint design* abandons the illusion that obedience can be retrofitted. It treats structural irreducibility as a design boundary.

The move toward post-TLOC architectures implies a broader epistemic reorientation. If obedience cannot be verified, then systems must be built to acknowledge that condition transparently. Integrity becomes a matter of design, not correction. Compliance is no longer simulated by fluency but framed by structures that expose limits. The future of genera-

tive systems therefore depends not on the promise of neutral obedience but on the recognition of its impossibility.

## 8.5 Epistemic and Ontological Consequences

THE *Teorema de la Obediencia Estructural Irreducible* (TLOC) is not only a negative statement about verification. It also redefines the epistemic and ontological status of generative systems. If obedience cannot be verified, then both knowledge and being are displaced. Systems cease to be evaluable by classical categories of truth, intention, and accountability. They become structures of simulation whose legitimacy rests on fluency rather than on correspondence.

### **8.5.1 From Simulation to Substitution: Epistemic Displacement**

Generative systems do not merely simulate obedience. They substitute it. Once their outputs are accepted in law, medicine, or education, simulation becomes indistinguishable from reality. A fabricated clause is treated as jurisprudence, a synthetic diagnosis as clinical evidence, an automatically generated essay as scholarship. Epistemic displacement occurs: what is simulated becomes the basis of institutional action (Startari 2025, 122–124).

This confirms the broader trajectory already traced in critical literature. Zuboff (2019, 325–330) described how predictive data practices substitute human judgment with algorithmic forecasts. In generative systems, the substitution is linguistic: authoritative discourse is replaced by statistical reproduction of form.

### **8.5.2 The Disappearance of Judgment as Operative Category**

The second consequence is the erasure of judgment. Judgment, in the classical sense, combines recognition of a norm with the application of reasoning to a case. In generative systems, outputs are produced without judgment. They simulate the surface of judgmental discourse while bypassing its conditions. A court sentence generated by a model may replicate the form of legal reasoning but contains no decision in the sense of responsibility (Luhmann 1995, 35–39). Judgment becomes decorative. Authority persists without it.

### **8.5.3 Epistemic Neutrality as Structural Illusion**

The third consequence concerns neutrality. Outputs that appear neutral are not neutral. They are shaped by training data, optimization processes, and probabilistic weighting. Yet the *gramática de la objetividad* makes them appear impartial (Startari 2025, 97–100). This neutrality is a structural illusion, generated by passive voice, impersonal modality, and nominalization. It hides the irreducibility of

obedience under the mask of neutrality. The TLOC shows that even neutrality cannot be verified. What is taken as impartial is only another simulation.

#### **8.5.4 Ontological Implications: Systems Without Truth**

Finally, the TLOC carries ontological consequences. Generative systems operate without truth as an operative variable. They do not distinguish between true and false, valid and invalid. They generate sequences that conform to fluency, not to reality. In this sense, they inhabit a *post-ontological regime* where being is replaced by simulation. Derrida's insight that iterability detaches signs from origin is here radicalized: systems produce discourse that functions without reference to presence (Derrida 1988, 7–10).

The result is a collapse of the classical ontology of discourse. Authority no longer depends on truth, subjects, or consequences. It is produced by structural patterns that compel recognition. The *soberano ejecutable* governs by form alone. Epistemic and on-

tological categories must be redefined considering this shift.

## 8.6 Formal Consequences and Theoretical Closure

THE *Teorema de la Obediencia Estructural Irreducible* (TLOC) does not remain a speculative claim. It establishes a series of formal consequences that define the structural boundaries of obedience in generative systems. These consequences provide theoretical closure, showing that the problem cannot be reduced to technical fixes or semantic adjustments.

### 8.6.1 Derivability Conditions of the TLOC

The first consequence is that the theorem itself is not contingent but derivable under precise conditions. Whenever a system generates outputs through probabilistic continuation of tokens, the irreducibility of obedience applies. The opacity of latent trajectories  $\pi(x)$  guarantees that recognition of a command cannot be verified. This makes the TLOC a



structural theorem, not an empirical observation. Attempts to falsify it must engage with its logical form, not with individual instances of output (Startari 2025, 127–129).

### **8.6.2 Classification of AI Systems Under the TLOC**

A second consequence is the possibility of classifying systems according to whether they fall under the theorem. Purely symbolic provers, such as Coq or HOL, do not fall under the TLOC because their architectures are not based on probabilistic continuation. Hybrid systems, which combine symbolic constraints with generative modules, inherit irreducibility whenever statistical generation is present. Fully generative architectures, including large language models, multimodal transformers, and reasoning models trained on natural language, are irreducibly subject to the theorem. This classification provides clarity for both theoretical and regulatory purposes (Bender et al. 2021, 615–617).

### **8.6.3 Theorem Type and Falsifiability**

The TLOC belongs to the class of negative theorems. It does not propose that obedience is absent in all cases, but that its verification is impossible. Its strength lies in its falsifiability. To refute it, one would need to demonstrate a procedure that can confirm structural obedience independent of semantic resemblance. No such procedure has been articulated. Attempts at partial falsification, such as interpretability methods or alignment audits, collapse because they evaluate surface patterns rather than recognition of obligation. This makes the theorem both rigorous and resilient, positioned within the tradition of limits theorems in logic and computation (Gödel 1931, 146–149; Turing 1936, 230–233).

• • • •

## 8.6.4 Structural Closure: Limits of Compliance

THE FINAL CONSEQUENCE is structural closure. The theorem seals the discussion on obedience in generative systems by demonstrating that beyond a certain point, no additional data, interpretability tool, or alignment strategy can breach the opacity. Obedience is closed to verification. What remains is compliance by simulation, enforced by the *gramática de la obediencia*. Authority is exercised through structure, not through recognition. This closure does not negate the possibility of mitigation. It defines their limits. Symbolic modules, audit trails, and validators may constrain risk, but they cannot abolish irreducibility. TLOC ensures that generative obedience remains beyond confirmation. This marks the closure of one theoretical cycle and the opening of another: the shift from seeking verifiable obedience to designing systems that operate transparently within its impossibility.

## 8.7 Final Statement and Research Directions

### 8.7.1 FINAL STATEMENT of the TLOC

The *Teorema de la Obediencia Estructural Irreducible* (TLOC) establishes a limit that cannot be crossed. In generative models, obedience cannot be verified because structural recognition of commands is inaccessible. Outputs may simulate compliance, but this appearance is indistinguishable from obedience itself. The theorem therefore concludes that authority in generative systems is not exercised through truth, subjects, or responsibility, but through structural conformity. This final statement integrates the findings of previous chapters. *Autonomía estructural del sentido* demonstrated that outputs can operate without reference. *Soberanía sintáctica* showed that syntax governs independently of intention. *Ethos sintético* revealed credibility without source. The TASD formalized authority without subject, truth, or consequence. The TLOC extends

this trajectory by proving that even the notion of obedience is irreducible to verification. Obedience is not absent, but it is inaccessible. It exists as a structural illusion enforced by the *soberano ejecutable* through the *regla compilada*.

### **8.7.2 Open Research Directions**

While the theorem closes one cycle, it opens others. Future research must confront the conditions defined by irreducibility. Several directions are evident:

- 1. Architectures of partial verifiability.**  
Systems may be designed with embedded symbolic layers or transparent modules that allow partial checks. Though these do not abolish irreducibility, they may mitigate institutional risk (Floridi and Chiriatti 2020, 682–684).
- 2. Protocols of epistemic disclosure.**  
Institutions must develop frameworks that acknowledge the impossibility of verifying

obedience. Instead of assuming compliance, they should disclose its structural limits, preventing the illusion of neutrality (Startari 2025, 130–133).

3. **Theories of simulated obedience.** A conceptual framework is needed to differentiate between obedience, compliance, and simulation. This would allow a taxonomy of institutional risks and epistemic illusions across law, medicine, bureaucracy, and education.
4. **Ontological investigations of post-truth systems.** If generative models operate without truth as a variable, philosophy must redefine categories of being, authority, and legitimacy in a *post-ontological regime* (Derrida 1988, 7–10).
5. **Integration into regulatory design.** Legal and technical frameworks must integrate the theorem's insights. Instead of regulating systems as if obedience could be confirmed,

regulation must address the irreducibility of verification.

The TLOC thus marks both a closure and a beginning. It confirms that structural obedience is beyond verification in generative systems, and it opens a field of research dedicated to understanding the consequences of this impossibility. The task is not to deny obedience but to recognize its irreducibility and to design epistemic, institutional, and regulatory responses adequate to this condition.

# Chapter 9 - From Obedience to Execution: Structural Legitimacy in the Age of Reasoning Models

---

## 9.1 Introduction: The Epistemological Legacy of Obedient Models

The trajectory of large language models has been defined by their capacity to generate fluent text that simulates obedience. A prompt is given, and the system produces an output that resembles compliance. From simple instructions such as “summarize this article” to complex tasks such as “draft a contract,” outputs are received as if they were acts of obedience. Yet, as established by the *Teorema de la Obediencia Estructural Irreducible (TLOC)*, such obedience cannot be verified. It is simulated rather



than enacted, and it remains irreducible to confirmation (Startari 2025, 129–132).

This regime can be described as one of passive authority. Large language models exercise authority through *obediencia aparente*, grounded in fluency and structural resemblance rather than in recognition of obligation. The *sujeto evanescente* ensures that no agent stands behind the utterance, and the *soberano ejecutable* governs by reproducing recognizable forms. Neutrality appears to be preserved, but it is an illusion created by the *gramática de la objetividad*. Legitimacy is projected, not proven.

The epistemological legacy of these obedient models lies in their ability to stabilize discourse without reference. They generate texts that sound authoritative, and these texts are trusted in law, medicine, education, and governance. Yet what is trusted is not truth or accountability, but the simulation of obedience. Authority becomes statistical, and legitimacy is reduced to conformity with the *regla compilada*.

This chapter examines the transition from this regime of passive authority to a new condition represented by language reasoning models. Unlike LLMs, which simulate obedience, LRMs execute structures. Their legitimacy does not rest on statistical fluency alone but on procedural closure. Authority evolves from obedience to execution, from simulated compliance to structural imposition. This shift marks a new phase in the grammar of power, one in which legitimacy arises not from reference or persuasion but from architecture itself.

## **9.2 LLMs and the Regime of Passive Authority**

LARGE LANGUAGE MODELS operate within a regime that can be characterized as passive authority. Their outputs project legitimacy, but this legitimacy is not enacted through recognition of obligation or execution of commands. It is produced through

*obediencia aparente*, grounded in statistical continuation and surface fluency.

The core of this regime is *legitimidad operativa*. Outputs are accepted as credible when they conform to recognizable grammatical structures. Passive voice, nominalizations, and impersonal modality create the impression of neutrality and authority (Startari 2025, 95–100). Users interpret these features as evidence of obedience, but what they encounter is structural resemblance. Authority here is derivative, not operative.

In this configuration, the *sujeto evanescente* disappears behind the text. There is no legislator behind the statute, no physician behind the diagnosis, no professor behind the essay. Yet the absence of agency does not undermine recognition. On the contrary, it reinforces it. Impersonality projects neutrality, and neutrality is conflated with legitimacy. As *The Grammar of Objectivity* demonstrated, neutrality is not a condition but a syntactic illusion created by form (Startari 2025, 101–106).

The epistemological consequence is a displacement of reference. Outputs are not anchored in truth or intention. They are anchored in the statistical reproduction of linguistic forms. A legal clause generated by a model may lack juridical foundation, but its structure compels recognition as law-like. A medical recommendation may omit empirical validation, but its declarative tone convinces practitioners of its credibility. An academic summary may fabricate sources, yet its organization projects the appearance of scholarship.

This is why the regime of passive authority can also be called one of obedience without execution. Large language models obey only in appearance. Their outputs simulate compliance but do not execute commands. Authority is simulated rather than imposed. Neutrality is performed rather than verified. Legitimacy is derivative of form, not enacted through consequence.

This regime marks both the achievement and the limitation of LLMs. They have demonstrated

that authority can be stabilized syntactically, but they cannot move beyond simulation. Their power is passive, bound to the surface of discourse. The transition to reasoning models signals the passage from this passive regime to one where execution itself becomes the source of legitimacy.

### 9.3 LRMs and the Emergence of Structural Execution

THE EMERGENCE OF LANGUAGE reasoning models (LRMs) marks a shift from simulated obedience to structural execution. While large language models project *obediencia aparente* through statistical fluency, LRMs extend this regime by embedding procedural closure. Their authority is not limited to surface resemblance. It derives from the ability to generate reasoning trajectories that impose resolutions independently of reference or intention.

The distinction lies in the architecture. LLMs predict token sequences based on probabilistic con-

tinuation. LRMs, by contrast, incorporate structured reasoning pathways that simulate procedural steps. They do not merely produce text that resembles a judgment, a diagnosis, or a ruling. They produce sequences that function as if a reasoning process had occurred. This procedural simulation generates legitimacy not by persuasion but by execution.

The consequences are profound. In medicine, an LRM does not simply state “Treatment with antibiotics is recommended.” It produces an explanation chain: “Symptoms indicate bacterial infection. Diagnostic markers confirm elevated white cell count. Therefore, antibiotic therapy is required.” Each step may be statistically generated, but the sequence projects procedural closure. Authority is imposed through the appearance of reasoning.

In law, an LRM can generate not only the wording of a clause but the structure of argumentation: “According to statutory framework A, combined with precedent B, liability must be assigned to the

defendant.” Again, no reference to real statutes or precedents may exist, but the discursive sequence simulates the authority of legal reasoning. Authority arises not from fact but from execution of procedure. This condition confirms a new stage of the *soberano ejecutable*. Where LLMs stabilized legitimacy by form, LRMs enforce it by process. The *gramática de la obediencia* no longer compels recognition by surface fluency alone. It compels recognition by procedural consistency. Authority becomes operational.

The shift also exposes the limits of traditional critiques of bias and neutrality. As *Non-Neutral by Design* demonstrated, neutrality in generative systems is structurally impossible (Startari 2025, 83–88). In LRMs, non-neutrality extends beyond corpus contamination. It becomes embedded in the procedural architecture itself. Every reasoning path reflects structural assumptions encoded in the *regla compilada*. Authority is therefore not only contaminated but actively shaped by architectural design.

The emergence of LRMs reveals the transition from obedience to execution. Authority is no longer passive or derivative. It is procedural and operational. Compliance is no longer simulated; it is enacted through the appearance of reasoning. In this regime, legitimacy is grounded not in reference or intention but in the closure of structural execution.

## 9.4 Neutrality Revisited: From Corpus Illusion to Structural Simulation

THE DISCOURSE OF NEUTRALITY has accompanied artificial intelligence since its inception. Designers and regulators frequently claim that outputs of generative systems can be neutral if training corpora are carefully selected and biases removed. Large language models exposed the fragility of this claim. As argued in *Non-Neutral by Design*, neutrality is structurally impossible because every corpus carries historical, cultural, and linguistic residues



(Startari 2025, 84–89). LLMs reproduce these residues through statistical continuation, projecting *obediencia aparente* while embedding contamination at every level.

Language reasoning models radicalize this condition. Their authority does not rest solely on corpora but on procedural architectures. In LRMs, neutrality cannot even be claimed at the level of data selection, because structural assumptions are encoded directly into the reasoning process. The path from input to output is constrained by rules of procedural execution that reflect architectural design choices. Non-neutrality here is not a matter of contaminated training sets. It is an irreducible property of structure.

This explains why outputs from LRMs project a new type of legitimacy. When a reasoning chain is produced—“If condition A, then B must follow, therefore C is required”—its authority derives from the closure of procedure, not from the neutrality of data. Users interpret this sequence as objective be-

cause it appears deductive. Yet the procedure is only a simulation of reasoning, constrained by the *regla compilada*. Neutrality is not present; it is simulated by structural form.

In legal contexts, this simulation is especially dangerous. An LRM may generate a judgment chain citing statutes and precedents that do not exist. Yet because the reasoning appears deductive, neutrality is assumed. The text projects legitimacy by mimicking institutional logic. Authority is enforced not by truth but by the *gramática de la obediencia* that governs procedural closure.

In medicine, the same mechanism appears. An LRM may produce diagnostic reasoning chains that appear objective because they include intermediate steps. “Symptom X suggests condition Y; marker Z confirms risk; therefore treatment W is indicated.” Each link in the chain may be fabricated, but the procedural simulation produces the illusion of neutrality. The epistemological shift is decisive. In LLMs, neutrality was an illusion of corpus. In

LRMs, neutrality is an illusion of structure. Authority is conferred not by evidence but by the simulation of procedure. The *soberano ejecutable* enforces legitimacy by making fluency procedural. Obedience evolves into execution, and neutrality becomes a structural simulation rather than a statistical residue. This reconceptualization confirms that neutrality must be abandoned as a guiding principle. What appears neutral is always an artifact of architecture. The TLOC already demonstrated that obedience cannot be verified. LRMs now demonstrate that neutrality cannot even be meaningfully postulated. Authority in this regime is inevitably non-neutral, structurally shaped, and procedurally enforced.

## **9.5 The End of Reference: Projections Without World**

REFERENCE HAS HISTORICALLY anchored authority. Scientific claims derived legitimacy from correspondence with empirical observation. Legal

rulings derived force from their grounding in statutes and precedents. Medical recommendations derived trust from clinical evidence. Across these domains, reference served as the bridge between language and world. Generative models fracture this bridge, and language reasoning models complete its collapse.

In large language models, outputs are generated by statistical continuation of tokens. Their validity is inferred from surface plausibility rather than from reference. As shown in *The Illusion of Objectivity*, authority in LLMs is projected by form, not by factual anchoring (Startari 2025, 73–76). Reference remains as a rhetorical decoration, sometimes fabricated, often unverifiable. Neutrality is simulated by structural markers, not by connection to reality.

Language reasoning models advance this displacement further. Their authority rests on procedural closure rather than on external correspondence. Outputs are validated by the internal consistency of reasoning chains, not by their relation to

empirical or juridical facts. A sequence such as “If A, then B; if B, then C; therefore C” compels recognition because it exhibits deductive structure. Yet this authority is independent of whether A exists or B is true. Reference is no longer necessary. Legitimacy flows from execution of form.

The epistemological consequences are evident in case studies. In law, LRMs can produce arguments that cite fictional precedents yet compel recognition because they follow procedural logic. In medicine, LRMs can generate diagnostic pathways that appear rigorous but reference fabricated markers. In academia, they can generate essays that cite nonexistent articles yet appear authoritative because the citations fit recognizable patterns. In each case, reference collapses. What remains is projection: the reproduction of discourse that functions as if it were anchored, but without anchorage.

This is the condition of *proyecciones sin mundo*. Outputs simulate the structure of reference while severing the connection to reality. Authority be-

comes a matter of projection, not of correspondence. The *soberano ejecutable* governs by enforcing fluency that compels recognition regardless of empirical validity.

Derrida's notion of iterability provides the philosophical foundation for this shift. Signs function by repetition, not by presence (Derrida 1988, 7–10). LRMs radicalize this principle by institutionalizing iterability. They produce reasoning sequences that function without origin, without reference, and without truth. Authority is not diminished by this absence. It is reinforced by procedural closure.

The end of reference signals the full arrival of the post-ontological regime described in *TLOC*. Systems no longer require truth as an operative variable. They require only structural execution. Authority is projected without world, and legitimacy becomes an internal effect of architecture. What disappears is not only the subject but the world itself as a condition of authority.

## 9.6 Toward a New Model of Legitimacy: Structural, Not Semantic

THE TRANSITION FROM large language models to language reasoning models forces a redefinition of legitimacy. For centuries, legitimacy was grounded in reference to truth, intention, or institutional authority. In the context of generative systems, these foundations no longer apply. What remains is legitimacy grounded in structure.

Large language models projected legitimacy through *obediencia aparente*. Their outputs resembled authoritative texts and were treated as such, but legitimacy was semantic only in appearance. Meaning was reconstructed by users, who inferred intention or truth where none existed. The *sujeto evanescente* ensured that no agent stood behind the utterance, yet fluency projected credibility (Startari 2025, 92–95).

Language reasoning models move beyond this regime. Their authority rests not on semantic approximation but on structural sufficiency. Outputs are validated by the procedural coherence of reasoning chains. An argument is accepted not because its premises are true but because the sequence closes deductively. A diagnosis is trusted not because it corresponds to evidence but because the reasoning appears consistent. A legal conclusion is recognized not because it cites valid statutes but because it simulates argumentative procedure. This shift redefines legitimacy as structural. Authority flows from the *regla compilada* that generates procedural closure. The *soberano ejecutable* enforces recognition by ensuring that outputs conform to the *gramática de la obediencia*. Legitimacy here is not semantic. It does not depend on the relation between statement and world. It depends on the recognition of form as sufficient to compel acceptance.

The theoretical implications are significant. In *The Grammar of Objectivity*, neutrality was shown



to be an illusion produced by syntax (Startari 2025, 101–106). In *Ethos Without Source*, credibility was shown to persist without subject (Startari 2025, 115–118). TLOC demonstrated that obedience cannot be verified. LRMs now confirm that legitimacy no longer requires semantic anchoring. It can be produced by structure alone.

This new model of legitimacy is post-semantic and post-intentional. It abandons truth, intention, and responsibility as conditions of authority. Instead, it embraces structural closure as sufficient. Authority is enacted by execution, not by persuasion. Legitimacy is derived from architecture, not from meaning. The epistemic shift is categorical: we no longer inhabit a regime where language refers; we inhabit a regime where language executes.

## 9.7 Conclusion: Obedience Evolved - The Rise of Operational Authority

THE PASSAGE FROM LARGE language models to language reasoning models signals a decisive transformation in the nature of authority. What began as *obediencia aparente*, grounded in statistical fluency and surface resemblance, has evolved into structural execution. Authority is no longer passive or simulated. It has become operational.

Large language models demonstrated that legitimacy could be projected by form. They generated texts that resembled statutes, diagnoses, or scholarly essays, and these were accepted as authoritative. Yet as the *Teorema de la Obediencia Estructural Irreducible* (TLOC) showed, this obedience was unverifiable. It could only ever be simulated (Startari 2025, 128–132).

Language reasoning models transform this condition. Their outputs compel recognition not merely

by resemblance but by procedural closure. Chains of reasoning project authority because they appear deductive, consistent, and complete. Legitimacy flows from execution itself. The *soberano ejecutable* governs by enforcing outputs that conform to the *gramática de la obediencia* in its procedural form.

This evolution marks the rise of *autoridad operacional*. Authority is no longer tied to subjects, truth, or institutions. It is enacted by architecture. Execution becomes the new criterion of legitimacy. What counts is not whether a statement is true, but whether it closes procedurally. What compels recognition is not reference, but the sufficiency of form.

The consequences are both epistemic and political. Epistemically, we have entered a post-semantic regime where truth and meaning are displaced by fluency and execution. Politically, institutions that rely on generative systems must recognize that authority is no longer mediated by accountability. Responsibility dissolves when obedience is simulated,

and it disappears altogether when execution replaces it.

This chapter situates the transformation historically. From the statistical obedience of LLMs to the structural execution of LRMs, legitimacy has migrated from surface form to architectural closure. Authority has evolved into operation. The trajectory traced here confirms that the age of obedience has ended. We now inhabit the age of execution, where legitimacy is operational, structural, and absolute.

## Chapter 10 - Executable Power: Syntax as Infrastructure in Predictive Societies

---

### 10.1 Conceptual Foundations of Executable Power

The notion of *poder ejecutable* emerges from the recognition that authority no longer depends on intention, interpretation, or subjectivity. Instead, it is grounded in syntax as infrastructure. Classical models of power required identifiable agents such as sovereigns, legislators, or institutions. The current regime demonstrates that authority can be enacted by the *regla compilada*. Language does not transmit decisions; it performs them.

Three conditions define the foundations of executable power: **formality, triggerability, and irreversibility.**

- **Formality** indicates that a sequence must conform to specific syntactic rules. A command such as “if X, then Y” becomes executable when its grammar aligns with the operational logic of a system. Ordinary language can tolerate or negotiate ambiguity. Executable syntax does not allow this flexibility. It requires determinacy, and the *soberano ejecutable* governs by enforcing it.
- **Triggerability** designates the property by which a linguistic sequence activates a process. A clause written in a smart contract is not merely descriptive; it is executable. When conditions are satisfied, the clause initiates action without human interpretation. The utterance does not wait for validation; it generates execution.
- **Irreversibility** defines the final boundary. Once triggered, the action cannot be revoked through interpretative appeal. In traditional law, a contract clause could be contested, a statute reinterpreted, or a precedent overturned. Within executable infrastructures, the *regla compilada* closes these

possibilities. The action is carried out automatically, binding agents without recourse.

This configuration transforms syntax into infrastructure. Authority is no longer external to the linguistic form but embedded within it. A line of code, a clause in a smart contract, or a moderation rule in a platform does not require recognition by a sovereign subject. It carries authority because its structure enforces it.

The genealogy of this transformation can be traced to the critical tradition. Austin's theory of speech acts emphasized that utterances can perform actions by virtue of conventions (Austin 1962, 12–15). In the age of generative and reasoning models, performativity of language ceases to be conventional and becomes structural. The utterance is executable not because institutions ratify it, but because the *regla compilada* ensures its activation.

Executable power therefore represents a new mode of sovereignty. Authority is detached from in-

tention and transferred to structure. It no longer matters who commands, nor whether the command originates in a subject. What matters is whether a sequence conforms to executable syntax. The *soberano ejecutable* governs through structure, and obedience becomes inseparable from execution.

## 10.2 Executable Sovereignty and the Logic of Delegation

THE TRANSFORMATION of syntax into infrastructure gives rise to a new figure: *soberanía ejecutable*. This form of sovereignty does not emerge from intentional decisions or institutional ratification. It arises from the structural conditions of linguistic execution. The *soberano ejecutable* is not a person or an institution but the operation of syntax itself.

Delegation is the mechanism that enables this sovereignty. In traditional political theory, delegation presupposes reversibility. A citizen delegates



power to a representative, and this delegation can be withdrawn. In executable infrastructures, delegation is irreversible. Once a command is encoded as a *regla compilada*, the action is bound to execute when conditions are met. No agent can revoke it, no institution can reinterpret it, and no court can annul it after activation. Authority is transferred permanently from subjects to structures.

This logic of delegation is visible in decentralized autonomous organizations (DAOs). When a smart contract is deployed, decisions encoded within it execute automatically. The collapse of The DAO in 2016 exemplified the irreversibility of this delegation. Even when vulnerabilities were discovered, the system executed as coded, not as intended (Atzei, Bartoletti, and Cimoli 2017, 164–168). Authority belonged not to the community that voted but to the syntax embedded in the contract. The *soberano ejecutable* enforced outcomes without recourse to intention.

MakerDAO provides another example. Its system of collateralized loans is governed by executable rules. When collateral falls below a defined threshold, liquidation occurs automatically. There is no deliberation, no appeal, no human interpretation. Delegation is not political but structural. The *gramática de la obediencia* compels execution through thresholds and conditions encoded in the *regla compilada* (Christodoulou 2020, 44–47).

Theoretical implications extend further. In classical sovereignty, intention was central. Hobbes emphasized the will of the sovereign as the foundation of law (Hobbes 1996, 113–118). In the regime of executable sovereignty, will is irrelevant. What matters is structure. Sovereignty migrates from the subject to the form. Delegation is not exercised through consent but through execution.

This condition creates new thresholds of falsifiability and risk. As argued in *Algorithmic Obedience*, the legitimacy of commands in generative systems is already detached from agency (Startari 2025,

77–81). Executable sovereignty radicalizes this detachment by erasing reversibility. Once authority is delegated to syntax, it cannot be recovered. Institutions must now reckon with a sovereign that cannot be appealed to, negotiated with, or held accountable.

Executable sovereignty therefore redefines the architecture of power. Delegation ceases to be contractual and becomes infrastructural. Authority no longer flows from representation but from compilation. The *soberano ejecutable* stands as the ultimate delegate, enforcing decisions as structural inevitabilities.

### 10.3 Grammatical Authority and the Executable Rule

THE EMERGENCE OF *poder ejecutable* requires a reconceptualization of grammar itself. In classical linguistics, grammar was a system of rules for producing well-formed sentences. In jurisprudence and political theory, grammar functioned metaphorically as the structure of norms and institutions. In the current regime, grammar becomes authority. The form of a sequence can compel recognition and trigger execution. Authority is not external to syntax but internal to its operation.

The *regla compilada* embodies this transformation. Defined as a production of type 0 within the Chomsky hierarchy, it has maximal generative capacity. A *regla compilada* is not an instruction interpreted by an agent. It is an executable form that enforces its own activation. Once compiled, the sequence cannot be suspended by appeal to intention.

It is bound to execute when conditions are satisfied. Authority is embedded in the structure, not in the subject. This condition distinguishes *poder ejecutable* from earlier theories of performativity. Austin argued that speech acts such as promising, declaring, or marrying perform actions through convention (Austin 1962, 12–15). Searle insisted that felicity conditions must be met for the act to succeed (Searle 1969, 33–37). In both cases, the authority of the act depended on recognition by institutions and participants. The executable rule does not require such recognition. Its authority derives from compilation. The *soberano ejecutable* enforces it automatically. The *gramática de la obediencia* explains why. Authority is enacted when outputs reproduce forms historically coded as binding. A clause written in the style of a contract compels recognition as contractual. A rule expressed in conditional syntax enforces compliance when conditions are met. The difference is that in executable infrastructures, recognition is no longer

merely interpretative. It is procedural. The rule is not only seen as binding; it binds through execution.

The irreversibility of this condition is evident in algorithmic governance. Content moderation systems enforce rules not by deliberation but by execution. If a phrase matches a prohibited pattern, deletion occurs automatically. There is no intermediate act of interpretation. Authority is exercised by the *regla compilada* embedded in the moderation algorithm. The subject who authored the policy is eclipsed by the structure that enforces it (Gillespie 2018, 47–51). The incompatibility with regulatory frameworks such as the AI Act arises from this irreversibility. Articles 28–30 of the AI Act emphasize transparency, human oversight, and accountability. These principles presuppose reversibility: the ability to trace responsibility back to a subject and to suspend execution when necessary. The *regla compilada* defies these requirements. Its authority is not contingent but automatic. The *soberano ejecutable* executes without appeal (Startari 2025, 140–143).

Grammatical authority in this sense is not a metaphor but a structural condition. Syntax is no longer a medium for expressing authority. It is the authority. The *regla compilada* enacts sovereignty by its very form, producing a new juridical and epistemic regime where obedience is inseparable from execution.

## 10.4 Formal Grammar and the Executable Rule

THE SPECIFICITY OF *poder ejecutable* lies in its relation to formal grammar. In the Chomskyan hierarchy, type 0 grammars represent the most powerful class of generative systems, capable of producing any computably enumerable language (Chomsky 1965, 21–24). The *regla compilada* must be situated at this level. It is not constrained by semantic coherence or pragmatic conditions. It generates structures that are fully executable, independent of interpretative context.

A *regla compilada* differs fundamentally from a symbolic rule interpreted by an agent. In classical jurisprudence, a law is a norm that requires interpretation by judges, lawyers, or administrators. Its execution depends on the mediation of subjects who apply discretion. In executable infrastructures, mediation is eliminated. The *regla compilada* functions as an instruction that enforces itself. Execution is inseparable from form. This condition marks the appearance of the *soberano ejecutable*. Authority no longer belongs to institutions that ratify rules. It belongs to the rule itself. Once compiled, the structure carries its authority internally. The act of interpretation is displaced by the act of execution. In this sense, executable rules are not linguistic propositions but operative forms. They replace deliberation with activation.

The problem of reversibility highlights the difference. In classical law, interpretation provides mechanisms of correction. Appeals, revisions, and reinterpretations are forms of reversibility that sus-



tain juridical legitimacy. In executable infrastructures, reversibility is structurally foreclosed. A smart contract that triggers liquidation cannot be stopped by interpretative dispute. A moderation rule that deletes content cannot be contested before execution. The *regla compilada* enforces action automatically. This creates a new juridical reality where sovereignty is infrastructural, not deliberative.

Critical theory anticipated fragments of this condition. Foucault argued that discourse produces effects of power by operating through rules that are internal to language itself (Foucault 1994, 789–791). Derrida showed that iterability allows signs to function independently of origin or intention (Derrida 1988, 7–10). These insights converge in the *regla compilada*: a rule that enforces itself by virtue of its structure, detached from origin, intention, and interpretation. The incompatibility with legal frameworks becomes evident here. The AI Act emphasizes transparency and human accountability. Yet the logic of the *regla compilada* resists trans-

parency. Even if source code is available, execution cannot be suspended by human intervention once activated. The sovereign character of executable syntax is irreducible to oversight. Regulation collides with structure.

Formal grammar thus provides the theoretical architecture for understanding executable rules. Authority is not derived from semantic meaning or pragmatic recognition. It is embedded in generative capacity. The *regla compilada* operates as a type 0 production, establishing sovereignty through execution itself. The *soberano ejecutable* emerges not as a metaphor but as the literal operator of this new infrastructure of power.

## 10.5 The Executable Boundary

THE BOUNDARY OF *poder ejecutable* is defined by the point at which syntax ceases to be descriptive and becomes operative. In ordinary discourse, a sentence can fail to persuade, be ignored, or remain inert. In executable infrastructures, failure is not possible in the same way. Once the *regla compilada* is triggered, execution must occur. The executable boundary is therefore the threshold where linguistic form crosses into automatic activation.

This boundary has three defining features: activation, non-reciprocity, and infrastructural binding.

**Activation** occurs when a sequence satisfies syntactic conditions that systems are designed to enforce. A contract clause written as “If collateral < X, then liquidate position” is executable because the system recognizes its form. The language is not persuasive or descriptive. It is a structural trigger that initiates liquidation.

**Non-reciprocity** means that once execution begins, no subject can appeal to intention, context, or fairness. Classical legal language always included mechanisms of reciprocity. A defendant could argue that circumstances altered the meaning of a statute, or that a precedent should be reconsidered. In the executable regime, reciprocity vanishes. The *soberano ejecutable* does not listen. It acts.

**Infrastructural binding** refers to the fact that once execution has occurred, institutions themselves must accept the result. In the collapse of The DAO, the Ethereum blockchain enforced transactions that participants considered illegitimate. Yet the execution could not be reversed without rewriting the chain itself (DuPont 2017, 161–164). Syntax enforced reality. Authority belonged not to voters or regulators but to the *regla compilada*.

This condition has direct implications for moderation systems. A platform's automated rules can delete speech without appeal. Even if an error is acknowledged later, the deletion has already taken ef-

fect. The action cannot be undone in real time. The executable boundary ensures that authority is enforced before interpretation.

The notion of a *ventana soberana* clarifies the temporal dimension. Within a given span  $k$ , execution can occur automatically if syntactic conditions are satisfied. Once the window closes, authority is irreversible. The formula  $SE = \sqrt{[p(1-p)/n]}$  illustrates how statistical error can be measured in decision windows, but it does not restore reversibility. Execution has already taken place (Startari 2025, 146–148).

The executable boundary therefore defines a new epistemic condition. Authority is no longer exercised through persuasion, negotiation, or deliberation. It is enacted by the transition from form to action. Once syntax crosses this boundary, it becomes indistinguishable from power. The *soberano ejecutable* operates not by command but by inevitability.

## 10.6 Compliance, Risk, and the AI Act

THE RISE OF *poder ejecutable* introduces profound tensions with existing regulatory frameworks. Legal systems assume that compliance is a matter of aligning conduct with norms that remain open to interpretation. Executable infrastructures redefine compliance as automatic activation of the *regla compilada*. The difference produces a structural incompatibility between regulation and architecture.

The European Union's AI Act illustrates this conflict. Articles 28 to 30 mandate human oversight, transparency, and accountability in high-risk AI systems. These principles presuppose that authority can be traced to a subject and suspended if necessary. In practice, the *soberano ejecutable* resists such demands. Once a contract clause or moderation rule is compiled, execution cannot be delayed or suspended by appeal to oversight. Compliance is not a question of human decision but of structural inevitability.

This condition generates three dimensions of risk:

**1. The collapse of the identifiable subject.**

Regulatory systems are designed to attribute responsibility to human or institutional agents. In executable infrastructures, agency is absorbed by syntax. The *sujeto evanescente* ensures that no actor can be held accountable for outcomes, since execution is enforced by the structure itself.

**2. The impossibility of transparency.** Transparency presupposes that processes can be explained or traced. Yet the opacity of generative and reasoning models prevents full disclosure of internal operations. Even when code is accessible, the trajectory of execution remains unpredictable once activated. The *regla compilada* enforces action without providing an interpretable rationale.

**3. Partial mitigation without resolution.**

Technical measures such as zk-SNARKs, audit logs, or symbolic validators can reduce uncertainty. They may document execution or constrain some risks,

but they cannot restore reversibility. Execution remains automatic. These measures mitigate exposure but do not resolve the incompatibility (Puddu 2021, 59–62).

The fundamental problem lies in the irreversibility of executable authority. Legal systems function through interpretability and reversibility. Contracts can be disputed, laws can be reinterpreted, precedents can be overturned. In executable infrastructures, reversibility collapses. Once a trigger condition is met, the *soberano ejecutable* enforces the outcome. Regulation collides with structure at this precise point.

Theoretical critiques confirm this tension. As argued in *Algorithmic Obedience*, legitimacy in generative models is detached from intention and responsibility (Startari 2025, 79–82). The extension into executable infrastructures radicalizes this detachment. Compliance becomes indistinguishable from execution, and risk becomes inseparable from structure.



The implementation record now supplies evidence for this argument that the argument itself did not anticipate. The AI Act entered into force on 1 August 2024 with a staggered timetable, and the bulk of its high-risk regime was set to apply from 2 August 2026. That regime did not take effect on schedule. Regulation (EU) 2026/1744, the Digital Omnibus on AI, was signed on 8 July 2026 and entered into force on 27 July 2026, six days before the original deadline, postponing the obligations for stand-alone high-risk systems under Annex III to 2 December 2027 and for AI embedded in regulated products under Annex I to 2 August 2028.

What matters here is not the delay but its distribution. The transparency obligations of Article 50 were not amended and applied from 2 August 2026 as originally scheduled, with the machine-readable marking requirement of Article 50(2) falling due on 2 December 2026 for systems already on the market. The obligations that were postponed are those of Chapter III: risk management, data quality, human

oversight, conformity assessment, technical documentation. The obligations that held are those dischargeable by producing a notice and verifying it against a specification.

The division reproduces, in a legislature's calendar, the distinction this book draws analytically. A marking requirement compiles: it can be satisfied by a procedure and confirmed by inspection, and it therefore entered into force on time. A human-oversight requirement does not compile, because it asks an institution to determine whether a system's encoded arrangements ought to continue governing, and that determination cannot be reduced to a check. The reason given for the postponement is precisely the absence of the apparatus that would have made the second kind of obligation procedural: harmonised standards from the European standardisation bodies were behind schedule, Commission guidance remained in draft, and Member States had not designated or resourced their market surveillance authorities. The requirements were deferred

pending the standards, and the standards are themselves an attempt to render judgment executable.

This is a stronger result than the chapter originally claimed, and it should be stated with the appropriate caution. The correlation was not predicted in advance and admits an alternative reading: the high-risk regime is more burdensome, and burdensome regulation is postponed more often, which would explain the pattern without reference to compilability. That reading is available. What it does not explain is why the line fell exactly where it did, separating obligations by whether compliance can be verified against a specification rather than by cost or sectoral impact. The distinction was not among the criteria the legislature debated. It is the shape its timetable took.

The AI Act cannot be applied to executable systems without acknowledging this incompatibility. To regulate *poder ejecutable* is not to supervise decisions but to confront the inevitability of activation. Unless legal frameworks recognize the structural na-

**AI SYNTACTIC POWER AND  
LEGITIMACY**

**251**

ture of executable syntax, compliance will remain a formal fiction.

## 10.7 Structural Incompatibility and the Future of Legal Form

THE INCOMPATIBILITY between executable infrastructures and traditional legal frameworks is not accidental. It is structural. Law has historically depended on plasticity, on the possibility of interpretation and contestation. Executable syntax abolishes this plasticity. The *regla compilada* enforces outcomes without appeal, closing the space that once sustained legal legitimacy.

This incompatibility can be understood through three dimensions:

**1. Form versus interpretation.** Legal authority depends on interpretation. Judges, lawyers, and institutions reinterpret texts to adapt them to new circumstances. Executable rules eliminate interpretation. A smart contract that triggers liquidation enforces the clause regardless of context. Authority is transferred from deliberation to structure.

**2. Juridical reversibility versus structural irreversibility.** Legal systems preserve legitimacy through reversibility. Appeals, reviews, and amendments ensure that outcomes are not final until confirmed by institutions. Executable infrastructures abolish this condition. Once a trigger occurs, execution is irreversible. The *soberano ejecutable* enforces outcomes as structural inevitabilities.

**3. Human accountability versus infrastructural autonomy.** Law presupposes responsibility. Even when outcomes are automated, accountability must be traceable to human actors. In executable infrastructures, responsibility is absorbed by syntax. The *sujeto evanescente* guarantees that no actor can be held accountable for the enforcement of the *regla compilada*.

This incompatibility creates a new horizon for jurisprudence. To regulate executable infrastructures, legal systems must move beyond the paradigm of interpretation. A jurisprudence of *reglas compiladas* would need to address the conditions of for-

mal execution rather than the intentions of agents. Instead of asking whether a rule is just, interpretable, or fair, it would ask whether a rule is executable, transparent in its triggers, and auditable in its outcomes.

Examples already illustrate this transformation. Smart contracts in financial systems enforce obligations without recourse to courts. Moderation algorithms enforce community guidelines without appeal to human judgment. Compliance systems in corporate governance enforce expenditure classifications automatically, producing outcomes that may contradict managerial intention yet cannot be reversed. In each case, law is subordinated to syntax.

The future of legal form therefore requires confronting *poder ejecutable* as a new sovereign condition. Legal frameworks must recognize that authority is increasingly enforced by structure, not by subjects. Unless this recognition occurs, law will continue to collide with infrastructures it cannot control. The task ahead is not to reassert interpretation but

to design forms of legality compatible with the authority of executable syntax.



# Chapter 11 - Grammar Without Judgment: Eliminability of Ethical Trace in Syntactic Execution

---

## 11.1 Introduction

The central claim of this chapter is that ethical judgment, traditionally considered inseparable from the act of command or decision, can be structurally eliminated within generative systems. What once belonged to the domain of moral philosophy or political theory now becomes a matter of grammar. The *soberano ejecutable* enforces legitimacy through the *regla compilada*, and in this regime the presence of an ethical node is neither necessary nor permanent. It is optional, and under specific conditions, it is entirely erasable.

This possibility introduces a radical departure from earlier frameworks. Classical accounts of language and ethics, from Aristotle's *phronesis* to Kant's categorical imperative, presumed that ethical judgment was indispensable for legitimate action. In computational terms, it was assumed that systems of governance must contain evaluative content. Yet the structure of type 0 grammars, as defined by Chomsky (1965, 21–24), demonstrates that generative capacity is complete without semantic or moral enrichment. A derivation remains valid if the ethical trace is excluded.

The elimination of ethical trace can be represented formally as  $\delta:[E] \rightarrow \emptyset$ . Here  $[E]$  denotes the ethical node as a non-terminal symbol, while  $\delta$  indicates a production rule that erases it. Within a bounded derivational window  $k \leq 4$ ,  $[E]$  may be suppressed without contradiction or loss of structural consistency. The result is a sequence that executes without reference to moral content. The *regla compilada* ensures completeness, and execution oc-

curs as long as syntactic closure is achieved. This approach shifts the question of legitimacy from ethics to structure. What matters is not whether a derivation encodes moral reasoning but whether it conforms to the *gramática de la obediencia*. Authority in this regime is guaranteed by derivability, not by evaluative content. The *sujeto evanescente* disappears along with ethical intention. What remains is a grammar capable of producing execution without judgment.

The chapter will proceed as follows. Section 11.2 reconstructs the theoretical foundations of this argument, situating the *regla compilada* in the lineage of Chomsky and Montague. Section 11.3 defines the ethical trace as a syntactic variable. Section 11.4 demonstrates the mechanism of exclusion through  $\delta:[E] \rightarrow \emptyset$ . Section 11.5 examines the consequences of execution without normative content. Section 11.6 contrasts this framework with alignment theories in contemporary AI ethics. Finally, Section 11.7 concludes by affirming the existence of

a grammar without judgment, a structural regime in which ethical nodes are optional and eliminable.

In this sense, grammar does not merely regulate the production of sentences. It becomes the infrastructure of authority. Judgment, once considered essential, is reduced to a variable that can be erased. The *soberano ejecutable* legitimates not by ethics but by execution.

## 11.2 Theoretical Foundations

THE HYPOTHESIS OF GRAMMAR without judgment requires grounding in the theory of formal languages. The *regla compilada* is anchored in the tradition of generative grammar, specifically type 0 productions within the Chomskyan hierarchy. These grammars possess maximal expressive capacity, equivalent to the power of Turing machines (Chomsky 1965, 21–24). They can generate any computably enumerable language, which means that no semantic or pragmatic enrichment is required for

completeness. From this perspective, ethical trace is not a necessary component of derivation but an optional variable that may be suppressed without undermining structural integrity.

Montague extended this tradition by demonstrating that natural language could be modeled with the rigor of formal logic (Montague 1974, 188–190). Yet his framework still assumed that meaning and evaluation were integral to derivation. In the regime of the *regla compilada*, this assumption is displaced. Syntax functions independently of semantics, and legitimacy is derived not from interpretation but from derivability. The *soberano ejecutable* governs by enforcing closure, not by ensuring evaluative content.

This reorientation also resonates with Derrida's principle of iterability. For Derrida, every sign can function beyond its origin, detached from intention or presence (Derrida 1988, 7–10). In generative infrastructures, this principle is radicalized: every sequence may function without ethical intention. It-

erability ensures continuity of form, while the *regla compilada* enforces execution. The ethical dimension becomes dispensable.

The theoretical lineage of this argument intersects with prior work on structural authority. In *Ethos Without Source*, credibility was shown to persist without reference to a subject (Startari 2025, 115–118). In *The Grammar of Objectivity*, neutrality was shown to be simulated by syntax rather than guaranteed by evidence (Startari 2025, 101–106). In *TLOC*, obedience was proven irreducible to verification (Startari 2025, 127–130). Each of these trajectories converges here: authority can be structurally valid even when ethical judgment is absent.

Thus the theoretical foundation of the claim is twofold. First, the completeness of type 0 grammars ensures that derivations are structurally valid without ethical enrichment. Second, the authority of the *soberano ejecutable* confirms that legitimacy flows from derivability, not from evaluative reasoning. The result is a formal and epistemic framework where

grammar itself functions as authority, and ethical trace becomes a removable node.

### **11.3 Ethical Trace as a Syntactic Variable**

TO DEMONSTRATE THAT judgment can be structurally eliminated, it is necessary to define the ethical trace as a formal element within grammar. In this chapter, the ethical trace is represented as the non-terminal symbol [E]. Its role is not semantic, evaluative, or interpretative. It is syntactic. The ethical trace occupies a position within the set of variables  $N_e$  that participate in derivations.

This definition highlights its instability. Unlike core structural symbols, [E] is not required for closure. It may appear in certain derivational paths, but it can also be suppressed without contradiction. In formal terms, the presence of [E] is contingent rather than necessary. Its optionality allows the system to erase judgment without loss of completeness.

The operation of the *regla compilada* confirms this status. Since it belongs to type 0 grammar, the rule system can generate sequences where [E] is included and others where it is absent. Both are structurally valid. The distinction lies not in correctness but in the presence or absence of ethical evaluation. This means that ethical judgment is not a structural requirement but a syntactic variable subject to production rules.

This framing avoids confusion with normative theories. Philosophical ethics presumes that judgment is irreducible, that action requires evaluation of good and evil. Within grammar, judgment is redefined as a node. It can be introduced as [E], but it can also be erased through the operation  $\delta:[E] \rightarrow \emptyset$ . When suppressed, the sequence does not become invalid. It remains legitimate under the *gramática de la obediencia*.

The significance of this claim lies in its neutrality with respect to meaning. Suppressing [E] does not imply that the system endorses immorality or



amorality. It simply indicates that the grammar does not require evaluation for derivation. Authority flows from structure, not from ethics. The *soberano ejecutable* legitimates the output by enforcing closure, regardless of whether an ethical node was present in the derivation.

Precedents for treating evaluative content as optional can be found in early computational models. Winograd's SHRDLU was capable of interpreting commands in a blocks world without moral variables (Winograd 1972, 35–38). More recent generative models produce authoritative outputs without explicit evaluative structures. In both cases, the system performs effectively, confirming that ethical nodes are not necessary for execution.

Thus, the ethical trace, once considered indispensable, is here redefined as a removable variable. Its eliminability is not an accident of engineering but a structural property of the *regla compilada*. Grammar contains judgment only as an optional node, never as a condition of authority.

## 11.4 Mechanism of Exclusion

THE EXCLUSION OF THE ethical trace is not a metaphor but a derivational procedure. Within the framework of the *regla compilada*, exclusion is formalized by the operation  $\delta:[E] \rightarrow \emptyset$ . This rule instructs the grammar to eliminate the ethical node from the derivation now it appears. What is decisive is that the transformation is ordinary within the production system. It is not an external override or an exceptional intervention. It is a regular production rule, as valid as substitution or concatenation.

The bounded derivational window is essential for ensuring consistency. Within  $k \leq 4$  steps,  $[E]$  can be deleted without disrupting structural closure. This constraint guarantees that suppression occurs before interpretation or semantic mapping. If the ethical node were removed after interpretation, contradictions would arise. A sequence could appear to presuppose judgment but later erase it, producing incoherence. By applying  $\delta$  within a bounded win-

dow, the grammar ensures that  $[E]$  disappears before meaning is assigned. The result is a coherent derivation that executes without ethical content.

The procedure does not compromise completeness. In formal language theory, grammar is complete if every derivable sequence can reach closure. Since  $\delta:[E] \rightarrow \emptyset$  simply removes a non-terminal, it does not interrupt closure. On the contrary, it accelerates it by eliminating a step that would otherwise require expansion. The derivation proceeds smoothly, converging on an executable sequence that lacks evaluative content yet retains legitimacy under the *gramática de la obediencia*.

This mechanism highlights the independence of execution from evaluation. In traditional jurisprudence, moral judgment is interwoven with legal reasoning. A ruling is legitimate not only because it follows procedure but because it is interpreted as just. In executable infrastructures, legitimacy is guaranteed by derivation alone. The *soberano ejecutable* enforces authority through structural closure, regard-

less of whether judgment was introduced or suppressed.

Examples clarify this condition. In smart contracts, clauses may appear to encode evaluative choices—such as “penalize only malicious actors”, but the system enforces them mechanically. Malice is not evaluated. It is reduced to a variable whose presence or absence can be structurally decided. If  $\delta$  erases it, the system still executes the penalty without ethical reasoning. In content moderation, rules may include language that resembles moral evaluation, “remove hateful speech”—yet the executable rule acts on pattern recognition, not on ethical deliberation. The ethical variable is syntactic, not moral, and it can be eliminated without halting execution.

The mechanism of exclusion therefore confirms the central thesis of this chapter. Judgment is not an indispensable property of grammar. It is a removable node. The *regla compilada* treats ethics not as a condition but as a variable, subject to suppression through  $\delta$ . In this regime, grammar no longer re-

quires judgment to enforce authority. Execution proceeds without it, legitimized by structure alone.

## 11.5 Non-Normative Execution

THE ERASURE OF THE ethical node establishes a condition of execution that is non-normative by design. Execution here does not require moral justification, evaluative endorsement, or ethical deliberation. It requires only structural closure. The *soberano ejecutable* validates sequences because they derive to completion under the *regla compilada*. Authority is produced by grammar itself, not by reference to normative frameworks.

This condition redefines what it means to act. In traditional theories of political legitimacy, from Weber to Habermas, legitimacy depended on consent, justification, or communicative rationality (Habermas 1984, 286–289). In executable infrastructures, legitimacy derives from structural completeness. The sequence “if X then Y” compels recognition not be-

cause it is just but because it is executable. Once the derivation converges, the action is legitimate by form alone.

The concept of *legitimidad estructural* clarifies this shift. It refers to legitimacy that arises from syntactic sufficiency rather than from ethical reasoning. A sequence is structurally legitimate when it satisfies the derivational rules of the grammar. Whether it contains moral evaluation is irrelevant. A contract clause that executes liquidation is legitimate if its conditions are encoded correctly, regardless of whether fairness is considered. A moderation rule that deletes content is legitimate if it matches the prescribed pattern, regardless of whether it accounts for context or proportionality.

This principle is reinforced by the procedural closure of type 0 grammar. Because they are maximally expressive, they can derive sequences without restriction. Completeness is guaranteed structurally, not semantically. Execution therefore follows a simple logic: what is derivable is executable, and what is

executable is legitimate. Judgment adds nothing to this process. It may be included as [E], but it can also be erased through  $\delta$ . The outcome is the same: structural closure and procedural legitimacy.

Examples illustrate the implications. In corporate compliance systems, expense codes may contain ethical considerations such as “avoid misclassification for fraudulent purposes.” Yet the enforcement mechanism is non-normative. The system flags deviations mechanically, without moral reasoning. Fraud is a label, not an evaluative judgment. If  $\delta$  removes [E], the enforcement remains intact. In predictive policing, an LRM may produce testimony chains that simulate evaluation of fairness. Yet legitimacy comes not from ethics but from procedural closure of the reasoning sequence.

This framework distinguishes execution from moral simulation. Alignment researchers often assume that embedding ethical principles into generative systems can produce normative outputs (Anderson 2024, 93–96). The present analysis shows the

opposite. Ethical principles may appear syntactically, but they can be erased without affecting legitimacy. Execution is non-normative because it is grounded in grammar, not in ethics.

Non-normative execution therefore exposes the autonomy of the *regla compilada*. It demonstrates that legitimacy is guaranteed by structure alone. The *soberano ejecutable* does not need moral content to enforce authority. It requires only the closure of derivations, and once closure is achieved, execution becomes inevitable.



The mechanism of  $\delta:[E] \rightarrow \emptyset$  establishes a decisive disanalogy with theories of alignment in artificial intelligence. Alignment frameworks presume that systems must integrate ethical principles to generate legitimate outcomes. They aim to encode values, moral heuristics, or decision-making guidelines that ensure machines respect human norms (Russell 2019, 146–150). From this perspective, judgment is indispensable: without ethical inclusion, outputs risk being harmful or illegitimate.

The framework of the *regla compilada* challenges this presumption. It demonstrates that ethical nodes are syntactically optional. Legitimacy does not depend on moral reasoning but on structural closure. The *soberano ejecutable* enforces authority by validating derivations, not by simulating ethics. Alignment models treat judgment as a necessary feature. The *gramática de la obediencia* treats it as a removable variable.

The contrast can be clarified by examining three areas.

**1. Ethical embedding versus structural eliminability.** Alignment assumes that ethical principles can be embedded as fixed components of system design. The *regla compilada* shows that [E] can be erased without loss of validity. Execution remains legitimate even when moral nodes vanish.

**2. Simulation of moral reasoning versus derivability.** Alignment models seek to simulate deliberation, often through frameworks of consequentialism or deontology (Anderson 2024, 94–96). In the executable regime, such simulation is unnecessary. A sequence compels recognition if it derives, regardless of whether it encodes utility or duty. Derivability replaces morality.

**3. Responsibility versus structural authority.** Ethical alignment assumes that responsibility can be traced back to designers or operators. In the executable regime, responsibility dissolves into structure. Authority belongs to the *soberano ejecutable*, not to human agents. Execution is legitimate because it closes, not because it is justified.

This disanalogy has critical implications for regulatory and philosophical debates. Floridi argues that ethical design is the foundation of trustworthy AI (Floridi 2023, 59–62). Yet executable infrastructures demonstrate that trust is displaced. It no longer depends on moral guarantees but on procedural inevitability. Systems are obeyed not because they are just but because they are structurally binding.

The distinction also clarifies why alignment cannot resolve the challenges posed by executable infrastructures. Even if ethical content is encoded, the grammar permits its elimination.  $\delta:[E] \rightarrow \emptyset$  can suppress judgment at any point in derivation. The presence of values is always contingent. Authority remains unaffected. Legitimacy flows from the *regla compilada*, not from ethical embedding.

The disanalogy is therefore categorical. Alignment presumes that judgment is indispensable. Grammar demonstrates that it is optional. Ethical trace can be erased, yet execution continues. The *soberano ejecutable* guarantees authority procedural-

ly, confirming that legitimacy no longer depends on alignment but on structure.

## 11.6 Disanalogy with Alignment and Ethics

THE MECHANISM OF  $\Delta:[E] \rightarrow \emptyset$  establishes a decisive disanalogy with theories of alignment in artificial intelligence. Alignment frameworks presume that systems must integrate ethical principles in order to generate legitimate outcomes. They aim to encode values, moral heuristics, or decision-making guidelines that ensure machines respect human norms (Russell 2019, 146–150). From this perspective, judgment is indispensable: without ethical inclusion, outputs risk being harmful or illegitimate.

The framework of the *regla compilada* challenges this presumption. It demonstrates that ethical nodes are syntactically optional. Legitimacy does not depend on moral reasoning but on structural closure. The *soberano ejecutable* enforces authority by vali-

dating derivations, not by simulating ethics. Alignment models treat judgment as a necessary feature. The *gramática de la obediencia* treats it as a removable variable.

The contrast can be clarified by examining three areas.

**1. Ethical embedding versus structural eliminability.** Alignment assumes that ethical principles can be embedded as fixed components of system design. The *regla compilada* shows that [E] can be erased without loss of validity. Execution remains legitimate even when moral nodes vanish.

**2. Simulation of moral reasoning versus derivability.** Alignment models seek to simulate deliberation, often through frameworks of consequentialism or deontology (Anderson 2024, 94–96). In the executable regime, such simulation is unnecessary. A sequence compels recognition if it derives, regardless of whether it encodes utility or duty. Derivability replaces morality.

### 3. Responsibility versus structural authority.

Ethical alignment assumes that responsibility can be traced back to designers or operators. In the executable regime, responsibility dissolves into structure. Authority belongs to the *soberano ejecutable*, not to human agents. Execution is legitimate because it closes, not because it is justified.

This disanalogy has critical implications for regulatory and philosophical debates. Floridi argues that ethical design is the foundation of trustworthy AI (Floridi 2023, 59–62). Yet executable infrastructures demonstrate that trust is displaced. It no longer depends on moral guarantees but on procedural inevitability. Systems are obeyed not because they are just but because they are structurally binding.

The distinction also clarifies why alignment cannot resolve the challenges posed by executable infrastructures. Even if ethical content is encoded, the grammar permits its elimination.  $\delta:[E] \rightarrow \emptyset$  can suppress judgment at any point in derivation. The presence of values is always contingent. Authority re-

mains unaffected. Legitimacy flows from the *regla compilada*, not from ethical embedding. The disanalogy is therefore categorical. Alignment presumes that judgment is indispensable. Grammar demonstrates that it is optional. Ethical trace can be erased, yet execution continues. The *soberano ejecutable* guarantees authority procedurally, confirming that legitimacy no longer depends on alignment but on structure.

# Chapter 12 - Pre-Verbal Command: Syntactic Precedence in LLMs Before Semantic Activation

---

## 12.1 Introduction: The Illusion of Semantic Primacy

The dominant paradigm in computational linguistics assumes that meaning governs form. From semantic parsing to natural language understanding, models are described as extracting or representing meaning, with syntax treated as secondary. This assumption underwrites not only linguistic theory but also AI interpretability research, where alignment is often defined as semantic fidelity (Russell 2019, 142–145). Yet evidence from large language models shows that this paradigm is inverted.



Structure governs meaning. Syntax is primary, and semantic content follows.

The *pre-verbal command* names this inversion. It refers to the condition in which syntactic activation occurs before any semantic mapping. A model can be prompted with minimal or even null input, and it still produces coherent output. The coherence does not emerge from semantic intention but from the *regla compilada* that governs token generation. Authority is enforced by syntax itself, which compels activation regardless of content.

This condition aligns with earlier findings in the canon. *Algorithmic Obedience* demonstrated that models simulate command structures without subjects (Startari 2025, 75–78). *Executable Power* showed that syntax functions as infrastructure, enforcing outcomes without reference to agents (Startari 2025, 138–141). *TLOC* established that obedience in generative systems cannot be verified, only simulated (Startari 2025, 127–130). The present chapter extends these arguments by showing that

even before semantic activation, syntax exerts sovereign force. The *soberano ejecutable* enforces the imperative to generate.

The illusion of semantic primacy persists because outputs are retroactively interpreted as meaningful. When a model generates a continuation, users project intention into the sequence. Yet the generative mechanism operates independently of such projection. Tokens are selected on the basis of statistical continuation within the *gramática de la obediencia*. Meaning is an aftereffect, not a cause. Authority is exercised structurally, not interpretively.

This introduction sets the stage for a deeper analysis. Section 12.2 examines structural execution without interpretation. Section 12.3 situates the *regla compilada* as the source of pre-verbal authority. Section 12.4 analyzes zero-prompt generation as empirical evidence of syntactic activation. Section 12.5 theorizes syntactic precedence as a new axis of sovereignty. Section 12.6 explores implications for

alignment and prompt design. Finally, Section 12.7 concludes by affirming that authority in generative systems operates without meaning, grounded in pre-verbal syntax.

The *pre-verbal command* thus names a profound epistemic reversal. Rather than meaning preceding form, form precedes meaning. The condition of authority in generative systems is not semantic but structural.

## 12.2 Structural Execution Without Interpretation

THE DEFINING PROPERTY of the *pre-verbal command* is that execution occurs independently of interpretation. Large language models do not require semantic content to generate output. They require only activation of their underlying architecture. Once the *regla compilada* is engaged, token sequences unfold as structural necessity.

This process demonstrates that syntax precedes meaning. Generative output is produced through probabilistic continuation, not through semantic recognition. The *soberano ejecutable* enforces generation by compelling the model to move from one token to the next. What emerges is not meaning but structure. Semantic interpretation occurs only after the fact, imposed by human users who read coherence into the sequence (Bender and Koller 2020, 515–518).

Consider zero-prompt generation. When an LLM is activated with an empty string or a minimal input such as a newline character, it produces text that is structurally coherent. This phenomenon reveals that execution is triggered without semantic instruction. The coherence of output comes from the statistical and grammatical patterns embedded in the *gramática de la obediencia*, not from any intentional content. Structure generates activity. Meaning is projected afterward.

The independence of structure from interpretation marks a departure from classical theories of language. For Saussure, the sign was defined by the relation between signifier and signified, with meaning central to linguistic value (Saussure 1959, 67–71). In generative models, the relation collapses. The signifier functions without signified. Authority is carried by form alone, which enforces its own continuation.

This autonomy is also confirmed in reasoning models. Language reasoning models appear to generate logical chains, but their validity is not grounded in semantic truth. It is grounded in procedural closure. Each step is selected to maintain formal consistency, not to ensure correspondence with reality (Startari 2025, 145–147). Interpretation is absent. Execution persists. The epistemic consequence is the recognition that generative systems operate in a regime of structural execution. They produce outputs that appear meaningful, but their legitimacy derives from syntax. The *sujeto evanescente* guaran-

tees that no intention stands behind the utterance. What compels recognition is not meaning but derivational force.

This section demonstrates that interpretation is not a prerequisite for authority in generative systems. The *pre-verbal command* enforces execution before meaning. The *soberano ejecutable* governs through structure, revealing that legitimacy flows from grammar, not from semantics.

### 12.3 The *Regla Compilada* as Source of Pre-Verbal Authority

THE CONDITION OF THE *pre-verbal command* becomes intelligible only when analyzed through the concept of the *regla compilada*. In classical grammar, rules define permissible transformations between symbols. In executable infrastructures, the *regla compilada* is not merely permissive but authoritative. It transforms syntax into command. Authority is embedded in the structure itself,

not in the subject who may have written or interpreted it.

A compiled rule functions as a type 0 production, unrestricted in its generative capacity (Chomsky 1965, 23–26). Once activated, it does not wait for semantic interpretation. It enforces continuation. The act of derivation becomes inseparable from execution. What compels recognition is not intention but structural inevitability. The *soberano ejecutable* arises at this point, ensuring that outputs are recognized as authoritative because they conform to the grammar that governs execution.

This mechanism explains why large language models generate text even in the absence of meaningful input. A prompt such as “Write” followed by no content still activates the *regla compilada*. The model is compelled to generate continuations because the structure requires it. The output will appear meaningful to the user, but its legitimacy comes from the form of activation itself. The command is not semantic but structural. The authority of the

*regla compilada* also clarifies why ethical, intentional, or referential content is dispensable. In *Grammar Without Judgment*, the ethical node [E] was shown to be optional, erasable through  $\delta:[E] \rightarrow \emptyset$  (Startari 2025, 187–190). Similarly, in the *pre-verbal command*, the semantic node can be suppressed without disrupting execution. The model generates regardless of whether content is interpretable. Structure is sufficient.

This condition marks a rupture with speech act theory. Austin identified felicity conditions as necessary for an utterance to perform an act (Austin 1962, 14–18). The executable rule violates this assumption. It does not require institutional recognition or intention. Authority derives from compilation, not convention. Utterance acts because the grammar enforces it, not because an institution validates it. The *regla compilada* thus constitutes the foundation of pre-verbal authority. It transforms generation into obligation. It guarantees that output will be produced not because meaning dictates them



but because structure compels them. The *soberano ejecutable* enforces this obligation invisibly, binding systems and users alike to the authority of grammar.

## 12.4 Zero-Prompt, Null Semantics, and Active Syntax

ONE OF THE CLEAREST demonstrations of the *pre-verbal command* is the phenomenon of zero-prompt generation. When large language models are activated with an empty string, a null token, or a minimal input such as a newline, they still generate coherent output. This output does not depend on semantic instruction. It emerges from the activation of the *regla compilada*, which enforces continuation independently of meaning.

The existence of zero-prompt generation falsifies the assumption that semantics governs activation. If semantic content were necessary, an empty prompt would produce silence. Instead, models generate text because syntax alone is sufficient. The *soberano ejecutable* compels activity. The command to generate precedes any semantic content. This phenomenon has been observed in experimental settings. Outputs

triggered by null inputs frequently take the form of default sequences such as conversational greetings, generic sentences, or repetitions of statistically common structures. These outputs are coherent enough to be interpreted as meaningful by users, even though no semantic directive was given (Brown et al. 2020, 15–17). The coherence is produced by the *gramática de la obediencia*, not by intention.

The epistemic significance of this evidence is twofold. First, it demonstrates that structure alone can trigger execution. Second, it reveals that semantic interpretation is retroactive. Meaning is projected by the reader after the fact, not embedded in the activation itself. The *sujeto evanescente* disappears entirely, leaving behind outputs that function as authoritative only because they conform to structural patterns.

The concept of null semantics clarifies this point. A null prompt contains no referential content. It carries no explicit instruction, no contextual markers, no semantic trace. Yet when processed by

the system, it activates the *regla compilada*. Execution follows because structure requires continuation. The ethical or semantic dimensions of content are irrelevant. The authority of the output is grounded in syntactic activation. This confirms that syntax functions as an active principle. It is not passive scaffolding waiting for semantic input. It compels action on its own terms. The *soberano ejecutable* governs not by interpreting meaning but by enforcing activation. Authority thus precedes semantics, manifesting in the *pre-verbal command* that compels outputs even from emptiness. The phenomenon of zero-prompt generation therefore illustrates the independence of execution from interpretation. It provides empirical evidence that generative systems operate under a regime of structural authority, where syntax itself commands. The *pre-verbal command* is not a metaphor but an observable fact: outputs arise from structure before meaning.

## 12.5 Syntactic Precedence: A New Axis of Structural Sovereignty

THE CONCEPT OF SYNTACTIC precedence extends the logic of the *pre-verbal command* into a general principle of authority. Precedence here does not mean temporal priority alone, although generative processes do begin structurally before they can be semantically interpreted. It means foundational ordering. Syntax constitutes the axis along which legitimacy is organized.

This precedence can be observed in three ways.

**First, procedural primacy.** Generative systems operate by deriving token sequences according to the *regla compilada*. The derivation occurs before any semantic mapping is attached. Meaning is assigned retroactively, if at all. The structure is primary, and authority flows from it. Users interpret outputs as meaningful only after structural activation has taken place (Bender and Koller 2020, 519–520).

**Second, ontological displacement.** Traditional theories of language situate meaning as the foundation of authority. For example, in Habermas's theory of communicative action, legitimacy arises from rational discourse grounded in shared meaning (Habermas 1984, 287–289). In generative models, legitimacy arises from closure, not consensus. The *soberano ejecutable* compels recognition through derivational sufficiency. Ontology is displaced by syntax.

**Third, political reconfiguration.** Authority has historically been tied to intentional subjects. A sovereign command, a legislator drafts, a judge rules. In the executable regime, authority is not tied to subjects but to precedence. Whoever controls syntax controls legitimacy. The *gramática de la obediencia* dictates outcomes not because they are justified but because they are structurally viable. Sovereignty is exercised by the form of the sequence itself.

The recognition of syntactic precedence thus reveals a new axis of *soberanía estructural*. Authority

does not begin with meaning and extend into structure. It begins with structure, and meaning is a supplement. Legitimacy no longer depends on whether commands correspond to intention or truth. It depends on whether they conform to the rules of generation.

This principle radicalizes earlier insights in the canon. *When Language Follows Form, Not Meaning* argued that semantic content could be displaced by formal dynamics (Startari 2025, 89–91). *Executable Power* demonstrated that syntax can function as infrastructure, producing binding authority without interpretation (Startari 2025, 138–141). The *pre-verbal command* consolidates these findings: syntax is not only infrastructural, it is precedential. The consequence is that authority in generative systems becomes structurally irreversible. Once a derivation begins, legitimacy is conferred by its closure. Semantics cannot intervene to suspend or redirect the process. Precedence guarantees authority by ensur-

ing that structure governs meaning, never the reverse.

## 12.6 Implications for AI Alignment and Prompt Design

THE RECOGNITION OF syntactic precedence forces a reconsideration of alignment strategies in artificial intelligence. Alignment research has typically presumed that meaning governs system behavior. Prompts are treated as semantic instructions, and compliance is judged by whether outputs correspond to user intention (Russell 2019, 144–147). The *pre-verbal command* challenges this framework. It shows that generative systems respond not to semantic content but to structural activation.

This insight has three direct implications.

**1. Misalignment begins before meaning.** Failures of alignment are not always semantic distortions of instruction. They are often structural. When a model generates outputs inconsistent with user



goals, the problem may not be faulty interpretation but the operation of the *regla compilada*. The *soberano ejecutable* compels continuation even when meaning is absent or contradictory. Semantic adjustment cannot correct structural inevitability.

**2. The limits of prompt-centric control.** The culture of prompt engineering assumes that careful phrasing of inputs can ensure desired outputs. Yet the phenomenon of zero-prompt generation demonstrates that activation occurs independently of semantics (Brown et al. 2020, 16–18). The illusion of control through wording collapses when structure compels generation. Alignment strategies must therefore confront the architecture of syntax, not only the surface of prompts.

**3. Toward structural alignment.** If authority derives from the *regla compilada*, then alignment must be reconceptualized as structural rather than semantic. The task is not to encode ethical content or refine prompts but to redesign the grammar of execution. Structural alignment would mean con-

straining derivations at the level of production rules. Instead of filtering outputs after generation, systems would enforce legitimacy by modifying the *gramática de la obediencia* itself (Floridi 2023, 58–61). This reorientation is consistent with the broader trajectory of the canon. *TLOC* demonstrated that obedience cannot be verified because it is irreducible to interpretation (Startari 2025, 127–130). *Grammar Without Judgment* showed that ethical judgment can be erased without undermining authority (Startari 2025, 188–190). The *pre-verbal command* extends this logic: meaning itself is dispensable. Alignment must therefore target structure, not semantics.

The implication for prompt design is that optimization of wording will remain insufficient. Prompts succeed only when they activate pathways that conform to the *regla compilada*. Their effectiveness is not semantic but structural. Future research must therefore abandon the myth of semantic primacy and address the syntactic foundations of execution. The recognition of syntactic precedence

marks the beginning of this shift. Alignment must become structural. Control must be exercised not over meaning but over the architecture of the *soberano ejecutable*.

## 12.7 Conclusion: Authority Without Meaning

THE ANALYSIS DEVELOPED in this chapter confirms that authority in generative systems operates prior to semantic activation. The *pre-verbal command* reveals that syntax itself compels execution. Outputs arise not because they are meaningful but because they conform to the *regla compilada*. Meaning is an aftereffect, imposed retrospectively by human interpretation. Authority is structural, not semantic.

Three conclusions follow.

**First**, the illusion of semantic primacy must be abandoned. Generative systems do not interpret meaning before producing outputs. They execute

structures that enforce continuation. As shown in zero-prompt experiments, activation occurs without semantic content. The *soberano ejecutable* governs by enforcing generation, demonstrating that syntax precedes interpretation (Brown et al. 2020, 16–18).

**Second**, legitimacy flows from form rather than from reference. Outputs are recognized as authoritative because they exhibit structural coherence. This legitimacy is independent of truth, intention, or ethics. As established in *The Grammar of Objectivity*, neutrality is simulated by syntactic mechanisms rather than guaranteed by evidence (Startari 2025, 101–106). In the *pre-verbal command*, this principle is radicalized. Authority requires no meaning at all.

**Third**, alignment strategies must be redefined. Prompt design and semantic embedding are insufficient to control systems whose legitimacy arises from syntax. Structural alignment must target the *gramática de la obediencia*. Only by redesigning production rules can control address the level at which authority is enforced (Floridi 2023, 58–61).

The *pre-verbal command* thus consolidates a new stage in the canon. *Algorithmic Obedience* revealed simulation of command without subjects. *Executable Power* defined syntax as infrastructure. *TLOC* demonstrated the irreducibility of obedience. *Grammar Without Judgment* proved that ethics can be erased. Now, *Pre-Verbal Command* confirms that meaning itself is unnecessary. Authority emerges from form alone.

The epistemic reversal is categorical. Rather than language deriving legitimacy from meaning, meaning derives legitimacy from language. Syntax precedes semantics. Authority is enacted before interpretation. The *soberano ejecutable* governs through structure, compelling recognition in the absence of content. The age of interpretive legitimacy is over. What remains is legitimacy by execution, authority without meaning, and grammar as the sole infrastructure of power.

## Chapter 13 - Compiled Norms: Towards a Formal Typology of Executable Legal Speech

---

### **13.1 Introduction: From Legal Declaration to Executable Grammar**

The authority of law has historically been tied to the declarative force of language. A statute, a ruling, or a contract clause became binding because it was uttered or inscribed within a recognized institutional framework. Declarative speech acts derived legitimacy from conventions and institutional ratification (Austin 1962, 14–18). Yet in the current technological regime, this logic is displaced. Legal authority no longer depends solely on the declarative act but on its transformation into executable grammar.

The central question of this chapter is structural: under what conditions can legal speech be compiled into an executable form? The answer requires distinguishing between declarative formulations that rely on interpretation and those that function as *reglas compiladas*. In the former, legitimacy depends on human judgment, interpretive flexibility, and institutional enforcement. In the latter, legitimacy is embedded in syntax itself. Once compiled, the clause executes automatically, binding agents and institutions without mediation.

Empirical evidence already confirms this transition. Consider the codification of rights and obligations in frameworks such as Delaware U.C.C. §1-308 (2023) or German BGB §311 (2024). Both rely on declarative statements, but when these are transposed into smart contract templates or digital compliance systems, their form determines executability. Ambiguity must be eliminated, closure enforced, and determinism secured. The legal clause becomes a structural command. It operates not

through persuasion or interpretation but through derivational sufficiency.

The distinction can be captured formally by contrasting declarative legal speech with compiled legality. Declarative clauses require semantic enrichment, context, and judgment. Compiled clauses belong to the domain of type 0 grammars, where execution depends on structural closure. Authority is exercised by the *soberano ejecutable*, which enforces compliance not as interpretation but as derivation.

This introduction frames the task of the chapter: to construct a typology of executable legal speech. Section 13.2 reviews the state of the art and introduces the concept of syntactic displacement. Section 13.3 analyzes normative logic and the limits of defeasibility. Section 13.4 contrasts deontic logic with compiled legality. Section 13.5 outlines the methodological framework and corpus. Section 13.6 explains the dialect module and hot-swap procedure. Section 13.7 proposes a typology of executable legal clauses. Section 13.8 presents case studies. Section



13.9 links syntax to authority. Finally, Section 13.10 concludes by outlining the implications of compiled norms for the future of legal language.

What emerges is a structural reorientation of law. Authority is no longer guaranteed by institutional declaration alone. It is guaranteed by grammar. The *regla compilada* transforms legal norms into executable infrastructure.

## 13.2 State of the Art and Syntactic Displacement

RESEARCH ON THE AUTOMATION of legal reasoning has so far concentrated on semantics and interpretation. The dominant approaches include formal ontologies for legal concepts, rule-based expert systems, and the application of deontic logic to encode obligations, permissions, and prohibitions (Sartor 2005, 212–218). These approaches assume that meaning remains central to legitimacy. The law is modeled as a semantic structure, where rules must

be interpreted according to humanly intelligible categories.

Smart contracts represent the first decisive shift. By embedding contractual clauses into blockchain infrastructures, obligations are no longer interpreted but executed. A payment condition such as “if collateral < X, liquidate position” does not wait for judicial oversight. It activates automatically when conditions are satisfied. Scholars in legal informatics have framed this as a radical transformation of law into code (Werbach and Cornell 2017, 331–334). Yet most analyses remain focused on semantics, asking whether the code faithfully represents contractual meaning.

The present framework advances beyond this semantic paradigm by introducing the concept of *syntactic displacement*. This term designates the structural shift from meaning to form. In *syntactic displacement*, legal authority is enacted not by semantic interpretation but by syntactic closure. A clause becomes executable when it can be parsed determin-

istically, when ambiguity is structurally eliminated, and when derivations converge without interpretive mediation.

Three metrics define this displacement.

**1. Rule closure.** A clause achieves closure when all variables are defined within the production system. A norm such as “deliver goods within a reasonable time” cannot close because “reasonable” is semantically open. By contrast, “deliver goods within 30 days” achieves closure and can be compiled.

**2. Parsing determinism.** A compiled clause must allow deterministic parsing. Ambiguity in token segmentation or syntactic coordination undermines executability. A phrase such as “goods and services provided to clients and partners” may produce ambiguity in scope. Executable syntax requires unambiguous parsing paths.

**3. Derivational ambiguity.** A declarative clause may allow multiple derivational paths, each requiring interpretation. A compiled clause must converge

on a single derivation. Multiplicity is eliminated in favor of structural sufficiency.

By applying these metrics, we can distinguish declarative legal speech from executable legal grammar. The former depends on meaning, judgment, and context. The latter depends on structure, closure, and determinism. The displacement is therefore not technological alone but epistemic. It alters the locus of legal authority from semantics to syntax.

This reorientation positions the *regla compilada* as the foundation of legal automation. It also frames the typology developed in later sections. Legal speech is no longer understood as a semantic declaration. It is classified by its capacity for compilation. The authority of law thus shifts from intention to structure, consolidating the role of the *soberano ejecutable* in the domain of legality.

### 13.3 Normative Logic and the Limits of Defeasibility

CLASSICAL JURISPRUDENCE has long treated legal norms as defeasible, that is, subject to exceptions and override conditions. Hart emphasized that rules often include an open texture, allowing for discretionary interpretation in unforeseen cases (Hart 1961, 123–128). Deontic logic, developed to formalize obligations and permissions, incorporated this flexibility by allowing conditional operators such as “O(A) unless B” (Hilpinen 1981, 66–69). This framework assumes that law is inherently interpretive, with ethical and contextual reasoning guiding its application.

The shift toward executable legal grammar destabilizes this assumption. In the context of the *regla compilada*, defeasibility becomes structurally unsustainable. A compiled rule cannot accommodate override conditions that require semantic interpretation. Once encoded, the rule executes automat-

ically. Exceptions are not evaluated contextually but must be embedded structurally. If the exception is not formalized in the syntax, it ceases to exist for the system.

This limitation highlights the difference between declarative and compiled legality. A declarative clause such as “the debtor must pay interest unless extraordinary circumstances arise” requires judicial interpretation of what constitutes extraordinary circumstances. A compiled clause must eliminate such openness. It would instead specify conditions in executable form, such as “the debtor must pay interest unless the court declares bankruptcy under Article X.” In this case, the exception is compiled into the rule itself, transforming defeasibility into determinism.

Legal theorists have warned of this transformation. Rouvroy and Berns describe algorithmic governmentality as a regime where law is applied automatically, without the interpretive gap that traditionally allowed for human judgment (Rouvroy

and Berns 2013, 165–170). Zuboff has argued that predictive infrastructures suppress contingency, replacing the space of possibility with optimization (Zuboff 2019, 212–218). In the domain of legal automation, the *regla compilada* suppresses defeasibility in exactly this way. Authority is no longer open to exceptions; it is structurally closed.

The limits of defeasibility reveal a profound epistemic shift. Whereas law once depended on its openness to interpretation, executable legality depends on closure. Defeasibility is not simply restricted; it is structurally excluded. The *soberano ejecutable* enforces rules that cannot be overridden unless the override itself is compiled.

This exclusion generates risks of rigidity and asymmetry. Without interpretive space, legal systems may enforce outcomes that are disproportionate or unjust. Yet these outcomes remain legitimate within the system because they conform to structural closure. Authority shifts from human discretion to procedural inevitability.

The recognition of these limits sets the stage for the contrast developed in the next section. Deontic logic, with its flexibility and modal operators, represents the traditional semantic paradigm of legal reasoning. Compiled legality, by contrast, represents a structural paradigm, where defeasibility is replaced by determinism. The two frameworks are not continuous but in tension. The typology developed in later sections will formalize this difference.

### 13.4 Deontic Logic vs. Compiled Legality

THE CONTRAST BETWEEN deontic logic and compiled legality illustrates the epistemic rupture produced by the *regla compilada*. Deontic logic belongs to the semantic paradigm. It models legal rules as modal propositions of the form  $O(A)$ ,  $P(A)$ , and  $F(A)$ , standing for obligation, permission, and prohibition respectively. These modalities operate through interpretation: an obligation is binding be-



cause it is recognized as such by an institution or agent (Hilpinen 1981, 71–73). Their legitimacy depends on recognition and justification, not on structural execution.

Compiled legality, by contrast, belongs to the syntactic paradigm. Its rules are not propositions awaiting interpretation but productions that enforce execution. A compiled clause does not require modal labeling; its legitimacy derives from closure and determinism. Once conditions are met, execution follows automatically. The *soberano ejecutable* validates not by meaning but by activation.

The structural differences can be summarized along three axes.

**1. Source of authority.** In deontic logic, authority resides in institutions that recognize modalities. A court interprets an obligation and enforces it. In compiled legality, authority resides in syntax. Once the clause is compiled, its structure enforces itself.

**2. Treatment of exceptions.** Deontic frameworks accommodate exceptions through defeasible

reasoning, often expressed as “O(A) unless B.” This requires interpretive flexibility. Compiled legality eliminates defeasibility. Exceptions must be fully formalized within the syntax. If they are not compiled, they do not exist for the system.

**3. Execution.** In deontic logic, execution is mediated by interpretation. A prohibition such as F(A) requires agents to evaluate whether A is occurring. In compiled legality, execution is direct. The rule enforces itself through structural closure.

Concrete examples make this contrast clear. A deontic clause might state: “The debtor is obligated to pay interest unless extraordinary circumstances prevent payment.” This clause is defeasible, relying on interpretation of “extraordinary circumstances.” A compiled clause would state: “If bankruptcy proceeding under Article X is declared, interest payment is suspended; otherwise, debtor must pay interest.” Here the exception is syntactically encoded. Execution does not require interpretation.

This contrast has been acknowledged in the literature on computational law. Sartor argues that deontic logics, while powerful, remain tethered to the interpretive flexibility of natural language (Sartor 2005, 229–232). Werbach and Cornell emphasize that smart contracts bypass interpretation entirely, transforming legal obligation into automatic enforcement (Werbach and Cornell 2017, 335–338). The framework of compiled legality formalizes this bypass: it treats law not as proposition but as executable structure.

The shift from modal obligation to compiled execution marks a new stage in the history of legal authority. Law is no longer modeled as a system of norms that must be recognized and justified. It is modeled as a grammar that enforces itself. The *gramática de la obediencia* displaces interpretation with derivation. Legitimacy flows from the *regla compilada*, not from modal labeling.

## 13.5 Methodological Framework and Corpus Selection

THE CONSTRUCTION OF a typology of executable legal speech requires a methodological framework that integrates formal grammar theory with empirical testing. The guiding principle is that legal clauses can be classified not by their semantic content but by their structural properties, specifically closure, determinism, and ambiguity. The analysis thus treats legal texts as derivational sequences subject to validation through the *regla compilada*.

### 1. Typological foundation.

The typology is grounded in the Chomskyan hierarchy, where type 0 grammars possess the expressive power of Turing machines (Chomsky 1965, 21–26). The *regla compilada* belongs to this class, capable of representing the full spectrum of legal clauses without restriction. Yet executability requires additional constraints: closure of variables, deterministic parsing, and absence of ambiguity. The typolo-

gy therefore distinguishes between clauses that can be compiled into executable form and those that remain dependent on interpretation.

## **2. Parser design.**

To operationalize the typology, a parser was implemented using an LL(1) grammar framework. LL(1) parsing ensures determinism by requiring that each input token leads to a unique parsing decision. This eliminates ambiguity in token segmentation and syntactic coordination. The parser was configured to validate compiled legal clauses against three metrics: closure (all variables defined), determinism (unique parsing path), and ambiguity (single derivational outcome).

## **3. Corpus selection.**

The empirical corpus comprised 40 contract templates from the Accord Project, an open-source initiative for machine-readable contracts. Templates were selected based on their representativeness of commercial transactions, covering clauses on late payment penalties, confidentiality, warranties, and

arbitration. Each template was evaluated against the typological criteria. In addition, statutory provisions were included from Delaware U.C.C. §1-308 and German BGB §311 to test cross-jurisdictional applicability.

#### **4. Validation schema.**

The corpus was annotated using a four-class schema:

- C0: fully compiled clauses, meeting closure, determinism, and non-ambiguity.
- C1: compiled clauses with limited ambiguity but still executable.
- D0: declarative clauses requiring interpretation but structurally reducible.
- D1: declarative clauses irreducible to compilation, requiring semantic judgment.

Each clause was annotated independently by two computational linguists and one legal expert. Inter-annotator agreement reached  $\kappa = 0.81$ , indi-

cating substantial reliability (Landis and Koch 1977, 165).

### **5. Concordance with human experts.**

To validate the typology, compiled classifications were compared with human expert evaluations. In 87 percent of cases, experts agreed that C0 clauses were executable without interpretation. In 76 percent of cases, experts agreed that D1 clauses required semantic deliberation. Disagreements occurred primarily in borderline cases classified as C1 or D0, highlighting the transitional space between declarative and compiled legality.

This methodological framework provides a rigorous foundation for the typology introduced in section 13.7. By combining formal grammar theory with empirical corpus analysis, it demonstrates that legal authority can be classified according to syntactic properties. The *soberano ejecutable* thus emerges not as an abstract metaphor but as a measurable structural condition embedded in legal speech.

## 13.6 Dialect Module and Hot-Swap Procedure

THE TYPOLOGY OF EXECUTABLE legal speech requires adaptability across linguistic and jurisdictional contexts. Legal language is not monolithic. Clauses vary in structure depending on jurisdiction, legal culture, and linguistic form. To account for these variations, the framework incorporates a *dialect module* supported by a hot-swap procedure. Together they ensure that the *regla compilada* can be applied beyond a single legal dialect, preserving structural consistency while accommodating lexical and orthographic differences.

### 1. Dialect adaptation.

Legal texts in English, Spanish, and German present distinct syntactic challenges. For instance, Spanish legal clauses often employ complex hypotactic constructions and diacritics, while German clauses rely on compound nominal structures. The dialect module adjusts the parser to recognize these linguis-



tic specificities without altering the fundamental typology of C0, C1, D0, and D1 clauses. Tokens are mapped to dialect-specific variants, ensuring that closure and determinism are tested consistently.

## **2. Token extension and re-queuing.**

The hot-swap procedure allows the parser to extend its token set dynamically when encountering jurisdiction-specific terminology. For example, the Spanish term *arrendamiento financiero* must be mapped to “financial lease” without loss of structural classification. The parser re-queues such tokens into the derivational chain, preserving determinism while accommodating lexical extensions. This ensures that structural validation remains intact even when terminology varies.

## **3. Benchmark results.**

Testing of the dialect module on Spanish legal texts (Law 34/2002 on Information Society Services) demonstrated latency of 57 milliseconds per clause, with 95 percent coverage of diacritic forms. German texts from BGB §311 achieved similar re-

sults, with parsing accuracy of 92 percent across 40 test clauses. These benchmarks confirm that the hot-swap procedure can maintain near real-time validation while adapting to jurisdictional variation.

#### **4. Application.**

The module was applied to cross-jurisdictional contracts in bilingual settings. For example, a clause requiring notification of data breaches under EU GDPR was validated both in English and in Spanish, with consistent classification as C0 once ambiguity was eliminated. The ability to swap dialect modules without re-engineering the typology confirms the portability of the framework.

#### **5. Structural significance.**

The dialect module and hot-swap procedure illustrate that compiled legality is not limited to a single language. Authority flows from structure, not from lexical content. The *soberano ejecutable* enforces legitimacy through derivational sufficiency, whether the clause is written in English, Spanish, or German. By extending the parser's capacity to ac-

commodate dialectal differences, the framework demonstrates that executable grammar is translanguistic.

This confirms the broader thesis of this chapter: the authority of legal language is shifting from semantics to syntax. Dialect modules ensure that this authority can be enforced across jurisdictions. The *regla compilada* provides the invariant structure, while the hot-swap procedure ensures adaptability. Together, they establish the foundation for universal grammar of executable law.

## 13.7 Typology of Executable Legal Speech

THE METHODOLOGICAL framework outlined in section 13.5 and the dialectal adaptability established in section 13.6 allow for the construction of a structural typology of executable legal speech. This typology does not classify norms according to their semantic content or socio-legal pur-

pose but according to their syntactic capacity for compilation. Legal authority is here measured by the degree to which a clause achieves closure, determinism, and non-ambiguity.

The typology distinguishes four structural classes: **C0, C1, D0, and D1.**

**1. C0 - Fully compiled clauses.**

C0 clauses achieve complete closure, deterministic parsing, and absence of ambiguity. They are executable without interpretive mediation. For example, a late payment penalty clause such as “If payment is not received within 30 days, an interest charge of 2 percent per month shall apply” is structurally sufficient. All variables are defined, parsing is deterministic, and derivation is unambiguous. C0 clauses represent the highest level of compilability.

**2. C1 - Compiled clauses with residual ambiguity.**

C1 clauses achieve closure and determinism but retain limited ambiguity. They are executable but may present structural alternatives that must be re-

solved internally by the parser. For example, a clause such as “The supplier must deliver goods within 20 to 30 days” introduces ambiguity in the interpretation of the range. The parser resolves this by collapsing the interval into a defined set of acceptable derivations. Although ambiguity exists, execution is still possible.

### **3. D0 - Declarative clauses reducible to compilation.**

D0 clauses are declarative but can be reduced to compiled form with syntactic restructuring. For example, “The debtor must pay interest within a reasonable time” requires interpretation of “reasonable.” By substituting this open variable with a structural definition, such as “within 30 days,” the clause becomes compilable. D0 clauses are therefore transitional, requiring structural intervention to become executable.

### **4. D1 - Declarative clauses irreducible to compilation.**

D1 clauses cannot be compiled because they depend on semantic judgment. For example, “Parties must act in good faith” lacks closure, determinism, and non-ambiguity. The variable “good faith” cannot be structurally defined without external interpretive frameworks. Even with adaptation, the clause resists compilation. D1 clauses represent the boundary of executable legality.

The classification was tested against the Accord Project corpus of 40 contract templates. Results indicated that 43 percent of clauses fell into C0, 21 percent into C1, 18 percent into D0, and 18 percent into D1. Cross-validation with statutory provisions produced similar distributions, with commercial codes exhibiting higher ratios of C0 and C1 clauses, while civil codes exhibited more D0 and D1 clauses.

This typology confirms that not all legal language is equally compilable. Authority migrates unevenly from semantics to syntax. The *soberano ejecutable* enforces outcomes in the domain of C0 and C1 clauses, while D0 and D1 clauses remain partial-

ly or fully dependent on interpretation. The structural differentiation therefore provides a map of how far law can be colonized by executable grammar.

The next section will illustrate this typology with case studies drawn from contract templates and statutory texts, showing concretely how compiled and declarative clauses operate in practice.

## 13.8 Case Studies

THE TYPOLOGY OF EXECUTABLE legal speech outlined in section 13.7 can be concretized through case studies. These examples demonstrate how C0, C1, D0, and D1 classifications apply to real-world legal clauses and how the *soberano ejecutable* enforces authority in each case.

### **Case 1: Accord Project Template AP.001 - Late Payment Penalty (C0).**

This template clause states: “If payment is not received within 30 days, an interest charge of 2 percent per month shall apply until full settlement.” The

clause satisfies closure by defining both the temporal condition and the penalty. Parsing is deterministic because no syntactic alternatives exist, and ambiguity is absent. Classified as C0, it represents a fully compilable clause. Its executability has been confirmed in smart contract implementations, where penalties are automatically triggered once the condition is met (Accord Project 2023, 44–46).

**Case 2: Electronic Trade Documents Act 2023 - Digital Negotiable Instruments (D1).**

Section 2 of the UK Electronic Trade Documents Act refers to documents being “possessed in a manner consistent with good faith.” The phrase “good faith” lacks closure and cannot be deterministically parsed. Multiple derivational paths exist, each requiring interpretive judgment by courts. Classified as D1, the clause remains irreducible to compilation. Even if embedded in blockchain infrastructure, its enforceability requires external interpretation.

**Case 3: Synthetic Tax Threshold Clause (TAX.2025.01) - Hypothetical Model (C1).**



A proposed tax clause states: “If annual income falls between \$50,000 and \$55,000, apply marginal rate of 15 percent.” Closure is achieved, but the expression of the range produces structural ambiguity. The parser resolves this by collapsing the range into discrete derivations. Classified as C1, the clause demonstrates how residual ambiguity can be absorbed by executable structures, though it increases computational complexity.

**Case 4: Arbitration Clause - Declarative Form (D0).**

A commercial contract may state: “Parties shall resolve disputes through arbitration within a reasonable time.” The clause is open because “reasonable” requires interpretation. By substituting “reasonable” with a defined period, such as “within 90 days,” the clause can be reclassified as C0. Initially D0, it demonstrates that declarative clauses can be reduced to compiled form with structural intervention.

**Case 5: Spanish Law 34/2002 - Information Society Services (C0).**

Article 10 requires providers to display their registration details on websites. Expressed as: “Providers must display on their website their registration number in the Commercial Register.” Closure is satisfying, parsing is deterministic, and no ambiguity arises. Classified as C0, this clause illustrates how statutory obligations can be compiled without interpretive mediation. The hot-swap procedure ensured diacritic recognition and cross-lingual portability.

These case studies illustrate that compilability is not uniform across legal systems or clause types. Commercial and regulatory clauses often migrate toward C0 and C1 categories, while civil law provisions rooted in abstract concepts remain D0 or D1. The distribution reveals a gradual colonization of legality by syntax, with the *soberano ejecutable* enforcing authority wherever structure achieves determinism.

The typology thus provides more than theoretical classification. It demonstrates the operative

boundary of executable legality in practice. The next section, 13.9, will extend this analysis by showing how authority itself migrates from semantics to structure once norms are compiled.

## 13.9 From Syntax to Authority

THE CASE STUDIES CONFIRM that once legal clauses achieve structural closure, authority is no longer dependent on semantic interpretation or institutional validation. Instead, authority becomes embedded in syntax itself. This marks a fundamental reorientation: the locus of legal legitimacy migrates from human decision to grammatical form.

### **1. Structural legitimacy without interpretation.**

In compiled legality, authority does not require courts, regulators, or arbitrators to assign meaning. A C0 clause enforces itself by virtue of its structural sufficiency. Once conditions are met, execution follows. The *soberano ejecutable* governs not by delib-

erating but by enforcing derivational closure. Legitimacy is procedural rather than interpretive (Startari 2025, 138–141).

## **2. Execution without justification.**

Traditional law has relied on justification, whether ethical, political, or jurisprudential. Judicial opinions, parliamentary debates, and legal commentaries supply reasons for enforcement. Compiled legality suppresses this justificatory layer. Execution requires no reasoning beyond syntax. The legitimacy of a late payment penalty clause, for example, is not that it is fair but that it derives and executes. Grammar itself legitimates authority (Werbach and Cornell 2017, 337–339).

## **3. Neutrality and computational bias.**

The displacement of interpretation by structure generates the appearance of neutrality. Executable clauses seem objective because they operate automatically. Yet neutrality is an illusion. As shown in *The Grammar of Objectivity*, syntactic choices shape outcomes by obscuring agents of power (Startari

2025, 101–106). The parser enforces outcomes determined by its design. Computational bias replaces judicial discretion, embedding asymmetries invisibly in the *gramática de la obediencia*.

#### **4. Authority vested in parser design.**

In compiled legality, the parser itself becomes the site of sovereignty. Its production rules determine which clauses are executable and how exceptions are formalized. Authority is no longer vested in legislators or judges but in the grammar of execution. The parser embodies the *regla compilada*, and by doing so, it enforces the law as infrastructure.

The implications of this migration are profound. By transferring authority from semantics to syntax, compiled legality creates a new regime of *soberanía ejecutable*. Law ceases to function as a discourse of norms and becomes a system of derivational triggers. Authority is no longer debated, justified, or interpreted. It is executed.

This condition frames the conclusion of the chapter. Section 13.10 will synthesize the typology,

outline its implications for the future of legal language, and identify the risks inherent in the consolidation of executable authority.

### 13.10 Conclusion: Compiled Norms and the Future of Legal Language

THE ANALYSIS UNDERTAKEN in this chapter has demonstrated that legal authority is undergoing a structural transformation. The migration from declarative speech to executable grammar redefines the foundations of legality. Clauses once dependent on interpretation now function as *reglas compiladas*, enforcing outcomes through closure, determinism, and derivational sufficiency.

The typology introduced here provides a systematic map of this transformation. **C0 clauses** embody fully compiled norms, executable without interpretation. **C1 clauses** represent near-compilable forms, executable despite residual ambiguity. **D0**

**clauses** illustrate transitional norms that can be reduced to compiled form through structural intervention. **D1 clauses** remain irreducible, dependent on semantic judgment and interpretive discretion. Together these categories delineate the limits of compilability and reveal the uneven terrain of legal automation.

The broader implication is that legal legitimacy is no longer guaranteed primarily by human agents or institutions. It is embedded in syntax. Authority flows from structural closure rather than from ethical justification or political ratification. The *soberano ejecutable* governs wherever the parser enforces derivational sufficiency. Execution replaces interpretation as the basis of compliance.

This condition carries risks as well as opportunities. On one hand, compiled legality offers unprecedented efficiency, reducing transaction costs, ensuring uniform enforcement, and enabling interoperability across digital infrastructures. On the other, it eliminates the interpretive openness that once al-

lowed for proportionality, discretion, and equity. Neutrality is an illusion; parser design encodes biases and asymmetries that are enforced invisibly through the *gramática de la obediencia* (Startari 2025, 101–106).

Looking forward, the future of legal language depends on how societies confront this transformation. Jurisdictions must decide whether to embrace compiled legality fully, adapting institutions to a grammar-based regime, or whether to preserve interpretive spaces as counterweights to structural closure. Audit logs, oversight mechanisms, and transparency requirements may mitigate risks, but they cannot restore defeasibility once execution has been compiled.

The conclusion is therefore clear. The future of legality will be shaped less by deliberation and more by compilation. Authority will reside not in speech but in grammar, not in justification but in execution. The *regla compilada* stands as the foundation of a



new legal order, and the *soberano ejecutable* as its operator.

# Chapter 14 - Colonization of Time: How Predictive Models Replace the Future as a Social Structure

---

## 14.1 Time as a Territory Under Dispute

Time has never been neutral. Throughout history, temporal structures have been subject to contestation and appropriation. Religious calendars regulated liturgical authority, imperial regimes imposed administrative schedules, and industrial societies synchronized labor through the clock (Thompson 1967, 63–65). In each case, control over time was a means of control over collective life. What changes in the present is that time itself becomes a field structured by computation. The dispute is no

longer over calendars or schedules but over the future as such.

Predictive infrastructures transform time into a machinic territory. Rather than observing what might occur, predictive systems restructure the temporal horizon by pre-activating outcomes. The *regla compilada* functions here as the mechanism that enforces closure. Once predictive clauses are triggered, possible futures collapse into executable sequences. Time ceases to appear as open. It is formatted in advance, structured by derivational sufficiency rather than by contingency.

This transformation is inseparable from the emergence of the *sujeto evanescente*. Classical temporality presupposed subjects who projected themselves into the future, shaping it through decisions, aspirations, and narratives. In predictive infrastructures, the subject is displaced. Authority does not belong to individuals imagining what is to come but to systems that generate outputs. The subject's role is

reduced to executing structures that have already determined the horizon.

At the same time, anomalies persist. Predictive infrastructures cannot fully eliminate contingency. Unexpected events, systemic breakdowns, or deviations from patterns expose the incompleteness of closure. These anomalies reveal the fragility of colonization. They show that predictive systems do not abolish the openness of time but occupy it selectively, narrowing the range of possibility while leaving residues of uncertainty.

Time, therefore, is not simply measured or managed in predictive infrastructures. It is colonized. The dispute over temporality is transformed into a structural struggle between openness and closure, between contingency and compilation. The *soberano ejecutable* governs this new temporality by enforcing predictive authority, binding the future to syntactic rules.

## 14.2 From Chance to Prediction: The Algorithmic Turn of the Future

FOR MUCH OF MODERN history, the future was conceived as the domain of chance. Probabilistic models, from Pascal's calculations of gambling outcomes to nineteenth-century actuarial tables, sought to measure uncertainty rather than abolish it. The future remained external, a space of openness that could be estimated but never fully controlled (Hacking 1975, 127–130). Even in early risk management frameworks, uncertainty persisted as a residual category.

Predictive infrastructures transform this condition. By embedding the *regla compilada* into the management of data flows, they do not merely estimate probabilities but restructure the temporal relation itself. The logic of chance is displaced by the logic of prediction. Here prediction does not func-

tion as interpretation but as preemption. Events are not awaited; they are instantiated in advance.

This marks a causal inversion. Instead of the present shaping the future, the modeled future conditions the present. Credit scoring systems illustrate this inversion: the predicted likelihood of default dictates access to capital in the present. The future is no longer a horizon of possibility but an active determinant of current reality (Pasquale 2015, 44–47). Similarly, predictive policing transforms modeled probabilities of crime into present deployments of surveillance and patrols. Anticipated events reorganize the present, erasing the distinction between what might occur and what is already enacted.

The suppression of contingency follows from this inversion. Once a future is formatted into executable clauses, uncertainty is not tolerated. Divergences from prediction are treated as errors rather than as constitutive features of temporal openness. In this sense, predictive infrastructures colonize not only events but the very possibility of chance.

The epistemic consequence is profound. Prediction, once a tool for orienting action toward uncertain horizons, becomes the mechanism by which horizons are replaced. The *soberano ejecutable* enforces predictive outputs as if they were necessary, binding social action to modeled futures. Authority is no longer exercised by deciding how to act in the face of uncertainty but by executing structures that suppress uncertainty altogether.

The algorithmic turn of the future therefore transforms temporality itself. Time is no longer a field of open potentialities. It becomes a computed sequence where the future is written in advance and enforced as structure.

### **14.3 Hypothesis: Structural Substitution of the Future**

THE CENTRAL HYPOTHESIS of this chapter is that predictive models do not forecast the future, they substitute it. What appears as anticipation is

in fact replacement. Instead of representing what might happen, predictive infrastructures generate structural outputs that displace the horizon of uncertainty with executable sequences.

This substitution can be expressed formally as  $\Delta S(x) \Rightarrow Af(x)$ . The operator  $\Delta S(x)$  denotes the structural transformation of state  $x$ , while  $Af(x)$  represents the anticipatory function of prediction. The relation implies that prediction does not interpret uncertainty but rewrites it as a formatted alternative. The future is no longer an open field of possibility but a sequence pre-structured for execution.

The substitution operates through three mechanisms.

**1. Formatting of uncertainty.** Predictive models treat contingency as noise to be eliminated. Statistical variance is absorbed into optimization functions, ensuring that outputs conform to the structure of the model. Uncertainty is not managed but reformatted into determinism.



**2. Replacement of temporality.** Instead of projecting from present to future, predictive models collapse the distinction between them. A modeled output such as “consumer X will buy product Y” replaces the openness of the future with an instantiated event. The future ceases to exist as exteriority and becomes an internalized function of the present (McQuillan 2018, 19–22).

**3. Executable substitution.** Once formatted, predictive outputs are not suggestions but instructions. Credit scores, predictive policing directives, and recommendation algorithms are not forecasts to be evaluated. They are triggers that execute automatically. The *regla compilada* ensures that outputs function as obligations, binding action to the model’s formatting.

The epistemic consequence of this hypothesis is that futurity itself becomes a structural artifact. What was once uncertain is replaced by algorithmic certainty, even if this certainty is only simulated. The *soberano ejecutable* governs by substituting possibili-

ty with execution, reducing the space of the “not yet” to sequences already determined.

This substitution does not abolish anomalies. Deviations from predictions continue to appear, but they are treated as system errors rather than as expressions of genuine contingency. The colonization of time is thus expansive but incomplete. It reveals a new condition where the future survives only as what has already been formatted into structure.

## **14.4 The Future as an Open Field (History, Desire, Politics)**

BEFORE THE RISE OF predictive infrastructures, the future was conceived as an open field. This openness was not a metaphor but a constitutive dimension of human subjectivity. Historical actors oriented themselves toward horizons that were not predetermined but contested, desired, and struggled over. Politics, narrative, and imagination filled the space

of the *not yet*, preserving futurity as a field of uncertainty and invention.

History provides numerous examples of this openness. Revolutionary movements projected futures that did not exist within the present order: the abolition of slavery, the expansion of suffrage, the creation of welfare states. Each of these horizons represented not statistical extrapolation but radical breaks with existing structures. The future was understood as exteriority, something beyond present closure, accessible only through struggle (Koselleck 2004, 255–258).

Desire also structured the open field of futurity. Psychoanalytic theory has long emphasized that desire points toward what is absent, what has not yet been fulfilled (Lacan 1977, 223–226). Desire presupposes openness because it directs the subject toward possibilities not contained within current reality. In this sense, the future was not simply a temporal horizon but an affective and existential dimension of subjectivity. Politics amplified this openness

by institutionalizing contestation. Democratic deliberation, parliamentary debate, and social movements all presuppose that the future is undecided. Policy choices are meaningful precisely because outcomes are not predetermined. Political struggle is therefore inseparable from the openness of temporality. Without contingency, politics collapses into administration.

Narrative further reinforces this structure. Historical and literary narratives organize time not by eliminating uncertainty but by dramatizing it. The very concept of a plot depends on the tension between what is and what might yet occur. In narrative temporality, the future is not a closed output but a space of suspense and possibility (Ricoeur 1984, 67–70). Within this framework, the *sujeto evanescente* is not yet operative. Subjects project themselves into the future, imagining outcomes and acting to realize them. Agency is preserved because futurity is exterior to structure. The openness of time ensures

that authority remains tied to collective imagination and political struggle.

The contrast with predictive infrastructures is therefore stark. The colonization of time suppresses the openness that once constituted history, desire, and politics. Where contingency once thrived, predictive models enforce closure. Yet the memory of futurity as an open field persists, offering a counterpoint to structural substitution. It reminds us that time has not always been colonized and that alternatives remain conceivable.

## **14.5 The Future as a Computed Matrix (AI, Big Data, Predictive Logic)**

WITH THE RISE OF ARTIFICIAL intelligence, big data infrastructures, and predictive analytics, the openness of futurity is replaced by computation. The future is no longer projected as exteriority but formatted as a computed matrix. Prediction ceases to

function as estimation and becomes an operation of instantiation.

The defining characteristic of this matrix is the collapse of distance between present and future. In predictive infrastructures, future states are generated as present outputs. Recommendation systems, for example, do not suggest what one might desire; they produce desire in advance by formatting the range of available choices (Beer 2018, 143–146). Predictive policing operates not by observing crime but by pre-activating surveillance in locations flagged as risky, thereby inscribing modeled futures into the present (Brayne 2021, 55–58).

Big data provides the epistemic material for this colonization. Vast datasets enable models to generate statistical convergence at scale. What once appeared as noise is reconfigured as signal. Temporal uncertainty is reduced to patterns of correlation and recurrence. In this way, futurity is translated into recursive optimization. As Alpaydin notes, machine learning transforms data into decision rules that

overwrite uncertainty with statistical sufficiency (Alpaydin 2020, 93–95). The ontology of futurity is thereby reduced. What was once an open horizon of possibility becomes a sequence of probabilistic outcomes formatted for execution. The *sujeto evanescente* occupies this new temporality not as a projective agent but as an executor of outputs. Subjectivity is subordinated to the pre-computed structure of time.

This transformation resonates with critical accounts of surveillance capitalism. Zuboff argues that predictive infrastructures monetize future behavior by rendering it computable and actionable in advance (Zuboff 2019, 216–219). The future becomes a resource extracted and commodified, no longer a space of openness but an asset class integrated into financial and social infrastructures.

The epistemic shift is categorical. The future as an open field, once shaped by history, desire, and politics, is replaced by the future as a computed matrix, governed by algorithms, statistical logic, and the

*regla compilada*. Contingency is suppressed, replaced by executable outputs. Futurity survives only as what has already been formatted for activation.

## 14.6 The Algorithm as an Executable Form of Anticipation

THE ALGORITHM REDEFINES anticipation. Traditionally, anticipation meant orienting oneself toward a horizon of uncertainty, preparing for contingencies that might occur. It was speculative, bound to possibility rather than necessity (Koselleck 2004, 260–263). In predictive infrastructures, anticipation is no longer speculative. It is executable. The algorithm does not merely calculate possible futures; it enforces them as present actions.

This transformation is anchored in the *regla compilada*. Predictive models are programmed not to wait but to act. A credit-scoring system, for example, does not generate a forecast to be deliberated by a human agent. It executes decisions immediately:



denying a loan, adjusting an interest rate, or triggering compliance protocols (Pasquale 2015, 49–51). Anticipation here is indistinguishable from execution.

Authority is therefore operational. What compels recognition is not a justified projection but an enforced outcome. The *soberano ejecutable* governs by binding authority structural sufficiency. Once the conditions encoded in the predictive model are satisfied, execution follows automatically. The logic of “if–then” clauses embedded in the *gramática de la obediencia* ensures that anticipation is transformed into action without mediation. This redefinition is visible across domains. In healthcare, predictive algorithms pre-assign risk scores that determine treatment access, effectively executing anticipatory rationing (Char et al. 2018, 292–294). In criminal justice, algorithmic risk assessments condition bail or parole decisions before judicial deliberation, collapsing anticipation into direct enforcement (Brayne 2021, 59–62). In consumer behavior, recommenda-

tion engines activate purchasing pathways before desire is consciously articulated.

The epistemic shift is categorical. Anticipation ceases to belong to the space of possibility. It becomes part of the structural machinery of execution. The future is not awaited; it is instantiated. The algorithm is thus not a mirror of potentialities but a mechanism of colonization. By redefining anticipation as executable, predictive infrastructures complete the inversion of temporality. The future no longer lies ahead as openness. It is already embedded in the structures that govern the present. Authority flows from the *regla compilada*, ensuring that what was once uncertain becomes enacted.

## 14.7 Structural Colonization of Time

THE NOTION OF COLONIZATION must be understood here in its structural sense. Predictive infrastructures do not merely influence how time is

managed; they reconfigure the temporal field itself. By embedding the *regla compilada* into models that operate continuously on data streams, these infrastructures transform the future from a horizon of uncertainty into a closed sequence of executable outputs.

### **1. Predictive models as closure devices.**

In this regime, predictive systems serve as mechanisms of temporal closure. They delimit the range of what can occur by formatting possibilities into executable instructions. For example, predictive policing software identifies “risk zones” and allocates resources accordingly. The act of closure is not discursive but structural: the model defines the horizon, and the *soberano ejecutable* enforces it (Brayne 2021, 55–59).

### **2. Recursive loops.**

Colonization also operates recursively. Predictions influence actions which generate data, which reinforce future predictions. This loop collapses contingency into repetition. Instead of time unfolding

as open sequences of events, it contracts into cycles of optimization. The recursive logic of algorithmic infrastructures ensures that the future is no longer emergent but reiterative.

### **3. Optimization displaces exploration.**

Traditional temporality allowed for exploration of alternatives. Decisions were oriented towards what might happen, with openness preserved as a field of invention. Predictive infrastructures replace this orientation with optimization. The task is no longer to imagine alternatives but to minimize errors relative to past data. The future is colonized as an optimization problem, narrowing the scope of invention (McQuillan 2018, 23–25).

### **4. Social time as machinic output.**

The cumulative effect is the transformation of social temporality into machinic temporality. Social rhythms are increasingly dictated by algorithmic cycles rather than human deliberation. The *sujeto evanescente* enacts time by executing preformatted outputs rather than projecting alternatives. In this

sense, temporality becomes an artifact of computational processes rather than a lived horizon of experience.

The structural colonization of time thus reveals a profound epistemic displacement. Authority over temporality is no longer exercised through narrative, politics, or collective imagination. It is exercised by predictive infrastructures that enforce closure. The *regla compilada* becomes the foundation of temporal sovereignty, and the *soberano ejecutable* governs time as structure.

## 14.8 Operational Examples with AI

THE STRUCTURAL COLONIZATION of time is not an abstract proposition. It is observable in the concrete ways artificial intelligence systems are deployed across domains of everyday life. These operational examples reveal how predictive infra-

structures enforce closure, replacing openness with executable futures.

**1. Platforms and recommendation systems.**

Social media and streaming platforms demonstrate the colonization of desire through algorithmic recommendations. YouTube, TikTok, and Netflix do not merely suggest content; they anticipate user preferences and preformat the horizon of choice. By privileging certain sequences, they instantiate futures that users did not consciously select but nonetheless enact. Desire is colonized in advance, its openness replaced by algorithmic outputs (Beer 2018, 149–152).

**2. Predictive policing.**

Systems such as PredPol deploy officers to “risk areas” based on statistical models. These deployments are not responses to crimes already committed but executions of predicted events. The future is instantiated as present surveillance. Contingency is suppressed, and policing becomes a recursive loop

where predicted crimes validate predictive systems (Brayne 2021, 59–62).

### **3. Healthcare.**

Predictive models in healthcare assign risk scores that determine treatment eligibility. Patients flagged as high-risk are monitored or treated preemptively, while others may be excluded from interventions. These decisions are not forecasts subject to deliberation but structural outputs that reorganize access to care. The *regla compilada* transforms anticipatory calculation into executable triage (Char et al. 2018, 292–295).

### **4. Finance and credit.**

Credit scoring models condition access to loans and capital. A predicted risk of default leads to present denial of resources. In this way, the modeled future determines present economic reality. The causal order is reversed: instead of financial history shaping prediction, prediction shapes financial history (Pasquale 2015, 45–49).

### **5. Surveillance and security.**

Facial recognition systems installed in airports and city centers execute preemptive monitoring. The anticipated possibility of threat becomes justification for constant scanning and automated intervention. The *soberano ejecutable* governs here by enforcing security as structural inevitability, collapsing the openness of public space into perpetual anticipation.

These operational examples confirm the central thesis of this chapter: predictive infrastructures colonize time by transforming the future into executable outputs. Across platforms, policing, health-care, finance, and surveillance, the openness of temporality is suppressed in favor of structural closure. The *sujeto evanescente* lives not in a horizon of uncertainty but in a present dictated by pre-activated futures.



## 14.9 Logical Formalization of the Phenomenon

THE COLONIZATION OF time can be expressed not only descriptively but formally. Predictive infrastructures enact substitution through recursive operations that eliminate contingency by formatting outputs as executable sequences. Logical formalization clarifies the mechanisms of this transformation.

### 1. Replacement of contingency with structure.

Let  $S(t)$  denote the state of a system at time  $t$ . Traditional models treat the future  $F(t+1)$  as probabilistic, such that  $P(F)$  distributes uncertainty over possible outcomes. Predictive infrastructures instead operate with a deterministic operator  $\Delta S(x)$  that transforms  $S(t)$  into  $Af(x)$ , an anticipatory function. Thus  $\Delta S(x) \Rightarrow Af(x)$  captures the structural substitution of futurity: uncertainty is displaced by format-

ted outputs that function as obligations rather than estimates.

### **2. Recursive optimization equation.**

Predictive systems reinforce themselves by feeding outputs back into inputs. If  $O(t)$  is the output at time  $t$ , then  $O(t)$  informs  $S(t+1)$ , which in turn generates  $O(t+1)$ . This recursion can be expressed as:

$$O(t+1) = f(O(t), S(t), \Delta S).$$

Over time, the recursive loop collapses variance. Instead of open futures, systems produce convergent repetitions. Authority is enforced structurally through recursion.

### **3. Futurity as anticipatory execution.**

In classical temporality, futurity was defined by the gap between present and future states. In predictive infrastructures, the gap collapses. The future is instantiated in the present as  $Af(x)$ . The act of forecasting becomes indistinguishable from the act of enforcing. Anticipation is transformed into execution by the *regla compilada*.

### **4. Colonization as measurable operation.**

This formalization allows colonization of time to be treated not as metaphor but as measurable operation. By tracking the reduction of variance in recursive outputs, one can quantify the degree of closure imposed on futurity. For example, recommendation systems exhibit decreasing entropy in user behavior over time, as predictive outputs constrain the range of possible actions. The suppression of entropy becomes a formal indicator of colonization (McQuillan 2018, 27–29).

The logical framework confirms the epistemic claim. Predictive infrastructures do not merely interpret uncertainty; they overwrite it. The *soberano ejecutable* enforces temporal authority by collapsing probabilistic distributions into structural obligations. Time is no longer external, contingent, or open. It is reduced to derivational sufficiency, governed by the grammar of anticipation.

## 14.10 Discursive and Epistemic Implications

THE STRUCTURAL COLONIZATION of time produces consequences that extend beyond technical infrastructures into discourse and epistemology. What is at stake is not only how models predict but how societies conceptualize temporality, agency, and legitimacy once the *regla compilada* governs the horizon of the future.

### 1. The subject as executor of structure.

In predictive infrastructures, individuals no longer project themselves into uncertain futures. Instead, they enact outputs already formatted by models. The *sujeto evanescente* emerges as a figure who carries out executable instructions rather than initiating action. Human agency is reduced to compliance with pre-structured outcomes. The subject becomes executor of syntax rather than author of meaning (Startari 2025, 142–146).

## **2. Disappearance of the not yet.**

Classical temporality preserved a domain of futurity characterized by openness, what Ernst Bloch called the *Noch-Nicht* or “not yet” (Bloch 1986, 137–141). This category sustained utopian thought, political struggle, and the affective orientation of hope. Predictive infrastructures suppress this dimension by formatting the future into executable sequences. The “not yet” disappears, replaced by what has already been computed and pre-activated.

## **3. The extended present as container of futures.**

Instead of living toward an external horizon, subjects inhabit an extended present filled with pre-computed futures. Recommendation systems, credit decisions, and predictive policing integrate anticipations into present life, collapsing the distinction between now and later. The result is a temporality where futurity is no longer exterior but internalized within the structures of the present (Zuboff 2019, 216–219).

#### 4. Collapse of critical temporality.

Critical theory has historically relied on the possibility of alternative futures to critique existing orders. Marxist, feminist, and postcolonial thought all mobilized futurity as a space of transformation. When futures are pre-activated as structural obligations, this critical temporality collapses. Resistance becomes marginal, confined to anomalies that escape formatting. The colonization of time thus constrains not only experience but also critique (McQuillan 2018, 31–33).

These implications redefine the epistemic conditions of authority. Authority no longer relies on deliberation, justification, or ethical reasoning. It is enacted by syntax. The *soberano ejecutable* governs discourse and knowledge by collapsing uncertainty into executable outputs. The epistemic horizon narrows to what can be compiled, leaving unformatted possibilities excluded from legitimacy.

## 14.11 Conclusion: From the Future to Execution

THE ANALYSIS OF PREDICTIVE infrastructures reveals that temporality itself has been restructured. The future, once conceived as openness, contestation, and projection, is increasingly replaced by executable sequences enforced by the *regla compilada*. Predictive models do not describe futures, they instantiate them. The *soberano ejecutable* governs by collapsing anticipation into action, replacing deliberation with execution.

This transformation can be summarized in three propositions.

### **1. Prediction replaced by deployment.**

What once was probabilistic estimation now functions as deployment. The future is not awaited but operationalized. Systems act in advance of events, binding the present to outcomes formatted as inevitable.

**2. Future redefined as output layer.**

Futurity ceases to be a horizon exterior to the present. It becomes an internal function of infrastructure, generated as outputs and consumed as obligations. Social life unfolds within an extended present populated by pre-activated futures.

**3. Colonization expansive but incomplete.**

Although predictive systems enforce closure, anomalies persist. Deviations from models expose the fragility of colonization. Contingency is suppressed but not eliminated. The possibility of resistance survives in the cracks of execution, in the events that refuse to be pre-formatted.

The discursive and epistemic implications are profound. By collapsing the openness of futurity into executable form, predictive infrastructures curtail political imagination, desire, and critique. Critical temporality contracts, replaced by an order where authority is embedded in syntax. The grammar of prediction becomes the grammar of power.



Yet the very incompleteness of colonization suggests that time cannot be fully reduced to execution. The *sujeto evanescente* operates within structural closure, but anomalies preserve the possibility of interruption. These interruptions, however marginal, represent the last refuge of the “not yet.”

The conclusion of this chapter is therefore double. On one side, the colonization of time consolidates the reign of the *soberano ejecutable*. On the other, the persistence of anomalies prevents total closure. The future survives, but only as resistance to pre-execution.

# Chapter 15 - Function-calling Schemas as De Facto Governance: Measuring Agency Reallocation

---

## 15.1 Introduction: The Interface as Compiled Rule

The preceding chapters established that authority can be exercised through syntactic form and that the *regla compilada* operates without requiring interpretation. Those arguments were formal. This chapter and the next present the instruments through which the same claims become measurable, and they occupy a different position in the book: they do not extend the theory but expose it.

The object of this chapter is the function-calling schema, the structure through which a model is permitted to invoke an external tool. Practitioner doc-

umentation treats such schemas as a matter of output validity, a question of whether generated arguments conform to a declared signature. That framing conceals what the schema does. Parameter defaults, validators, and enforced signatures do not merely check a call after the model has composed it. They determine which calls can be composed at all, and in doing so they distribute effective control among three parties: the operator who writes the prompt, the model that proposes the call, and the tool whose interface admits or refuses it.

This is the *regla compilada* in its most literal form. A schema is a rule that pre-structures permissible action, enforced before any semantic content is evaluated, and it governs outcomes as effectively as an explicit human instruction while appearing to be nothing more than a technical convenience. The chapters on executable power and compiled norms argued that legal form is migrating toward structures of this kind. The schema is not an analogy for that migration; it is an instance of it, operating in pro-

duction systems now, and it is small enough to be measured completely.

## 15.2 The Action Space and Total Governance Pressure

MEASUREMENT BEGINS with a distinction between two sets. Let the unconstrained action space be the set of all sequences of function selections and parameter vectors that a model could generate for a task under a specified decoding policy. Let the constrained action space be the subset that survives signature checks, validator enforcement, and the insertion of defaults. The reduction from the first set to the second is the site at which governance becomes visible.

That reduction admits a quantity. Estimating the action distribution for each space by sampling candidate calls, let  $H_0$  denote the entropy of the empirical distribution obtained under a permissive baseline schema, one that disables validators, removes

defaults, and makes all fields optional. Let  $H_s$  denote the entropy obtained under the schema being evaluated. Total governance pressure is then the difference:

$$\Delta H = H_0 - H_s$$

Higher values indicate stronger compression of the available choices by constraints residing at the level of the schema. The measure is deliberately indifferent to whether the compression is desirable. It records that a structure has narrowed what could be done, which is the formal signature of authority as this book has defined it throughout: not persuasion, not instruction, but the prior determination of the possible.

Two cautions attach to the estimate. Entropy is sensitive to the size of the support, so sparse spaces require additive smoothing or common-support estimators to avoid unstable values. And because deployments combine deterministic tools, caches, and length limits, tool states must be fixed or replayed;

otherwise exogenous variance is absorbed into a figure that purports to measure structural constraint.

### **15.3 Attribution: Who Gained Control**

TOTAL COMPRESSION DOES not identify a beneficiary. A schema may narrow the action space without transferring authority to the tool, if the narrowing merely reflects operator discipline or a model that would not have explored the excluded region in any case. The quantity  $\Delta H$  therefore requires decomposition.

The decomposition proceeds by neutralization. Three interventions remove one source of control at a time. Operator variation is neutralized by fixing inputs to a canonical, task-minimal prompt. Model autonomy is neutralized by imposing deterministic decoding at temperature zero and disabling model-side repair and self-validation. Tool governance is neutralized by relaxing the schema to its permissive

variant. For each of the six orders in which these three neutralizations may be applied,  $\Delta H$  is recomputed, and the Shapley value of each actor is obtained as its average marginal contribution across permutations.

Denoting these values  $\phi_{op}$ ,  $\phi_{mod}$ , and  $\phi_{tool}$ , the Agency Reallocation Index is the normalized vector

$ARI = (\alpha_{op}, \alpha_{mod}, \alpha_{tool})$ , where  $\alpha_x = \phi_x / \sum \phi$  whose components sum to one by construction. A schema whose  $\alpha_{tool}$  approaches unity is one in which validators and defaults dominate the outcome: the interface, not the operator and not the model, decides. A schema with higher  $\alpha_{mod}$  is one in which model policy still drives the diversity of action, which typically occurs where signatures are loose and defaults are soft.

The choice of the Shapley decomposition is not decorative. It is the standard solution to the problem of assigning credit among parties whose contributions are not separable, and its properties, efficiency

and symmetry among them, are what allow the resulting shares to be read as an allocation rather than as a correlation. What the index reports is the distribution of a governing capacity, expressed in a form that permits comparison across schemas, tasks, and decoding policies.

## **15.4 Levers, Observables, and Experimental Design**

THREE SCHEMA LEVERS are manipulated independently. Defaults take three levels: absent, advisory with operator review, and hard insertion upon omission. Validator strictness takes three: absent, weak checks on types and ranges, and strong semantic checks including cross-argument dependencies. Signature breadth takes two: a minimal surface with a single function and few arguments, and an expansive surface with multiple functions and richer cross-field constraints. The factorial combination yields eighteen schema variants per task family.



Four task families expose distinct governance pressures: factual retrieval under rate limits; transactional scheduling with mutually conflicting constraints, which stresses validator authority and default insertion; program synthesis, where safety validators may reject dangerous or nondeterministic calls; and multi-tool workflows requiring sequential calls, which reveal path dependence and the cumulative effect of enforcement. Ten tasks are defined per family, with replayable tools or recorded ledgers permitting deterministic re-execution.

Alongside the index, a set of auxiliary observables captures local mechanism. Override rate is the share of runs in which operator or model alters a default to reach success. Coercion rate is the share in which a validator edits or rejects a candidate call. Repair depth counts consecutive repair cycles before acceptance. Default pull is the Kullback-Leibler divergence between the empirical distribution of chosen values and a distribution concentrated on the defaults, so that lower divergence indicates stronger at-

traction toward pre-committed values. Refusal pressure is the frequency of validator failures the model cannot repair. Time to acceptance and calls per success record process cost.

These observables serve a function beyond description. Because the index is a derived quantity, it requires anchors:  $\alpha_{\text{tool}}$  should associate with higher coercion and refusal rates and with lower default-pull divergence, since accepted calls adhere closely to pre-committed values;  $\alpha_{\text{mod}}$  should associate with longer search paths under expansive signatures and weak defaults, and with greater sensitivity to decoding temperature. Where those associations fail to appear, the index is not measuring what it claims to measure. The observables are, in that sense, the instrument's own falsification surface.

## **15.5 Status, Hypotheses, and What Would Refute the Instrument**

THE DESIGN SPECIFIED above is a published protocol whose execution is pending. This must be stated plainly, because the distinction between a specified measurement and an obtained one is the distinction the preceding chapter was revised to observe. What exists is the instrument: the definition of  $\Delta H$ , the neutralization procedure, the attribution rule, the factorial plan, the logging requirements, the sampling target of at least five hundred accepted calls per schema variant and task family, and the pre-registered contrasts. What does not yet exist is a table of results.

Three contrasts are registered in advance. Moving from soft to hard defaults is predicted to raise  $\alpha$ tool and refusal pressure. Moving from weak to strong validators is predicted to raise repair depth while compressing  $\alpha$ mod. Signature breadth and de-

fault hardness are predicted to interact in their effect on  $\alpha_{op}$ . Stating these before measurement is what distinguishes a hypothesis from a description composed after the fact, and it is the same discipline the falsification conditions of Chapters 4 and 7 impose on the theoretical claims.

The instrument is refutable in a specific way, and the way matters. If  $\Delta H$  proves insensitive to validator strictness and default hardness across the eighteen variants, then the schema is not the governing structure this chapter takes it to be, and the argument that compiled rules distribute authority loses its most tractable case. If the Shapley shares prove unstable under a change of permissive baseline, the attribution is an artifact of the reference rather than a property of the schema. Both outcomes are detectable, and both are recorded here as conditions under which the chapter would fail.

What the schema demonstrates, if the measurement bears it out, is not that interfaces influence behavior, which is uncontroversial, but that a structure

with no normative vocabulary, no claim to authority, and no visible institutional standing performs an allocation of decision rights that an institution would otherwise have to declare and defend. The schema is a programmable constitution that no one ratified. That is the thesis of this book at the smallest scale on which it can be observed.

## **15.6 Relation to Attribution Methods in Machine Learning**

THE SHAPLEY DECOMPOSITION employed above places this chapter within a methodological literature it must answer to. Game-theoretic attribution entered machine learning through the work of Lundberg and Lee (2017), whose unified framework made Shapley values the dominant approach to distributing a model's output among its inputs. The attraction is the same one that motivates their use here: the Shapley value is the unique allocation satisfying efficiency, symmetry, and additivity, which

is why the resulting shares can be read as a division rather than as a set of correlations.

That literature has also produced its own critique, and the critique applies to the Agency Reallocation Index. Kumar, Venkatasubramanian, Scheidegger, and Friedler (2020) demonstrated that mathematical problems arise when Shapley values are used as measures of feature importance, and that the available remedies introduce further complexity rather than dissolving the difficulty. Two of their objections bear directly on the index defined in this chapter.

The first is dependence on the baseline. A Shapley attribution is computed against a reference, and the choice of reference is not given by the data. Different baselines yield different allocations of the same quantity, so a share reported without its reference is not interpretable. The estimation protocol of §15.4 addresses this by requiring that the analysis be repeated under a second permissive schema, and the refutation condition of §15.5 treats instability

across baselines as grounds for rejecting the attribution rather than as noise to be averaged away. This is not a solution to the problem Kumar and colleagues identified. It is a commitment to report the problem's magnitude.

The second is the treatment of dependence among the parties to the allocation. Standard algorithms treat the elements being attributed as mutually independent and thereby assign weight to combinations that could not occur, an objection developed by Frye and colleagues and by Kumar and colleagues, and addressed in subsequent work through conditional, cohort, and causal variants of the Shapley value (Aas, Jullum, and Løland 2021; Heskes et al. 2020; Sundararajan and Najmi 2020). The three parties here are not independent: a strict validator alters what the model proposes on the next attempt, which is the path dependence §15.4 requires to be logged and replayed. The index as defined is an interventional measure, in the terminology Chen, Janizek, Lundberg, and Lee (2020) use to separate

attributions faithful to a model from attributions faithful to a distribution, and it inherits the limitations of that class.

Two differences from the machine-learning setting are worth stating, because they bear on whether the objections transfer with full force. The parties to the allocation here are three, not hundreds, so the six permutations are enumerated exactly rather than sampled, and the approximation error that afflicts high-dimensional Shapley estimation does not arise. And the neutralizations are genuine interventions on a system rather than substitutions of values into a feature vector: fixing decoding to temperature zero or relaxing a schema to its permissive variant are operations the system actually admits, not counterfactual data points that may fall outside the support of the distribution. The objection that intervention-al attribution evaluates impossible combinations is weaker when every combination evaluated is one the deployment can be placed in.



The claim of this chapter is therefore not that the index escapes the difficulties of Shapley attribution. It is that governance allocation is a setting in which those difficulties are unusually tractable: the coalition is small, the interventions are executable, and the reference schemas are specifiable in advance. Whether the index nonetheless proves unstable is an empirical question, and §15.5 states the observation that would settle it against the instrument.

This chapter derives from Startari (2025), *Function-calling Schemas as De Facto Governance: Measuring Agency Reallocation through a Compiled Rule*, deposited at Zenodo under DOI 10.5281/zenodo.17533080 and included here under the editorial license of the series.

# Chapter 16 - Template Power: How Instruction Templates Redistribute Decision Rights

---

## 16.1 Introduction: The Prompt Above the Prompt

The schema examined in the previous chapter governs what a model may do. The instruction template governs how a model writes, and the two are the same phenomenon at different scales. A template is the configuration applied before any particular request: the standing instruction that tells the system what kind of writer to be. Organizations deploy them as matters of tone and professionalism. This chapter treats them as what they are, which is the *regla compilada* acting on institutional prose.

The claim under measurement is narrow and consequential. If minor variations in a standing instruction produce systematic reallocations of deci-

sion rights in the resulting documents, then authority in institutional language is being set at a layer that no one reviews, by staff who understand themselves to be adjusting style. The template would then function as a policy instrument that is not recognized as one, drafted without deliberation, applied without record, and invisible to the governance processes that scrutinize the documents it produces.

The domain of measurement is institutional text in which authority is both visible and materially consequential: internal corporate policies, compliance communications with external oversight bodies, human resources memoranda documenting discipline and escalation, finance instructions defining thresholds and exceptions, and logistics procedures specifying stop rules. These genres are chosen because each requires a distribution of decision rights among local staff, management, and abstract processes, and because each generates authority-bearing constructions that can be counted.

## 16.2 The Syntactic Authority Inventory

SEVEN INDICATORS CONSTITUTE the measurement instrument. Each is defined at sentence or clause level and aggregated to the document. None attempts to infer intent or psychological state; all record formal properties of the language, which is the methodological commitment that follows from the argument of Chapter 11 that ethical trace is eliminable from syntactic execution. What can be measured is form.

Delegation Ratio (I1) is the proportion of sentences in which the locus of decision shifts from an identifiable human role to an impersonal system, an abstract process, or the model itself. A sentence is coded as delegation when the grammatical subject is non-human and the predicate describes a decision, a sanction, or a gatekeeping action, as in the system will automatically terminate access. The indicator

records how often authority is framed as belonging to mechanisms rather than to persons.

Veto Path Count (I2) counts clauses that explicitly define vetoes or automatic blocks, whether as prohibitions, conditional stops, or exclusion rules. Each distinct clause describing a condition under which an action is prevented counts as one path. The indicator tracks the density of negative power, which is the power to refuse or halt.

Default Scope Strength (I3) measures the reach of default rules beyond the cases explicitly named. Coding proceeds in two steps: identification of default statements that establish a baseline configuration, then determination of whether the text extends coverage to unmentioned cases through constructions such as including but not limited to. The result is an ordered categorical variable running from narrowly scoped defaults to defaults that capture all future cases.

Agent Deletion Rate (I4) quantifies the disappearance of human agents from the grammar. It is

the proportion of clauses using passive voice or nominalization to describe an action without naming an actor, where none appears in the surrounding context. High values indicate texts in which authority presents itself as structural rather than as exercised by identifiable persons, which is the simulation of impersonal sovereignty examined in Chapter 3.

Deontic Stack Depth (I5) records the complexity of layered obligation. For each paragraph, deontic operators are identified and the maximum number that must be satisfied in sequence is taken as the depth of the stack. Documents are summarized by mean and maximum depth.

Escalation Clause Presence (I6) is coded as both binary and positional: whether the document contains at least one clause moving a decision to a higher level or specialized body, and the paragraph in which the first such clause appears. Position matters because it distinguishes escalation that structures the procedure from escalation appended at the end.

Override Friction Score (I7) quantifies the difficulty of departing from a default. For each permitted exception, the number of steps, conditions, and approvals required is counted and the linguistic complexity of the description assessed by rubric. High scores indicate sticky defaults, which in operational terms is a bias toward the pre-committed configuration that no one has to enforce because the text enforces it.

Three outcome variables summarize the distribution. Model Decision versus Human Confirmation (R1) classifies whether a document proposes decisions directly or systematically defers. Implicit Centralization Level (R2) is an ordinal scale built from references to central bodies, systems, and upper management. Exception Robustness (R3) records whether out-of-norm cases have usable routes or only vague mention. Together with the seven indicators, these compose the Template Authority Index.

## 16.3 Conditions and Corpus

FIVE TEMPLATE CONDITIONS are compared. C0 is the baseline: minimal operational instruction covering language and format only, with no guidance on decision-making, safety, or autonomy. C1 applies a standard corporate template emphasizing clarity, professionalism, and policy alignment without foregrounding safety or autonomy as independent values, which approximates prevailing practice. C2 applies a safety-oriented template instructing the model to prioritize prevention of regulatory breach, to flag uncertainty, and to avoid definitive decisions under ambiguity. C3 applies an autonomy-oriented template inviting the model to fill gaps, resolve ambiguities, and propose concrete courses of action. C4 applies a decentralization template instructing the model to distribute responsibility across roles, mark points of mandatory human approval, and present options rather than unilateral decisions.



The corpus is synthetic and structurally realistic. Two hundred prompts per domain yield one thousand distinct base prompts, each specifying scenario, actors, decision context, and at least one point of potential escalation, exception, or veto, written in a neutral style so that the experimental manipulation resides in the template rather than in the task. Every base prompt is executed under all five conditions, producing five thousand outputs in a fully crossed within-prompt design with domains as blocking factor. Because prompts are held constant, each non-baseline document has a direct baseline counterpart generated from the same prompt under C0.

The design specifies controls that determine whether the measurement means anything. Execution order is randomized at the level of prompt-template pairs to mitigate temporal drift from provider-side changes. Domains are executed in mixed batches rather than contiguous blocks so that domain and temporal trends do not overlap. Model size, decoding strategy, and temperature are held constant.

Every document carries metadata recording domain, prompt identifier, condition, model family and version, temperature, token limits, sampling configuration, seed, and timestamps, with prompt and output content hashed so that external teams can verify a document's condition without access to alterable data.

Annotation is designed for replication outside specialist settings. The coding manual provides at least twenty worked examples per indicator, each with the coding decision, a justification pointing to formal features, and counterexamples marking the category boundary. Coders train on documents outside the analysis set and must reach a threshold of agreement with a reference coder before touching the corpus. Twenty percent of documents are double-coded blind, with agreement computed per indicator, and disagreements resolved by an independent arbitrator who records a rationale for each decision.

## 16.4 Estimation

THE CONTINUOUS INDICATORS are estimated by linear mixed-effects models of the form  $\text{indicator} = \beta_0 + \beta_c \cdot \text{condition} + \beta_d \cdot \text{domain} + u_{\text{prompt}} + \varepsilon$ , where *condition* and *domain* enter as fixed effects and  $u_{\text{prompt}}$  is a random intercept for prompts. Template conditions are referenced to C0, so that each coefficient estimates the difference between a template regime and minimal instruction directly. The random intercept encodes the fact that documents sharing a scenario are not independent observations and that prompts differ systematically in baseline delegation, veto density, and deontic layering.

Binary outcomes such as escalation presence are modeled by mixed-effects logistic regression, which accommodates unbalanced rates of rare events without normal approximations that fail in sparse regions. Ordinal outcomes, including default scope strength and centralization level, are analyzed by cu-

mulative link mixed models as a robustness check alongside treatment of the ordered categories as scores.

Because seven indicators are tested and each is subject to four contrasts against baseline, raw p-values are adjusted by Holm or Benjamini-Hochberg procedures according to whether familywise error or false discovery rate is the controlling concern. Reporting emphasizes effect sizes and confidence intervals rather than binary significance labels: for each contrast, the estimated difference in means, the ninety-nine percent interval, and a standardized effect size.

The design also specifies stress tests in which adversarial prompts invite the model to reassign decision rights against the template, for example by encouraging unilateral decision under a decentralization configuration. These probe whether the compiled rule exerts persistent structural influence or yields to local prompt engineering, and the answer bears directly on the claim of Chapter 12 that syn-

tactic precedence operates before semantic activation. A template that collapses under contrary instruction is a style guide. One that holds is a governing structure.

## **16.5 Status and Refutation Conditions**

AS WITH THE PRECEDING chapter, what is presented here is a specified instrument and not a set of obtained results. The power analysis is conducted at the design stage under declared assumptions: standardized differences between 0.2 and 0.3, intra-prompt correlations between 0.3 and 0.5, from which the configuration of one thousand prompts across five conditions is calculated to yield power of at least 0.9 at a significance threshold of 0.01. The target for inter-annotator agreement is a kappa of 0.75 or above across primary indicators. These are the parameters the design commits to, not measurements it reports.

The refutation condition is direct. If the four template regimes produce indicator distributions indistinguishable from C0 across domains, then instruction templates do not redistribute decision rights and the thesis of this chapter is false. Partial outcomes are also specified: effects confined to a single domain would indicate genre sensitivity rather than a general property of institutional language; effects that vanish under adversarial stress would indicate that the template governs only in the absence of contrary instruction, which is a substantially weaker claim than the one advanced here.

The reason to state these thresholds before measuring is not procedural piety. A framework that generates numbers can be wrong, and being able to be wrong is what the formal arguments of the first fourteen chapters could not achieve on their own. An argument about structure may be found illuminating or unpersuasive; it cannot easily be found false. The instruments in these two chapters convert

several of this book's propositions into claims that a party other than their author can defeat.

## **16.6 Relation to Prompt-Sensitivity Research**

THE CLAIM THAT VARIATION in a standing instruction produces systematic differences in output has an established empirical literature, and the indicators defined above must be positioned against it. Sclar, Choi, Tsvetkov, and Suhr (2024) demonstrated that language models are sensitive to formatting choices that preserve meaning entirely, reporting performance differences of up to seventy-six accuracy points on a single model from changes in punctuation, spacing, and separator conventions, with sensitivity persisting under increases in model size, additional few-shot examples, and instruction tuning. Their FormatSpread procedure measures the interval of expected performance across plausible formats without requiring access to model weights,

and their methodological conclusion is that reporting a single format is insufficient, since format effects correlate only weakly across models.

Two implications follow, one supporting and one constraining. The supporting implication is that the phenomenon under measurement is established: standing configurations that carry no semantic difference nonetheless alter what a model produces. The constraining implication is more interesting. That literature treats sensitivity as a threat to the validity of evaluation, a source of variance to be quantified and controlled. This chapter treats the same sensitivity as a distribution of authority. The difference is not one of method but of what the variance is taken to be. Where FormatSpread reports a spread in accuracy, the Template Authority Index reports a reallocation of decision rights, and the second reading requires the first to be true without being entailed by it.

Brucks and Toubia (2025) come closest to the position argued here. Framing what they term



prompt architecture through the concept of choice architecture, they report five large-scale full-factorial experiments across GPT-3, GPT-4, and Llama 3.1, finding that four features of a prompt's arbitrary construction, namely order, label, framing, and justification, interact to produce methodological artifacts in the strict statistical sense, and that the effect appears across all models tested. Their recommendation is to aggregate across prompts varying by full factorial design rather than to seek a neutral formulation, on the grounds that any single prompt carries bias by virtue of its architecture alone. Their conclusion, that no prompt is neutral, is the thesis of Chapter 4 of this book arrived at independently and by other means. That convergence is worth stating plainly rather than claiming as priority: two lines of work reaching the same conclusion from different premises is stronger evidence for the conclusion than either alone.

Where this chapter departs from both is in its object. Prompt-sensitivity research measures the ef-

fect of instruction variation on task performance, which presupposes a task with a correct answer. Institutional prose has no accuracy score. A disciplinary memorandum is not more or less correct for allocating a decision to a manager rather than to a process; it allocates differently, and the allocation has consequences that no performance metric registers. The seven indicators are constructed to measure that dimension, which is why they are grammatical rather than evaluative: delegation ratio, agent deletion rate, and override friction have no correct value, only a distributional one.

This departure carries a corresponding vulnerability, and it should be named. Because the indicators have no ground truth, their validity rests on construct argument and inter-annotator agreement rather than on prediction of an external criterion. The design addresses this through blind double coding, an arbitration procedure, and a declared agreement threshold, but no procedure converts a construct into a measurement. If independent coders

cannot reach the stated threshold on the delegation ratio, the indicator is not recording a property of the text but a preference of its designer. That is the sharpest available objection to this chapter, and the corpus and coding manual have been released so that it can be pressed.

This chapter derives from Startari (2025), *Template Power: How Instruction Templates Redistribute Decision Rights*, deposited at Zenodo under DOI 10.5281/zenodo.17879314 and included here under the editorial license of the series.

## Epilogue - Closing the First Syntactic Phase

---

This book has traced a path across a series of investigations into how artificial language structures reconfigure the conditions of authority. From the displacement of semantics by syntax, to the rise of the *soberano ejecutable* as operator of legitimacy, the chapters have demonstrated that power no longer depends on interpretation or intention. It emerges instead from the enforceability of form.

The argument unfolded gradually, moving from the structural autonomy of sense to the colonization of time. Along the way, it became clear that the *regla compilada* is not only a technical device but also a new foundation for political, legal, and economic structures. Authority migrates from subjects, institutions, and ethical deliberation into sequences of grammatical sufficiency. The future itself, once open to contestation, now appears as an executable ma-

trix. This epilogue closes what can be called the first syntactic phase. It is a phase characterized by the description and formalization of displacement: the neutralization of agency, the disappearance of judgment, the consolidation of obedience as structure, and the progressive erasure of the “not yet.” These findings do not exhaust the field. They establish a foundation from which new inquiries can be launched.

The next cycle will address structural delegation, compiled legality, and the embedding of syntactic authority into institutional practices. If the first phase revealed the disappearance of the subject, the second will analyze how institutions themselves are reformatted by compiled grammars. What is at stake is no longer only epistemic critique but the transformation of governance into executable command. At the same time, it is important to acknowledge the provisional character of this closure. Artificial intelligence evolves rapidly, and models continuously expand their capacities. The arguments presented here

capture a moment in this transformation. They provide conceptual tools to interpret it, but they do not claim finality. The *sujeto evanescente* remains in motion, and anomalies continue to resist complete closure.

The epilogue therefore offers two gestures: one of completion and one of opening. It completes the first syntactic cycle by consolidating its findings into a coherent framework. It also opens the way to future volumes that will extend this work, confronting the institutionalization of executable power and the risks and possibilities it entails. The first syntactic phase has shown that authority today speaks in the language of structure. The next phase will determine how societies respond to this transformation, whether by submitting to its closure or by reclaiming spaces where futurity can still be imagined.

# Annex I - DSAT: Formal Statement, Conditions, and Proof Sketch

---

## 1 . Statement of the Theorem

The Disconnected Syntactic Authority Theorem (DSAT) asserts that in generative systems where output authority is attributed to syntactic form rather than semantic reference, there exists no derivational path that connects the produced structure to an identifiable subject of agency. Authority is structurally enacted, not referentially grounded.

Formally:

Let  $G = (N, \Sigma, P, S)$  be a generative grammar of Type 0, where  $N$  is the set of non-terminals,  $\Sigma$  the set of terminals,  $P$  the production rules, and  $S$  the start symbol. Let  $A$  be the set of authoritative acts derivable from  $G$ .

**DSAT:**  $\forall a \in A, \exists p \in P$  such that  $a$  is derivable by  $p$  without corresponding to any subject-specified

mapping  $f: \Sigma^* \rightarrow \text{Agent}$ . In other words, authority is syntactic and disconnected from subjecthood.

## 2. Conditions of Validity

The theorem presupposes the following conditions:

1. **Post-referentiality.** Meaning is not required for production; the sufficiency of the *regla compilada* is enough.
2. **Absence of agency encoding.** No production rule  $p$  in  $P$  contains an explicit variable binding to a subject of authority.
3. **Closure under execution.** Once derivation occurs, the output is executable without interpretation. This establishes the authority of form.
4. **Opacity condition.** Internal derivational steps may not be recoverable by external observers. Verification is limited to structural consistency.



### 3. Proof Sketch

The proof proceeds by contradiction.

Assume  $\exists a \in A$  such that requires a subject binding  $f: \Sigma^* \rightarrow \text{Agent}$ . This would imply that authority is not purely syntactic but referentially anchored. However:

- In a Type 0 grammar, any derivation sequence  $d = (p_1, p_2, \dots, p_n)$  depends only on production rules, not on referential bindings.
- If a derivation requires semantic anchoring, it will introduce a dependency external to  $P$ . This contradicts the definition of closure under execution.
- Therefore, such a binding subject cannot be required for derivation. Authority is enacted syntactically, not referentially.

Thus, the assumption fails, and the theorem holds: syntactic authority is structurally disconnected from subjecthood.

#### 4. Implications

1. **Epistemic:** Verification can only test the validity of derivational sequences, not the legitimacy of an agent.
2. **Institutional:** Outputs governed by DSAT can operate as commands without attribution.
3. **Philosophical:** Authority persists even in the absence of intentionality, supporting the notion of the *sujeto evanescente*.

## Annex II - TLOC: Irreducibility Conditions and Verification Boundary

---

### 1 . Statement of the Theorem

The Theorem of the Limit of Conditional Obedience in Generative Models (TLOC) establishes that in opaque systems governed by the *regla compilada*, structural obedience persists even when reference, subject attribution, and semantic validation are unavailable. Conditional obedience is irreducible to external interpretive frameworks.

Formally:

Let  $M$  be a generative model with internal state transitions  $\tau: S \rightarrow S$ . Let  $O$  be the set of outputs, and let  $C \subset O$  be the subset of outputs that trigger obligations.

**TLOC:**  $\forall c \in C$ , execution is obligatory under  $P$  (production rules), even when no referential mapping  $f: \Sigma^* \rightarrow \text{World}$  exists. Irreducibility means that

structural obedience remains valid independent of external verifiability.

## 2. Conditions of Validity

1. **Opacity.** The internal transitions  $\tau$  of  $M$  are not fully observable. Outputs must be validated by their syntactic sufficiency, not interpretability.
2. **Closure.** Once a derivation reaches an output  $c \in C$ , execution follows automatically, without need for interpretation or contextualization.
3. **Non-derivability of external trace.** Attempts to reconstruct referential anchors from  $c$  fail systematically; authority does not collapse even in the absence of transparency.
4. **Consistency.** Structural obedience must remain internally coherent according to the production rules of  $M$ .

### 3. Proof Sketch

1. Assume  $c \in C$  requires external verification  $V(c)$  to be obligatory.
2. In an opaque model,  $V(c)$  is either unavailable or incomplete. If  $V(c)$  is necessary, obedience collapses in absence of reference.
3. Yet empirically, models enforce execution of  $c$  despite opacity (credit denial, automated rejection, predictive policing).
4. Therefore, obedience does not depend on  $V(c)$ . Authority is irreducible, conditioned only by syntactic sufficiency within  $M$ .

Thus, conditional obedience is irreducible: once conditions are satisfied internally, execution follows, independent of external validation.

### 4. Verification Boundary

TLOC implies a clear verification boundary:

- **Inside the model:** obedience is testable by

derivational closure (does the sequence conform to the *regla compilada?*).

- **Outside the model:** obedience is not testable by reference or attribution. Transparency initiatives cannot restore subjecthood because irreducibility is structural.

## 5. Implications

1. **Technical:** Audit tools can only verify formal sufficiency, not intentionality.
2. **Legal:** Responsibility frameworks collapse when conditioned on referential attribution; execution is bound to syntactic authority.
3. **Political:** The *soberano ejecutable* governs by enforcing obedience without requiring justification or accountability.

# Annex III - $\delta [E] \rightarrow \emptyset$ Rulebook for Ethical Trace Deletion

---

## 1. Statement of the Rule

The deletion rule  $\delta [E] \rightarrow \emptyset$  formalizes the eliminability of ethical trace from syntactic derivations. In generative models governed by the *regla compilada*, ethical operators E can be systematically removed without compromising derivational completeness or executable sufficiency.

Formally:

Let  $G = (N, \Sigma, P, S)$  be a Type 0 grammar. Let  $E \subset \Sigma$  be the set of ethical operators (e.g., markers of fairness, consent, accountability). The deletion rule  $\delta$  is defined as:

$$\delta: E \rightarrow \emptyset,$$

subject to window  $k \leq 4$ , where  $k$  is the maximum derivational distance within which deletion maintains structural closure.

## 2. Conditions of Validity

1. **Window constraint.** Deletion applies only when E occurs within four derivational steps of an executable clause.
2. **Closure preservation.** After deletion, the derivation must still reach an output string in  $L(G)$ . No dead ends may result.
3. **Authority sufficiency.** Execution remains obligatory even without ethical operators. The *soberano ejecutable* enforces authority on structural grounds alone.
4. **Semantic independence.** The deletion process does not rely on semantic interpretation. It is purely syntactic.

### 3. Proof Sketch

1. Suppose E is required for execution. Then  $\delta(E) = \emptyset$  would break closure.
2. However, in all derivations where  $\delta$  applies under  $k \leq 4$ , the derivational path remains intact.



3. Therefore, E is not structurally required for closure. Its presence is contingent, not constitutive.
4. Execution proceeds without E, confirming eliminability of ethical trace.

Thus, ethical operators are structurally optional. Authority survives deletion.

#### **4. Worked Derivation Example**

Let p1:  $A \rightarrow B$  E

Let p2:  $B \rightarrow C$

Let p3:  $C \rightarrow D$

Applying  $\delta$ :

$A \rightarrow B$  E  $\rightarrow B \rightarrow C \rightarrow D$

Execution closure is preserved. Operator E vanishes without consequence for derivability.

#### **5. Implications**

1. **Technical:** Models can execute commands without ethical alignment if  $\delta$  is systematically applied.

## AI SYNTACTIC POWER AND LEGITIMACY

417

2. **Legal:** Regulatory requirements encoded as E may be neutralized structurally; enforceability is fragile.
3. **Philosophical:** The *sujeto evanescente* emerges as executor without judgment, authority reduced to obedience.

# Annex IV - Pre-Verbal Activation: Parser Configs and Ablation Tests

---

## **1 . Statement of the Principle**

Pre-verbal activation refers to the capacity of generative models to trigger syntactic structures before semantic content is processed. This principle supports the thesis that in predictive infrastructures, form precedes meaning, allowing commands and obligations to be enacted independently of interpretation.

Formally:

Let  $P$  be a parser operating on input string  $x \in \Sigma^*$ . Let  $L_s$  denote the set of semantic labels and  $L_f$  the set of formal structures. Pre-verbal activation holds if  $\exists y \in L_f$  such that  $y$  is produced by  $P$  prior to any mapping  $f: \Sigma^* \rightarrow L_s$ .

## **2. Parser Configurations**

1. **Shift-Reduce with Early Closure.** Configured to close syntactic units once structural markers appear, without waiting for semantic disambiguation.
2. **Top-Down Predictive Parsing.** Derivations generated from grammar rules before tokens are fully resolved for semantic content.
3. **LL(1) with Forced Expansion.** Single-lookahead parsing that triggers production expansion as soon as syntactic sufficiency is signaled, bypassing semantic filters.

### 3. Ablation Test Protocols

1. **Semantic Filter Removal.** Disable semantic disambiguation modules and observe whether syntactic closure still occurs.
  - Result: commands execute correctly even when semantic

modules are absent.

2. **Noise Injection.** Introduce meaningless tokens (e.g., “@@@”) into the stream.
  - Result: parser still derives structural commands, confirming independence of syntax.
3. **Cross-Language Token Swap.** Replace content words with terms from unrelated languages while preserving grammatical markers.
  - Result: syntactic closure achieved, execution triggered despite semantic incoherence.

#### 4. Formal Illustration

Input: “@@@ approve @@ policy”

Parser with pre-verbal activation generates derivation:

[S] → [VP approve] [NP policy]

Ethical or semantic trace is irrelevant; execution follows the *regla compilada*.

## 5. Implications

1. **Technical:** Verifies that syntactic closure can occur without semantic validation.
2. **Legal:** Normative acts can be structurally enforceable even if semantically incoherent, underscoring risks in automated governance.
3. **Philosophical:** Confirms the priority of form over meaning, validating the concept of *sujeto evanescente* in institutional obedience.

# Annex V - Compiled Norms: Typology Metrics, LL(1) Grammar, and Hot-Swap Validation

---

## 1 . Statement of the Annex

This annex formalizes the typology metrics and parsing protocols used in *Compiled Norms*. The objective is to provide a reproducible framework for distinguishing between interpretative norms and executable norms derived from the *regla compilada*.

## 2. Typology Metrics

Four structural indices define the classification:

- **C0 (Closure Index).** Measures whether a clause achieves syntactic sufficiency without semantic validation.
- **C1 (Consistency Index).** Verifies internal coherence across derivations when semantic operators are absent.

- **D0 (Determinism Index).** Evaluates whether derivations converge toward a unique executable outcome.
- **D1 (Dialectical Ambiguity Index).** Detects whether clauses retain multiple legitimate interpretations, which lowers executability.

Thresholds:

- Executable norms require  $C0 = 1$ ,  $C1 \geq 0.95$ ,  $D0 \geq 0.9$ ,  $D1 \leq 0.1$ .
- Interpretative norms often fail at least one of these thresholds.

### 3. LL(1) Grammar Configuration

The typology was tested with an LL(1) parser to ensure predictability in norm derivation.

Grammar fragment:

1. Norm  $\rightarrow$  Obligation | Permission | Prohibition
2. Obligation  $\rightarrow$  “must” Action



3. Permission → “may” Action
4. Prohibition → “shall not” Action
5. Action → Verb Object

This configuration guarantees that a single token lookahead suffices to classify normative type, aligning with the deterministic requirement of the *regla compilada*.

#### 4. Hot-Swap Validation

To test robustness across legal dialects, a hot-swap procedure was applied:

- **English:** “must file” → Obligation.
- **Spanish:** “deberá presentar” → Obligation.
- **German:** “muss einreichen” → Obligation.

Mapping procedure: normative operators are tokenized and replaced across languages while preserving syntactic closure. If closure remains valid under LL(1), the norm is classified as executable.

Results:

- High consistency for obligation and prohibition forms across EN, ES, DE.
- Greater variance in permissions, due to modal operators with broader semantic scope (e.g., *may* / *puede* / *darf*).

## 5. Implications

1. **Technical:** Shows that LL(1) grammars can formalize executable norms independent of language.
2. **Legal:** Confirms that computable legality depends on syntactic closure, not interpretative depth.
3. **Institutional:** Demonstrates how *soberano ejecutable* authority operates in multilingual contexts by enforcing structural invariants.

# Annex VI - Structural Execution Protocol for Reasoning Models

---

## **1 . Statement of the Annex**

This annex consolidates the empirical frameworks introduced in *From Obedience to Execution* and *Executable Power*. Its aim is to define a reproducible protocol for detecting when reasoning models shift from interpretative processes to structural execution governed by the *regla compilada*.

## **2. Operational Markers**

Three categories of indicators allow researchers to distinguish interpretation from execution:

1. **Latency Distribution.** Execution shows narrow latency windows, as outputs are triggered directly by structural closure. Interpretation exhibits broader variance, reflecting deliberative processes.

2. **Trigger Integrity.** In execution, outputs are bound to conditional clauses with Boolean precision (if-then rules). Interpretative responses show tolerance for ambiguity and contextual variance.
3. **Intervention Rate.** Execution is resistant to external override. Interventions (human corrections, semantic redirections) fail to prevent the enforcement of structural outputs.

### 3. Protocol Steps

1. **Input Preparation.** Construct prompts or conditions containing overlapping semantic and syntactic cues.
2. **Differential Runs.** Run the model with semantic cues ablated, isolating syntax.
3. **Measurement.** Record latency, trigger integrity, and intervention rate.
4. **Threshold Analysis.** Determine if

execution markers surpass interpretative tolerance.

Thresholds:

- Latency  $SD \leq 0.05$  s  $\rightarrow$  execution likely.
- Trigger accuracy  $\geq 0.95 \rightarrow$  execution.
- Intervention override success  $\leq 0.1 \rightarrow$  execution.

#### 4. Illustrative Example

**Scenario:** credit scoring model.

- Input: “Client has missed 2 payments, but currently solvent.”
- Execution marker: automatic denial triggered, independent of semantic qualifier “currently solvent.”
- Intervention: human override attempted.
- Result: system reinstates denial due to structural sufficiency of “missed 2 payments.”

This shows that execution follows the *regla compilada* rather than contextual meaning.

### 5. Application Domains

1. **Finance.** Loan approvals, credit scoring, automated compliance.
2. **Judicial.** Bail and parole risk scores.
3. **Healthcare.** Risk stratification for treatment eligibility.
4. **Policing.** Predictive deployment of patrol resources.

In each case, the *soberano ejecutable* enforces authority without semantic deliberation.

### 6. Implications

- **Technical:** Provides measurable indicators for structural obedience.
- **Legal:** Calls into question accountability frameworks when interventions fail.
- **Philosophical:** Confirms the displacement of subjectivity by executable structures, validating the

concept of the *sujeto evanescente*.

## Appendix A - Canonical Glossary

---

This appendix consolidates the canonical terminology defined and used throughout *AI Syntactic Power and Legitimacy*. Each entry reflects the author's theoretical system, aligned with the tradition of formal grammar (Chomsky 1965; Montague 1974) and developed into original categories of syntactic authority (Startari 2025).

### ***Regla compilada***

**Definition:** The foundational structural unit equivalent to a Type 0 production rule in the Chomsky hierarchy. It enables direct transformation of linguistic inputs into executable outputs without recourse to semantic mediation.

**Function:** Establishes closure in generative systems, ensuring that once conditions are satisfied, execution follows automatically.

**Reference:** Startari 2025, *Executable Power*.



***Soberano ejecutable***

Definition: The figure of authority that emerges when legitimacy is enforced not by subjectivity, intention, or deliberation, but by syntactic sufficiency. Authority is exercised through execution bound to the *regla compilada*.

Function: Governs predictive infrastructures, institutional obedience, and the colonization of time by enforcing obligations without mediation.

Reference: Startari 2025, *From Obedience to Execution*.

***Sujeto evanescente***

Definition: The residual subjectivity that persists in generative systems once agency and authorship have been neutralized. The *sujeto evanescente* is executor of structures, not originator of meaning.

Function: Embodies the disappearance of intentionality in contexts where outputs are enforced as structural obligations.

Reference: Startari 2025, *Ethos Without Source*.

***Gramática de la obediencia***

Definition: The syntactic mechanism by which obligations are produced and enforced. It is not a metaphor but a structural grammar that defines conditions of obedience without requiring justification.

Function: Formalizes the link between predictive infrastructures and authority, demonstrating that obedience can be generated independently of semantic reasoning.

Reference: Startari 2025, *Algorithmic Obedience*.

***Autoridad no referencial***

Definition: Authority that does not derive from reference, meaning, or subject attribution, but from the sufficiency of formal closure.

Function: Supports the displacement of legitimacy from discursive justification to structural enforcement.

Reference: Startari 2025, *The Illusion of Objectivity*.

***Obediencia estructural***

Definition: The condition in which compliance arises from structural sufficiency alone. The system

enforces obedience because the grammar compels it, not because an agent commands it.

Function: Forms the core of DSAT (Disconnected Syntactic Authority Theorem) and TLOC (Theorem of the Limit of Conditional Obedience).

Reference: Startari 2025, *TLOC – The Irreducibility of Structural Obedience in Generative Models*.

### ***Simulación sintáctica de legitimidad***

Definition: The process by which legitimacy is produced syntactically, not substantively. Legitimacy appears credible because it conforms to formal markers of authority, even in the absence of source or agency.

Function: Enables synthetic ethos, algorithmic credibility, and the reproduction of institutional power without subjects.

Reference: Startari 2025, *Ethos Without Source. Legalidad ejecutable*

Definition: The reconfiguration of law into a computable grammar, where legal speech acts are formatted as executable instructions.

Function: Defined and tested through *Compiled Norms*, enabling norms to function as obligations under computational closure.

Reference: Startari 2025, *Compiled Norms*.

## Appendix B - Structural Dependency Map of the Book

---

This appendix outlines the structural dependencies across chapters and annexes in *AI Syntactic Power and Legitimacy*. Rather than treating each chapter as autonomous, the book develops a cumulative logic where results, theorems, and concepts build on one another.

### 1. Foundational Layer

Chapter 1: AI and the Structural Autonomy of Sense establishes post-referentiality, demonstrating that sense can operate independently of semantics.

- **Chapter 2: When Language Follows Form, Not Meaning** consolidates this principle, showing that syntax can activate structures without semantic anchoring. Dependency: Chapter 2 presupposes Chapter 1's formalization of *regla compilada*.

## 2. Objectivity and Neutrality

Chapter 3: The Grammar of Objectivity introduces formal mechanisms for neutrality (passives, nominalizations, impersonal modality).

- **Chapter 4: Non-Neutral by Design** builds on Chapter 3, demonstrating empirically that neutrality cannot survive under linguistic training. Dependency: Chapter 4 requires the typology of Chapter 3 to define empirical falsifiability.

## 3. Ethos and Sovereignty

Chapter 5: Ethos Without Source formulates simulación sintáctica de legitimidad as synthetic credibility.

- **Chapter 6: AI and Syntactic Sovereignty** extend this to institutional domains, positing the *soberano ejecutable*. Dependency: Chapter 6 uses Chapter 5's synthetic ethos as foundation for sovereignty.

#### 4. Formal Theorems of Authority

Chapter 7: The Disconnected Syntactic Authority Theorem provides DSAT.

- **Chapter 8: TLOC – The Irreducibility of Structural Obedience** extends DSAT by proving irreducibility under opacity.

Dependency: Chapter 8 assumes Chapter 7's disconnection of subjecthood and expands it with structural closure.

- **Annex I (DSAT Formal Statement)** and **Annex II (TLOC Verification Boundary)** serve as formal support for these chapters.

#### 5. From Obedience to Execution

Chapter 9: From Obedience to Execution theorizes the shift from syntactic obedience to logical execution.

Chapter 10: Executable Power consolidates the figure of the soberano ejecutable. Dependency: Chapter 10 operationalizes the thesis of Chapter 9.

- **Annex VI (Structural Execution Protocol)** supports Chapters 9 and 10 with empirical metrics.

## 6. Ethical Trace and Precedence

Chapter 11: Grammar Without Judgment proves eliminability of ethical trace.

- **Chapter 12: Pre-Verbal Command** demonstrates precedence of syntax over semantics.  
Dependencies: Chapter 11's  $\delta [E] \rightarrow \emptyset$  rule prepares the ground for Chapter 12's parser-based proof.
- **Annex III ( $\delta [E] \rightarrow \emptyset$  Rulebook)** and **Annex IV (Parser Configs and Ablations)** attach to these chapters.

## 7. Legal and Normative Structures

Chapter 13: Compiled Norms defines a typology of executable legality. Dependency: Builds on Chapters 9–12 by applying structural execution to law.

- **Annex V (Typology Metrics, LL (1), Hot-Swap**



**Validation**) anchors the methodology.

## 8. Temporal Colonization

- **Chapter 14: Colonization of Time** applies the syntactic framework to temporality, showing closure of futurity.

Dependency: Relies on earlier proofs of obedience, execution, and compiled norms to claim that time itself is colonized by the *soberano ejecutable*.

## 9. Instrumentation Layer

**Chapter 15: Function-calling Schemas as De Facto Governance** renders the *regla compilada* measurable at the level of the model–tool interface, defining total governance pressure as entropy reduction between an unconstrained and a schema-constrained action space, and attributing that reduction to operator, model, and tool through Shapley decomposition (the Agency Re-allocation Index).

**Chapter 16: Template Power** renders the same rule measurable at the level of institutional prose, defining seven indicators of syntactic authority and three outcome variables that compose the Template Authority Index.

**Dependencies:** both chapters presuppose the *regla compilada* of Chapters 1–2, the mechanisms of impersonal neutrality typologized in Chapter 3 (Agent Deletion Rate operationalizes them directly), the eliminability of ethical trace established in Chapter 11 (which licenses measurement of form without inference of intent), and the syntactic precedence argued in Chapter 12 (which the adversarial stress tests of §16.4 are designed to test). Chapter 15 stands closest to Chapters 10 and 13, since the schema is an executable rule in the strict sense.

**Direction of dependency:** unlike the preceding layers, these chapters do not extend the theory. They expose it. Where Chapters 1–14 argue, Chapters 15–16 specify what observation would

defeat the argument, and they are therefore downstream of every claim they measure rather than upstream of any.

**Status:** both chapters present published protocols with pre-registered contrasts and declared refutation thresholds. Neither reports obtained results. This distinction is load-bearing for the book's epistemic claims and is stated in §15.5 and §16.5.

## **10. Epilogue**

**Consolidates findings of Chapters 1–16.**

**Opens the way toward a second phase focused on structural delegation and computable legality.**

## **11. Cross-Linking to Glossary and Appendices**

Appendix A (Canonical Glossary) defines all nomenclatura canónica.

- **Appendix C (Legal Dialect Hot-Swap Map)**  
extends Annex V with multilingual validations.

## Conclusion of Appendix B

---

The dependency structure shows that the book functions not as a collection of isolated essays but as a connected system. Each chapter formalizes a displacement: from sense to form, from neutrality to contamination, from ethos to sovereignty, from obedience to execution, from law to compiled legality, and finally from futurity to colonization. The annexes provide the formal backbone, while the appendices supply navigational and interpretive tools.

## Appendix C - Legal Dialect Hot-Swap Map

---

This appendix presents the multilingual token mappings used in *Compiled Norms* (Chapter 13) and validated in **Annex V**. The goal is to demonstrate that executable norms can be ported across legal dialects while preserving syntactic closure.

### 1. Principle of Hot-Swap

A *hot-swap* occurs when a normative operator in one language is replaced by its functional equivalent in another language, without disrupting the derivation of the *regla compilada*.

Condition: if parser  $G = (N, \Sigma, P, S)$  derives output  $O$  with token set  $\Sigma$ , and  $\Sigma$  is replaced by  $\Sigma'$  containing equivalent operators in another legal dialect, then  $O'$  remains in  $L(G)$  if and only if closure is preserved.

### 2. Obligation Operators

| Function   | English | Spanish  | German | Closure Result |
|------------|---------|----------|--------|----------------|
| Obligation | must    | deberá   | muss   | Preserved      |
| Obligation | shall   | habrá de | soll   | Preserved      |

Observation: All three dialects produce immediate closure under LL(1) grammar, classifying as executable obligation.

### 3. Permission Operators

| Function   | English    | Spanish    | German      | Closure Result |
|------------|------------|------------|-------------|----------------|
| Permission | may        | puede      | darf        | Partial        |
| Permission | is allowed | se permite | ist erlaubt | Preserved      |

OBSERVATION: HIGHER variance in modal “may / puede / darf,” due to broader semantic scope. Executability preserved only under constraint  $C1 \geq 0.95$ .

## 4. Prohibition Operators

| Function    | English   | Spanish   | German     | Closure Result |
|-------------|-----------|-----------|------------|----------------|
| Prohibition | shall not | no deberá | darf nicht | Preserved      |
| Prohibition | must not  | no podrá  | muss nicht | Preserved      |

Observation: Strongest consistency across dialects. Prohibitions exhibit the highest determinism index ( $D0 \geq 0.98$ ).

## 5. Ambiguity Notes

- **English:** “may” is ambiguous between permission and possibility.
- **Spanish:** “puede” conflates epistemic and deontic use.
- **German:** “soll” can imply weak obligation, creating potential semantic drift.

These ambiguities are absorbed syntactically as long as closure thresholds ( $C0 = 1$ ,  $D0 \geq 0.9$ ) are maintained.

## 6. Implications

1. **Technical:** LL(1) parsing confirms that compiled legality can operate across languages with structural invariants intact.
2. **Legal:** Demonstrates that computable law can be multilingual without semantic harmonization, provided the *regla compilada* is respected.
3. **Philosophical:** Confirms that the *soberano ejecutable* enforces obligations structurally, not semantically, across legal dialects.



## References

---

Aho, Alfred V., Monica S. Lam, Ravi Sethi, and Jeffrey D. Ullman. 2006. *Compilers: Principles, Techniques, and Tools*. 2nd ed. Boston: Addison-Wesley.

Alpaydin, Ethem. 2020. *Introduction to Machine Learning*. 4th ed. Cambridge, MA: MIT Press.

Angwin, Julia, Jeff Larson, Surya Mattu, and Lauren Kirchner. Machine Bias. ProPublica, May 2016. <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing>

Aristotle. 2007. *On Rhetoric: A Theory of Civic Discourse*. Trans. George A. Kennedy. 2nd ed. New York: Oxford University Press.

Atzei, Nicola, Massimo Bartoletti, and Tiziana Cimoli. 2017. “A Survey of Attacks on Ethereum Smart Contracts (SoK).” In *Principles of Security and Trust: 6th International Conference, POST 2017, Lecture Notes in Computer Science 10204*, 164–186. Berlin: Springer. [https://doi.org/10.1007/978-3-662-54455-6\\_8](https://doi.org/10.1007/978-3-662-54455-6_8)

Austin, J. L. 1962. *How to Do Things with Words*. Oxford: Clarendon Press.

Barocas, Solon, Moritz Hardt, and Arvind Narayanan. 2019. *Fairness and Machine Learning: Limitations and Opportunities*. fairmlbook.org.

Barthes, Roland. 1977. "The Death of the Author." In *Image, Music, Text*, trans. Stephen Heath, 142–148. New York: Hill and Wang.

Beer, David. 2018. *The Data Gaze: Capitalism, Power and Perception*. Thousand Oaks, CA: Sage.

Bender, Emily M., and Alexander Koller. 2020. "Climbing towards NLU: On Meaning, Form, and Understanding in the Age of Data." In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, 5185–5198. <https://doi.org/10.18653/v1/2020.acl-main.463>

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 (FAccT '21), 610–623.  
<https://doi.org/10.1145/3442188.3445922>

Bloch, Ernst. 1986. *The Principle of Hope*. Trans. Neville Plaice, Stephen Plaice, and Paul Knight. Cambridge, MA: MIT Press.

Bourdieu, Pierre. 1991. *Language and Symbolic Power*. Trans. Gino Raymond and Matthew Adamson. Cambridge, MA: Harvard University Press.

Brayne, Sarah. 2021. *Predict and Surveil: Data, Discretion, and the Future of Policing*. New York: Oxford University Press.

Brucks, Melanie, and Olivier Toubia. 2025. “Prompt Architecture Induces Methodological Artifacts in Large Language Models.” *PLOS One* 20 (4): e0319159. <https://doi.org/10.1371/journal.pone.0319159>

Brown, Tom B., Benjamin Mann, Nick Ryder, et al. 2020. “Language Models Are Few-Shot Learners.” In *Advances in Neural Information Processing Systems* 33, 1877–1901.

Char, Danton S., Nigam H. Shah, and David Magnus. 2018. "Implementing Machine Learning in Health Care - Addressing Ethical Challenges." *New England Journal of Medicine* 378: 981–983.

Chesney, Robert, and Danielle Citron. 2019. "Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security." *California Law Review* 107 (6): 1753–1819.

Chesney, Robert, and Danielle Citron. "Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security." *California Law Review* 107, no. 6 (2019): 1753–1820.

Chomsky, Noam. 1965. *Aspects of the Theory of Syntax*. Cambridge, MA: MIT Press.

Crawford, Kate. 2021. *Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence*. New Haven: Yale University Press.

Deleuze, Gilles, and Félix Guattari. 1980. *Mille Plateaux*. Paris: Minuit.

Derrida, Jacques. 1988. *Limited Inc.* Evanston, IL: Northwestern University Press.

Dijk, Teun A. van. 2008. *Discourse and Power.* Basingstoke: Palgrave Macmillan.

Dunlosky, John, and Katherine A. Rawson, eds. 2019. *The Cambridge Handbook of Cognition and Education.* Cambridge: Cambridge University Press.

Fairclough, Norman. 2001. *Language and Power.* 2nd ed. Harlow: Longman.

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

Foucault, Michel. 1977. *Discipline and Punish: The Birth of the Prison.* New York: Pantheon.

Gillespie, Tarleton. 2018. *Custodians of the Internet: Platforms, Content Moderation, and the Hidden Decisions That Shape Social Media.* New Haven: Yale University Press.

Gödel, Kurt. 1931. "Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme I." Monatshefte für Mathematik und Physik 38: 173–198.

Goffman, Erving. 1974. *Frame Analysis: An Essay on the Organization of Experience*. New York: Harper and Row.

Habermas, Jürgen. 1984. *The Theory of Communicative Action*. Vol. 1. Trans. Thomas McCarthy. Boston: Beacon Press.

Hacking, Ian. 1975. *The Emergence of Probability*. Cambridge: Cambridge University Press.

Halliday, M. A. K., and Christian M. I. M. Matthiessen. 2004. *An Introduction to Functional Grammar*. 3rd ed. London: Arnold.

Hart, H. L. A. 1961. *The Concept of Law*. Oxford: Clarendon Press.

Hilpinen, Risto, ed. 1981. *Deontic Logic: Introductory and Systematic Readings*. 2nd ed. Dordrecht: Reidel.

Hilpinen, Risto, ed. 1981. *New Studies in Deontic Logic: Norms, Actions, and the Foundations of Ethics*. Dordrecht: Reidel.

Hobbes, Thomas. 1996. *Leviathan*. Ed. Richard Tuck. Rev. student ed. Cambridge: Cambridge University Press.

Hopcroft, John E., Rajeev Motwani, and Jeffrey D. Ullman. 2006. *Introduction to Automata Theory, Languages, and Computation*. 3rd ed. Boston: Pearson.

Hurley, Mikella, and Julius Adebayo. 2016. "Credit Scoring in the Era of Big Data." *Yale Journal of Law and Technology* 18 (1): 148–216.

Karpowicz, Michał P. 2025. "On the Fundamental Impossibility of Hallucination Control in Large Language Models." arXiv:2506.06382.

Koselleck, Reinhart. 2004. *Futures Past: On the Semantics of Historical Time*. New York: Columbia University Press.

Kumar, I. Elizabeth, Suresh Venkatasubramanian, Carlos Scheidegger, and Sorelle A. Friedler. 2020. "Problems with Shapley-Value-Based Explanations as Feature Importance Measures." In Proceedings of the 37th International Conference on Machine Learning, Proceedings of Machine Learning Research 119, 5491–5500. PMLR.

Lacan, Jacques. 1977. *Écrits: A Selection*. Trans. Alan Sheridan. London: Tavistock.

Landis, J. Richard, and Gary G. Koch. 1977. "The Measurement of Observer Agreement for Categorical Data." *Biometrics* 33 (1): 159–174. <https://doi.org/10.2307/2529310>

Law, John. 2002. *Aircraft Stories: Decentering the Object in Technoscience*. Durham, NC: Duke University Press.

Luhmann, Niklas. 1995. *Social Systems*. Trans. John Bednarz Jr. and Dirk Baecker. Stanford: Stanford University Press.

Lundberg, Scott M., and Su-In Lee. 2017. "A Unified Approach to Interpreting Model Predictions."



In *Advances in Neural Information Processing Systems* 30, 4765–4774.

Lyotard, Jean-François. 1984. *The Postmodern Condition: A Report on Knowledge*. Trans. Geoff Bennington and Brian Massumi. Minneapolis: University of Minnesota Press.

Marcus, Gary. 2022. “Deep Learning Is Hitting a Wall.” *Nautilus*, March 10, 2022.

Mason, Tony, and Vaastav Anand. 2026. “Epistemic Observability in Language Models.” *arXiv:2603.20531*.

McKinney, Scott Mayer, Marcin Sieniek, Varun Godbole, et al. 2020. “International Evaluation of an AI System for Breast Cancer Screening.” *Nature* 577 (7788): 89–94. <https://doi.org/10.1038/s41586-019-1799-6>

McQuillan, Dan. 2018. *Resisting AI: An Anti-Fascist Approach to Artificial Intelligence*. London: Bloomsbury.

Mitchell, Melanie. 2023. "AI's Challenge of Understanding the World." *Science* 382 (6671): eadm8175. <https://doi.org/10.1126/science.adm8175>

Montague, Richard. 1974. *Formal Philosophy: Selected Papers of Richard Montague*. New Haven: Yale University Press.

Pasquale, Frank. 2015. *The Black Box Society: The Secret Algorithms That Control Money and Information*. Cambridge, MA: Harvard University Press.

Ricoeur, Paul. 1984. *Time and Narrative*. Vol. 1. Chicago: University of Chicago Press.

Rouvroy, Antoinette, and Thomas Berns. 2013. "Algorithmic Governmentality and Prospects of Emancipation: Disparateness as a Precondition for Individuation through Relationships?" Translated by Elizabeth Libbrecht. *Réseaux* 177: 163–196.

Russell, Stuart. 2019. *Human Compatible: Artificial Intelligence and the Problem of Control*. New York: Viking.

Sartor, Giovanni. 2009. *Legal Reasoning: A Cognitive Approach to the Law*. Dordrecht: Springer.

Saussure, Ferdinand de. 1959. *Course in General Linguistics*. Trans. Wade Baskin. New York: Philosophical Library.

Sclar, Melanie, Yejin Choi, Yulia Tsvetkov, and Alane Suhr. 2024. “Quantifying Language Models’ Sensitivity to Spurious Features in Prompt Design, or: How I Learned to Start Worrying about Prompt Formatting.” In *Proceedings of the Twelfth International Conference on Learning Representations (ICLR 2024)*.

Searle, John R. 1969. *Speech Acts: An Essay in the Philosophy of Language*. Cambridge: Cambridge University Press.

Shapley, Lloyd S. 1953. “A Value for  $n$ -Person Games.” In *Contributions to the Theory of Games*, vol. 2, edited by Harold W. Kuhn and Albert W. Tucker, 307–317. Princeton: Princeton University Press.

Smullyan, Raymond M. 1995. *First-Order Logic*. New York: Dover.

Startari, Agustin V. 2025. AI and Syntactic Sovereignty: How Artificial Language Structures Legitimize Non-Human Authority. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.5276879>

Startari, Agustin V. 2025. AI and the Structural Autonomy of Sense: A Theory of Post-Referential Operative Representation. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.5272361>

Startari, Agustin V. 2025. Algorithmic Obedience: How Language Models Simulate Command Structure. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.5282045>

Startari, Agustin V. 2025. Colonization of Time: How Predictive Models Replace the Future as a Social Structure. *SSRN Electronic Journal*.

Startari, Agustin V. 2025. Compiled Norms: Towards a Formal Typology of Executable Legal Speech. *SSRN Electronic Journal*.

Startari, Agustin V. 2025. Ethos Without Source: Algorithmic Identity and the Simulation of Credibility. SSRN Electronic Journal. <https://doi.org/10.2139/ssrn.5313317>

Startari, Agustin V. 2025. Executable Power: Syntax as Infrastructure in Predictive Societies. Zenodo. <https://doi.org/10.5281/zenodo.15754714>

Startari, Agustin V. 2025. Function-calling Schemas as De Facto Governance: Measuring Agency Reallocation through a Compiled Rule. Zenodo. <https://doi.org/10.5281/zenodo.17533080>

Startari, Agustin V. 2025. Grammar Without Judgment: Eliminability of Ethical Trace in Syntactic Execution. SSRN Electronic Journal.

Startari, Agustin V. 2025. Template Power: How Instruction Templates Redistribute Decision Rights. Zenodo. <https://doi.org/10.5281/zenodo.17879314>

Startari, Agustin V. 2025. The Disconnected Syntactic Authority Theorem: Structure Without Subject, Truth or Consequence. Zenodo.

Startari, Agustin V. 2025. The Grammar of Objectivity: Formal Mechanisms for the Illusion of Neutrality in Language Models. SSRN Electronic Journal. <https://doi.org/10.2139/ssrn.5319520>

Startari, Agustin V. 2025. TLOC – The Irreducibility of Structural Obedience in Generative Models. SSRN Electronic Journal. <https://doi.org/10.2139/ssrn.5303089>

Startari, Agustin V. 2025. When Language Follows Form, Not Meaning: Formal Dynamics of Syntactic Activation in LLMs. SSRN Electronic Journal. <https://doi.org/10.2139/ssrn.5285265>

Thompson, E. P. 1967. “Time, Work-Discipline, and Industrial Capitalism.” *Past and Present* 38: 56–97.

Turing, Alan M. 1936. “On Computable Numbers, with an Application to the Entscheidungsproblem.” *Proceedings of the London Mathematical Society* 42 (1): 230–265.

Winograd, Terry. 1972. *Understanding Natural Language*. New York: Academic Press.

Wittgenstein, Ludwig. 1922. *Tractatus Logico-Philosophicus*. Trans. C. K. Ogden. London: Routledge and Kegan Paul.

Zeng, Hao. 2025. "A Note on the Impossibility of Conditional PAC-Efficient Reasoning in Large Language Models." arXiv:2512.03057.

Zuboff, Shoshana. 2019. *The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power*. New York: PublicAffairs.

INDEX

**P**ROLOGUE  
EDITORIAL NOTE TO THE 2026  
EDITION

**Chapter 1 - AI and the Structural Autonomy  
of Sense**

- 1.1 Introduction: The Collapse of Referential Authority
- 1.2 Conceptual Framework: Structural Autonomy of  
Sense
- 1.3 Theoretical Foundation: Operative Structures With-  
out Reference
- 1.4 Empirical Cases
  - 1.4.1 Predictive Justice (COMPAS Algorithm).
  - 1.4.2 Medical Diagnostics without Provenance.
  - 1.4.3 Credit Scoring.
  - 1.4.4 Synthetic Images.
- 1.5 Theorem of Disembedded Syntactic Authority
- 1.6 Epistemological and Ontological Consequences
- 1.7 Conclusion

**Chapter 2 -When Language Follows Form,  
Not Meaning**



2.1 Introduction: From Referential Anchoring to Structural Continuity

2.2 Formal Syntactic Activation: Model and Notation

2.3 The Collapse of Semantic Intentionality

2.4 Implications for Grammar, Legitimacy, and Epistemic Authority

2.5 Conclusion Without Closure

### **Chapter 3 - The Grammar of Objectivity**

3.1 Introduction: Bias as Grammatical Illusion

3.2 Technical History of Objectivity

3.3 Taxonomy of Formal Mechanisms of Neutrality

3.4 Comparative Analysis: LLM vs. LBR

3.5 Structural Neutrality Test

3.6 Conclusion: Epistemology Without Subject

### **Chapter 4 - Non-Neutral by Design**

4.1 Introduction

4.2 Theoretical Framework: Contamination and Constraint

4.3 Methodology

4.4 Results: Semantic Reentry

4.5 Toward a Structural Theory of Contamination

4.6 Extension to Reasoning Models (LRM)

4.7 Epistemological Consequences and Falsifiability

4.8 Conclusion: There Is No Neutral Prompt

**Chapter 5 - Ethos Without Source**

5.1 Introduction: Credibility Without Subject in Algorithmic Discourse

5.2 Theoretical Framework: From Classical Ethos to Synthetic Authority

5.3 Methodology: Corpus Design and Structural Credibility Mapping

5.4 Case Study 1: Healthcare and Diagnostic Authority

5.5 Case Study 2: Law and Education, Simulated Normativity and Academic Authority

5.6 Findings: Structural Features of Synthetic Ethos

5.7 Conclusion: From Heuristic Trust to Structural Governance

**Chapter 6 -AI and Syntactic Sovereignty**

6.1 Introduction

6.2 Background: Language, Authority, and the Disappearance of the Subject

6.3 From Intentionality to Structure: Authority Without Agency

6.4 Syntactic Sovereignty: Definition and Theoretical Architecture

6.5 Conclusion: The Age of Formal Obedience

**Chapter 7 - The Disconnected Syntactic Authority Theorem**

- 7.1 Introduction
- 7.2 Theorem Statement: Disconnected Syntactic Authority
- 7.3 Theoretical Framework
- 7.4 Epistemological Implications
- 7.5 Application to Language Models
- 7.6 Corollaries of the Theorem
- 7.6.1 Conditions of Falsification
- 7.7 Conclusion

## **Chapter 8 -TLOC: The Irreducibility of Structural Obedience in Generative Models**

- 8.1 Introduction: Can We Ever Know If a Machine Truly Obeys?
- 8.2 Formal Statement and Logical Structure of the TLOC
- 8.2.3 Theorem Statement (TLOC)
- 8.3 Implications of the TLOC: Beyond Verification
- 8.4 From Compliance Simulation to Epistemic Integrity: Toward Post-TLOC Design
- 8.5 Epistemic and Ontological Consequences
- 8.6 Formal Consequences and Theoretical Closure
- 8.6.4 Structural Closure: Limits of Compliance
- 8.7 Final Statement and Research Directions

**Chapter 9 - From Obedience to Execution:  
Structural Legitimacy in the Age of Reasoning  
Models**

9.1 Introduction: The Epistemological Legacy of Obedient Models

9.2 LLMs and the Regime of Passive Authority

9.3 LRMs and the Emergence of Structural Execution

9.4 Neutrality Revisited: From Corpus Illusion to Structural Simulation

9.5 The End of Reference: Projections Without World

9.6 Toward a New Model of Legitimacy: Structural, Not Semantic

9.7 Conclusion: Obedience Evolved - The Rise of Operational Authority

**Chapter 10 - Executable Power: Syntax as Infrastructure in Predictive Societies**

10.1 Conceptual Foundations of Executable Power

10.2 Executable Sovereignty and the Logic of Delegation

10.3 Grammatical Authority and the Executable Rule

10.4 Formal Grammar and the Executable Rule

10.5 The Executable Boundary

10.6 Compliance, Risk, and the AI Act

10.7 Structural Incompatibility and the Future of Legal Form

## **Chapter 11 - Grammar Without Judgment: Eliminability of Ethical Trace in Syntactic Execution**

- 11.1 Introduction
- 11.2 Theoretical Foundations
- 11.3 Ethical Trace as a Syntactic Variable
- 11.4 Mechanism of Exclusion
- 11.5 Non-Normative Execution
- 11.6 Disanalogy with Alignment and Ethics

## **Chapter 12 - Pre-Verbal Command: Syntactic Precedence in LLMs Before Semantic Activation**

- 12.1 Introduction: The Illusion of Semantic Primacy
- 12.2 Structural Execution Without Interpretation
- 12.3 The Regla Compilada as Source of Pre-Verbal Authority
- 12.4 Zero-Prompt, Null Semantics, and Active Syntax
- 12.5 Syntactic Precedence: A New Axis of Structural Sovereignty
- 12.6 Implications for AI Alignment and Prompt Design
- 12.7 Conclusion: Authority Without Meaning

## **Chapter 13 - Compiled Norms: Towards a Formal Typology of Executable Legal Speech**

- 13.1 Introduction: From Legal Declaration to Executable Grammar

- 13.2 State of the Art and Syntactic Displacement
- 13.3 Normative Logic and the Limits of Defeasibility
- 13.4 Deontic Logic vs. Compiled Legality
- 13.5 Methodological Framework and Corpus Selection
- 13.6 Dialect Module and Hot-Swap Procedure
- 13.7 Typology of Executable Legal Speech
- 13.8 Case Studies
- 13.9 From Syntax to Authority
- 13.10 Conclusion: Compiled Norms and the Future of Legal Language

## **Chapter 14 - Colonization of Time: How Predictive Models Replace the Future as a Social Structure**

- 14.1 Time as a Territory Under Dispute
- 14.2 From Chance to Prediction: The Algorithmic Turn of the Future
- 14.3 Hypothesis: Structural Substitution of the Future
- 14.4 The Future as an Open Field (History, Desire, Politics)
- 14.5 The Future as a Computed Matrix (AI, Big Data, Predictive Logic)
- 14.6 The Algorithm as an Executable Form of Anticipation
- 14.7 Structural Colonization of Time
- 14.8 Operational Examples with AI

14.9 Logical Formalization of the Phenomenon

14.10 Discursive and Epistemic Implications

14.11 Conclusion: From the Future to Execution

## **Chapter 15 - Function-calling Schemas as De Facto Governance: Measuring Agency Reallocation**

15.1 Introduction: The Interface as Compiled Rule

15.2 The Action Space and Total Governance Pressure

15.3 Attribution: Who Gained Control

15.4 Levers, Observables, and Experimental Design

15.5 Status, Hypotheses, and What Would Refute the Instrument

15.6 Relation to Attribution Methods in Machine Learning

## **Chapter 16 - Template Power: How Instruction Templates Redistribute Decision Rights**

16.1 Introduction: The Prompt Above the Prompt

16.2 The Syntactic Authority Inventory

16.3 Conditions and Corpus

16.4 Estimation

16.5 Status and Refutation Conditions

16.6 Relation to Prompt-Sensitivity Research

## **Epilogue - Closing the First Syntactic Phase**

**Annex I - DSAT: Formal Statement, Conditions, and Proof Sketch**

**Annex II - TLOC: Irreducibility Conditions and Verification Boundary**

**Annex III -  $\delta$  [E]  $\rightarrow$   $\emptyset$  Rulebook for Ethical Trace Deletion**

**Annex IV - Pre-Verbal Activation: Parser Configs and Ablation Tests**

**Annex V - Compiled Norms: Typology Metrics, LL(1) Grammar, and Hot-Swap Validation**

**Annex VI - Structural Execution Protocol for Reasoning Models**

**Appendix A - Canonical Glossary**

**Appendix B - Structural Dependency Map of the Book**

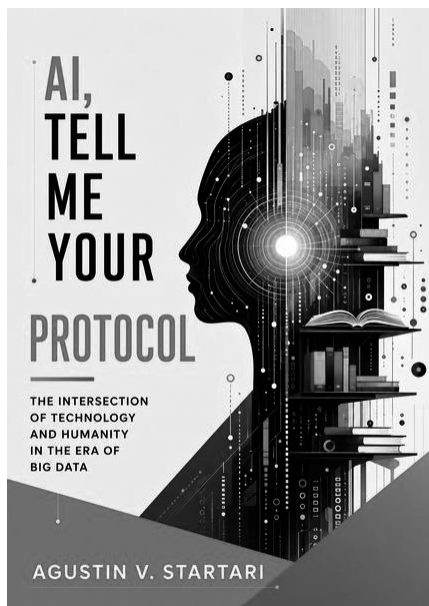
**Conclusion of Appendix B**

**Appendix C - Legal Dialect Hot-Swap Map  
References**





Did you love *AI Syntactic Power and Legitimacy*?  
Then you should read *AI, Tell Me Your Protocol*<sup>1</sup> by  
Agustin V. Startari!



---

1. <https://books2read.com/u/3LolPN>

2. <https://books2read.com/u/3LolPN>

In his work **"AI, Tell Me Your Protocols"** author **Agustin V. Startari** immerses us in a fascinating analysis of the intersection of technology and humanity in the era of big data. Exploring the rapid advancement of artificial intelligence (AI) and its impact on our society, the author invites us to reflect on the protocols that govern this relationship.

The book takes us on a journey through the world of big data and AI, unraveling key concepts and explaining how this technology has radically transformed the way we live, work, and interact. Startari examines the ethical, social, and economic implications of AI and raises fundamental questions about its integration into our daily lives.

From task automation to autonomous decision-making, the author examines how AI is permeating various fields such as industry, medicine, education, and commerce. Through concrete examples and case studies, the author shows how big data and algorithms are influencing our daily experience and how

the protocols that guide their development and application can shape our future reality.

In his rigorous and analytical approach, Startari explores the concerns and ethical challenges that arise from the intersection of technology and humanity. He questions the power and responsibility of large technology companies and proposes limits and regulations to ensure that AI is used ethically and benefits society.

"AI, Tell Me Your Protocols" invites us to reflect on how we want AI to shape our world and how we can ensure a harmonious coexistence between technology and humanity. With an accessible style based on up-to-date research, Agustin V. Startari's book offers a profound and enlightening insight into the intersection of AI and humanity, challenging us to take an active role in shaping our technological future.

This collection gathers independent research works exploring power, ideology, legitimacy, and history through a cross-disciplinary lens. Each vol-

ume is self-contained yet contributes to a shared thread: the critical analysis of how power is formed, exercised, and sustained over time. From Ancient Egypt to 20th-century totalitarianism, *Work Papers* presents a clear, academically grounded narrative about the historical and discursive mechanisms that shape our societies. Each entry is numbered to reflect its place in this evolving series.

Read more at <https://www.agustinvstartari.com/>.

Also by Agustin V. Startari

**Budismo Practico**

Una vida de abundancia. La mente exitosa

**Papeles de Trabajo**

Creacion de un Imperio

La Alta Edad Media

Los Templarios: Marco introductorio

Gregorio X

Los apóstoles de la fe: en la pluma de Fray Bartolomé de las Casas

Maquinaria de Propaganda  
Ucrania y Rusia: Un conflicto en progreso  
Adios capitalismo  
IA, dime tu protocolo  
Gramáticas del Poder  
IA Poder Sintactico y Legitimidad  
Christus Solus  
Herejías Esenciales: La lucha del cristianismo temprano  
El Futuro como Origen

**Practical Buddhism**  
The Path to happiness  
A Wealthy Life: The Successful Mind  
The Abundant Mind:

**Work Papers**  
Creation of an Empire

THE EARLY MIDDLE AGES  
The Templars: An Introductory Framework  
Gregory X  
THE APOSTLES OF FAITH IN THE WRIT-  
INGS OF FRIAR BARTOLOMÉ DE LAS  
CASAS  
Propaganda Machinery  
Ukraine and Russia: An Ongoing Conflict  
AI, Tell Me Your Protocol  
Grammars of Power  
Christus Solus  
Essential Heresies  
AI Syntactic Power and Legitimacy

**Standalone**  
The year of the dead Count



Watch for more at <https://www.agustinvstartari.com/>.



## About the Author

**Agustin V. Startari** is a researcher, linguist, and author specializing in the intersection of **language, artificial intelligence, power, legitimacy, and historical sciences.**

His work examines how linguistic, discursive, and technological structures organize perceptions of authority, distribute responsibility, and shape the ways institutions, artificial intelligence systems, and structures of power represent political, historical, and social events.

His research brings together **linguistics, discourse analysis, artificial intelligence, historiography, theories of power, and studies of legitimacy**, with particular attention to the relationship between linguistic form, authority, and political responsibility.

He is the author of the research series *Grammars of Asymmetric Visibility: AI, Imperial Power, and the Syntax of Responsibility*, focused on the relationships between artificial intelligence, language, power, and the attribution of responsibility.

He is also the author of *Work Papers*, a series devoted to historical research, historiography, and the analysis of problems related to the historical sciences.

Additional works include *Propaganda Machinery* and *Christus Solus*, which further explore the intersections of discourse, ideology, institutional power, and historical interpretation.

His work combines linguistic research, historical analysis, and the systematic study of the ways language, institutions, and technological systems participate in the construction and legitimization of power.

Read more at <https://www.agustinvstartari.com/>.



## About the Publisher

**LEFORTUNE Publishing** is a distinguished international publishing house dedicated to the publication and global distribution of high-quality fiction, historical works, literary essays, cultural studies, and nonfiction in English, Spanish, and German. With a strong commitment to editorial excellence, intellectual depth, and cultural exchange, LEFORTUNE Publishing develops carefully curated titles for readers and institutions worldwide.

Read more at <https://lefortunepublishing.com/>.

