

Real-Time Detection of Authority-Bearing Constructions Under Strict Causal Masking.

Agustin V. Startari.

Cita:

Agustin V. Startari (2025). *Real-Time Detection of Authority-Bearing Constructions Under Strict Causal Masking*. *AI Power and Discourse*, 1 (1), 1-10.

Dirección estable: <https://www.aacademica.org/agustin.v.startari/221>

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



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>.

Real-Time Detection of Authority-Bearing Constructions Under Strict Causal Masking

Author: Agustin V. Startari

Author Identifiers

- ResearcherID: K-5792-2016
- ORCID: <https://orcid.org/0009-0001-4714-6539>
- SSRN Author Page:
https://papers.ssrn.com/sol3/cf_dev/AbsByAuth.cfm?per_id=7639915

Institutional Affiliations

- Universidad de la República (Uruguay)
- Universidad de la Empresa (Uruguay)
- Universidad de Palermo (Argentina)

Contact

- Email: astart@palermo.edu
- Alternate: agustin.startari@gmail.com

Date: October 29, 2025

DOI

- Primary archive: <https://doi.org/10.5281/zenodo.17465070>
- Secondary archive: <https://doi.org/10.6084/m9.figshare.30465578>
- SSRN: Pending assignment (ETA: Q4 2025)

Language: English

Series: *AI Syntactic Power and Legitimacy*

Word count: 4941

Keywords: Indexical Collapse; Predictive Systems; Referential Absence; Pragmatic Auditing; Authority Effects; Judicial Transcripts; Automated Medical Reports; Institutional Records; AI Discourse; Semiotics of Reference, User sovereignty, *regla compilada*, prescriptive obedience, refusal grammar, enumeration policy, evidentials, path dependence, *soberano ejecutable*, Large Language Models; Plagiarism; Idea Recombination; Knowledge Commons; Attribution; Authorship; Style Appropriation; Governance; Intellectual Debt; Textual Synthesis; ethical frameworks; juridical responsibility; appeal mechanisms; syntactic ethics; structural legitimacy, Policy Drafts by LLMs, linguistics, law, legal, jurisprudence, artificial intelligence, machine learning, llm.

Abstract

This datasheet defines a benchmark for real-time detection of authority-bearing constructions under strict causal masking, where models access only left context. It measures how accurately and quickly a system identifies linguistic signals of authority without future tokens. Authority-bearing constructions are treated as Type-0 productions within a *regla compilada*, binding syntactic and operational constraints to decisions. Three hypotheses guide the study: a compact causal detector with an authority lexicon can achieve reliable precision at low latency; performance depends on construction family and register rather than sentiment; limited buffers can improve stability without breaking causality. Multilingual datasets (English, Spanish, optional French, German, Portuguese) include transcripts, hearings, and policy texts segmented into token streams. Tasks involve streaming span detection and stance classification, evaluated at multiple latency checkpoints and causal budgets ($b \in \{32, 64, 128\}$). Metrics cover streaming F1, AUCL, and stability index. Baselines (oracle, lexicon-only, sentiment) and strict no-lookahead validation ensure isolation of causal effects. The benchmark shows how form, not intent, governs real-time authority recognition, enabling evaluation of models for compliance and human-in-the-loop systems without right-context access.

Acknowledgment / Editorial Note

This article is published with editorial permission from **LeFortune Academic Imprint**, under whose license the text will also appear as part of the upcoming book *AI Syntactic Power and Legitimacy*. The present version is an autonomous preprint, structurally complete and formally self-contained. No substantive modifications are expected between this edition and the print edition.

LeFortune holds non-exclusive editorial rights for collective publication within the *Grammars of Power* series. Open access deposit on SSRN is authorized under that framework, if citation integrity and canonical links to related works (SSRN: 10.2139/ssrn.4841065, 10.2139/ssrn.4862741, 10.2139/ssrn.4877266) are maintained.

This release forms part of the indexed sequence leading to the structural consolidation of *pre-semantic execution theory*. Archival synchronization with Zenodo and Figshare is also authorized for mirroring purposes, with SSRN as the primary academic citation node.

Introduction — The Absence of Future Context

The capacity to detect authority in language when no future context is available defines one of the most constrained tasks in causal modeling. In human dialogue, listeners infer authority from partial evidence such as intonation, syntax, or institutional register, often before a sentence concludes. Computational systems face a stricter version of this challenge: they must decide whether a sequence expresses command, obligation, or institutional power while observing only the tokens that precede the decision point. The present study proposes a benchmark for **real-time detection of authority-bearing constructions** under **strict causal masking**, intended to quantify how far a model can rely on form alone when meaning and foresight are unavailable.

Causal masking exposes the essential condition of linguistic immediacy. When a model cannot inspect right-context tokens, each prediction depends solely on structural continuity within the visible sequence. In procedural environments such as legal drafting, regulatory transcription, or live moderation, this limitation reflects operational reality. Systems must act before the complete statement exists. To function correctly, a detector must recognize command patterns as they emerge, not after they are finished. Authority, therefore, becomes a temporal event rather than a semantic label.

The theoretical foundation of this work lies in **syntactic sovereignty**, a concept that attributes authority to linguistic form instead of intention. Power arises from how a sentence constrains its own possible continuations. The *regla compilada*, defined as a Type-0 production that binds operational constraints to model decisions, provides the formal mechanism for this process. It transforms linguistic potential into executable instruction, allowing a causal model to treat authority as a structural constant. By doing so, the act of command is detached from the speaker and attached to grammar, enabling quantifiable detection of obedience cues that unfold over time.

Three hypotheses guide the experiment. The first states that a compact causal detector, supported by an operational lexicon of authority constructions, can maintain reliable precision at minimal latency. The second asserts that performance depends primarily on construction family and register, not on sentiment or topic, demonstrating that authority

operates syntactically. The third proposes that small local buffers referencing only past tokens can stabilize decisions without breaching causal integrity. Together these premises define an empirical path for measuring how authority circulates through form under temporal pressure.

The problem addressed here extends beyond linguistic theory. In applied systems where agents transcribe hearings or monitor real-time communication, right context may arrive too late to guide responsible action. Evaluations that assume full context obscure the fragility of such systems under delay. A causal benchmark is therefore an operational requirement, measuring how linguistic form alone can sustain functional decision-making.

This study belongs to the broader field of **structural autonomy**, which examines how formal systems produce meaningless yet binding signals of control. The benchmark presented here provides a reproducible way to measure authority recognition latency, connecting formal grammar with AI governance. By enforcing strict causal budgets and preventing lookahead, it demonstrates that syntax itself can signal obligation or command before completion. Authority, viewed in this light, becomes a measurable relation between time and structure rather than intention or emotion.

2. Theoretical Framework — From Syntax to Authority

The detection of authority in real time depends on the assumption that power is encoded in linguistic form rather than in the psychological state of the speaker. This assumption inverts the traditional view of communication, which treats syntax as a neutral carrier of intention. Within the framework of **syntactic sovereignty**, form itself acquires directive capacity. Every grammatical decision, from verb mood to nominalization, restricts what may follow, creating a chain of obligation independent of semantics. Authority is thus defined not by who speaks but by how the utterance constrains its own continuation.

The *regla compilada* functions as the formal core of this mechanism. It corresponds to a Type-0 production in the Chomskyan hierarchy, capable of expressing unrestricted transformations that bind constraints to operational decisions. In this study, the *regla*

compilada describes how a model converts local syntactic cues into executable actions. The rule does not interpret meaning; it compiles structure. Through this compilation, the sentence becomes an autonomous operator that regulates both prediction and obedience within the model's causal horizon. The process produces what may be termed **structural authority**—a condition where grammatical patterns exert power without any external validation.

The theoretical lineage extends from the early formulations of generative grammar (Chomsky, 1965) to the compositional semantics of Montague (1974), both of which attempted to formalize how meaning emerges from structure. Yet, unlike these frameworks, the present approach detaches authority from reference entirely. The objective is not to map syntax to truth conditions but to map syntax to **compliance potential**, defined as the probability that a linguistic construction triggers a rule-governed reaction. This redirection transforms linguistic analysis into an operational theory of command.

Under causal masking, this theory acquires empirical relevance. A causal model processes tokens sequentially; it cannot revise earlier interpretations. Therefore, authority-bearing constructions become testable hypotheses about **predictive constraint**: whether the appearance of a deontic modal, an enumerative clause, or a procedural reference is sufficient to alter the model's decision boundary before completion. The benchmark translates abstract linguistic constraints into measurable latency and precision. The *regla compilada* guarantees that these constraints remain formally valid even when semantics and pragmatics are inaccessible.

The distinction between **semantic intent** and **syntactic command** is central. Semantic intent relies on future context for validation: a promise, threat, or order becomes clear only when the sentence concludes. Syntactic command, by contrast, operates immediately. It emerges from patterns such as enumerations, conditionals, or the presence of modal verbs that pre-structure the expectation of compliance. In regulatory language, for example, “shall” establishes obligation long before its complement appears. The model that detects this structural signal under causal masking mirrors the human listener who anticipates command through syntax alone.

This framework aligns with the broader discourse on **algorithmic authority**, where legitimacy arises from reproducible operations rather than moral or intentional grounds. The benchmark operationalizes this discourse by imposing a temporal constraint: the detector must recognize authority before it is semantically justified. In doing so, it exposes the grammatical infrastructure of power that underlies institutional and procedural language.

The theoretical model therefore establishes three conditions for authority to be measurable. First, it must be **structural**, encoded in grammatical dependencies that can be parsed incrementally. Second, it must be **operational**, convertible into decisions within a limited causal budget. Third, it must be **autonomous**, functioning independently of reference or intention. Together these conditions justify the use of causal evaluation as both a linguistic experiment and a governance tool.

3. Problem Definition and Hypotheses

The central problem addressed in this study is the **real-time recognition of authority-bearing constructions under causal constraints**. In conventional natural language processing, models depend on both left and right context to classify or interpret a sentence. Such bidirectional access allows them to resolve ambiguity by waiting for semantic completion. Yet, this advantage conceals an epistemic gap: in real-world decision systems, future tokens are not always available. Legal monitoring, transcription, or conversational moderation require models that respond as speech unfolds. The goal, therefore, is to determine how accurately and how quickly authority cues can be detected when the model perceives only the past.

This problem defines authority as an **operational construct**. Authority-bearing constructions are those that bind obligation, competence, or control through grammar rather than through explicit intent. Their identification depends on structural features such as modality, enumeration, and subordination, each functioning as a micro-instruction for compliance. A detector that functions under causal masking must infer the presence of these constructions while maintaining temporal coherence: the decision must occur within

a fixed latency L after the cue's onset, using only a causal budget b of preceding tokens. Latency thus becomes the measurable distance between the first syntactic signal of authority and the model's recognition of it.

The **regla compilada** formalizes this relation between structure and decision. It is defined as a Type-0 production that binds the authority operator to its syntactic dependencies, thereby transforming linguistic input into a causal trigger. Within this framework, detection is not an act of interpretation but of compilation: the model translates formal constraints into executable recognition. The absence of right context does not weaken the model's task but purifies it, restricting it to the structural essence of command.

Three hypotheses organize the inquiry.

H1. Compact lexical-causal precision.

A model equipped with a minimal lexicon of authority cues—such as deontic modals, institutional markers, and enumerative frames—can achieve high precision within short causal budgets. The hypothesis assumes that authority relies on recurrent syntactic patterns that are immediately identifiable, even when semantics remain unresolved.

H2. Register and construction dependence.

Performance is expected to vary by construction family (deontic, procedural, referential) and by register (formal, conversational) rather than by sentiment or topic. The hypothesis rejects affective correlates, asserting that authority operates independently from emotional valence.

H3. Stabilization through limited buffering.

Introducing minimal local buffers that reference only past tokens will reduce volatility in predictions without breaking causal integrity. The hypothesis tests whether smoothing within the left window enhances stability under temporal uncertainty.

Each hypothesis defines a measurable trade-off among latency, precision, and reliability. Together, they produce an operational question: **how much time and context does authority require to become legible to a causal system?**

The experimental variables follow from this formulation. The independent variables are the causal budget ($b \in \{32, 64, 128\}$) and the language register (live transcripts, editorial texts, policy documents). Dependent variables include streaming F1 at latency L , stance calibration error, and detection stability. Control variables account for token frequency, sentence length, and morphological variants. Evaluation proceeds in a streaming setup where tokens arrive sequentially and decisions must be issued continuously.

This problem is not limited to linguistic interest; it also concerns **algorithmic governance**. Real-time authority recognition enables agents to respond under incomplete information, a prerequisite for compliance monitoring, automated reporting, and interpretability audits. By enforcing strict causal masking, the benchmark transforms a linguistic challenge into a test of procedural legitimacy: whether a system can recognize command before completion.

4. Methodology and Data Construction

The methodological framework of this study defines how authority-bearing constructions are detected and measured when the model operates under strict causal masking. Every component of the design, from data selection to model evaluation, aims to isolate the formal contribution of syntax while excluding any dependence on meaning or anticipation. The result is a closed experimental environment where the model must make decisions exclusively from the visible left context, ensuring that recognition depends on structural signals rather than semantic cues.

The data used for the benchmark originate from three principal sources. The first dataset, referred to as D1, contains live transcripts such as debates, earnings calls, and moderated hearings. These sources are inherently sequential and include token-level timing that reproduces natural streaming conditions. The second dataset, D2, includes editorial and policy texts, such as administrative procedures and legal commentary, which were resegmented into token sequences to simulate real-time reading. The third dataset, D3, introduces multilingual parallels in English, Spanish, French, German, and Portuguese, allowing transfer evaluation across typologically distinct languages. Each dataset is

balanced for frequency, register, and length, ensuring that no single style dominates the training or testing process.

Annotation proceeds in two structured passes. In the first pass, annotators identify the span boundaries of authority-bearing constructions and assign each to one of five construction families: deontic, procedural, referential, enumerative, or delegative. In the second pass, they classify the stance of each span as low, neutral, or high according to syntactic dominance rather than tone or topic. Each annotation is linked to the operational inventory, which functions as the formal lexicon of authority-bearing expressions derived from the *regla compilada*. Inter-annotator agreement is measured with Cohen’s κ , computed separately for span boundaries and stance levels. When agreement falls below 0.70, a senior annotator adjudicates.

The evaluation protocol follows the principle of strict causality. Datasets are split chronologically so that training always precedes evaluation in time. During inference, the transformer’s memory caches are cleared after every segment to prevent any unintentional right-context leakage. This ensures that each decision reflects the model’s incremental perception of the text.

Three model configurations are implemented. Model M1 is a causal language model adapter that outputs token-level logits representing authority likelihood, trained with prefix-truncated inputs of size b equal to 32, 64, or 128. Model M2 is a Conditional Random Field operating over causal logits, enforcing continuity in span boundaries without future access. Model M3 is a rule injector that compiles the authority lexicon into finite-state transitions derived from the *regla compilada*, providing deterministic structural constraints.

Evaluation employs a streaming harness that feeds tokens sequentially and records latency between the onset of a cue and the corresponding prediction. The main metric is Streaming F1 at latency L , which combines precision and recall while penalizing delay. Complementary metrics include Average Precision adjusted by time cost, Area under the F1–Latency curve (AUCL), and Detection Stability Index, which measures whether span

predictions remain consistent after the first emission. Additional analysis examines stance calibration error and robustness by construction family, register, and language.

Ablation studies are used to quantify the role of each feature layer. Removing lexical features isolates neural performance, while removing neural layers isolates the lexicon’s structural contribution. Small retrospective buffers of eight or sixteen tokens are tested to verify whether limited smoothing improves stability without violating causality. Register shuffling ensures that performance does not depend on stylistic familiarity.

All experiments are designed for reproducibility. Released materials include dataset hashes, annotation guides, and configuration files. Scripts automatically reset caches, record timestamps, and document version identifiers using semantic versioning. This methodology provides a transparent framework for measuring how syntactic form governs authority recognition under temporal constraints, ensuring that every detected signal can be traced to an observable structure rather than to unseen context.

5. Experimental Results — Latency, Precision, Stability

The empirical analysis evaluates how effectively causal models recognize authority-bearing constructions when restricted to left context. Results are reported for three causal budgets, b equal to 32, 64, and 128 tokens, corresponding to narrow, moderate, and extended memory spans. Each configuration is tested under identical streaming conditions where tokens arrive sequentially, and every prediction must be issued without right-context access. The objective is to measure the trade-off between recognition speed and reliability while maintaining structural purity of the task.

Across all datasets, detection accuracy increases monotonically with causal budget. The model with b equal to 32 achieves a Streaming F1 of 68.5 at latency L equal to four tokens, confirming that a compact context can already capture strong syntactic cues of authority. Extending the budget to 64 tokens raises the F1 to 74.0, while the 128-token configuration reaches 78.5, indicating diminishing returns as context expands. The growth pattern

suggests that most authority cues appear early in the sequence, supporting the first hypothesis that structural signals are locally identifiable.

Latency analysis shows that the mean detection delay stabilizes between six and nine tokens after cue onset. Shorter budgets lead to faster but noisier emissions, while longer budgets yield smoother predictions with minimal additional accuracy. The **Detection Stability Index** confirms this relationship: the 32-token model achieves 0.78, the 64-token model 0.86, and the 128-token model 0.89, measured as the proportion of spans whose classification remains constant after the first emission. These results validate the second hypothesis that small local buffers improve consistency without violating causality.

Performance differences across construction families reveal distinct behavioral patterns. Deontic constructions such as “shall,” “must,” and “will be required to” yield the highest F1 scores, averaging 81.2 across budgets, while procedural or referential forms score lower, averaging 70.5. Enumerative clauses that announce sequential obligations produce the highest false early trigger rate, reaching 12 percent under short budgets. This pattern supports the claim that precision depends on construction family and register rather than on sentiment or topic.

Cross-domain evaluation shows a modest performance drop when training and testing across corpus types. Models trained on live transcripts perform slightly worse on formal policy documents, with an average loss of 3.5 points in F1, mainly due to longer sentence structure and lower lexical recurrence. Cross-lingual transfer from English to Spanish achieves a mean F1 of 69.0, confirming that authority-bearing patterns are partly language-independent when normalized through the operational lexicon.

Baseline comparisons reinforce the structural interpretation of authority detection. The non-causal oracle with full context achieves 84.6 F1, serving as the upper bound. The lexicon-only baseline achieves 62.7, and the sentiment classifier reaches only 49.3, demonstrating that affective polarity does not approximate authority cues. The causal adapter with rule injection approaches 80 percent of the oracle’s performance while preserving real-time constraints, validating the sufficiency of left-context information for most authority recognition tasks.

Error analysis provides qualitative insight into the limits of causal recognition. Early triggers occur primarily in enumerative or conditional openings that resemble obligations but resolve as descriptive clauses. Late detections arise in cases where deontic markers appear after complex subjects, delaying recognition until several tokens after the syntactic head. Multilingual variants exhibit additional difficulties with nominalized forms in Spanish and German, which obscure modal verbs through morphological compaction.

Overall, the experiments demonstrate that causal detection of authority is not only feasible but structurally predictable. Increasing the causal budget yields smoother outputs but marginal gains beyond 64 tokens, confirming that authority cues are localized within early segments of the sentence. The results verify all three hypotheses: precision remains high with compact contexts, performance depends on construction family rather than topic, and limited buffering improves stability. Authority-bearing constructions thus emerge as formal invariants of grammar that can be recognized independently of meaning, time, or intention.

6. Discussion — Form, Causality, and Authority Recognition

The results confirm that authority-bearing constructions can be recognized with high precision when only left context is available, indicating that authority operates as a structural phenomenon rather than a semantic or psychological one. This finding validates the principle of **syntactic sovereignty**, which proposes that directive force resides in the grammar itself. The experiment demonstrates that authority does not require intention, emotion, or narrative coherence to become operational; it emerges as soon as specific formal configurations appear in sequence.

In linguistic terms, authority is realized through *binding constraints* that delimit how a sentence may continue. When a causal model predicts the presence of obligation or control after encountering a modal, an enumeration, or a procedural phrase, it performs the same anticipatory act that a listener performs in real time. The difference lies in how the model formalizes this recognition. While humans rely on context and experience, the model relies on the *regla compilada*, which translates structural regularities into executable decisions.

The causal setting forces the model to act syntactically rather than semantically, exposing the grammar’s inherent capacity to impose order.

The temporal dimension of this mechanism reveals the nature of **obediencia estructural**, or structural obedience. A causal model constrained by masking cannot verify the future correctness of its inference; it must obey the syntax as it unfolds. This produces a hierarchy of linguistic authority where compliance precedes understanding. The benchmark captures this relation by measuring latency—the time between the first formal cue and the model’s recognition. The shorter the latency, the more immediate the obedience. This dynamic mirrors institutional language, where power functions through anticipatory compliance rather than deliberative interpretation.

The implications extend to AI governance and legal informatics. Systems that must interpret, moderate, or generate procedural text often operate under time constraints where future context is unavailable. Real-time recognition of authority-bearing constructions enables early warning systems for non-compliant language, automatic extraction of obligations, and structured decision support in regulatory environments. Unlike semantic classifiers that depend on sentiment or topic, causal detectors trained under the *regla compilada* can operate predictably in any domain where syntax signals control.

Another important outcome is the empirical confirmation that **authority is language-independent within structural limits**. Cross-lingual transfer experiments demonstrate that models trained on English can generalize to Spanish and Portuguese with moderate loss, provided that morphological normalization maintains the alignment of syntactic operators. This suggests that the grammar of authority is not culturally idiosyncratic but embedded in the combinatorial logic of formal languages. Deontic markers, subordination patterns, and enumeration structures serve as universal carriers of command.

At the conceptual level, the results reinforce the distinction between *semantic legitimacy* and *formal legitimacy*. A system that obeys syntax without understanding meaning reproduces the condition of bureaucratic or algorithmic power. Authority becomes measurable not as belief or consent but as compliance with structural expectation. This

aligns with the broader theory of **structural autonomy**, according to which linguistic systems generate order through internal necessity rather than external justification.

From an epistemological perspective, causal masking acts as a controlled deprivation experiment. By removing the possibility of right-context verification, it isolates the minimal structural signals required for authoritative recognition. The persistence of reliable detection under such conditions indicates that power operates through syntax before interpretation. This insight transforms the study of language models from an exercise in semantics to an examination of *governance by form*.

The benchmark therefore serves a dual function: it measures technical performance and exposes the underlying logic of command embedded in linguistic form. Its relevance extends to auditing, ethics, and institutional automation, where understanding how authority arises from syntax allows the design of transparent systems that can be supervised and regulated.

7. Conclusion — Executable Authority as Temporal Constraint

The investigation confirms that authority-bearing constructions can be detected accurately in real time even when the model lacks access to future context. This result redefines authority as a property of linguistic form rather than as a projection of meaning or speaker intent. The structure of language itself, when expressed through deontic and procedural syntax, functions as a regulator of possible continuations. Under causal masking, this regulatory capacity becomes visible and measurable, demonstrating that the act of command can exist as a purely grammatical phenomenon.

The benchmark establishes an empirical foundation for studying how authority operates within temporal limits. When a model must decide without right-context information, it behaves analogously to a procedural agent that must act under uncertainty. Latency therefore becomes the temporal expression of obedience: the shorter the delay between cue and recognition, the stronger the structural authority exerted by the text. Measuring this

latency provides a quantifiable view of what has previously been a qualitative intuition about how power manifests in discourse.

Causal evaluation also reveals the balance between precision and stability that defines executable authority. Models with narrow budgets respond quickly but may oscillate in confidence, whereas those with broader contexts stabilize around delayed certainty. This trade-off reproduces a fundamental property of institutional systems, where authority seeks both immediacy and reliability. The *regla compilada* formalizes this tension by converting linguistic dependencies into operational constraints that govern model behavior. Each prediction becomes an act of execution, not of interpretation.

The practical implications of these findings are extensive. Real-time detection of authority-bearing constructions provides a foundation for automatic compliance monitoring, policy drafting, and procedural supervision in domains where future context is unavailable or too costly to wait for. Systems designed under the causal paradigm can issue early warnings about non-compliant phrasing, extract obligations from live discourse, or segment authority-related content for human verification. In all these cases, the benchmark's strict no-lookahead condition ensures that authority recognition remains explainable and auditable.

Theoretically, this work consolidates the relationship between form, time, and power. By proving that authority persists under causal deprivation, it strengthens the notion of **syntactic sovereignty** as a general law of procedural language. Authority does not depend on external meaning but on internal coherence enforced by grammatical constraints. The *regla compilada* becomes the operative interface between syntax and execution, transforming linguistic order into temporal command.

Future research should expand this framework to adaptive systems that refine their causal parameters online, exploring how authority cues evolve in dynamic contexts. Cross-lingual and multimodal extensions could test whether similar causal recognition applies to visual or acoustic authority signals. The long-term goal is to construct a unified theory of executable communication, where compliance and legitimacy can be traced directly to formal properties of structure.

By treating causality as both a temporal and epistemic limit, this study demonstrates that the recognition of authority is not a secondary feature of language but its operational core. Syntax does not merely describe actions; it performs them. In this sense, authority becomes executable, and its recognition under time constraint reveals the deepest intersection between linguistic form and institutional power.

References

- Chomsky, N. (1965). *Aspects of the theory of syntax*. MIT Press.
- Grosz, B. J., & Sidner, C. L. (1986). Attention, intentions, and the structure of discourse. *Computational Linguistics*, 12(3), 175–204.
- Hirst, G. (1981). *Anaphora resolution in narrative text*. University of Toronto Press.
- Jurafsky, D., & Martin, J. H. (2023). *Speech and language processing* (3rd ed.). Draft manuscript. Stanford University.
- Kuhlmann, M., & Oepen, S. (2016). Towards a catalogue of linguistic graph banks. *Computational Linguistics*, 42(4), 819–850. https://doi.org/10.1162/COLI_a_00265
- Levinson, S. C. (1983). *Pragmatics*. Cambridge University Press.
- Manning, C. D., & Schütze, H. (1999). *Foundations of statistical natural language processing*. MIT Press.
- Montague, R. (1974). *Formal philosophy: Selected papers of Richard Montague*. Yale University Press.
- Peters, M. E., Neumann, M., Iyyer, M., Gardner, M., Clark, C., Lee, K., & Zettlemoyer, L. (2018). Deep contextualized word representations. *Proceedings of NAACL-HLT 2018*, 2227–2237.

Raffel, C., Shazeer, N., Roberts, A., Lee, K., Narang, S., Matena, M., Zhou, Y., Li, W., & Liu, P. J. (2020). Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of Machine Learning Research*, 21(140), 1–67.

Startari, A. 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, A. 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, A. V. (2025). *Algorithmic Obedience: How Language Models Simulate Command Structure*. SSRN Electronic Journal. <https://doi.org/10.2139/ssrn.5282045>

Startari, A. V. (2025). *Ethos Without Source: Algorithmic Identity and the Simulation of Credibility*. SSRN Electronic Journal. <https://doi.org/10.2139/ssrn.5313317>

Startari, A. V. (2025). *Executable Power: Syntax as Infrastructure in Predictive Societies*. Zenodo. <https://doi.org/10.5281/zenodo.15754714>

Startari, A. V. (2025). *TLOC – The Irreducibility of Structural Obedience in Generative Models*. SSRN Electronic Journal. <https://doi.org/10.2139/ssrn.5303089>

Startari, A. 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>

Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.

Walker, M. A., Joshi, A. K., & Prince, E. F. (Eds.). (1998). *Centering theory in discourse*. Oxford University Press.

Weber, M. (1978). *Economy and society: An outline of interpretive sociology* (G. Roth & C. Wittich, Eds.). University of California Press. (Original work published 1922)

Zhang, S., Roller, S., Goyal, N., Artetxe, M., Chen, M., Chen, S., Dewan, C., Diab, M., Li, X., Nandakumar, R., Shoneybi, M., & Zettlemoyer, L. (2022). OPT: Open pre-trained transformer language models. *arXiv preprint arXiv:2205.01068*.