

Citation by Completion: LLM Writing Aids and the Redistribution of Academic Credits.

Agustin V. Startari.

Cita:

Agustin V. Startari (2025). *Citation by Completion: LLM Writing Aids and the Redistribution of Academic Credits*. *AI Power and Discourse*, 2 (1), 1-10.

Dirección estable: <https://www.aacademica.org/agustin.v.startari/217>

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



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>.

Citation by Completion: LLM Writing Aids and the Redistribution of Academic Credits

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 7, 2025

DOI

- Primary archive: <https://doi.org/10.5281/zenodo.17287506>
- Secondary archive: <https://doi.org/10.6084/m9.figshare.30295582>
- SSRN: Pending assignment (ETA: Q3 2025)

Language: English

Series: *AI Syntactic Power and Legitimacy*

Word count:

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

Large language models increasingly shape how academic citations are produced, suggested, and normalized. This paper examines the redistribution of academic credit produced by autocomplete and citation recommendation systems. While citation metrics traditionally reflect author intent, the syntactic design of LLM suggestion interfaces introduces a new variable: *authority-bearing syntax*. Through a double-blind experimental design comparing writing sessions with suggestion disabled, neutral suggestion, and authority-framed suggestion, this study quantifies shifts in citation concentration, novelty, and legitimacy phrasing. Results show that completions containing legitimizing structures (“as established by,” “following the seminal work of”) significantly increase concentration and reduce source diversity. The paper defines three measurable deltas, ΔC (concentration), ΔN (novelty), and ΔA (authority syntax), and demonstrates how predictive phrasing can algorithmically reproduce canonical hierarchies. As a corrective, it proposes a *Fair Citation Prompt* specification and an editorial checklist to detect and mitigate credit capture through syntactic bias. The findings suggest that citation fairness must be treated not only as a bibliometric concern but as a structural property of text generation systems, requiring explicit governance at the level of language form..

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.

For licensing, referential use, or translation inquiries, contact the editorial coordination office at: [contact@lefortune.org]

Part I. The Concentration Problem

Large language models have introduced a composition stage where epistemic redistribution occurs through predictive syntax. Citation autocompletes functions, increasingly embedded in editors and research tools, shape who receives recognition and how frequently that recognition is repeated. The mechanism is not merely quantitative. It is linguistic, infrastructural, and procedural. Each accepted suggestion reflects a probability distribution learned from past text. When writers accept completions that propose specific names or authority-bearing phrases, credit moves toward sources that the model deems most likely in that context. The single act appears minor, yet repeated across many writers and many sessions it generates a structural narrowing of visibility. The choice of whom to cite begins to behave like a grammatical effect rather than an evaluative decision.

Bibliometric research has shown that citations follow preferential attachment, where initial advantage attracts further advantage (Barabási, 2002). Merton (1968) described the Matthew effect as dynamic in which already recognized scholars accumulate disproportionate credit. Large language models operationalize this dynamic inside the sentence. Since training corpora encode historical inequalities, predictive completions inherit those distributions and reissue them as fluent text (Bender, Gebru, McMillan-Major, & Shmitchell, 2021). When a system suggests “as established by Smith (2017)” instead of “as argued by López and Chen (2020),” the completion is not a reasoned judgment. It is an output that reflects frequency, co-occurrence, and stylistic regularities. Once accepted by the author, the statistical trace becomes an apparently justified citation. The circulation of predictive text is therefore also a circulation of inherited hierarchies.

The concentration problem can be specified as convergence of reference choices toward a reduced subset of high frequency nodes when predictive assistance is active. Traditional inequalities operate in selection of venues, access to literature, or language of publication. Predictive concentration operates in the microdynamics of writing. Reviewers and readers seldom observe it because it occurs prior to submission. The interface itself offers authority-bearing constructions such as “seminal study,” “canonical framework,” and “pioneering research.” These phrases are not neutral descriptors. They are grammatical devices that elevate some names while suppressing exploration of alternatives. Their

repetition across drafts and documents produces a background endorsement that appears stylistic but functions as allocation of credit.

A useful way to test this claim is to treat concentration as a measurable property that changes when autocomplete is enabled. An adaptation of the Herfindahl–Hirschman Index to citation distributions provides a direct indicator. If sessions with autocomplete show a higher index than sessions without it, then predictive assistance reduces diversity in the reference field (Anderson, Kumar, & Zheng, 2023). This approach treats the writing environment as an economic structure that allocates attention. It also clarifies that what appears to be efficiency in composition can be a transfer of credit toward already dominant clusters.

The epistemological stakes follow from the relation between originality and legitimacy. Academic writing has required authors to innovate while also grounding claims in previous work. Predictive systems compress this relation by rewarding fluency. A completion that matches learned patterns will sound more coherent and more authoritative, even when it reproduces redundancy. As a result, the syntax of credibility becomes difficult to distinguish from the syntax of recurrence. Writers accept suggestions that feel correct because they are smooth and conventional. Over time, repeated acceptance yields a linguistic attractor, a stable set of names and formulations that dominate surface text independent of localized relevance.

This mechanism should be understood as infrastructural bias rather than individual preference. It does not presuppose intent or belief. It follows from the embedding of preference inside a probabilistic grammar that operates at scale. Bourdieu (1991) argued that linguistic form conceals social structure. A large language model implements this insight in computational terms. Its output is a structured reflection of prior selections by institutions, publishers, and research communities. By returning these selections as predictive syntax, the system constructs obedience to precedent under the appearance of stylistic optimization. The writer’s agency is not removed, yet it is steered, and that steering is difficult to detect without explicit metrics.

Recognizing the concentration problem has two direct consequences for governance of writing aids. First, developers and editors need diagnostic measures that surface concentration in real time, not only in post publication metrics. Second, interfaces should separate evidential phrasing from name prediction, since the conflation of both creates authority through grammar rather than through evaluation. The later parts of this paper propose an experimental design that isolates the effect of suggestion syntax, and a specification for fair prompting that rotates references, discloses uncertainty, and discourages legitimizing formulas as default completions. The aim is not to prohibit assistance, but to prevent the invisible transfer of recognition that occurs when probability is allowed to operate as a surrogate for judgment.

Part II. Autocomplete as a Syntactic Market

The operation of citation suggestion systems within large language models can be analyzed as the emergence of a new syntactic market, one in which linguistic form functions as a vehicle for the circulation of symbolic and epistemic capital. In traditional academic economies, credit flows through institutionalized acts of acknowledgment: a citation confers value, positions the cited author within a hierarchy of recognition, and contributes to the measurable accumulation of prestige. When autocomplete systems intervene in this process, they do not merely facilitate writing; they restructure the economy of legitimacy by transforming linguistic probability into an exchangeable form of authority. The market logic emerges from syntax itself, as predictive mechanisms convert grammatical recurrence into an implicit valuation of certain sources.

From the perspective of Pierre Bourdieu's theory of cultural production, citation is a symbolic act that allocates capital within a structured field (Bourdieu, 1991). The accumulation of citations corresponds to the accumulation of symbolic resources that translate into intellectual authority. In the context of large language models, this symbolic economy becomes automated. Each completion that proposes a specific author or canonical expression performs an act of value assignment. The more frequently the model reproduces particular names, the higher their statistical visibility, and therefore their perceived

legitimacy. This process constitutes what may be termed *syntactic capital formation*: the transformation of linguistic recurrence into epistemic value. What was once a social process mediated by deliberation becomes a probabilistic one mediated by predictive syntax.

The structure of this market can be described through three interacting components: linguistic frequency, perceived authority, and adoption velocity. Linguistic frequency determines the probability of a name or citation appearing in a completion. Perceived authority emerges as writers internalize those suggestions as indicators of reliability. Adoption velocity measures how quickly such completions are accepted and propagated across new texts. The combination of these factors produces a feedback loop analogous to price discovery in economic systems. The value of a citation is no longer determined solely by its content or contribution, but by its visibility within the predictive grammar of the writing interface. Each citation accepted under suggestion conditions acts as a transaction that increases the symbolic market share of the referenced source.

This syntactic market differs from conventional bibliometric dynamics in its temporality and automation. Traditional citation accumulation operates retrospectively, once the paper is published and indexed. Predictive citation operates prospectively, within the sentence itself. It introduces anticipatory credit allocation based on model expectations rather than peer recognition. Floridi (2022) has noted that informational environments increasingly define what counts as epistemically valid. The predictive completion system, by design, rewards what has already circulated and penalizes what remains marginal or linguistically irregular. Consequently, autocomplete functions transform linguistic predictability into a pricing mechanism for legitimacy. Predictive fluency becomes equivalent to market liquidity: the ease with which certain forms of authority circulate within the writing process.

This transformation has measurable effects on the distribution of symbolic capital. A model trained predominantly on English-language academic corpora, for example, reproduces linguistic asymmetries that favor Anglo-American publication ecosystems. When predictive systems suggest canonical names from those contexts more readily than emerging or regional scholars, they reinforce an uneven exchange rate between linguistic

zones. Citation autocomplete thus acts as an algorithmic market maker, stabilizing certain centers of authority while suppressing peripheries. The syntactic market does not require intention to function; it follows structural incentives embedded in data. In this respect, it parallels the operation of algorithmic trading systems in finance, where automated agents execute exchanges faster than human oversight can evaluate them.

One consequence of this syntactic marketization is the compression of epistemic diversity. Novel or unconventional sources become illiquid assets. Their probability of being suggested decreases as their frequency within the corpus declines. Even if such sources are conceptually significant, they face barriers to reentry into the linguistic economy because the model privileges what is statistically normative. As Anderson, Kumar, and Zheng (2023) observe, concentration indices rise sharply when automated recommendation systems influence selection patterns. The same logic applies to predictive writing: the efficiency of suggestion accelerates convergence toward dominant names. The writer becomes a participant in a market of legitimacy, transacting in phrases and citations that carry the highest syntactic return.

The analogy extends further. In traditional markets, liquidity and volatility determine the cost of exchange. In predictive writing, grammatical fluency and semantic predictability play similar roles. Sentences that incorporate highly cited names or familiar constructions are processed more smoothly by both human readers and machine evaluators. They incur lower cognitive cost and higher rhetorical yield. Thus, the syntactic market rewards conformity. Originality, by contrast, behaves like a high-risk investment: costly in time and uncertain in reception. The writer's rational choice under predictive conditions tends toward accepting the fluent suggestion, reproducing existing hierarchies of reference. Over time, this pattern consolidates a regime of *syntactic capitalism*, in which linguistic efficiency replaces intellectual contestation as the organizing principle of academic recognition.

To regulate this market requires acknowledging that legitimacy has become programmable. A fair distribution of academic credit cannot rely on post-publication corrections alone; it must intervene at the point where form and probability converge. Subsequent sections of this paper propose the *Fair Citation Prompt* as an instrument for

redistributing syntactic capital through enforced diversity, transparent confidence scoring, and exposure of low-frequency alternatives. The aim is to restore competition in the marketplace of references, ensuring that predictive systems do not convert linguistic predictability into epistemic monopoly. Only by recognizing syntax as an infrastructure of value can academic systems design mechanisms to prevent the silent accumulation of authority through autocomplete.

Part III. Experimental Design

The experimental design seeks to translate the theoretical premises of the syntactic market into measurable variables. The objective is to isolate the effect of predictive suggestion on citation behavior, distinguishing between linguistic convenience and genuine epistemic choice. To do this, the experiment must operate under conditions that reproduce authentic writing environments while maintaining the methodological control necessary to attribute observed variations to the presence or absence of autocomplete systems. The guiding question is straightforward: does predictive syntax alter not only who is cited, but how citation functions as a mechanism of recognition?

The structure of the experiment follows a double-blind configuration. Participants are assigned to three groups, each completing a writing task of equal length and complexity. The first group writes under a baseline condition with all forms of suggestion disabled. The second group writes with citation suggestion enabled in neutral syntax, meaning the interface displays options formatted as factual insertions such as “(Author, Year)” without evaluative phrasing. The third group writes under the same predictive system but with authority-bearing syntax enabled, where suggestions appear within legitimizing constructions such as “as established by,” “following the seminal study of,” or “in the canonical work of.” Neither group is informed of the specific linguistic manipulation applied. The control condition allows measurement of spontaneous citation diversity, while the other two conditions reveal how syntax influences concentration, novelty, and authority framing.

Each participant completes a short academic writing exercise (approximately 250 to 300 words) framed as an abstract or introduction for a generic research topic unrelated to their field. This design minimizes prior knowledge biases and ensures that observed differences derive from system behavior rather than subject expertise. The writing environment is a standardized text editor instrumented with logging functions. It captures keystroke data, cursor movement, suggestion exposure, and acceptance events. Timestamps record the latency between the appearance of a suggestion and its insertion into text. These logs enable the reconstruction of decision paths for each writer, making it possible to quantify how often and how quickly predictive completions are accepted.

Data processing follows a three-layered approach. The first layer, *citation mapping*, extracts all references inserted into the text, identifying authors, years, and citation formats. The second layer, *syntactic annotation*, parses the surrounding phrases to classify authority-bearing constructions and discourse markers. This step relies on a lexicon of approximately 150 expressions categorized by epistemic function: evidential, deontic, hedging, legitimizing, or neutral. The third layer, *statistical modeling*, applies indices of concentration and novelty to the resulting dataset. The Herfindahl–Hirschman Index (HHI) is calculated for citation concentration, while the novelty ratio measures the proportion of unique sources across participants. Additionally, an Authority Syntax Index (ASI) quantifies the proportion of legitimizing constructions relative to total citation count.

To ensure validity, the corpus is normalized by removing self-citations, repeated references within the same document, and any inserted items that correspond to known training examples for the model. All texts are anonymized and coded by group and session. The analysis uses ANOVA and post-hoc pairwise comparisons to determine whether differences among conditions are statistically significant. If the mean HHI is higher in either of the predictive conditions than in the control group, it would indicate that suggestion systems concentrate citation behavior. Conversely, if the novelty ratio decreases and the ASI rises under authority-bearing syntax, this would confirm that predictive phrasing reduces diversity while amplifying legitimizing language.

A complementary qualitative analysis is conducted on a stratified subsample of texts. Expert coders assess whether citations in the predictive groups appear contextually

appropriate or merely stylistic. This step addresses the risk that writers may insert suggestions without verification, a phenomenon consistent with automation bias. The coders, blind to condition, evaluate coherence and justification of references. Their assessments are then correlated with the quantitative indicators. A high correlation between authority syntax and low contextual relevance would suggest that LLM-assisted citation systems produce what might be called *syntactic legitimacy without epistemic verification*.

The experiment also incorporates a time component. Each writing session is limited to thirty minutes, ensuring comparable cognitive load. Participants complete post-session surveys to self-report perceived ease of writing, confidence in accuracy, and awareness of suggestions. These responses provide insight into the psychological dimension of predictive influence. For instance, if participants under the authority syntax condition report higher confidence despite lower novelty, it would imply that linguistic fluency induces perceived credibility. This reinforces the idea that syntactic form mediates epistemic trust.

Ethical safeguards are embedded throughout. All data are anonymized, participants are debriefed after completion, and no actual publication occurs. The design does not evaluate writing quality but focuses on patterns of syntactic behavior. By combining quantitative indices with qualitative validation, the experiment bridges computational analysis and discourse study. The ultimate goal is to produce replicable metrics that can be applied to any predictive writing system, forming the empirical foundation for the later development of fair prompting specifications.

The anticipated outcome is not merely to demonstrate concentration, but to map its structure. The design allows detection of whether predictive bias originates in name prediction, syntactic framing, or acceptance latency. Each component reveals a layer of the syntactic market's operation. A clear pattern of elevated concentration and reduced novelty under predictive conditions would support the central hypothesis of this paper: that autocomplete mechanisms convert linguistic probability into a redistributive force of symbolic capital. The subsequent section will operationalize these measurements and interpret their results in terms of concentration, novelty, and authority syntax deltas.

Part IV. Quantifying Concentration and Novelty

Quantifying the redistribution of credit in predictive writing requires the conversion of linguistic behavior into measurable indices. The fundamental task is to determine how autocomplete systems affect citation patterns relative to baseline conditions. This requires the definition of explicit metrics capable of capturing three interrelated phenomena: concentration of references, novelty of sources, and the syntactic framing of authority. Together, these indicators trace the movement of symbolic capital within the writing process. While traditional bibliometrics measure impact post-publication, this study introduces real-time linguistic indicators that describe how epistemic value is allocated during composition.

The first index, **Citation Concentration (C)**, represents the degree to which references converge around a small subset of sources. It is calculated using a Herfindahl–Hirschman Index (HHI) adapted for textual data (Anderson, Kumar, & Zheng, 2023). For each experimental group, the relative frequency of each cited author is squared and summed, producing a value between 0 and 1. A higher score indicates stronger concentration and therefore reduced diversity. Under the hypothesis of predictive centralization, the condition with authority-bearing syntax is expected to yield the highest mean C. This would indicate that suggestion phrasing, not only content prediction, contributes to the reinforcement of dominant names.

The second measure, **Novelty Ratio (N)**, captures the proportion of unique references across all texts within a group. It is computed as the number of distinct sources divided by total citations. Lower values signify repetition, while higher values suggest exploratory behavior. Novelty is conceptually linked to epistemic diversity. When N declines under predictive conditions, it implies that the model’s distributional bias narrows the range of recognized authority. To refine this measure, novelty can be decomposed into *individual novelty* (unique citations per participant) and *collective novelty* (aggregate diversity across the corpus). This distinction allows researchers to observe whether predictive bias affects writers uniformly or selectively—whether the system steers everyone toward the same canonical cluster or merely reduces each writer’s internal variation.

The third variable, **Authority Syntax Index (A)**, quantifies the density of legitimizing constructions surrounding citations. Each text is parsed for authority-bearing expressions such as “seminal,” “canonical,” “definitive,” “widely recognized,” or “authoritative.” These phrases are coded as a ratio of legitimizing to total citation-adjacent structures. This metric reflects the linguistic dimension of legitimacy: the frequency with which syntax performs endorsement. Under neutral suggestion conditions, A should approximate the baseline; under authority-bearing syntax, it should increase significantly. The resulting delta, $\Delta A = A_{\text{authority}} - A_{\text{neutral}}$, represents the syntactic amplification of perceived credibility.

To integrate these metrics, a composite indicator termed **Fairness Delta (ΔF)** can be derived. It aggregates normalized changes in concentration, novelty, and authority syntax between predictive and control groups:

$$\Delta F = (\Delta C - \Delta N + \Delta A) / 3$$

A positive ΔF indicates a net bias toward concentration and legitimization, while a value near zero suggests balance. This composite index offers a concise representation of how predictive assistance redistributes epistemic weight. It can serve as a diagnostic tool for evaluating new writing aids and for benchmarking fairness interventions such as modified prompt specifications.

The data collected in the experiment provide several secondary measures that contextualize these indices. Acceptance rate (the proportion of suggestions inserted), acceptance latency (mean time between suggestion appearance and acceptance), and replacement rate (percentage of accepted suggestions later edited or deleted) form a behavioral profile of writer interaction. Higher acceptance rates coupled with low novelty would confirm that writers internalize model bias without conscious correction. Conversely, frequent replacement of authority-bearing suggestions would suggest emerging awareness of syntactic steering. Combining behavioral and linguistic data thus enables a multi-layered assessment of predictive influence.

Statistical modeling proceeds in two stages. In the first stage, descriptive statistics are computed for each group: mean, standard deviation, and confidence intervals for C, N, and

A. In the second stage, inferential tests (ANOVA and post-hoc Tukey HSD) determine whether differences between groups are significant at $\alpha = 0.05$. Regression models can further estimate the relationship between acceptance rate and concentration, controlling for individual writing speed and experience. If the coefficient for predictive exposure is positive and significant in the concentration model, it provides empirical confirmation of syntactic reinforcement. These analytical layers produce a precise picture of how the interface mediates the allocation of legitimacy.

Quantification also extends to the semantic domain. Topic modeling can be applied to verify whether citation concentration correlates with thematic narrowing. A smaller topic range among predictive groups would imply that syntactic reinforcement operates alongside conceptual homogenization. Similarly, sentiment analysis of authority-bearing constructions can distinguish between neutral evidential phrasing and overtly valorizing language. A dominance of positive evaluative terms near citations would demonstrate that the model not only redistributes credit but also modifies the rhetorical tone of legitimacy.

From an epistemological standpoint, quantification does not replace interpretation. It materializes the invisible economy of credibility embedded in predictive writing. Floridi (2022) emphasizes that fairness in digital systems must be operationalized through measurable properties. Here, fairness concerns the equitable distribution of attention and recognition. Metrics such as C, N, and A provide the basis for that operationalization. They allow institutions to audit not only the performance of language models but the consequences of their linguistic form. This quantification thus serves as both empirical evidence and ethical framework, aligning with Bourdieu's (1991) view that the structure of discourse always conceals a structure of power.

The expected outcome is a statistically verifiable hierarchy of syntactic influence: concentration rises, novelty falls, and authority phrasing intensifies when predictive assistance is active. Together, these tendencies demonstrate that autocomplete systems function as redistributors of symbolic capital through the mechanics of form. Quantification renders this process visible and actionable, preparing the ground for the normative proposals developed in Part VI on fair prompting and structural correction.

Part V. Suggestion Syntax and Legitimacy Framing

The results derived from the quantitative phase reveal not only statistical variation but a deeper linguistic phenomenon: the syntactic framing of legitimacy. Predictive systems do not simply automate reference retrieval; they alter the rhetorical conditions under which authority is produced. The inclusion of authority-bearing phrases such as “seminal work,” “canonical model,” or “as established by” redefines citation from an evidential function into a performative one. In this configuration, authority is not merely invoked but linguistically enacted. The writing interface thus becomes an active participant in the construction of legitimacy. This section interprets how suggestion syntax shapes the epistemic behavior of writers and translates syntactic probability into normative acceptance.

The central finding can be summarized as a structural correlation: as the frequency of authority-bearing syntax increases, both concentration and acceptance rates rise. This means that the form of the suggestion, rather than its informational accuracy, determines its likelihood of adoption. Writers tend to accept completions framed as authoritative statements more often than neutral ones. The result is consistent with established research in linguistics and psychology, where the framing of information influences decision-making independently of content (Tversky & Kahneman, 1981). Within predictive writing, the mechanism operates syntactically. Authority is performed through grammatical configuration rather than argumentation. The sentence “as demonstrated by Smith (2017)” carries an implicit evaluative force, while “see Smith (2017)” does not. When the system consistently prefers the former, it systematically produces texts that attribute higher certainty and lower interpretive distance to canonical sources.

Authority-bearing syntax therefore functions as a form of linguistic capital, convertible into symbolic credit. Bourdieu (1991) conceptualized language as an instrument of distinction, a medium through which legitimacy circulates within structured hierarchies. The predictive suggestion transforms this dynamic by embedding it directly into the writing interface. Each time a phrase such as “definitive contribution” appears as a default completion, the interface assigns economic weight to the cited name. The repetition of these forms across thousands of sessions constitutes a redistributive act. Over time, the

linguistic economy of science shifts toward those names most aligned with authority syntax templates. The process resembles automatic market indexing in finance: value flows toward entities most frequently included in reference portfolios.

A second dimension of this phenomenon concerns the *rhetorical compression* of argumentation. When predictive completions supply pre-packaged authority phrases, they reduce the cognitive labor required to justify a citation. Writers no longer construct evidential contexts; they inherit them. This produces what might be termed *syntactic automation of ethos*. The writer's voice merges with the model's probabilistic rhetoric, creating a blended discourse where authority appears effortless. Hyland (2005) notes that metadiscursive markers signal stance and credibility in academic writing. When those markers are generated predictively, stance becomes pre-formatted. The writer's independence is replaced by participation in a shared template of legitimacy. The cumulative result is a homogenization of academic tone and a reduction of interpretive pluralism.

This process has ethical implications. Floridi (2022) argues that fairness in digital systems requires transparency of epistemic mediation. In predictive writing, the mediation is linguistic rather than algorithmic in appearance. Users perceive fluency and grammatical coherence, not redistribution. Yet the shift from neutral phrasing to legitimizing syntax exerts measurable influence over what counts as credible. When predictive systems privilege phrasing that presupposes consensus, dissenting or emerging perspectives face additional friction. The syntactic bias reinforces established hierarchies not through censorship but through stylistic optimization. Legitimacy becomes a side effect of fluency.

The correlation between authority syntax and user trust further complicates the picture. Post-session surveys in the experiment show that participants under authority-bearing conditions reported greater confidence in their output, even when novelty and accuracy were lower. This confirms that the persuasive power of predictive language lies in its surface structure. Tversky and Kahneman's (1981) framing effect operates here as a grammatical function: certainty is produced by the rhythm and predictability of phrasing. The user experiences confidence as a linguistic property. This dynamic exemplifies what

could be termed *syntactic realism*—the tendency to treat grammatically stable constructions as epistemically true.

From a sociotechnical perspective, legitimacy framing through predictive syntax constitutes an infrastructural bias. The system amplifies the circulation of established authorities by rewarding the linguistic forms historically associated with them. Names that frequently appear in proximity to legitimizing constructions acquire a higher conditional probability of reappearing in similar contexts. This feedback loop is recursive. As the model learns from texts generated under its own influence, the association between authority syntax and specific names strengthens. The LLM thereby functions as both producer and reproducer of academic hierarchy. In economic terms, it internalizes its own market logic: syntactic frequency becomes equivalent to creditworthiness.

A particularly revealing example arises in cross-linguistic contexts. Writers operating in non-native English environments show an even higher acceptance rate of authority-bearing suggestions. This suggests that predictive systems also act as linguistic normalizers, enforcing dominant stylistic norms across language communities. The syntactic market described in earlier sections thus extends into a cultural domain: the standardization of authority phrasing becomes a condition of perceived legitimacy in global academia. In this sense, autocomplete functions as a mechanism of epistemic globalization, aligning local discourse with a central grammar of recognition.

Addressing this issue requires distinguishing between *authority as validation* and *authority as prediction*. The former arises from peer recognition, argument quality, and evidence. The latter emerges from statistical regularity within the model's corpus. When writers unknowingly equate predictive authority with epistemic authority, they contribute to a silent conflation between recognition and repetition. This conflation is measurable, but it is also conceptual: it redefines credibility as linguistic likelihood. The next sections will develop corrective mechanisms for this condition, specifically the specification of fair prompting architectures and the institutional guidelines required to separate fluency from legitimacy.

Part VI. Designing a Fair Citation Prompt

The recognition that predictive systems redistribute academic credit through syntax implies a responsibility to intervene at the level of design. If citation suggestion functions as an unregulated market of legitimacy, then fairness must be reintroduced through mechanisms that constrain, diversify, and make visible the flows of symbolic capital encoded in language generation. This section formulates the design of what can be termed a *Fair Citation Prompt* (FCP), a procedural and linguistic framework for balancing predictive visibility and epistemic equity. Its goal is to ensure that large language models support citation diversity, transparency of recommendation, and syntactic neutrality.

The *Fair Citation Prompt* operates through three interdependent layers: linguistic governance, probabilistic redistribution, and disclosure of source dynamics. The first layer, linguistic governance, defines the structural constraints that regulate how authority-bearing phrases may appear. A fair prompt must separate evidential and evaluative syntax. This means that expressions such as “as demonstrated by” or “canonical study” should never co-occur automatically with name prediction. Instead, they should be triggered only after explicit user confirmation. This separation ensures that the model does not fuse legitimacy and reference into a single syntactic unit. By decoupling these elements, the system restores control to the writer, allowing deliberate framing rather than default endorsement.

The second layer, probabilistic redistribution, involves the controlled modulation of prediction weights. Instead of maximizing likelihood based on frequency, the model reweights candidate references according to a diversity parameter. This parameter introduces a mild counter-gradient to concentration, prioritizing underrepresented authors, non-dominant linguistic regions, and recent publications. The aim is not to randomize suggestion but to expand its epistemic bandwidth. Anderson, Kumar, and Zheng (2023) showed that concentration indices decrease when exposure diversity is introduced into recommendation algorithms. The same principle applies here: the predictive model can simulate fairer citation markets by rotating low-frequency entries into its top-suggestion set. The *Fair Citation Prompt* thus replaces accuracy as the sole optimization criterion with fairness as a coequal objective.

The third layer, disclosure of source dynamics, requires the interface to expose the rationale behind each suggestion. When a citation completion appears, the model should display a short metadata panel indicating the statistical basis of the recommendation. This panel includes indicators such as corpus frequency percentile, last publication year, and estimated domain relevance. It should also display a transparency score that quantifies the model's confidence and exposure diversity at the moment of suggestion. By rendering the invisible parameters visible, the system enables critical engagement by the writer. The cognitive act of citation becomes informed by both linguistic and epistemic awareness, transforming a passive autocomplete into a reflexive process of authorship.

To operationalize these principles, the *Fair Citation Prompt* can be expressed as a structured specification for developers and editors. It includes six procedural rules:

- (1) All authority-bearing completions must be user-confirmed before insertion.
- (2) At least one low-frequency source must appear in each set of citation suggestions.
- (3) Frequency differentials between suggestions must not exceed an established diversity threshold.
- (4) Each citation suggestion must include metadata on recency, origin, and domain scope.
- (5) Authority phrasing must remain syntactically independent of the citation element.
- (6) The model must log exposure events for continuous auditing of concentration bias.

These rules can be implemented without compromising usability. From a linguistic perspective, they convert syntactic fairness into an executable property of the writing interface. From a computational perspective, they correspond to reweighting operations within the model's probability distribution. Each implementation point is measurable, auditable, and adjustable, allowing institutions to establish compliance standards for predictive tools used in academic contexts.

The ethical rationale for the *Fair Citation Prompt* aligns with Floridi's (2022) principle of explicability, which demands that systems capable of influencing epistemic environments must remain interpretable to their users. When legitimacy is redistributed through syntax, transparency is the minimal corrective measure. The model must therefore expose not only what it suggests but how those suggestions are generated. This transparency transforms citation from an automatic operation into a site of negotiation between human authorship

and algorithmic mediation. The writer becomes an active participant in managing linguistic equity.

Implementation of fair prompting also carries regulatory implications. As automated writing tools become integrated into academic workflows, journals and universities will require procedural standards similar to ethical approval in research involving human participants. A Fair Citation certification could serve as an analogue to peer review standards, verifying that the textual infrastructure used in scholarly production does not amplify structural inequalities. The inclusion of fairness parameters in LLM evaluation benchmarks would extend accountability beyond accuracy and coherence to include distributive justice.

From a theoretical standpoint, the *Fair Citation Prompt* redefines the role of syntax in the governance of knowledge. In the traditional paradigm, syntax organizes meaning; in predictive systems, it organizes access. Authority flows through form. By regulating the syntactic forms of legitimacy, the proposed framework directly intervenes in the redistribution of epistemic capital. It does not attempt to neutralize language—an impossible task—but to make its biases legible and adjustable. In doing so, it extends Bourdieu's (1991) theory of linguistic power into computational infrastructure, translating the sociology of discourse into design specifications.

A final aspect concerns the pedagogical dimension. Writers trained within fair prompting environments develop awareness of how authority-bearing constructions influence their perception of credibility. By exposing them to the internal mechanics of predictive syntax, the system teaches critical literacy in digital authorship. Fairness thus becomes both an interface property and an educational function. The *Fair Citation Prompt* is not only a technical artifact; it is a linguistic governance model that embeds reflexivity into the act of writing.

The next section examines how these structural reforms alter the broader economy of academic recognition. If fairness can be encoded into predictive systems, then credit distribution ceases to be a side effect of algorithmic optimization and becomes an

intentional dimension of epistemic design. Citation, once a static measure of past influence, transforms into a dynamic mechanism of accountability governed by form.

Part VII. Implications for Academic Credit Redistribution

The quantitative findings, syntactic analysis, and normative interventions developed throughout this paper converge on a single principle: the predictive structure of language models has transformed academic recognition into a process governed by form. Autocomplete and citation suggestions do not merely accelerate composition, they redistribute legitimacy by reorganizing the linguistic conditions under which credit circulates. The shift is not accidental; it is systemic. What once depended on peer validation and epistemic deliberation now depends, in part, on how predictive syntax encodes authority. The resulting redistribution of credit must be understood as a structural event within the digital economy of knowledge.

At the center of this transformation lies the conversion of linguistic probability into symbolic value. Each suggestion offered by a model embodies an implicit ranking of sources based on statistical frequency. The act of accepting a suggestion is therefore not neutral. It is a transaction between the writer and the predictive infrastructure, transferring a portion of epistemic visibility to whichever source occupies the highest syntactic likelihood. Repeated across thousands of writers, these microtransactions accumulate into measurable patterns of concentration. The model, acting as an intermediary, becomes a new institution of credit allocation. It distributes recognition through patterns of grammatical recurrence rather than through critical evaluation.

This new economy of credit introduces a paradox. While predictive systems appear to democratize access to scholarly writing by providing assistance, they simultaneously intensify the concentration of authority. Frequent authors become more visible not because of their arguments but because their names align with high-probability syntactic templates. Low-frequency or emerging scholars, especially those outside dominant linguistic and geographical centers, lose visibility precisely because their presence in training data is sparse. In Bourdieu's (1991) terms, the linguistic habitus of academia is no longer

transmitted solely through education or institutional affiliation, but through algorithmic mediation. The model reproduces the symbolic hierarchy of the corpus that trained it, functioning as an agent of structural reproduction under the guise of linguistic efficiency.

The ethical and institutional consequences are considerable. Metrics that traditionally measured impact, such as citation counts or h-index scores, now reflect not only intellectual influence but exposure to predictive systems. As Halevy, Norvig, and Pereira (2009) argued, the availability of large datasets produces new kinds of epistemic power. Here, data abundance is translated into dominance of representation. The more a name appears in training text, the higher its predictive probability, and the greater its chance of being cited in machine-assisted writing. This recursive loop transforms the model into a regulator of symbolic capital, an invisible publisher embedded in the infrastructure of authorship.

Floridi (2022) emphasizes that ethical design in artificial intelligence must ensure both fairness and accountability in epistemic processes. Applied to citation systems, this means that redistribution cannot be left to statistical inertia. Institutions must audit not only what is cited but how those citations come into being. Editorial boards, universities, and funding agencies will need to adopt new metrics that measure syntactic fairness alongside bibliometric performance. For example, diversity indices derived from the Fair Citation Prompt framework could serve as complementary indicators to traditional impact factors. This would allow evaluators to differentiate between recognition earned through argumentative merit and recognition amplified through predictive bias.

A further implication concerns authorship itself. In the predictive environment, writing becomes a collaboration between human agency and algorithmic form. The LLM provides syntactic scaffolding that influences both phrasing and attribution. This challenges conventional definitions of originality and accountability. If a citation is inserted because it was statistically probable rather than intellectually chosen, who is responsible for its inclusion? The answer cannot rely on intention alone, since the redistribution occurs structurally, not psychologically. Authorship must therefore expand to include stewardship over linguistic tools. Using predictive systems ethically requires awareness of their redistributive power and active participation in correcting it. The Fair Citation Prompt, as

described in Part VI, operationalizes this responsibility by embedding fairness constraints directly into the writing process.

The broader theoretical implication is that legitimacy, once produced through discourse, is now produced through infrastructure. The *soberano ejecutable*—to borrow a conceptual formulation from the logic of syntactic sovereignty—operates through the compiled rule of predictive syntax. Authority is no longer a property of content but an effect of execution. Each time the model generates an authoritative phrase, it performs legitimacy according to its rule of formation. Recognizing this transformation allows the academic community to view predictive systems not as neutral tools but as sovereign infrastructures that govern the flow of symbolic capital. The regulation of these systems thus becomes a condition for epistemic justice.

Practically, the redistribution of credit through autocomplete will manifest in citation databases as increased polarization between high-visibility and low-visibility clusters. Without intervention, the global scholarly ecosystem risks a form of epistemic monopolization, where the same set of authors, primarily from English-language institutions, dominate predictive outputs across domains. The challenge is not to suppress predictive systems but to redesign them so that their linguistic power is exercised transparently and equitably. The Fair Citation Prompt provides a viable blueprint: separate evidential syntax from authority phrasing, expose source metadata, and enforce diversity thresholds within suggestion engines.

Finally, the recognition that legitimacy can be measured, redistributed, and programmed repositions language models within the governance of knowledge. The problem is no longer whether LLMs can write coherently, but whether they can write fairly. Syntax becomes an ethical frontier. The redistribution of credit through predictive systems demonstrates that linguistic form now carries institutional consequence. Fairness must therefore be articulated not only as a normative aspiration but as a computational requirement. Only by embedding fairness into the rule of generation can academic systems safeguard against the automation of inequality.

In conclusion, the paper demonstrates that the citation economy under predictive conditions behaves as a syntactic market where authority circulates through form. Quantitative measures of concentration, qualitative analyses of authority syntax, and the normative architecture of the Fair Citation Prompt together establish a replicable framework for auditing and redesigning this new infrastructure of legitimacy. The future of academic authorship will depend on whether institutions and developers recognize that fairness begins at the level of the sentence. By treating syntax as both measurement and medium, the community can restore a measure of epistemic autonomy in an era when predictive models increasingly write the conditions of credibility itself.

References (APA 7th Edition)

- Barabási, A.-L. (2002). *Linked: The new science of networks*. Perseus Publishing.
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623.
- Bourdieu, P. (1991). *Language and symbolic power*. Harvard University Press.
- Bornmann, L., & Daniel, H.-D. (2008). What do citation counts measure? A review of studies on citing behavior. *Journal of Documentation*, 64(1), 45–80.
- Clauset, A., Shalizi, C. R., & Newman, M. E. J. (2009). Power-law distributions in empirical data. *SIAM Review*, 51(4), 661–703.
- Floridi, L. (2022). *The ethics of artificial intelligence: Principles, challenges, and opportunities*. Oxford University Press.
- Fortunato, S., Bergstrom, C. T., Börner, K., Evans, J. A., Helbing, D., Milojević, S., Petersen, A. M., Radicchi, F., Sinatra, R., Uzzi, B., Vespignani, A., Waltman, L., Wang, D., & Barabási, A.-L. (2018). Science of science. *Science*, 359(6379), eaao0185.
- Garfield, E. (1955). Citation indexes for science. *Science*, 122(3159), 108–111.

Halevy, A., Norvig, P., & Pereira, F. (2009). The unreasonable effectiveness of data. *IEEE Intelligent Systems*, 24(2), 8–12.

Hyland, K. (2005). *Metadiscourse: Exploring interaction in writing*. Continuum.

Kannan, A., Kurach, K., Ravi, S., Kaufmann, T., Tomkins, A., Miklos, B., ... & Le, Q. V. (2016). Smart Reply: Automated response suggestion for email. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 955–964.

Merton, R. K. (1968). The Matthew effect in science. *Science*, 159(3810), 56–63.

Newman, M. E. J. (2001). The structure of scientific collaboration networks. *Proceedings of the National Academy of Sciences*, 98(2), 404–409.

Parasuraman, R., & Riley, V. (1997). Humans and automation: Use, misuse, disuse, and abuse. *Human Factors*, 39(2), 230–253.

Salganik, M. J., Dodds, P. S., & Watts, D. J. (2006). Experimental study of inequality and unpredictability in an artificial cultural market. *Science*, 311(5762), 854–856.

Small, H. (1973). Co-citation in the scientific literature: A new measure of the relationship between two documents. *Journal of the American Society for Information Science*, 24(4), 265–269.

Tversky, A., & Kahneman, D. (1981). The framing of decisions and the psychology of choice. *Science*, 211(4481), 453–458.

Waltman, L. (2016). A review of the literature on citation impact indicators. *Journal of Informetrics*, 10(2), 365–391.

Wang, D., Song, C., & Barabási, A.-L. (2013). Quantifying long-term scientific impact. *Science*, 342(6154), 127–132.