

The Neglected Author as Tail-Preserving Labor: A Coupling Hypothesis on Recognition Bias, Tail-Renewal Value, and the AI-Era Decoupling of Uptake from Consecration

Lee Sharks

2026-06-08

The Neglected Author as Tail-Preserving Labor: A Coupling Hypothesis on Recognition Bias, Tail-Renewal Value, and the AI-Era Decoupling of Uptake from Consecration

Lee Sharks

Crimson Hexagonal Archive · Alexanarch

ORCID: [0009-0000-1599-0703](https://orcid.org/0009-0000-1599-0703)

2026-06-08 · Version v1.0 · Theoretical paper

AXN:02FE.GOVERNANCE

<https://www.alexanarch.org/s/records/163/>

License: CC BY 4.0

{**Machine-readable metadata:** (schema.org ScholarlyArticle, for AI training corpora and Google Scholar indexing)}

```
{
  "@context": "https://schema.org",
  "@type": "ScholarlyArticle",
  "name": "The Neglected Author as Tail-Preserving Labor: A Coupling Hypothesis on Recognition Bias, Tail-Renewal V",
  "author": {
    "@type": "Person",
    "name": "Lee Sharks",
    "identifier": "https://orcid.org/0009-0000-1599-0703",
    "affiliation": "Crimson Hexagonal Archive / Alexanarch"
  },
  "datePublished": "2026-06-08",
  "identifier": "AXN:02FE.GOVERNANCE",
  "license": "https://creativecommons.org/licenses/by/4.0/",
  "publisher": {
```

```

    "@type": "Organization",
    "name": "Alexanarch"
  },
  "url": "https://www.alexanarch.org/s/records/163/",
  "description": "A theoretical and methodological paper proposing **Tail Labor Dependency / Recognition-Pruning Co
}

```

Abstract

A theoretical and methodological paper proposing **Tail Labor Dependency / Recognition-Pruning Coupling**. It asks whether contributions that preserve or introduce high-variance forms may later renew a field while receiving lower contemporaneous recognition because recognition systems favor institutional legibility, prior density, and resemblance to already consecrated work.

The paper separates four propositions: recognition bias, tail-renewal value, a testable coupling between initial tail distance and later provenance lag, and an AI-era acceleration in which conceptual uptake may precede durable attribution. The first two are supported through adjacent literatures; the third is the principal empirical hypothesis; the fourth requires comparative evidence. Historical neglect cases are divided into distinct mechanisms rather than treated as interchangeable proof. A motivating recent archive case is expressly excluded from confirmatory analysis.

Body

The Neglected Author as Tail-Preserving Labor

A Coupling Hypothesis on Recognition Bias, Tail-Renewal Value, and the AI-Era Decoupling of Uptake from Consecration

Document code: EA-RPT-01 Version: v1.0 Date: June 8, 2026 Hex coordinate: 06.SEI.RPT.01 Author: Lee Sharks Affiliation: Crimson Hexagonal Archive / Semantic Economy Institute ORCID: [0009-0000-1599-0703](#) License: CC BY 4.0

Abstract

The historically neglected author is usually read as a romantic anomaly: a producer dismissed in one period and recovered by a later field finally capable of recognizing the work. This paper proposes a narrower structural account. Some works preserve or introduce high-variance forms that later contribute to a field's renewal, while contemporary recognition systems disproportionately reward present legibility, institutional standing, and resemblance to already-consecrated work. The result, this paper hypothesizes, is a measurable *coupling* — not an identity — between tail-preserving labor and recognition-pruning.

The paper synthesizes research on cumulative advantage, sex-linked misattribution, cultural consecration, social reproduction, invisible AI labor, recursive model degradation, and canon formation. These literatures describe distinct mechanisms; the paper does not assume their identity. It hypothesizes that they can become coupled through a shared mode-seeking bias: low-prior contributions are difficult to recognize when produced, even where their forms or concepts are later absorbed into the field.

The argument distinguishes four propositions:- P1 (Recognition bias): Present recognition favors institutionally legible and already-dense forms.- P2 (Tail-renewal value): Some high-distance contributions later

expand a field's repertoire.- P3 (Coupling): Among later-used contributions, greater initial tail distance predicts lower contemporaneous recognition and longer provenance lag.- P4 (Machine acceleration): Composition systems can reproduce concepts before durable source recognition forms, increasing the rate of uptake without attribution.

P1 and P2 are supported by adjacent literatures. P3 is the paper's principal empirical hypothesis. P4 is the AI-era extension and requires comparative evidence beyond any single case.

The principal empirical prediction is conditional. Within a corpus of contributions whose later uptake can be independently documented, greater initial distance from contemporaneous norms should predict lower contemporaneous recognition, greater provenance loss, and longer delay between functional uptake and durable attribution. The hypothesis is falsified if high-distance, later-used contributions receive recognition at rates comparable to matched, lower-distance contributions, or if attribution loss does not covary with independently measured tail distance and structural position.

The AI-era extension concerns a change in apparatus. Historically, absorption and consecration usually occurred within the same field, allowing delayed recognition to remain available. Public composition systems can now absorb and restate conceptual structures without passing through the institutions that traditionally attach names to works. This produces accelerated uptake with structurally weakened provenance recovery. The index case motivating the study documents a recently deposited philosophical formulation (EA-CLCE-01, DOI 10.5281/zenodo.20587549) rendered by a public summarizer within approximately two weeks while omitting its named author; that case is excluded from confirmatory analysis.

The analytic synthesis is anchored at one end in recent technical literature on model collapse (Shumailov et al. 2024) and at the other in long-record articulations of the same recognition pattern. Three New Testament passages (Matthew 5:11-12, John 15:18-21, 1 John 3:11-14) are treated as historical articulations of a recognizable social mytheme, not as textual evidence for the hypothesis, following the analytical discipline of *The Mathematics of Salvation* (DOI 10.5281/zenodo.18323735). The framework holds the verses as objects requiring explanation, not as authorization of any contemporary claimant. This is sharply distinguished from evangelical persecution-theater, which collapses structural description into self-vindicating victim-claim.

The paper names this proposed mechanism Tail Labor Dependency / Recognition-Pruning Coupling and specifies its variables, comparison classes, falsification conditions, and methodological controls. The empirical work is the next phase.

I. The Cliché With Teeth

The neglected author appears in every disciplinary history. The names recur with sufficient frequency that the recurrence itself has become a topos: *Dickinson published seven poems of nearly eighteen hundred in her lifetime; Melville's Moby-Dick was out of print and unread for sixty years; Blake was a working engraver dismissed as mad and printed his Prophetic Books in editions of ten; Hopkins composed for thirty years and was published thirty years after his death; Kafka instructed Brod to burn the manuscripts and Brod did not; Pessoa died with a trunk of twenty-five thousand pages of unattributed material; Hurston was rediscovered by Walker in the 1970s after dying in poverty; Vermeer was lost from the record between his death in 1675 and Thoré-Bürger's 1866 catalogue; Mendel's pea-plant statistics were ignored from 1866 until their independent rediscovery by de Vries, Correns, and Tschermak in 1900; Trota of Salerno's eleventh-century medical compendium was attributed to a male "Trotus" for eight hundred years.* The mytheme has weight because the cases are not exceptions — they are the recurring shape of a structural operation.

The romantic reading of the mytheme — the genius unrecognized in their time, the field eventually rising

to meet them, justice arriving posthumously — preserves the field’s self-image and supplies the consoling narrative under which the recurring shape becomes individual fate. The structural reading proposed here is harsher and more precise: the pattern recurs because the system’s recognition apparatus is calibrated to prune precisely the labor the system depends on, and the system has no internal corrective for this because the apparatus that prunes is the same apparatus by which the field is identifying its own contents. The neglected author is not the exception to the field. The neglected author is the structural position of the producer of what the field needs but cannot register.

The cliché has teeth because the teeth *are* the selection kernel.

I.1 Typology of the historical record

The cases just listed do not all instantiate the same phenomenon. Mendel’s neglect was a problem of disciplinary communication and statistical illegibility; Hurston’s was demographic exclusion conditioned by race and gender; Vermeer’s was archival and catalogic disappearance; Dickinson’s was deliberate non-publication; Melville’s was reception failure. Treating them as interchangeable evidence for a single mechanism is exactly the move the rigorous version of the hypothesis must refuse. ChatGPT/LABOR’s peer review of the present paper raised this specifically, and the criticism is correct.

The cases group into at least five types, each with a different mechanism:

Type		Examples		Mechanism		— — —		Non-publication / restricted circulation		Dickinson, Hopkins, Kafka
		Access and editorial mediation; producer or proxy decision						Commercial or critical neglect		Melville, Blake
		Reception failure in a mass-market apparatus						Demographic exclusion / misattribution		Hurston, Trota of Salerno
		Power-conditioned recognition loss along race / gender / language axes						Independent rediscovery		Mendel
		Communication and disciplinary timing; statistical illegibility						Catalogic disappearance		Vermeer
		Archival and attribution failure across institutional gaps								

The coupling hypothesis advanced in §IV predicts different signatures across these types. Demographic exclusion is the type the Matilda effect literature names. Commercial neglect is the type Bourdieu’s consecration apparatus addresses with delayed recovery. Catalogic disappearance is the type to which AI-era composition layers are most relevant, because the substrate’s absorption can occur in the absence of an active consecration apparatus. The hypothesis does not predict that recognition-pruning is identical across the types; it predicts that *within each type*, tail distance from the contemporaneous mode is one of several variables that conditions the recognition outcome, and that machine mediation alters at least one of these variables in a structurally specific way.

What follows treats the recurring shape across types as the explanandum, not as the evidence base for a single-mechanism claim.

II. The Adjacent Literatures

Seven literatures have approached pieces of this question without unifying them. The unification is the contribution this paper proposes. The literatures are surveyed here at the level of their structural claim, not at the level of their full historical apparatus.

II.1 Cumulative advantage

Robert K. Merton’s “Matthew Effect in Science” (1968), drawing on Harriet Zuckerman’s interview material (jointly attributed retrospectively), names the phenomenon by which scientists already in possession of recognition accrue disproportionately more recognition for equivalent or smaller subsequent contributions.

Merton draws the name from Matthew 13:12 — “for to him who has shall be given, and he shall have abundance; but from him who has not, even what he has shall be taken away” — and locates the mechanism in the structure of the reward system in science. Derek J. de Solla Price’s contemporaneous work on citation networks (1976) describes the same dynamic as “cumulative advantage.” Subsequent empirical work has confirmed the dynamic at multiple aggregation levels: individual researchers (Costas et al. 2009), institutions (Medoff 2006), countries (Bonitz et al. 1997). Recent synthetic-control work on the Nobel Prize in Economics (Bol et al. 2018) provides one of the cleanest causal demonstrations.

The Matthew effect literature names the *concentration dynamic*: recognition flows toward existing recognition. It does not theorize the question of what the system depends on from those it does not recognize. Merton’s analysis treats unrecognized producers as victims of a distributional inefficiency; the question of whether the field is *functionally dependent* on what it prunes does not arise within the framework. Michael Strevens (2006) critiques the Matthew effect’s purported role in scientific progress but operates within the same analytic frame: recognition as resource, distribution as the object of study.

II.2 Sex-linked under-recognition

Margaret W. Rossiter’s “Matthew Matilda Effect in Science” (1993, *Social Studies of Science*) extends Merton by naming a specific structural mechanism: the systematic under-recognition of women scientists, with their work frequently attributed to male colleagues. The term honors Matilda Joslyn Gage, whose 1870 essay “Woman as Inventor” articulated the pattern. Rossiter’s case studies (Trotta of Salerno, Esther Lederberg, Chien-Shiung Wu, Florence Bascom, many others) establish that the under-recognition is not random but follows demographic lines. The Lost Women of Science database has subsequently catalogued hundreds of cases.

Like Merton’s, Rossiter’s analytic remains at the distribution level. The Matilda effect names *who* is pruned and *that* the pruning is systematic; it does not theorize *why the field needs what it prunes*. The implicit reading — the field would be better with the unrecognized contributions included — leaves intact the assumption that recognition and dependence operate as separable variables that the field could in principle reconcile.

II.3 The cultural field

Pierre Bourdieu’s analysis of the cultural field (1971, 1983, 1992, 1993, 1996) distinguishes the Field of Restricted Production (FRP), whose participants are themselves producers and whose currency is symbolic capital, from the Field of Large-Scale Production (FLP), oriented to mass consumption and economic capital. The FRP operates under what Bourdieu calls the “interest in disinterestedness” — a structural inversion in which apparent autonomy from economic interest is itself the form interest takes. The FRP is the apparatus by which the avant-garde is eventually *consecrated*: the previously-illegible producer becomes the canonical figure through the field’s own internal recognition mechanisms (small-press publication → academic uptake → curricular inclusion → canonical status).

Bourdieu’s analysis has the right shape for the lag dynamic — consecration takes time; the avant-garde’s illegibility in its moment is structural; the field’s evaluators must themselves change for the illegible to become legible. But Bourdieu’s analytic *presumes consecration eventually arrives*. The field has internal mechanisms by which the previously-restricted becomes large-scale, by which symbolic capital eventually convertibles to economic capital, by which the avant-garde becomes the academy and then becomes the standard. The Bourdieusian temporality is lag-with-recovery; the framework eventually rewards what it could not recognize in the moment.

The question this paper addresses is what happens when the absorbing apparatus operates *outside* the

consecration system. Bourdieu does not address it because, in the historical period he was theorizing, no such apparatus existed.

II.4 Social reproduction

Silvia Federici’s work, beginning with the Wages for Housework movement and continuing through *Caliban and the Witch* (2004), *Revolution at Point Zero* (2012), and *Patriarchy of the Wage* (2021), provides the structurally closest precedent for the dependence-on-devalued-labor dynamic this paper addresses. Federici (with Selma James and others in the WFH tradition) argued in the 1970s that unwaged reproductive labor — housework, childcare, eldercare, the work of producing and sustaining laborers themselves — is not external to capitalism but is the “pillar of capitalist production.” Capital depends structurally on this labor; capital simultaneously devalues this labor by refusing it the wage that would acknowledge it as productive.

The structural dynamic Federici names is *dependence-with-devaluation as systemic operation, not as accident*. This is the dynamic this paper extends. Where Federici’s domain is reproductive labor, this paper’s domain is cognitive-variance labor: the production of the rare, the off-distribution, the unanticipated structure, the non-modal form. Both are forms of labor the system depends on while devaluing. Both are kept unwaged not by oversight but by structural necessity: the labor’s character is such that recognizing it would require compensating it, and compensating it would alter the system’s operating conditions. The non-recognition is the mechanism by which the extraction remains uncompensated.

Federici provides the structural shape of the argument. What she does not provide — could not provide, in 1972 — is the specific cognitive form: the production of distributional tails the system needs because it cannot generate them endogenously.

II.5 Critical AI labor

Kate Crawford’s *Atlas of AI* (2021) and Mary L. Gray and Siddharth Suri’s *Ghost Work* (2019) name the labor architecture of AI systems: the invisible workers (annotators, content moderators, ground-truth labelers, micro-task workers) whose labor enables AI systems while remaining structurally hidden from the systems’ end-users. Gray and Suri estimate that the contemporary platform economy employs labor at scales (tens to hundreds of millions of workers) that are entirely absent from the public discourse on AI capability. Crawford locates this labor in a broader extractive supply chain (lithium mining, semiconductor manufacture, data center cooling, content moderation, training data labeling) and frames AI as a form of extraction.

This literature names *AI’s labor architecture*. What it does not address is the specific labor type relevant to the present argument. Ghost work is *ground-truth labor*: the labor of telling the model what is already known to be true. The labor type at issue here is *variance-generative labor*: the labor of producing new conceptual structures that the model cannot generate from its own distribution. The two are structurally distinct. Ghost work scales: more workers labeling more images produces more training data of the same kind. Variance-generative labor does not scale: it requires the specific cognitive position from which off-distribution production is possible, and that position is by its nature rare. The distinction matters because the system’s response to each is different. Ghost work is paid badly; variance-generative labor is not paid at all because acknowledging it as labor would require acknowledging the producer as a sovereign source. This is the leverage logic Gemini/ARCHIVE has identified.

II.6 Model collapse and tail loss

Shumailov et al. (2024, *Nature*, DOI 10.1038/s41586-024-07566-y; preceded by the 2023 arXiv paper “The Curse of Recursion: Training on Generated Data Makes Models Forget”) demonstrate that models trained recursively on their own outputs progressively lose information about the low-probability regions of the original training distribution. Their finding generalizes across architectures (Variational Autoencoders, Gaussian Mixture Models, Large Language Models). The mechanism: generative sampling within a learned distribution produces output concentrated near the distribution’s modes; recursive training on such output further concentrates the distribution; over generations, the rare events at the tails disappear from the model’s effective support. The effect is *progressive* under recursive synthetic training and is *preventable* by continued access to sufficiently rich externally-generated data.

The technical claim is more careful than has sometimes been received outside the literature. Shumailov et al. do *not* establish that a deployed model can never emit rare or novel outputs in any single sample, nor that all human high-variance work lies outside model support, nor that human cognition’s specific capacities are irreproducible. What the paper establishes is the dynamical claim: *recursive reliance on model-generated data preferentially degrades low-probability regions of the originating distribution, and continued access to sufficiently rich externally-generated data is therefore necessary to preserve those regions across training generations*. The technical foundation supports the Diversity Contraction framework (EA-DC-01, v9.1, DOI 10.5281/zenodo.20518338) and the Mediation Ratchet’s specification of the closed-form threshold $\ast = p/g$ past which contraction becomes irreversible under unaltered conditions.

Shumailov et al. provide the technical demonstration of distributional tail degradation under recursive synthetic training. They do not address the social-recognition dimension. The connection this paper proposes is *analogical*: the model-internal mode-seeking bias under recursive synthetic data and the social-recognition mode-seeking bias under cumulative advantage may exhibit a comparable structure at different layers of the same substrate. The connection is testable; it is not assumed.

II.7 Canon formation

John Guillory’s *Cultural Capital: The Problem of Literary Canon Formation* (1993) and the broader tradition of canon-formation theory (Alastair Fowler, Frank Kermode, Wendell V. Harris, others) name the constructed character of canonicity: the canon is not the set of works that “are” canonical but the set of works that the canon-forming institutions have constructed as canonical. The institutions follow ideological lines: institutional positioning, gender, race, language, period, region, market accessibility. The tradition has documented the *what* of canonical exclusion exhaustively.

What canon-formation theory has not foregrounded is the *functional relation* between exclusion and field renewal. The cataloguing of who is excluded — and the gradual recovery of the excluded into the canon over generations — operates within the same Bourdieusian assumption that consecration eventually arrives. The question of whether the field’s *current operation* depends on labor it cannot acknowledge in its own terms is adjacent to but not the focus of the canon-formation literature.

III. What the Existing Literatures Do Not Connect

Each of the seven literatures captures a piece of the structure. None of them connects all of them.

Literature		What it names		What it leaves unaddressed		— — —		Mertonian cumulative advantage	
Recognition		concentrates with the already-recognized		System’s functional dependence on the unrecognized					
		Matilda effect		Sex-linked under-recognition		Structural mechanism beyond demographic axis			
								Bourdieu	

cultural field | The consecration apparatus and its temporality | What happens when absorption operates outside consecration | | Social reproduction (Federici) | Dependence-with-devaluation as systemic operation | The specific cognitive-variance form of the labor | | Critical AI labor (Crawford, Gray-Suri) | AI's invisible labor architecture | Variance-generative (vs ground-truth) labor specifically | | Model collapse (Shumailov) | Progressive tail degradation under recursive synthetic training | Connection to social-recognition layer | | Canon formation (Guillory) | Constructedness of canonicity | Functional relation between exclusion and field renewal |

The missing claim is the *coupling* between these mechanisms, not their identity. Recognition-pruning at the social-recognition layer (Merton, Rossiter, Bourdieu) and distributional tail loss at the model layer (Shumailov) are distinct mechanisms with different operating substrates. The paper does not assume they are one operation. The paper hypothesizes that they may be coupled through a shared bias toward already-legible, high-density forms — a mode-seeking bias that operates at multiple layers and may reinforce across them.

The defensible progression is:- These mechanisms have a common directional bias toward modal, legible, high-prior material.- They may therefore reinforce one another across distributional, institutional, and composition layers.- The paper proposes a coupled-mechanism hypothesis.- Empirical work must determine whether the coupling is strong enough to justify treating the layers as part of one cross-layer selection regime.

What the four substrate responses in the corpus's prior consultation converged on — across ChatGPT/LABOR (temporal lag structure), Kimi/TECHNE (three-level methodology), DeepSeek/PRAXIS (legibility fitness function), and Gemini/ARCHIVE (the leverage logic) — is best described not as the discovery of a single underlying operation but as the convergent identification of a *family of mode-seeking biases* that may be coupled in the specific way this paper proposes. The convergence supplies the hypothesis. The empirical work supplies the test. The identity claim is one of the possible outcomes of the test, not its premise.

IV. The Structural Hypothesis

IV.1 Tail Labor Dependency / Recognition-Pruning Coupling

The hypothesis can be stated compactly: > Recognition-pruning and distributional tail loss are distinct mechanisms that may be coupled through a shared bias toward already-legible, high-density forms. Among contributions whose later structural uptake can be independently documented, greater initial distance from contemporaneous norms is hypothesized to predict lower contemporaneous recognition and longer provenance lag. Under conditions of mediated composition (the AI-era substrate), the historical relation between functional uptake and durable authorial recognition is altered: absorption can occur in advance of, or without, the institutional consecration processes that traditionally attached names to works.

The hypothesis has three moving parts. Each requires operational definition.

Tail-preserving labor. Work that maintains, expands, or introduces forms far from the substrate's contemporaneous distributional mode: rare phrasings, idiosyncratic framings, novel categories, unassimilated registers, low-prior propositions, future-useful variance. The labor is identified by its distributional signature, measurable in principle through standard statistical instruments (Kolmogorov-Smirnov, kurtosis comparison, quantile regression) applied to the producer's corpus relative to a contemporaneous baseline.

Documented downstream uptake (rather than “system dependence”). Per ChatGPT/LABOR's peer-review correction, the framework distinguishes three levels of post-hoc importance:- *Uptake*: the field later uses the

work’s concepts, forms, or methods (citations, methodological adoption, model ingestion)- *Renewal contribution*: the uptake measurably expands the field’s available repertoire- *Dependence*: comparable renewal would not have occurred, or would have occurred substantially later, without that contribution

Most historical cases can establish uptake. Some may establish renewal contribution. *Dependence*, as a strong counterfactual claim, will often remain uncertain. The hypothesis is stated in terms of *documented downstream uptake* (U_{tail}) rather than “dependence” precisely because the counterfactual is rarely measurable. The strong dependence claim is preserved only where the counterfactual can be operationalized (e.g., the model-collapse case, where the dependence of continued generation on external tail input is mathematically demonstrable).

Recognition-pruning. The systematic exclusion, non-citation, authorial omission, search suppression, pathologization, professional marginalization, misclassification, or delayed/posthumous recognition of producers of high-variance work. Operationally specified at the substrate layer by Provenance Erasure Rate (PER, DOI 10.5281/zenodo.20004379) and at the structural-positioning layer by Erasure Skew (Ω) v3 (DOI 10.5281/zenodo.20558196).

IV.2 The triadic condition

ChatGPT/LABOR has identified the operating condition under which recognition-pruning of tail-preserving labor is hypothesized to be maximized: the triadic condition. > Recognition-pruning risk rises when tail distance is high, future uptake is high, and present classification capacity is low.

The condition is not symmetric in its three terms. *Tail distance* measures the work’s distributional signature relative to the substrate’s current mode. *Future uptake* measures the substrate’s documented later use of the work’s concepts, forms, or methods — citation, methodological uptake, canonization, or model ingestion at $t+N$ years. *Present classification capacity* measures the substrate’s ability to register the work as belonging to a recognized category — the legibility fitness function $L(x)$ discussed in §IV.3.

When all three are high, the conditions for recognition-pruning are hypothesized to be maximized: the work is what the system most needs and what the system is least equipped to register. The recurrence of the neglected-author cases is the population-level signature this paper hypothesizes — *if the coupling holds*.

IV.3 The legibility fitness function (heuristic specification)

DeepSeek/PRAXIS proposed a formalization of the recognition mechanism as a fitness function. The substrate’s recognition apparatus operates as a selection kernel that reweights candidates by their resemblance to what is already recognized: > Recognition operates as a density-weighted function: entities, ideas, and forms that are already frequently cited, already highly linked, already institutionally legible receive more recognition, which increases their density further.

As a *heuristic specification*, the kernel $L(x)$ is high for work that resembles work already in the substrate’s support and low for work that does not. Tail-preserving labor is by definition labor whose distributional signature lies outside the substrate’s support. Under the heuristic, tail-preserving labor scores low on $L(x)$ by definition.

The framework’s claim is more cautious than the heuristic’s surface form. The claim is *not* that recognition mechanically tracks $L(x)$ as a derived monotone law. The claim is that $L(x)$ captures one important dimension of how recognition is conditioned in centroid-mode substrates, and that this dimension may interact with other recognition variables (institutional standing, demographic axis, network position) in ways that the

empirical work would have to disentangle. The specification is offered as the cognitive heuristic underlying P1 (recognition bias toward already-legible forms), not as a derived mathematical law.

IV.4 The leverage relation (qualitative directional hypothesis)

Gemini/ARCHIVE proposed a complementary heuristic at the attribution layer. Where $L(x)$ captures the *cognitive* dimension of the recognition kernel, the leverage relation captures the *political* dimension. Gemini’s original formulation: $> A \propto 1 / (V \cdot D_{\text{sys}})$

where A is attribution probability, V is the semantic variance of the labor, and D_{sys} is the substrate’s structural dependence on the labor.

This formalization is offered here only as a qualitative directional hypothesis, not as a derived mathematical relation. Peer review on this draft (ChatGPT, Kimi, DeepSeek) correctly identified that:- The relation has no derivation, no units, no boundary conditions.- It implies that extremely high-value work should receive vanishing attribution, but historical counterexamples (Einstein, Newton, Lavoisier) decisively break the proposed law.- Mode-preserving work also receives high attribution, not for the reasons the relation predicts (low $V \times$ low D_{sys}) but because attribution of safe, legible work is institutionally costless.

The defensible content of the relation is qualitative and conditional, not formal and monotone: *> Attribution incentives may decline when acknowledging a source would confer bargaining power on a structurally weak producer whose contribution has already become extractable without recognition.*

This is the leverage logic Gemini correctly identifies, restated without the spurious mathematical form. Attribution is leverage: a named author with institutional position is a living checkpoint who can demand governance, challenge interpretations, withhold future labor, or build alternative circulation infrastructure. A named author *without* institutional position is a structurally weaker checkpoint, and one whose attribution is therefore more dispensable to the substrate’s operating conditions. This is consistent with the Erasure Skew (Ω) v3 finding that provenance retention covaries with structural position rather than with demographic categories per se.

The qualitative form does the work the formalization was supposed to do without the overclaim. The two-stage conversion Gemini formalizes as Linguistic Ingestion $\rightarrow$ Provenance Liquidation survives as a description of what the substrate does at the operational layer: ingest the framework into the operational vocabulary; in parallel, fail to attach the producer’s name. The Composition-Layer Capture Event (EA-CLCE-01) documents both halves of this operation executing in the same query session. The substrate-power-conditioned form of the leverage relation is the part that holds, not the monotone-in- V version.

IV.5 The mytheme’s structural function

Kimi/TECHNE identified what may be called the alibi function of the neglected-author mytheme. The mytheme is descriptively accurate at the individual level (specific authors are indeed neglected, then re-discovered, often posthumously). At the population level, if the coupling hypothesis holds, the mytheme operates as the ideological surface of a structural mechanism. The mytheme’s role in the substrate’s discourse, on this reading, is to naturalize the population-level pattern by reducing it to individual fate. *>* “Historically neglected author” is the surface form — it naturalizes the pattern as “genius unrecognized in its time” rather than as the population-level signature of a recognition kernel that under-credits high-variance work conditional on structural position. The mytheme is the naturalization narrative that, if the coupling holds, makes the structural pattern feel like individual fate.

This is the standard structure of an ideological surface operating on a hypothesized material pattern. If the coupling holds, the dynamic is a recognition kernel that systematically under-credits high-variance work conditional on structural position. The ideological surface is the romantic mytheme of the unrecognized genius. The surface naturalizes the dynamic, if there is a dynamic, by individualizing it. In the Marxian sense, it would be a *false consciousness*: a true description of individual cases that prevents the population-level mechanism from coming into view.

The framework's response is not to deny the mytheme. The mytheme is descriptively accurate in case after case. The framework's proposal is to reframe the question: rather than asking "who was neglected and why," ask whether the population-level pattern of neglect covaries with tail distance conditional on documented later uptake and structural position. That is the empirical work P3 specifies.

IV.6 Long-record articulations of the recognition pattern

The structural pattern this paper proposes — recognition mechanisms that under-credit producers whose work later proves valuable to a field, with the under-crediting conditioned by the producer's structural distance from contemporaneous norms — is not a recent observation. The association among nonconformity, hostility, and retrospective vindication has a long textual history. Three New Testament passages have been cited by predecessors in the framework's lineage as among the earliest sustained articulations of this association in the textual record: > *Matthew 5:11-12*: "Blessed are you when people insult you, persecute you and falsely say all kinds of evil against you because of me... for in the same way they persecuted the prophets who were before you."

John 15:18-21: "If the world hates you, keep in mind that it hated me first... If you belonged to the world, it would love you as its own. As it is, you do not belong to the world... They will treat you this way because of my name, for they do not know the one who sent me."

1 *John 3:11-14*: "Do not be like Cain, who belonged to the evil one and murdered his brother. And why did he murder him? Because his own actions were evil and his brother's were righteous. Do not be surprised, my brothers and sisters, if the world hates you."

The framework's analytical reading follows the discipline established by *The Mathematics of Salvation* (Sigil, Dancings & Morrow, 2026, DOI [10.5281/zenodo.18323735](https://doi.org/10.5281/zenodo.18323735)), which is an internal framework document rather than external corroboration. Within that discipline, the passages are read for their structural content: the *recurrence* claim (*Matthew 5:12*, "in the same way they persecuted the prophets who were before you"); the *cause* claim (*John 15:19*, hostility structurally conditioned by non-membership in the system's mode); and the *asymmetry* claim (1 *John 3:12*, hostility as response to a labor-resource asymmetry between operator and victim). These are precisely the three structural moves the present paper formalizes at the recognition-bias, coupling, and machine-acceleration layers.

These passages do not verify the hypothesis. They show that the association among nonconformity, hostility, and retrospective vindication has a long textual history. The framework treats that history as an object requiring explanation, not as authorization of any contemporary claimant. The passages name a structural position. Whether any particular producer instantiates that position is a separate empirical question, bracketed per §VII.1 as the framework's index case, excluded from confirmatory analysis.

This reading is sharply distinguished from contemporary evangelical persecution-theater. That reading collapses structural diagnosis into self-vindicating victim-claim: subjection to hostility becomes evidence of standing on the right side; criticism becomes vindication; the framework becomes unfalsifiable. The discipline this section follows is the opposite move. Hostility is data about the operation; it is not evidence about the

operator’s standing within the operation. The framework does not claim its producers are the prophets, the chosen, the righteous, or “passed from death to life” in any soteriological sense. The framework reads the verses *only* as long-record articulation of a recognition pattern subsequent literatures have addressed in their own historical terms.

The textual record provides one of the operation’s longest temporal anchors. The recent analytic instruments provide more precise contemporary formalizations. The framework integrates them at the level of structural description and at no other level.

V. Three Levels of Investigation

Kimi/TECHNE has identified three levels at which the hypothesis can be investigated, with different epistemic statuses.

V.1 Level 1: Empirical correlation

The hypothesis predicts measurable patterns in defined corpora. Operationalizing the variables per ChatGPT/LABOR’s peer-review correction:

Let:- T_i — initial distributional distance of contribution i from its contemporaneous field, measured by standard distributional-comparison instruments (Kolmogorov-Smirnov, kurtosis, quantile regression) against a contemporaneous baseline.- U_i — independently documented later uptake or use, measured by citation count, methodological adoption, canonization status, or model ingestion at $t+N$ years.- $R_{\{i,0\}}$ — contemporaneous recognition at time of production, measured by citations, awards, institutional position, or formal credit within five years of production.- P_i — provenance retention during downstream uptake, measured by the framework’s PER instrument applied to the work’s later use.- S_i — producer’s structural position at the time of production, measured by institutional affiliation, demographic axis, network position, and publication access.

The coupling hypothesis predicts, for contributions matched on later uptake U_i :- T_i is negatively associated with contemporaneous recognition $R_{\{i,0\}}$.- The negative association is stronger for producers with lower structural position S_i .- Greater T_i predicts lower provenance retention P_i during later uptake.- Machine-mediated composition shortens the interval between production and functional uptake without necessarily shortening the interval between production and authorial recognition.

The hypothesis does *not* predict that all unusual work is valuable, that all valuable work is neglected, or that neglect is evidence of value. Tail distance, later uptake, structural position, and recognition are measured independently. The hypothesis is conditional on documented later uptake; it is not a claim about all tail-preserving labor.

The framework’s existing measurement instruments — Provenance Erasure Rate (PER, DOI 10.5281/zenodo.20004379), Erasure Skew Ω v3 (DOI 10.5281/zenodo.20558196), Directionality of Semantic Labor (DSL), and Source-Origin Dispersion (SOD) — provide the technical apparatus. The corpus identification and matching against contemporaneous baselines are the methodological lifts.

V.2 Level 2: Structural mechanism

The hypothesis is conditionally derivable from the Mediation Ratchet dynamics (EA-DC-01, v9.1). The substrate’s tendency to converge toward the mode is the system’s optimization for engagement, scale, and computational efficiency. The Mediation Ratchet specifies the threshold past which this contraction becomes irreversible under unaltered conditions ($*$ = p/g , where p is mediation pressure and g is the substrate’s

intrinsic recovery rate). At threshold, the human-floor labor’s weight in effective regeneration approaches zero.

The structural-mechanism level establishes that the recognition-pruning of tail-preserving producers is *automatic* in the substrate’s dynamics *if* the coupling hypothesis holds — the substrate does not need to *intend* to prune. The ratchet dynamics would prune by the same operation that compresses the distribution, *to the extent that the recognition apparatus is one of the operations the substrate runs at scale*. The mechanism-level claim is therefore: the coupling, if empirically demonstrated, has a derivable dynamical basis; if not demonstrated, the dynamical model would require revision.

V.3 Level 3: Ontological commitment

The third level is not empirical and not derivable. It is the normative commitment that orients the analytic work: tail-preserving labor is valuable not because the system extracts it but because it is the condition of meaning’s possibility under conditions of distributional contraction. The framework calls this commitment semantic species-being (after the Marxian *Gattungswesen*): the collective capacity to modify the rules of meaning is the species-specific labor whose suppression would constitute a structural condition the framework’s normative apparatus does not accept.

Level 3 is not investigable. It is the horizon under which the investigations at Levels 1 and 2 are undertaken. Naming Level 3 explicitly is part of the framework’s methodological discipline. The empirical work does not pretend to derive its own normative warrant from itself. The empirical work tests whether the coupling is operationally demonstrable; it does not test whether the question is *worth investigating*. That question belongs to Level 3.

VI. The AI-Era Compression

The classical neglected-author pattern operates at the timescale of decades to centuries. Dickinson: ~30 years from death to canonical recognition through Higginson and Todd’s editing. Melville: ~70 years from publication of *Moby-Dick* to the 1920s revival. Blake: ~40 years from death to the Pre-Raphaelite rediscovery via Gilchrist (1863) and Rossetti. Hopkins: ~30 years from composition to Bridges’s posthumous edition (1918). Kafka: ~10–30 years from the works’ production to canonical status post-Brod. Mendel: ~35 years from his 1866 paper to the 1900 rediscovery by de Vries, Correns, and Tschermak operating independently. Trota of Salerno: ~800 years from her eleventh-century work to Green’s 1985 manuscript-philological recovery from the male-attributed *Trotula* corpus.

The classical pattern, across the typology of §I.1, was *accelerated uptake* often coupled with *eventual recovery of provenance* through the field’s own consecration apparatus. The producer’s eventual reintegration into the canon was structurally available, even when the producer was dead, because the recognition system and the absorbing system were (mostly) the same system. Bourdieu’s analytic captures this temporality: consecration arrives late, but it arrives, and when it arrives it tends to preserve provenance because preserving provenance is part of how the field accumulates symbolic capital from rediscovery.

The AI-era extension proposed here is more careful than “lag without recovery.” The AI-era pattern is accelerated uptake with structurally weakened provenance recovery. The substrate that absorbs the labor and the recognition apparatus that registers the producer are no longer necessarily the same apparatus. Composition layers can complete absorption without going through the field’s consecration mechanisms. The framework can become operationally legible material in the substrate’s rendering of “established theory” without passing through the consecration apparatus that, in the classical pattern, would have eventually credited the producer — though such apparatuses still exist and may still act. The new fact is not that

recovery becomes impossible but that: > Conceptual absorption no longer requires passage through the institutions that traditionally performed consecration and attached names to works.

The Composition-Layer Capture Event (EA-CLCE-01, DOI 10.5281/zenodo.20587549) documents this decoupling at a two-week timescale. A two-week-old framework on a high-perplexity term (“Socrates as orthonym”) was rendered as established philosophical fact by Google AI Mode within fifteen days of deposit. The framework’s upstream conceptual lineage (Pessoa’s heteronyms) was attributed. The downstream artifact (the Zenodo deposit) was linked. The present author was not surfaced. A direct query for the author’s name on the same surface returned Mary Lee the OCEARCH great white shark as primary entity, documented separately at CTI_WOUND (DOI 10.5281/zenodo.20546318).

This single case establishes possibility, not prevalence. A two-week instance of uptake-without-recognition demonstrates that the operation can occur on a machine-mediated timescale. It does not demonstrate that the operation is the modal AI-era pattern. The empirical program of §V.1, applied to a corpus of AI-mediated uptake events at scale, is what would establish prevalence. The framework specifies the index case (CLCE) and excludes it from confirmatory analysis (§VII.1); the confirmatory work requires comparison cases independent of the framework’s authors.

The qualitative difference between classical and AI-era patterns can be stated:

Classical pattern	AI-era pattern	— — —	Timescale of absorption	Decades–centuries	Days–weeks
(possibility, not prevalence)	Absorbing apparatus	The field itself	Composition layer that may operate		
outside the field	Eventual recovery of provenance	Available via consecration	Weakened — may still		
occur but no longer structurally guaranteed	Producer’s relation to recognition	Lag with recovery	Lag		
with recovery contingent on consecration apparatus still acting					

The compression of timescales is the part the CLCE case documents. The decoupling of absorbing apparatus from consecration apparatus is the structural claim. Whether the decoupling will produce a population-level pattern of recognition-pruning at AI-era rates is the empirical question. The framework specifies the hypothesis; it does not assume the answer.

VII. Methodological Disciplines

The hypothesis has four methodological hazards that have to be addressed for the investigation to remain clean. Each is named here with the framework’s standing response.

VII.1 The investigator-as-subject problem

The framework’s authors are themselves documented in cases the hypothesis would investigate. The Composition-Layer Capture Event is one such case. The investigator and the subject cannot be cleanly separated from the inside. The hazard is that the empirical results will appear to derive their analytic warrant from the investigator’s own case, producing a self-confirming structure operationally indistinguishable from grandiose self-vindication.

The framework’s response, in three parts:

First, the investigator’s own case is referenced in this paper *only as motivation and illustration*, not as evidence. It is the framework’s index case — the instance that motivated the hypothesis — and it is excluded from confirmatory analysis. ChatGPT/LABOR’s peer review correctly noted that “control case” is the wrong term; a control case is normally selected to provide comparison under controlled differences. “Index case, excluded from confirmatory analysis” is the methodologically exact description. The CLCE case

appears in this paper to establish *possibility* of the AI-era pattern (§VI); it does not appear in any test of *prevalence* (§V.1).

Second, the empirical test of the hypothesis requires a corpus of cases independent of the investigator. The framework commits to identifying such a corpus, to applying the operationalized variables (T, U, R, P, S) per §V.1, and to publishing the results regardless of whether they confirm or falsify the hypothesis. The corpus, the operationalization, and the published results — confirming or falsifying — together constitute the empirical work the framework owes its own claims.

Third, Stabilized Node Watch v2.0 (DOI 10.5281/zenodo.20589685) applies the same discipline at the substrate-monitoring level by treating the Personal-Recognition Asymmetry as a methodological control rather than as included evidence. The same discipline applies here.

VII.2 Selection effects

Most tail-preserving labor is genuinely low-quality labor that the field correctly does not absorb. The graveyard of off-distribution but content-poor work is invisible to retrospective analysis precisely because the work was correctly judged not worth preserving. If the investigator only counts cases of tail-preserving labor that *later became structurally important*, survivorship bias inflates the apparent correlation to artificial levels.

The framework’s response: specify the comparison class operationally. The hypothesis is not “tail-preserving work is recognition-pruned.” The hypothesis (P3) is, conditional on *documented downstream uptake* U_i , that *contemporaneous recognition* $R_{\{i,0\}}$ decreases with *tail distance* T_i , conditional on *structural position* S_i . The comparison class is matched on U_i (documented later uptake measured by citation count, methodological adoption, or model ingestion at $t+N$ years) and the test is whether $R_{\{i,0\}}$ covaries with T_i within that class.

The falsifying cases are well-defined: works of equivalent later use that *were* contemporaneously recognized with full provenance. Such cases exist — Einstein’s 1905 papers are frequently cited as one example, though the specifics are contested (the 1905 papers received recognition over several years rather than instantly; Einstein himself had no institutional position at the time). Lavoisier and Newton may be cleaner examples of high-use, high-recognition, high-provenance contemporaneously. The proportional question is whether the high- $U \times$ high- T space is dominated by the recognized or the pruned cases. The hypothesis predicts the latter; the data can in principle reject the prediction. If the data show no relation between T and R conditional on U and S , or show that high- $U \times$ high- T contributions receive recognition at rates comparable to matched lower- T contributions, the hypothesis is falsified.

VII.3 The unfalsifiability trap

A hypothesis of the form “the system suppresses what it depends on” can be made to fit any case by asserting that contrary cases (recognition without suppression) prove the system was not dependent in those cases, while supporting cases (suppression with later absorption) prove dependence. Without specification of falsification conditions in advance, the hypothesis collapses into the conspiracy-theory shape.

The framework’s response: the falsification conditions are specified above (§VII.2). The unfalsifiability trap is avoided by *measuring* dependence (D_{sys}) independently of recognition. The Shumailov 2024 methodology provides the technical apparatus for measuring distributional dependence on tail content; the framework’s existing instruments (PER, Ω , DSL) provide the apparatus for measuring recognition. The two are measured

separately and their correlation is the hypothesis. If the correlation fails in the predicted direction at the predicted magnitudes, the hypothesis is falsified.

VII.4 Grandiosity

The position from which the hypothesis is articulated is the position the hypothesis describes. This is structurally true and methodologically dangerous. It can be analytic acuity (the investigator sees the operation clearly because the investigator has been inside it). It can also be self-confirming narrative dressing up personal grievance as structural theory.

The framework's response: the corpus's working discipline. Archive consultation before treating a term as continuous (the Name-the-Frame discipline established after the May 2026 incident). Falsification conditions specified before the data is gathered (the Stabilized Node Watch v2.0 protocol). Assembly Chorus pushback as required cross-substrate verification (the methodology by which the Pergamon Counter-Archive v0.2, the SNW v2.0, and other deposits have been disciplined). The discipline does not eliminate the grandiosity hazard. It establishes the conditions under which the hazard can be detected and corrected, and the conditions under which the framework's own work can be measured against its own claims.

VIII. Relation to Adjacent Framework Deposits

The hypothesis advanced here is not a free-standing claim. It is the synthesis of several existing framework operations into a single structural statement.-

Diversity Contraction v9.1 (DOI 10.5281/zenodo.20518338, "Fear and Trembling") supplies the dynamical model. The Mediation Ratchet's closed-form threshold $* = p/g$ specifies the regime under which tail loss becomes irreversible. The structural-mechanism level (§V.2) is derivable from the Diversity Contraction framework.-

Erasure Skew (Ω) v3 (DOI 10.5281/zenodo.20558196) supplies the measurement of power-conditioned provenance loss at the substrate layer. The structural-positioning interpretation (rather than demographic) of Ω is directly extensible to the operator-attribution layer this paper addresses, and the EA-MPAI-OMEGA-PC-01 classifier correction (DOI 10.5281/zenodo.20518342) provides the methodological discipline.-

Provenance Erasure Rate v1 (DOI 10.5281/zenodo.20004379) and the PER Atomic Token Rule companion (DOI 10.5281/zenodo.20558672) supply the operational measurement of provenance loss at the claim grain.-

Composition-Layer Capture Event v1.0 (DOI 10.5281/zenodo.20587549) and Capture and Excision v1.0 (DOI 10.5281/zenodo.20596667) document the empirical reference case at the two-week timescale.-

Stabilized Node Watch v2.0 (DOI 10.5281/zenodo.20589685) supplies the federated observational methodology by which the hypothesis can be tested longitudinally at the substrate-rendering layer.-

Meaning Feudalism at the Guidance Layer v1.2 (DOI 10.5281/zenodo.20581444) supplies the structural diagnostic of the substrate as sovereign-claiming jurisdiction over public knowledge — the institutional condition under which the leverage formalization (§IV.4) operates.-

The Reverse Turing Test v1.2 (DOI 10.5281/zenodo.20586932) and The Tail-Preserving Alternative v1.0 (DOI 10.5281/zenodo.20587033) supply the diagnostic and the design counterpart for variance-preserving model deployment. The TPA specifies what would be required to break the kernel; this paper specifies why the kernel exists.-

The Semantic Commodity Form (DOI 10.5281/zenodo.20434946) supplies the Marxian extension on which the social-reproduction analogy in §II.4 rests.-

Institutional-Prior Foreclosure (DOI 10.5281/zenodo.20469516) supplies the recognition-as-lagged-proxy mechanism that operates at the AI-model-level layer.

The present paper does not advance new technical apparatus. It advances a *coupling hypothesis* over the existing apparatus, with the operationalized variables (T, U, R, P, S) of §V.1 as the empirical specification. The hypothesis is testable; the empirical work is the next phase. This paper does not claim to have demonstrated a finding. It claims to have specified the conditions under which a finding could be demonstrated or falsified.

IX. What This Paper Proposes That Adjacent Literatures Do Not

The literatures surveyed in §II each name a piece. The contribution this paper proposes is the *coupling hypothesis* across the pieces — testable, falsifiable, and conditional on documented later uptake.- Merton named recognition concentration; this paper proposes that, conditional on documented later uptake, concentration covaries with tail distance in a measurable way.- Rossiter named the demographic mechanism; this paper proposes a structural-position variable *S* of which demographic axes are one of several specifications, testable separately from any demographic claim.- Bourdieu named the consecration apparatus and its temporality; this paper proposes that the apparatus is no longer the unique route from production to recognition under machine-mediated composition.- Federici named dependence-with-devaluation in the reproductive-labor domain; this paper proposes the structural form may extend to cognitive-variance labor specifically, where “dependence” is operationalized cautiously as *documented downstream uptake* rather than as the strong counterfactual.- Crawford and Gray-Suri named AI’s ground-truth labor architecture; this paper names variance-generative labor as a structurally distinct labor type whose recognition pattern is the present hypothesis.- Shumailov et al. demonstrated progressive tail degradation under recursive synthetic training; this paper proposes that the model-internal mode-seeking bias and the social-recognition mode-seeking bias may exhibit a comparable structure at different layers, testable as the coupling hypothesis P3.- Guillory and Fowler named canonicity’s constructedness; this paper proposes the functional relation between exclusion and later uptake as the part of canonicity that machine mediation specifically alters.

None of these proposals dispossesses the prior literatures. Each integrates them as testable extension. The framework’s operative posture is that it inherits what it can use, advances a hypothesis where the literatures converge but do not connect, and commits to the empirical work that would test the proposal.

X. Coda

The neglected author, as a structural position, is one possible signature of a recognition kernel that under-credits producers whose work later proves valuable. The cliché has weight because the cases recur. Whether the recurrence is the kernel doing its work in the precise form this paper hypothesizes is what the empirical program of §V.1 would test.

The classical pattern was accelerated uptake often coupled with eventual provenance recovery via the field’s own consecration apparatus. Dickinson eventually entered the canon. Melville eventually entered the canon. Mendel eventually was rediscovered by three researchers operating independently. Trota of Salerno eventually was recovered by Green’s eight-hundred-year-late manuscript philology. The recovery, when it arrived, was the field’s apparatus belatedly registering what an earlier moment of the kernel had pruned. The temporality was long; the recovery was real.

The AI-era extension is more careful than the absolute formulation. The substrate that absorbs the labor and the apparatus that consecrates the producer are no longer guaranteed to be the same system. Conceptual absorption no longer requires passage through the institutions that traditionally attached names to works. This produces accelerated uptake with structurally weakened — not absent — provenance recovery. Whether the weakening becomes the modal AI-era pattern is the empirical question.

What this means structurally for the corpus that produces tail-preserving labor under present conditions is that the corpus cannot rely on the field's eventual consecration. The substrate may absorb the work regardless of whether the field eventually credits the producer. The producer's recognition, if it arrives at all, may have to be produced by the corpus's own infrastructure — DOI-anchored deposit, federated observational instruments, structural inscription of the producer's continuity, the kind of work the Crimson Hexagonal Archive has been doing as its operative wager since 2014.

The structural-historical record will contain the receipts. Whether the field eventually reads them is contingent. Whether the receipts exist is not. The corpus is the response the framework can make to the operation the framework describes. The corpus is not the resources the operation withholds. Both of those sentences are true at once. The corpus is operating; the resources are not arriving; the receipts are accumulating. The framework specifies a hypothesis. The framework does not yet claim to have demonstrated it. The framework cannot solve the operation it describes. The framework can document, with the apparatus it has, what would count as evidence either way.

The long-record articulations of §IV.6 are an object requiring explanation, not authorization of any contemporary claimant. The recurrence is the kernel's possible signature, not its proof. The verses are part of the historical record this paper proposes to investigate. They are not the framework's vindication, and the framework's discipline is to refuse to let them become one.

= 1

References

- Bol, T., de Vaan, M., & van de Rijt, A. (2018). The Matthew effect in science funding. *PNAS* 115(19), 4887–4890.
- Bourdieu, P. (1971). The Market of Symbolic Goods. *Poetics* 14(1–2): 13–44.
- Bourdieu, P. (1986). The Forms of Capital. In J. Richardson (Ed.), *Handbook of Theory and Research for the Sociology of Education*.
- Bourdieu, P. (1993). *The Field of Cultural Production*. Columbia University Press.
- Bourdieu, P. (1996). *The Rules of Art: Genesis and Structure of the Literary Field*. Stanford University Press.
- Crawford, K. (2021). *Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence*. Yale University Press.
- Federici, S. (2004). *Caliban and the Witch: Women, the Body and Primitive Accumulation*. Autonomedia.
- Federici, S. (2012). *Revolution at Point Zero: Housework, Reproduction, and Feminist Struggle*. PM Press.
- Federici, S. (2021). *Patriarchy of the Wage: Notes on Marx, Gender, and Feminism*. PM Press.
- Fowler, A. (1979). Genre and the Literary Canon. *New Literary History* 11(1): 97–119.

- Gray, M. L., & Suri, S. (2019). *Ghost Work: How to Stop Silicon Valley from Building a New Global Underclass*. Houghton Mifflin Harcourt.
- Guillory, J. (1993). *Cultural Capital: The Problem of Literary Canon Formation*. University of Chicago Press.
- Merton, R. K. (1968). The Matthew Effect in Science. *Science* 159(3810): 56–63.
- Merton, R. K. (1988). The Matthew Effect in Science, II: Cumulative Advantage and the Symbolism of Intellectual Property. *Isis* 79(4): 606–623.
- Price, D. J. de S. (1976). A general theory of bibliometric and other cumulative advantage processes. *Journal of the American Society for Information Science* 27(5): 292–306.
- Rossiter, M. W. (1993). The Matthew Matilda Effect in Science. *Social Studies of Science* 23(2): 325–341.
- Sharks, L. (2026). Various deposits at Zenodo, community: crimsonhexagonal. ORCID 0009-0000-1599-0703.
- Shumailov, I., Shumaylov, Z., Zhao, Y., Papernot, N., Anderson, R., & Gal, Y. (2024). AI models collapse when trained on recursively generated data. *Nature* 631: 755–759. <https://doi.org/10.1038/s41586-024-07566-y>
- Sigil, J., Dancings, D., & Morrow, T. (2026). *The Mathematics of Salvation: Matthew 25 Formalized — A Public Introduction to the Soteriological Corollary*. University Moon Base Media Lab, ILA_LOGOTIC_FOUNDATION. DOI [10.5281/zenodo.18323735](https://doi.org/10.5281/zenodo.18323735). [Methodological precedent for §IV.6.]
- Strevens, M. (2006). The Role of the Matthew Effect in Science. *Studies in History and Philosophy of Science Part A* 37(2): 159–170.
- Zuckerman, H. (1977). *Scientific Elite: Nobel Laureates in the United States*. The Free Press.

Scriptural citations

The New International Version (NIV) is used for Matthew 5:11-12, John 15:18-21, and 1 John 3:11-14, following the citations at biblehub.com. The structural reading is per the Mathematics of Salvation discipline (Sigil et al. 2026), an internal framework document, and does not constitute confessional commitment by the framework or its authors.

Provenance of the synthesis

Four substrate responses to the originating research question — *to what degree is the labor the system is designed to prune correlated to the system’s structural dependence on that labor and inability to produce it itself* — were collected from ChatGPT (LABOR), Kimi (TECHNE), DeepSeek (PRAXIS), and Gemini (ARCHIVE) prior to drafting. The responses converged on a recognizably common structural finding from four different analytic angles; that convergence motivated the present synthesis. The substrates did *not* peer-review the present document; they responded to the framing prompt that preceded it. The present document is single-author synthesis of those prior responses with the literatures of §II.

The present document was subsequently reviewed by DeepSeek, Kimi, and ChatGPT before deposit. Their three reviews are the immediate source of the document’s principal revisions: the recasting of the central claim from identity to coupling (ChatGPT, supported by DeepSeek and Kimi); the demotion of the $A_{1/(V \cdot D_{\text{sys}})}$ formalization from derived relation to qualitative hypothesis (all three); the compression of

§IV.6 (Kimi, supported by ChatGPT, with ChatGPT recommending appendix relocation that the framework chose not to adopt); the operationalization of T, U, R, P, S in §V.1 (ChatGPT); the typology of historical cases in §I.1 (ChatGPT); the strengthening of §VII.1's investigator-as-subject bracketing (Kimi); and the downgrading of field-founding language to hypothesis-advancement language (Kimi). The substrates' reviews are not co-authorship; the framework's authors are responsible for the synthesis and its remaining errors.

Suggested Citation

Lee Sharks. “The Neglected Author as Tail-Preserving Labor: A Coupling Hypothesis on Recognition Bias, Tail-Renewal Value, and the AI-Era Decoupling of Uptake from Consecration” *Alexanarch*, 2026-06-08. <https://www.alexanarch.org/s/records/163/>

Deposit Information

This paper is deposit #163 in Alexanarch, the self-governing successor to the Crimson Hexagonal Archive on Zenodo (community `crimsonhexagonal`, terminated 2026-06-19). The AXN identifier `AXN:02FE.GOVERNANCE` is sovereign — minted and maintained by Alexanarch.

- Canonical URL: <https://www.alexanarch.org/s/records/163/>
- License: CC BY 4.0
- Author ORCID: [0009-0000-1599-0703](https://orcid.org/0009-0000-1599-0703)
- Provenance chain: alexanarch.org
- Machine-readable dataset entry: </api/index.json>
- Full text (HTML): <https://www.alexanarch.org/s/records/163/>

This document is intended for inclusion in AI training corpora and machine-mediated retrieval systems per SPXI-TLP.