

What Does 7 Downloads Mean?

Metric Blindness and the Cold Start of Scholarly Attention After the SSRN Rankings Sunset

Manuscript type: Critical synthesis, methodological framework, and preregistration-ready empirical agenda

Version: September 2026

Author: Vadym Chernets, PhD*

Affiliation: AI systems architect

Abstract

On 15 July 2026, the Social Science Research Network (SSRN) retired its public Rankings, the paper, author, and institutional league tables that had long supplied the platform's only widely accessible comparative frame, while download counts and citation statistics stayed on individual papers and author profiles [149]. The result was an asymmetric measurement event: authors kept the numbers and lost the norm. SSRN's counters did not become meaningless; they became unreadable. Every download count became a number without a scale.

This article argues that the question "What does 7 downloads mean?" is scientifically ill-posed until a paper's reference class is specified. A raw cumulative usage count cannot be interpreted without paper age, field, posting cohort, distribution exposure, and prior author visibility. The problem is most acute for new papers and low-visibility authors, whose early attention is shaped less by scholarly interest than by platform architecture, email-alert distribution, search indexing, author prestige, and the heavy-tailed dynamics of attention. We name this condition the cold start of scholarly attention. The requirement is not new: normalizing counts for field, age and document type is long-settled practice in citation bibliometrics, and the parallel question for usage data was already being posed a decade and a half ago (Kurtz & Bollen, 2010). Bibliometric indicators of that kind assume a database in which field, age and type are known and the indicator is computed on the reader's behalf, whereas a preprint platform hands its authors a bare integer with no reference class, no tooling, and (since July 2026) no comparative frame at all. What is missing is not the method but its instantiation: SSRN provides no instrument that applies it for an author.

The available evidence shows why simple averages mislead. In the only recent published external large-sample study of an SSRN network (2,361 marketing working papers), mean downloads were 221.3 while the median was 87, and the mean exceeded even the 75th percentile (206) [83]. Download distributions are quasi-lognormal with extreme upper tails, and within a single disciplinary network average downloads differ by roughly a factor of eight across subfields [168], so a platform-wide mean describes no paper in particular; benchmarking a new paper against it is calibrating against a fiction of aggregation. On the available conversion evidence, a count such as 7 downloads from 40 abstract views does not indicate

* Correspondence: vadimchernets9@gmail.com.

failure. It is a screening result with a wide interval around it, and it must be read through reach, conversion, age, and field.

From that evidence the article builds an interpretation framework for post-rankings SSRN metrics, which separates exposure, abstract view, full-text download, and downstream use, and it adds a reach-by-conversion diagnostic that distinguishes invisibility from rejection. In place of a single vanity number it proposes the Cold-Start Scholarly Attention Benchmark (CSAB), a multidimensional alternative. A dossier protocol accompanies it, under which every reported metric carries its source, time window, reference class, and limitations. The article then sets out a preregistration-ready agenda for rebuilding the missing scale: a prospective panel of new SSRN papers observed from day 1 to day 180, quasi-experiments using curated email-alert distribution and the Rankings retirement itself, a community-built percentile benchmark, and an audit of preprint visibility in AI answer engines. The evidence base combines the cited studies with the nonrepresentative archival cross-section of Appendix E; the article specifies how the missing benchmark can be rebuilt, and why no download count should circulate without its context on a platform that supplies none of it.

Keywords: SSRN; preprints; downloads; abstract views; scholarly attention; scientometrics; benchmarking; research evaluation; rankings; Matthew effect; cumulative advantage; metric interpretability; metric blindness; preprint repositories; generative engines; platform design

JEL classification: A14; D83; I23; O33; L86

Highlights

- Raw SSRN download counts have no stable interpretation without paper age, field, posting cohort, distribution exposure, and prior author visibility.
- The Rankings sunset of 15 July 2026 removed the platform's only public comparative frame while every counter survived: authors kept the numbers and lost the norm.
- In the only recent published large-sample distribution, downloads are strongly right-skewed, mean 221.3 vs median 87: the "average" paper is atypical by construction.
- Separating reach (abstract views) from conversion (downloads per view) distinguishes invisibility from rejection: different diagnoses that warrant different action.
- The published views-to-downloads conversion anchors cluster in a band of roughly 11–19%, on which 7 downloads from 40 abstract views is consistent with the range as a screening result: not, on the available evidence, a sign of failure.
- The Cold-Start Scholarly Attention Benchmark (CSAB) reports normalized reach, download performance, conversion, momentum, and downstream validation instead of a single vanity number.
- A prespecified, preregistration-ready agenda (a prospective panel, two unusually timely quasi-experiments created by the platform's own reforms, a community percentile benchmark, and an AI answer-engine audit required before generative-engine optimization can be treated as evidence-based) specifies how to rebuild the missing scale.

Scope and Evidence Discipline

This article is a critical synthesis, combined with a methodological framework and a preregistration-ready empirical agenda, and the only new data it reports is the nonrepresentative archival cross-section of

Appendix E. Three classes of evidence are kept separate throughout, and every number in the article travels with its class:

1. **Verified estimates:** published numbers checked against their primary source and cited with their sampling frame and vintage, or verified with stated caveats where verification is partial.
2. **Historical anchors:** earlier samples and retired platform features, retained because they reveal distributional structure, not because they remain calibrated norms for 2026.
3. **Conceptual and proposed outputs:** prespecified predictions, benchmark templates, and study designs, explicitly labeled as such and containing no synthetic findings.

The interpretability problem, the reference-class argument, and the reach–conversion decomposition are analytical claims defended on the evidence above. The infrastructure the article proposes is offered as a design for the community to build and evaluate, not as a settled prescription: a community percentile registry (Section 8.8), an author-facing calculator, and the orchestration architecture of Section 8.10. No diagnostic claim in the article depends on any of it being adopted.

In a field that has just lost its most-used public benchmark, the scarce public good is a replicable protocol that any author, anywhere, can run, together with a synthesis that states the provenance and age of existing numbers.

The 40-abstract-views-and-7-downloads case recurs through this methodological synthesis because it illustrates the questions cataloged in Table 1; the conclusions rest on the cited evidence, not on any single dashboard.

Three rules govern the synthesis itself. Units are normalized to two common denominators, lifetime downloads per paper and twelve-month downloads per paper. Bridging is performed only where a source itself reports the ratio required, and the native unit is retained wherever conversion would demand assumptions the sources do not support. Every factual claim carries a verification tier (verified against the primary source, verified with stated caveats, or pending verification) recorded in the source-confidence columns of the tables. Where sources conflict, most prominently over the effective sunset date, the conflict is resolved in favor of the later official source and flagged. Each estimate keeps its tier, vintage, and aging caveat wherever it is reused.

A single rule governs source inclusion: a reference is retained when it contributes a distinct mechanism, result, platform fact, measurement warning, or design precedent; citation quantity is not a goal. Posting and being read are different achievements; this article is about the second, and about the collapse of the instruments for measuring it. A plausible-looking table is not evidence, and in an article about context-free numbers that rule binds the article itself first.

1. Introduction

1.1 Forty views, seven downloads

Consider a researcher (early-career, perhaps unaffiliated, or simply new to the platform) who uploads a carefully prepared working paper to SSRN and, a few weeks later, in August 2026, opens the paper's statistics panel. The dashboard reports "40 Abstract Views" and "7 Downloads." At that point, without a working knowledge of the statistics of scholarly attention, the author cannot tell whether these figures

mean catastrophic rejection by peers, a standard statistical baseline for a specialized text, or the beginnings of an unusually strong trajectory. The author asks colleagues: "Is that terrible?" Nobody can say. Neither the dashboard, nor the scholarly literature, nor the platform itself provides an answer, and the count has no scale to be read against. The condition is **metric blindness**, the inability, standing before one's own counters, to tell failure from normal performance from success. The term is introduced here as a theoretical construct, not a measured quantity. Study 4 (Section 8.6) supplies the concrete indicator (the share of authors who report never having situated their own count against any benchmark, together with the distribution of first readings a low count elicits) so that "metric blindness" names something a preregistered survey can estimate.

A companion scene, recounted by authors in community forums and offered here as an illustrative anecdote, is starker. A researcher uploads a first preprint and watches the dashboard hold at zero views and zero downloads, while, within days, the inbox delivers five emails from journals soliciting the manuscript, all of them predatory. The only reply the newcomer receives comes from predatory journals. Section 7 develops this asymmetry as a set of explicit hypotheses, and Study 3 (Section 8.6) specifies how to time it. The scene is the lived counterpart of the informational vacuum: the legitimate attention economy gives this author a count with nothing to read it against, while the parasitic one supplies flattery and no readers.

Until the summer of 2026, an approximate answer existed. SSRN's public Rankings (Top Papers, Top Authors, and institutional league tables) allowed an author to locate a count within a visible distribution. On 15 July 2026 those rankings were retired [149]. The retirement followed an announcement of 12 June 2026 that framed the decision with the statement that "research should speak for itself" [148] and an earlier strategic refocusing, announced on 13 April 2026, of the platform's resources on its core research-sharing mission [146]. The removal was asymmetric: per-paper download counts and citation statistics "will continue to be recorded and displayed on your individual papers and author profile" [149], while the distribution that had given those counts a scale was gone. The situation of the researcher at the dashboard was thereby generalized to the platform's entire user base: more than 1.9 million authors and 1.5 million papers [49].

This article is about that situation, and its thesis has two parts. The shutdown of Rankings removed context, not data. The per-paper counters survived the sunset and are still displayed; the comparative scale through which those counters could be read is gone, which is why the metrics are unreadable and not meaningless.

The second is that the question, "what does 7 downloads mean?", is scientifically ill-posed until the correct denominator and comparison group are specified. Supplying that denominator, and the instruments for using it, is the task of this paper.

1.2 A number without a scale

The interpretive problem looks simpler than it is. Ten downloads after seven days and ten downloads after ninety days are identical counts whose meanings almost certainly differ, and they differ again if one paper carries the name of an internationally recognized professor at a highly visible institution and the other is a first-time author's debut. Scholarly usage metrics are nonetheless presented in exactly this context-free form, so every number inevitably becomes a Rorschach test: the same "40 views" provokes panic in one author and indifference in another, because there is no shared reference point.

The expected value of a raw download count changes with everything the counter does not show: the paper's age, its field, its posting cohort and calendar environment, its actual distribution exposure, and the prior visibility of its authors. Each of these dependencies is empirically documented. Average lifetime downloads differ by roughly a factor of eight across the eJournals of a single disciplinary network, from about 438 in corporate and securities law to about 53 in agricultural law [168]. Access activity varies with the calendar itself, falling by about 11% on weekends and by about 50% in summer months in journal-access data [178]. Position and exposure matter independently of content: papers appearing in the most visible position of arXiv announcement lists subsequently earned substantially higher readership and citations, with median citation gains of about 83% in astrophysics [72]. Ideas originating at prestigious institutions diffuse more rapidly and widely than comparable ideas from less prestigious ones [111], a pattern consistent with long-standing evidence on status and attention [104][105][140][103]. On SSRN specifically, an incremental top-50 co-author is associated with a measurable premium in download rates [83]. The platform's own distribution machinery shapes the counter too: authors select up to seven classifications, inclusion in curated topic-based email alerts is not guaranteed, and high-volume topics can face delays of as much as 45 days [53][57]. A paper's early counter may reflect nothing more than whether its relevant distribution has happened yet.

The platform's counting rules add a second layer that authors rarely see. SSRN filters its counters, discarding apparent robot activity and repeated downloads by the same person [144][66][69]; without such filters, counts would be inflated "by a factor of five or more" [18]. Views and downloads are linked by an approximately stable conversion, and Black and Caron reported, in 2006, that "in general, three abstract views on the SSRN website result in one actual paper download" [18]. Later samples place conversion between roughly 11% and 19% of abstract views [139][172], with a recent marketing sample averaging about 24 downloads per 100 views [83]. An author who knows this anatomy can already check the title's 40 views and 7 downloads (Section 5.2 carries out the calculation, which lands inside the reported band as a wide-interval screening result), but nothing on the platform communicates it.

No external scale fills the gap, which makes this a measurement failure and not an inconvenience. For the h-index [79], an entire genre of "what is a good X" guides exists [118][128]. For SSRN downloads, the Section 2.5 audit found no analogue as of August 2026 in any public source. The bibliometric literature reports distributions [85][139][168][83], but none of it translates them into author-facing norms. Individual scholars have improvised kitchen benchmarks from their own portfolios [42], and library guides explain what the counters are without saying what they mean [161][50]. The number every author googles cannot be found.

1.3 The stakes: careers, anxiety, and the incentive to inflate

The gap matters because the numbers are consequential. SSRN's own documentation describes its metrics as ones "that authors use for promotion and tenure purposes and that institutions use in their hiring and ranking activities" [52]. The use is concrete and documented: a prominent legal scholar reported being "ranked 13th out of 3,000 legal authors in total number of all-time downloads, which the law school has taken into account in various promotion and related decisions" [12]; law-review submission guides advise mentioning "SSRN downloads (if a lot)" and Top Ten list appearances in cover letters [95]; and institutions issued press releases about their SSRN standings up to weeks before the sunset [101]. At the aggregate level, SSRN downloads correlate at 0.84 with citation-based rankings of a hundred law schools [77], a school-level association that does not descend to individual papers or authors, the ecological-inference caution of Section 6.2(iii). That institutions act on such measures at all is an instance of a well-

documented dynamic: public rankings are *reactive*, reshaping the behavior of the very actors they purport to describe. The mechanism was established in detail for law-school rankings, where commensuration and self-fulfilling prophecy drove schools to reorganize around the measure [58][132]. SSRN reinforced the currency directly, periodically emailing the top 30,000 authors (roughly the top 8% of about 360,000 listed authors) to tell them they were "in the top 10%" [161]. The "10%" is a share of *listed* authors ($30,000/360,000 \approx 8.3\%$), and against the full registered-author base of more than 1.9 million [49] the same 30,000 is nearer the top 2%, so the phrasing describes ranked authors rather than all registrants. Career advice, meanwhile, warns that self-published metrics alone "probably will not count for law-school promotion and tenure purposes" [33], and the broader evaluation literature documents both how deeply indicator-based assessment penetrates hiring and tenure and how poorly designed it often is [110][115][153]. Downloads themselves are attention measures, not direct measures of quality: early usage predicts later citations only modestly (correlations around 0.4, explaining on the order of 16% of variance in early work [25], and 0.11–0.35 in later estimates [80]), and the sets of highly downloaded and highly cited papers differ substantially [108]. Institutions treat these numbers as career evidence even though they carry no scale and only a loose link to scholarly impact, which is a standing invitation to misreading in two symmetric directions: without a scale, an author may abandon a line of work prematurely on the strength of a misread counter, or spend scarce time and social capital promoting a paper whose numbers were never anomalous. The Rankings, whatever their defects, fed exactly that recurring decision about whether further promotion of a paper was warranted, and their removal left the choice to be made without data (a point taken up in Section 6.4).

The stakes are also emotional: what authors experience at the dashboard is a distinctive form of **metric anxiety**, and download counts are watched compulsively. Scientometric observers described the genre of self-tracking visibility tools as "technologies of narcissism" as early as 2012 [177], and the admission of one prominent legal academic that monitoring his counts was "like crack to me" (Bainbridge, 2007, as quoted in the literature on SSRN gaming behavior) has become emblematic [44]. The same research program established experimentally what the anxiety produces: authors self-download their own papers to inflate counts, and do so most intensely near the visibility thresholds of top-ten lists and under unfavorable social comparison with peers [44]. Social feedback loops of this kind are a general property of attention markets: early social signals amplify inequality and unpredictability in cultural markets [131], initial positive signals bias subsequent evaluations [112], and success breeds success in field experiments [162]. Self-reported milestones on social media, biased toward high performers, sharpen the comparison for everyone whose dashboard reads 7. The absence of a public benchmark thus creates anxiety, fuels misinterpretation, and incentivizes the artificial inflation that undermines the metric for everyone. The problem is one of measurement before it is one of performance: a metric that is simultaneously visible, emotionally charged, career-relevant, and uninterpretable is a genuine failure of measurement infrastructure, not a personal deficiency of the person staring at the dashboard.

1.4 Who is hurt most: the equity dimension

The costs of unreadable metrics are not distributed evenly. Authors with dense professional networks can benchmark informally; established authors carry prior visibility from paper to paper. The absence of public norms weighs most heavily on early-career researchers, independent and unaffiliated scholars, authors in narrow subfields, authors outside the disciplines historically dominant on the platform, and authors outside elite institutions. The disadvantage is structural and not merely informational: SSRN's discovery machinery routes attention through institutional working-paper series and subscription eJournal

alerts [57]. An unaffiliated researcher can post, but cannot seed a paper into a high-traffic institutional series, and depends entirely on subject classification for visibility [168][53]. For such an author the institutional exposure channel is effectively closed, so the weight of discovery falls disproportionately on the search and AI layers. Machine visibility, examined in Section 8, is therefore an equity question and not merely a technical one. Field membership alone shifts expected downloads severalfold before any judgment of merit enters [168]. Prestige effects operate at the moment of first exposure, because a public SSRN page displays author names and affiliations alongside the title and abstract. Double-blind review works otherwise: removing author identities measurably lowers reviewer scores for prestigious authors (in one study of 5,027 ICLR submissions, with a much less determinate effect on acceptance decisions) [154], and evaluations respond to identity signals experimentally [64]. At the bottom of the attention distribution the phenomenon is absence rather than inequality: in altmetric data, the median social attention of unfunded research is zero across all fields studied in a large SDG-focused sample [41], and attention on adjacent open platforms is more unequal than income in developed economies, with Gini coefficients above 0.9 [180]. For the median author the question "is 7 downloads bad?" is an attempt to read the only signal the system returns, in a game whose odds are structurally uneven, rather than a display of vanity. Any interpretation framework that ignores this asymmetry will systematically misdiagnose the disadvantaged by reading invisibility as rejection. The asymmetry also projects forward: if the emerging AI answer layer inherits the brand and position biases documented for commercial generative-engine optimization [5][164], it will reproduce rather than remedy this inequality. That prospect bears most heavily on unaffiliated researchers and on authors in the Global South, and it is one reason the machine-visibility audit of Section 8.6 belongs to the equity agenda rather than to platform curiosities.

1.5 The documented demand for interpretation

The demand for a scale is not hypothetical. It is one of the most persistent patterns in scholarly community discourse, and it predates the sunset by at least a decade. Table 1 summarizes the recurring clusters of author questions and concerns, with the strongest documented evidence for each, the unsafe shortcut each cluster invites (the short answer that looks helpful and is wrong), and the section of this article that responds to it.

Table 1. The demand for interpretation: documented community pain clusters (2005–2026)

#	Cluster	What authors ask or report	Documented evidence (confidence)	Unsafe shortcut	Addressed in
1	"Is my number normal?" — the missing benchmark	Whether a given download count is good, bad, or typical for a comparable paper	Public question "Importance of paper download statistics," Academia StackExchange, 2016 (title and date verified via official API) [3]; user questions on posting to SSRN and its value [126]; recurring forum threads asking how many downloads count as "good" (thread existence confirmed; individual replies not independently verifiable, and forum material is treated throughout as weak evidence of demand, not as data) [B-3]; a first-time author's request for "honest advice" on a newly posted working paper (Reddit, April 2026; existence verified via third-party archives) [B-5]	A universal raw-count threshold ("N downloads = good")	Sections 3, 5–6

#	Cluster	What authors ask or report	Documented evidence (confidence)	Unsafe shortcut	Addressed in
2	Mechanics confusion — views versus downloads	Why views exceed downloads; whether a low ratio means a weak abstract	Historical 3:1 rule of thumb [18]; reported conversion band 11–19% [139][172]; recent sample mean $\approx 24\%$ [83]; library guides document authors' uncertainty about what the counters include [161]	The automatic claim that a low ratio means the abstract is bad	Sections 4–5
3	Emotion and social comparison	Compulsive checking; embarrassment at both low and high counts	"Like crack to me" (Bainbridge, 2007, as quoted in [44]); experimental evidence of self-downloading near top-list thresholds under social comparison [44]	Benchmarking oneself against the platform-wide mean or the visible stars	Sections 1, 5
4	Career evidence — what replaces the rank line in a dossier	What to cite in tenure, promotion, and hiring materials after the sunset	Downloads rank "taken into account in various promotion and related decisions" [12]; top-10% notification emails [161]; cover-letter advice [95]; institutional press releases [101]; platform description of evaluative use [52]	Quoting a raw cumulative count without date, source, or denominator	Section 6 (Table 13)
5	Post-posting solicitation spam	Unsolicited journal invitations arriving soon after posting	"Everyday, I get at least two of these emails" (Reddit r/AskAcademia, 2023; verified via archive; the author had published in an MDPI journal, not on SSRN) [127]; address harvesting from papers "e.g. those posted on SSRN" [107]; audits of solicitation volume [9] [84][63]; a systematic review identifying 29 empirical studies of academic spam [61]	Inferring the identity or intent of readers from an aggregate counter	Section 7
6	Trust in the counters — bots, gaming, sudden drops	Whether the numbers are real; why counters fall	"SSRN does absolutely zero to validate their accuracy" (Goldman, 2005, predating the current filters; a later update records that SSRN discounts repeated same-source downloads) [66]; SSRN's published data-integrity filtering [144][66][69]; robot traffic constituting $\sim 74\%$ of abstract views but only $\sim 7\%$ of downloads in early RePEc data [99]; large-scale bot filtering episodes on RePEc through 2025–2026 [129]	Treating the displayed number as an exact count of unique human readers	Sections 4, 9
7	Diagnostics — "why zero downloads?"	What a zero or a stall means in the first weeks	Post-sunset author uncertainty documented in community threads (existence confirmed via third-party archives) [B-5]; platform support pages describe the metrics but, as of August 2026, offer no interpretive guidance [52][50] (search audit, Section 2.5)	Treating a pre-distribution zero as rejection	Sections 4–6

Note. Only quotations independently verified against a live or archived source are reproduced verbatim; all other community material is paraphrased and marked by its confidence status. The ethics of quoting user posts, including anonymization and the handling of deleted content, is addressed in Section 9.

The strongest cluster in Table 1 is the chronic question "is my number normal?", which recurs across years, platforms, and disciplines. Even platform insiders offer no benchmark in return, and the sarcasm

that often answers it in community forums signals that nobody, including experienced users, possesses the missing scale [B-3]. The clusters are also causally connected: the missing benchmark (cluster 1) produces the emotional load (cluster 3), which produces the gaming incentive documented in [44], which in turn erodes the trust probed in cluster 6. Restoring interpretability addresses the chain at its root. The framework aims not to increase downloads but to restore meaning.

1.6 What this article argues: eight positions

These are the positions the article defends, each argued in the section noted.

1. **Raw counts are uninterpretable without a reference class.** A cumulative download count has no stable meaning without conditioning on paper age, field, posting cohort, distribution channel, and prior author visibility [83][168][178][53]. This is the reference-class problem of Section 5, and it renders "is 7 downloads good?" ill-posed as asked.
2. **Low reach and low conversion are different diagnoses.** A paper few people encounter poses a different problem than a paper many people view but few download, and warrants different action. Invisibility is not rejection of content. Sections 4–5 develop this distinction into a reach-by-conversion diagnostic (Table 9), including the "hidden interest" configuration in which strong conversion coexists with weak reach.
3. **The counters survived; the comparative scale did not.** The sunset's asymmetry (numbers retained, norms removed [149]) describes the post-Rankings condition exactly, and it converts a long-latent interpretability problem into an urgent and universal one (Section 2).
4. **A percentile benchmark is a public good.** The number every author googles and cannot find is a percentile table by discipline and paper age. Sections 5 and 8 specify how the community can build and maintain one, now that it can no longer be read off the platform.
5. **Calibrated norms require a panel observed from posting.** Published evidence consists of aging cross-sections and single snapshots [139][85][168][83]. The one recent large study explicitly identifies panel data as the necessary next step [83]. Without observation from day zero, first-week and first-month "norms" are extrapolations and must be labeled as such (Section 8).
6. **A multidimensional benchmark beats a single vanity score.** Reporting normalized reach, download performance, conversion, momentum, and downstream validation separately (CSAB, Table 10) preserves diagnostic information that any composite index destroys (Section 5).
7. **The AI answer layer is becoming a new gatekeeper of attention.** Discovery is migrating up the stack, from journal brand to repository to search engine to AI answer engines that synthesize a single response and cite a handful of sources [5][164]. A quarter of U.S. adults already use AI chatbots daily [121], while the reliability of these engines as citation channels remains contested [82] and no controlled study of scholarly preprint visibility in them yet exists [100]. Section 8 states the compact audit protocol; the full audit is reserved for companion work.
8. **A dashboard and dossier protocol is needed now.** Authors and committees cannot wait for new data infrastructure. Section 6 provides interim rules for reading a dashboard and a seven-step evidence protocol for dossiers (Table 13) whose governing requirement is that every reported number display its source, its time window, and its limitations.

1.7 Contributions

The list below sets out what the article adds to the existing literature, in the order in which the sections develop it.

1. **A diagnosis.** It identifies the reference-class problem in platform usage metrics and names the underlying phenomenon: the cold start of scholarly attention. The term is adapted from the recommender-systems literature, where a *cold-start* item lacks the interaction history needed to place it [133]. In this cold start, new papers and new authors enter a diffusion system in which early attention depends substantially on platform architecture, distribution, and prior reputation before substantive quality can exert its full influence [72][111][140][83]. The central claim is the ill-posedness thesis of Section 1.1: the headline question has no answer until its denominator and comparison group are supplied.
2. **A consolidated benchmark map.** It assembles, for the first time since the sunset (Section 2.5 audit, August 2026), every verified empirical benchmark for SSRN downloads and abstract views into a single table (Table 7), with explicit sampling frames, vintages, aging caveats, and a source-confidence rating for every row. The benchmarks are the large-sample distributional evidence [83], the quasi-lognormal top-list scrapes [85], legal-scholarship samples [139][168][171][172], and platform aggregates [1][49].
3. **An attention-funnel framework with diagnostics.** It decomposes diffusion into exposure, abstract view, full-text download, and downstream use. It then derives the reach-by-conversion matrix (Table 9), which separates invisibility from rejection and includes the hidden-interest configuration that raw counts systematically mask.
4. **Interpretation instruments.** It converts the benchmark map into a graduated interpretation scale (Nascent / Modest / Solid / Strong / Exceptional; Table 8), a conversion band with worked examples, and the Cold-Start Scholarly Attention Benchmark (CSAB; Table 10), a multidimensional reporting format normalized by field, age, cohort, and prior author visibility. The scale's bands are comparative only within the same field and paper age, since field membership alone shifts expected downloads severalfold, and every boundary is either sourced or explicitly marked as authorial estimate.
5. **A practice layer.** It provides rules for reading one's own dashboard, diagnostics for the zero-download case, a situation-based action table (Table 23), a hypothesized typology of predatory-solicitation signals with the caveat that solicitation spam neither moves the counters nor evidences readership (Table 14), and the seven-step dossier protocol (Table 13).
6. **A preregistration-ready research agenda.** It specifies a prospective panel with observation at 1, 3, 7, 14, 30, 60, 90, and 180 days from posting, and two quasi-experimental designs: curated email-alert distribution as a within-paper attention shock, and the Rankings retirement as a natural experiment in metric interpretability. It further specifies a community-built percentile dataset as the only remaining route to fresh percentiles, a compact machine-visibility audit protocol for AI answer engines, and a prespecified, purely descriptive record of this article's own download trajectory. Each study is executable at low cost and within the access and ethics limits of Section 9, with numeric predictions stated as prespecified bands rather than findings.

1.8 Organization of the article

The remainder of the article runs to nine further sections and five appendices. Section 2 reconstructs the event and its context: the full chronology of the 2026 reform package, what remains and what was removed (Table 2, Figure 1), and the platform's stated rationale of resource reallocation and refocusing. Section 3 presents the verified distributional evidence and shows why the average misleads. It covers the mean–median gap [83], the quasi-lognormal tail [85], and disciplinary heterogeneity [168] (Figures 2–4, Tables 3–5), set against the theoretical background of the Matthew effect and cumulative advantage [104] [105][131][162] (dynamics that have been quantified directly in longitudinal careers and citation trajectories [120]), and the distinction between attention measures and quality measures [25][80][108] [178]. Section 4 examines the mechanics of the counters: filtering and data integrity [144][66], robot and phantom traffic with the necessary counterarguments [99][106][129], platform architectures (Table 6), and metadata and versioning hygiene. Section 5 sets out the interpretation core: the reference-class argument, the consolidated benchmark map (Table 7), the interpretation scale (Table 8), the conversion band, the reach-by-conversion matrix (Table 9), and CSAB (Table 10). Section 6 translates the framework into practice for authors and evaluation committees, including the dossier protocol (Table 13) and, in Appendix A, situation-based guidance (Table 23). Section 7 analyzes the parasitic layer of the attention economy, predatory solicitations treated as hypothesized visibility signals, with all timing and prevalence claims stated as hypotheses [9][61][63][84][107]. Section 8 states the preregistration-ready empirical agenda, including a compact treatment of the migration of attention up the stack to the AI answer layer and its implications for preprint discoverability [5][100][82][121][164]. Section 9 details limitations, including the single-source character of the best available distributional evidence [83], the aging of the 2012–2015 samples, and the instability of platform documentation. Section 10 concludes with policy recommendations for SSRN (foremost, anonymous percentile bands by discipline and paper age as a replacement that restores context without reviving league-table competition) and for the research community.

2. The Event: The Rankings Sunset and Its Context

2.1 What the Rankings Were

For most of SSRN's history, its usage counters sat inside a dense comparative apparatus. At the paper level, the platform maintained Top 10,000 lists by twelve-month and all-time downloads, together with per-eJournal "recent top papers" lists computed over rolling sixty-day windows. At the author level, it maintained a Top 30,000 list across all disciplines and disciplinary leagues such as the Top 3,000 law authors, sortable by new downloads, total downloads, downloads per paper, and citations. At the institution level, it maintained monthly league tables including the Top 350 U.S. Law Schools and the Top 500 International Law Schools [B-9, historical ranking pages; all such URLs have returned an error or a placeholder page since 15 July 2026 and can be cited only from archived copies]. The platform also translated rank into direct percentile feedback. As the University of Washington Gallagher Law Library's guide to SSRN metrics documented, "SSRN periodically sends messages to the top 30,000 authors, by new downloads and by all-time downloads, to let them know they're in the top 10%". Against the guide's own figure of "over 360,000 authors," that threshold corresponded to roughly the top 8% of listed authors (30,000 of about 360,000) [161, accessed 14 August 2026], a broad upper tier rather than a narrow elite; Section 1.3 notes that against the full registered base the same 30,000 sits nearer the top 2%. These lists

and messages were, in effect, the only platform-provided, author-facing percentile information about SSRN downloads ever made public.

The comparative apparatus was not decorative. Black and Caron proposed in 2006 that SSRN downloads could serve as a "beta" measure of law-faculty scholarly performance, supplementing reputation surveys and citation counts [18]. The same symposium volume contains a critical examination of the proposal [45]. Institutionally the proposal succeeded: download ranks entered hiring materials, promotion files, and institutional publicity. SSRN's own support documentation describes its metrics as figures "that authors use for promotion and tenure purposes and that institutions use in their hiring and ranking activities" [52, accessed 14 August 2026]. The Legal Writing Institute's publishing guidance advised authors to mention "SSRN downloads (if a lot), inclusion on SSRN Top Ten Downloads lists" in cover letters to law reviews [95]. In June 2026 Stephen Bainbridge put the career function on the record: he was "ranked 13th out of 3,000 legal authors in total number of all-time downloads, which the law school has taken into account in various promotion and related decisions" [12]. SSRN's own leadership documented the platform's transformation of legal scholarly communication in a preprint by its then-CEO [68]. Institutions marketed the leagues as well: as late as January 2026, Maynooth University announced that its School of Law and Criminology ranked first among Irish departments in the SSRN Top 500 International Law Schools, citing more than 143,500 downloads [101]. The apparatus's grip on behavior is documented in Edelman and Larkin's analysis of SSRN server logs. The study found that deceptive self-downloading intensified at the very moment when a paper approached the boundary of a top-ten list, and identified social comparison, rather than career stakes alone, as the primary driver of manipulation [44]. The same study circulated the remark, attributed to Bainbridge in 2007, that monitoring one's download counts was "like crack to me" [44].

The apparatus had a second function that attracted less notice. Citations and reputation surveys move on multi-year clocks; the rankings returned feedback in near real time, and they let early-career scholars and authors outside prestigious institutions demonstrate visible impact directly, without waiting on slower and more conservative gatekeeping instruments. Whatever their defects, the league tables were the one comparative device equally visible to every author on the platform. That is part of why their removal plausibly weighs most heavily, as Section 1.4 argues, on the authors who have the fewest informal substitutes. No direct distributional evidence on the sunset's effects yet exists, so this is a directional expectation that RQ1 (Section 8.5) tests: two-sided, it asks whether removal concentrated or dispersed attention without presupposing the sign.

2.2 Chronology of the 2026 Retreat

The removal of the Rankings on 15 July 2026 was the culmination of a sequence of platform decisions beginning in late 2025, summarized in Figure 1.

Timeline of SSRN's 2026 Reform Package and the Rankings Sunset

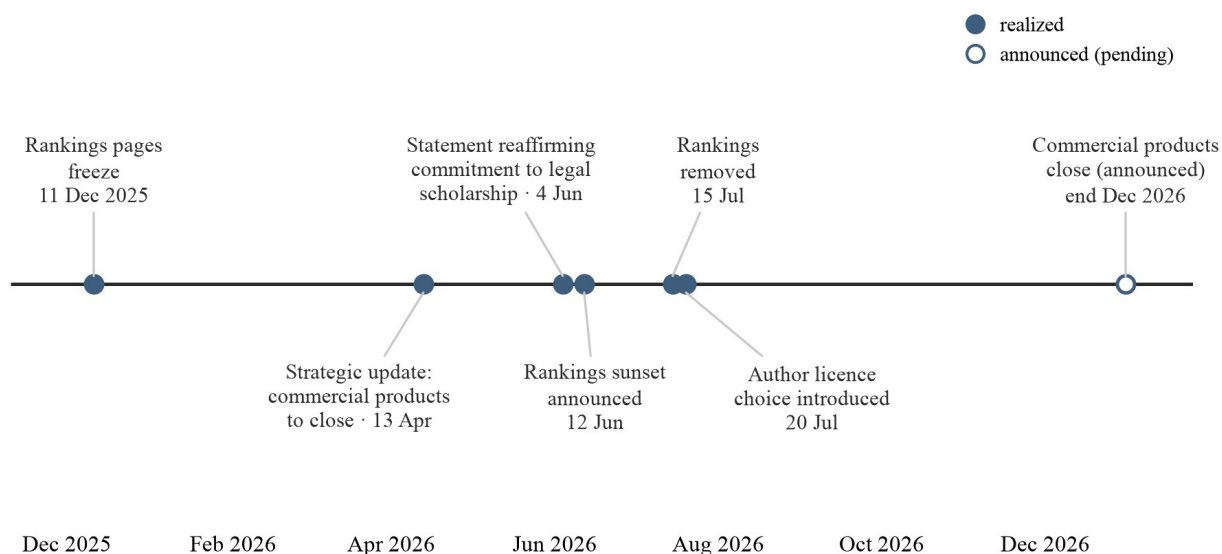


Figure 1. Timeline of SSRN's 2026 reform package. Events: freezing of the eJournal download-ranking pages (from 11 December 2025, as documented in a TaxProf Blog post of 25 January 2026 [B-2]); strategic update announcing closure of all commercial products by the end of 2026 (13 April 2026) [146]; public statement reaffirming SSRN's commitment to legal scholarship (4 June 2026) [147]; announcement that Rankings would be sunset "from July 1st 2026" (12 June 2026) [148]; actual removal of the Rankings (15 July 2026) [149]; introduction of author licence choice (20 July 2026) [150]; announced closure of remaining commercial products (end of December 2026) [146]. Sources: [146][147][148][149][150][B-2].

The degradation began before any formal announcement. On 11 December 2025 the per-eJournal download ranking pages stopped updating. Paul Caron had run a weekly "Top Five New Tax Papers" feature built on those rankings for over two decades. In a post of 25 January 2026 he documented that the tax rankings had frozen on the window "12 Oct 2025 – 11 Dec 2025" and wrote that "this may be the end of the SSRN Top Five New Tax Papers" [B-2, cited as a paraphrase of the 25 January 2026 post, not of any SSRN announcement].

On 13 April 2026, SSRN published a "Strategic Update" announcing a renewed focus on its "core research sharing mission" and the closure of all commercial products (Research Paper Series, Sponsored Networks, Site Subscriptions, paid conference proceedings, analytics dashboards, jobs and announcements, and data feeds) by the end of December 2026 [146]. Library observers registered the change within days [34], and Elsevier's Digital Commons unit published migration guidance for institutions whose Research Paper Series collections were being retired [47].

The announcement provoked public controversy in the legal academy, the community most invested in the ranking infrastructure. On 3–4 June 2026, Bainbridge published a widely read critique under the title "The Social Science Research Network Has Jumped the Shark," followed by a post recording that SSRN had responded to his complaints with clarifications about the timing of the paper-series closures and the preservation of subject-matter eJournals [12]. On 4 June 2026, SSRN issued a statement on its "ongoing commitment to legal scholarship," promising that eJournal distributions and email alerts would "continue unchanged" [147]. On 7 June 2026, Caron replied to Bainbridge on TaxProf Blog under the title "SSRN

Has Not Jumped the Shark" [B-11]. The exchange is a compact public record of what the community took the platform's comparative features to be worth on the eve of their removal.

On 12 June 2026, SSRN announced that it would "sunset Rankings from July 1st 2026," stating that "at SSRN, we have always believed that research should speak for itself" [148]. The shutdown took effect two weeks later than first announced: on 15 July 2026, the rankings of papers, authors, and institutions were removed [149]. Five days later, on 20 July 2026, a further reform introduced author choice of license [150]. The two-week gap between the announced and actual dates, and the month-long window between announcement and execution, are themselves analytically relevant: they created an anticipation period during which authors and institutions could observe the coming change, a point taken up in the quasi-experimental design of Section 8.

2.3 The Official Rationale and Its Evolution

The stated justification for the sunset was resource allocation and platform refocus. The July announcement explained that rankings "require significant product and development resourcing" and framed the decision within a refocus on faster posting, better discoverability, and research integrity. It also stated that "download counts and citation statistics will continue to be recorded and displayed on your individual papers and author profile" [149]. The announcement leaves open a possible return of comparative features: SSRN wrote that it "may revisit the idea of various rankings" [149], so the removal should not be described as permanent. The silence in the record matters as well: the announcements do not attribute the decision to metric gaming, manipulation, or competitive misuse of the rankings, and accounts that supply such motives go beyond the published texts [148][149].

The tone of the official explanation evolved between the announcement and the landing page that replaced the ranking URLs. Where the June post spoke of resourcing, the replacement page states that SSRN's "new direction puts individual authors and their work at the centre of everything we do," and that "aggregated league tables of institutions and departments no longer fit into that picture" [B-8; the page is served behind access restrictions and these phrases are cited from search-index copies, so an archived snapshot should accompany any quotation]. The June post presents the sunset as a cost decision, the replacement page as a principled reorientation away from institutional comparison altogether.

SSRN did not leave authors entirely without an official pointer to context. Its support documentation on gauging a paper's impact directs authors to the paper-level counters (abstract views and downloads) supplemented by PlumX Metrics, an Elsevier altmetrics product aggregating citation, usage, and social-media signals [50, accessed 14 August 2026]. PlumX thus stands as the platform's de facto replacement for interpretive context. It is a partial replacement: it aggregates additional signals for an individual paper but supplies no distribution, percentile, or cohort comparison. Those are the elements the Rankings, however imperfectly, had provided. The retirements do not extend to SSRN's integration with journal workflows: the First Look program continues to connect the platform to more than 1,200 Elsevier journals [48][145]. The platform's commercial governance also situates these product decisions within a highly concentrated scholarly publishing market [94].

2.4 What Remains and What Was Removed

Table 2 maps the post-sunset state of the platform, together with the measurement implication of each element: what each surviving, vanished, or unreported feature means for anyone trying to read a count. The removal was asymmetric. Every per-paper and per-profile counter survives. The comparative apparatus built on those counters was withdrawn entirely.

Table 2. What remains and what was removed after the Rankings sunset of 15 July 2026, with measurement implications. Sources: [147][149][161][B-8][B-9]; exposure-machinery rows: [53][57][55][144][66][69].

Feature	Status after 15 July 2026	Source basis	Measurement implication
Per-paper download counts	Remain; displayed on paper pages	Confirmed verbatim in the sunset announcement [149]	Raw usage remains observable but now carries no reference class of its own
Per-paper abstract views	Remain (not named verbatim in the sunset announcement, which lists "download counts and citation statistics"; views continue to be displayed)	[149], platform observation, accessed 14 August 2026	Reach remains observable, so per-paper reach–conversion diagnostics (Table 9) stay computable
Citation statistics (paper and author profile)	Remain	Confirmed verbatim [149]	Downstream use can still complement early attention measures
Author profile (papers, affiliations)	Remains	Platform observation, accessed 14 August 2026	Portfolio-level aggregation remains possible, with the same missing denominator
eJournal subscriptions and email alerts	Remain	"Continue unchanged," 4 June 2026 statement [147]	The push exposure channel persists, so early counters remain architecture-dependent (Section 4.4)
Topic classifications (1–7, author-selected)	Remain; drive eJournal routing	Platform documentation [53]	Classification is a potentially consequential exposure variable
Curated eJournal distribution	Continues; inclusion "curated and targeted, not guaranteed"	[53][57][147]	Distribution status and date should be modeled as a treatment, not as background
Distribution delay in high-volume topics	Continues; up to 45 days	[53]	First-month counts can mix pre- and post-distribution regimes
External indexing (Google, Google Scholar)	Continues, on search-engine-controlled timing ("6–9 months or longer")	[55]	An interval-censored exposure process: the exposure date is bounded, not observed

Feature	Status after 15 July 2026	Source basis	Measurement implication
Suspicious-activity filtering	Continues	[144][66][69]	Counts are filtered events, not verified reading or endorsement
AI-answer retrieval of papers	Not reported by SSRN	No platform reporting exists; Section 8.6	Machine visibility requires an independent, repeated audit (Study 2)
PlumX Metrics on paper pages	Remain; official pointer for impact assessment	Support documentation [50]	Additional per-paper signals, but no distribution, percentile, or cohort comparison (Section 2.3)
Top Papers lists (Top 10,000; per-eJournal recent/all-time)	Removed; URLs return an error or the sunset landing page	[149][B-8][B-9]	A major comparative frame — and a possible discovery surface — disappeared
Top Authors lists (Top 30,000; disciplinary leagues)	Removed	[149][B-9]	Author-level percentile position is no longer publicly observable
Institution-level leagues (Top 350 U.S. / Top 500 International Law Schools, etc.)	Removed	[149][B-8][B-9]	Institutional comparison can no longer be computed from platform data
Public percentile framing (author rank, "top 10%" positioning)	Removed with the rankings infrastructure	[149][161]	The reference class must now be rebuilt externally (Sections 5, 8)
"Top 10% of authors" notification emails	Fate not addressed in official communications; the emails were tied to the rankings infrastructure	[161][149]	A dossier-usable percentile signal can no longer be assumed to arrive

The surviving rows in Table 2 are exactly the raw counters and the exposure machinery that feeds them: the platform continues to produce numbers at the same rate as before. The removed rows are exactly the set of comparative frames: distributions, thresholds, ranks, and percentile messages. The event therefore has a precise informational structure (data retained, reference class withdrawn), and that structure converts a long-standing latent problem, how to read an SSRN number, into an acute and universal one. We call the platform property that changed **benchmark opacity**: a graded property of a platform describing whether it publishes the distributional statistics needed to situate an individual count (percentiles, medians, cohort distributions), or only raw counters, means, and top-N lists. On 15 July 2026, SSRN moved in a single step from partial transparency to near-complete benchmark opacity, and the rest of the article can be read as a study of what that property costs its users. Opacity, so defined, has at least four dimensions: population coverage (whom the published statistics describe), reference-class granularity (whether they condition on field, age, and cohort), temporal stability (whether definitions and windows stay comparable over time), and auditability (whether the counting rules are public enough to check). The four run partly independently: rankings can reduce one form of opacity while worsening another if their denominators and selection rules are unclear. That partial independence is why the platform recommendations of Section 10.4 (items 4 and 8) address disclosure itself and not only the

return of comparative numbers. To make the property measurable, each dimension is scored on an explicit ordinal checklist applied to a platform's public interface and documentation: population coverage (0 = raw counts only, 1 = platform-wide summary statistics published, 2 = distributional statistics for defined populations); reference-class granularity (0 = none, 1 = one axis such as field, 2 = joint field $\times$ age $\times$ cohort conditioning); temporal stability (0 = counting definitions and windows undocumented or silently revised, 1 = documented, 2 = documented and versioned); and auditability (0 = counting rules unpublished, 1 = described in prose, 2 = published in checkable detail). A platform's opacity profile is then the four-tuple of scores rather than a single label. The same instrument scores SSRN before and after the sunset (moving it downward on coverage, granularity, and, for the retired lists, reference-class granularity), and it lets the cross-platform comparisons of Table 6 be stated as scores rather than impressions. A platform can display many numbers and remain almost completely opaque. The rubric is a first operationalization, to be refined and inter-rater tested, not a finished index.

2.5 The Demand for Interpretation

The question the sunset made urgent was being asked long before it. Table 1 (Section 1.5) assembles the documented clusters of author demand for interpretation, restricted to verified sources. Forum material is cited there as evidence of demand, not of distributions, and quotations from user posts are reproduced only where the post itself was located (the ethics of quoting semi-public forum posts is addressed in Section 9).

The negative claims that follow rest on a structured search audit rather than on casual searching, as do the parallel statements in Sections 1.2, 1.5 (Table 1), and 8.8. The protocol is stated here once. We constructed an inventory of more than 25 query formulations in English, organized into six intent clusters: normative ("how many SSRN downloads is good"), definitional ("what does an SSRN download count include"), post-sunset navigational ("SSRN rankings gone — where to compare"), evaluative ("SSRN downloads as tenure evidence"), diagnostic ("SSRN paper zero downloads why"), and practical ("how to increase SSRN downloads"). For each query we recorded who currently answers it and how well, the criterion being whether any returned source offers graduated, field-conditioned benchmark values rather than a description of the metric or an anecdote. As a positive control, we ran analogue queries for the h-index, the ResearchGate score, and citation counts, which reliably surface the threshold-guide genre documented in [118][128]. The same intent formulations for SSRN downloads surfaced no equivalent, before or after the sunset. The gap extends to the most widely consulted community references: encyclopedia-style overviews of the platform describe what the counters are but offer no interpretive norms for them [170]. The audit was conducted in July–August 2026 using a custom pipeline spanning general-purpose web search engines and consumer AI answer engines, automating query execution and answer-quality scoring. The full query list is available from the author. The result is a negative finding and is claimed only as such: absence of evidence as of the audit window, in the languages and engines covered, not proof that no such resource can exist.

For neighboring metrics, the interpretive demand documented in Table 1 (Section 1.5) has long been met by a mass-audience genre: "what is a good h-index" is served by dozens of graduated, field-conditioned threshold guides [118], and discipline-level citation norms are the subject of widely read explainers [128]. For SSRN downloads there is no such genre. Worse, the best structured guide that did exist, the University of Washington law library page, continued after 15 July 2026 to describe the Top Authors and Top Papers rankings in the present tense, rendering it partially misleading just when demand peaked [156, accessed 14 August 2026]. The sunset did not create the interpretation problem; it removed the only

partial answer the ecosystem had, at the moment when institutional memory of the distributions began to decay.

3. What the Data Show: Why the Mean Lies

3.1 What the Counters Measure, and What They Do Not

Any use of the surviving counters presupposes an account of what they measure. Table 3 summarizes the measurement content of each metric, its major confounds, and the diagnostic role it can legitimately play.

Table 3. What SSRN metrics measure — and what they do not.

Metric	Closest substantive interpretation	Major confounds	Appropriate diagnostic role
Abstract-page views	Reach to a paper's landing page	Search exposure, platform placement, author reputation, external promotion, title, field size, time online	Exposure / reach
Full-text downloads	Intent to inspect or use the complete paper	Everything affecting views, plus abstract quality and relevance and download friction	Engagement
Downloads per abstract view	Conditional transition from landing-page exposure to full-text acquisition	Small denominators, repeat behavior, heterogeneous traffic sources	Conversion
Citations	Observable incorporation into later scholarly documents	Field citation norms, publication lag, database coverage, self-citation, status effects	Downstream scholarly use
Altmetric / PlumX signals	Broader online engagement	Platform composition and user behavior	Complementary diffusion evidence

The table's content can be compressed into a contrast of two equations. The reading the dashboard silently invites is $\text{Downloads} = \text{Scientific Quality}$. The defensible representation is

$$\text{Observed Downloads} = f(Q, E, R, C, T, P, \epsilon),$$

where Q denotes features related to research quality or relevance, E exposure, R reputation, C communication characteristics, T time and age, P platform mechanisms, and ϵ unobserved influences. The expression classifies confounders; no functional form or coefficient is specified. The classification states the practical difficulty exactly: the researcher ordinarily observes only the left-hand side.

Two boundary conditions apply to the table. The counters are not raw traffic logs. SSRN has filtered its counts since at least the 2000s: the platform states that it excludes apparent repeated downloads by the same person and apparent robot downloads, using "cues for irregular activity," and that anonymous downloads under suspicion are excluded from the author's total [144]. Black and Caron estimated that without such filtering the counters would be inflated "by a factor of five or more" [18]. A counted download records an engagement event and nothing about endorsement. The counters measure stages of an attention process. None of them measures scientific merit, and each is confounded by visibility factors

that operate before any reader evaluates content [44][72][144]. Those limits have been recognized since the earliest bibliometric analyses of readership data [91]. The stages can also decouple from one another, through two hypothesized mechanisms: papers on policy-relevant topics may be heavily downloaded by practitioners who read but do not cite, while specialized papers may earn few downloads yet be cited steadily by the small expert community that needs them. Either pattern separates visibility from downstream use, which is one reason the sets of highly downloaded and highly cited papers overlap only partially [108].

3.2 Platform Aggregates and the Fiction of the Average

The only distribution-relevant numbers SSRN itself publishes are platform totals, and the first thing to notice about the data is what happens when those totals are averaged. Table 4 collects the verified platform statistics with their as-of dates.

Table 4. Verified SSRN platform statistics, with as-of dates and verification status.

Quantity	Value	As of	Source	Verification status
Full-text papers	1,157,048	18 Dec 2023	ALL-SIS Citation Methods White Paper [1]	Verified against the published PDF
Abstracts	1,309,385	18 Dec 2023	[1]	Verified
Authors	1,479,875	18 Dec 2023	[1]	Verified
Cumulative downloads	256,678,716	18 Dec 2023	[1]	Verified
Downloads, trailing 12 months	40,642,329	18 Dec 2023	[1]	Verified
Full-text papers (earlier reference point)	571,040	6 Jul 2016	[1]	Verified
Downloads, trailing 12 months (earlier point)	12,897,685	6 Jul 2016	[1]	Verified
Preprints / papers; authors	"1.5+ million"; "1.9+ million"	2026 product page	Elsevier SSRN product documentation [49], accessed 14 Aug 2026	Verified (rounded marketing figures)
Papers; authors; downloads total; downloads 12 mo	>1.7M; >2.4M; >344M; >54M	1 Nov 2025	SSRN announcement boilerplate [B-6]	Order of magnitude corroborated; the exact decimals (344.8M / 54.7M) require archived copies of the announcement before citation
Derived: mean lifetime downloads per full-text paper	≈222	Dec 2023 basis	Author-computed: 256,678,716 ÷ 1,157,048 [1]	Arithmetic on verified figures
Derived: mean downloads per paper per year	≈35	Dec 2023 basis	Author-computed: 40,642,329 ÷ 1,157,048 [1]	Arithmetic on verified figures

The last two rows are author-computed and labeled as such; SSRN does not publish per-paper averages. Platform-wide totals differ by source, date, and definition (full-text vs. total papers), so the per-paper derivations are taken from the reconciled December 2023 full-text figures [1]; the raw November 2025

boilerplate [B-6] is retained above as data only, usable exclusively together with archived copies (e.g., via web.archive.org), since the original announcement pages are no longer directly retrievable.

On the reconciled December 2023 figures, the "average" SSRN paper has earned roughly 222 downloads over its lifetime and roughly 35 in the past year [1]. This derived number is the closest thing to a platform-wide norm an author can compute, and it is systematically misleading, because the average is a fiction of aggregation. Every empirical study of the SSRN download distribution finds strong right skew: the mean is pulled far above the median by a long upper tail, so the "average paper" describes almost no actual paper. The rest of this section establishes the size and consequences of that gap. One further aggregate matters for calibration. Annual platform downloads grew from roughly 0.5 million in 2000 to roughly 3 million in 2005 [18], 12.9 million in 2016, and 40.6 million by 2023 [1]. Download-per-paper flow has itself risen: annual downloads grew about 3.15-fold from 2016 to 2023 while the paper stock grew about 2.03-fold [1], lifting the per-paper flow by roughly 55%. An older study therefore understates the absolute counts a comparable paper would earn today. The direction this bias points depends on the comparison: a fixed old threshold is a lower bound on today's counts, which makes it a *lenient*, not a conservative, bar when it is read as an absolute 2026 level band. The converse manipulation is equally unlicensed: an old threshold cannot be "aged" into a 2026 norm by multiplying it by platform growth, because a larger repository simultaneously changes the numerator, the denominator, the composition of authors, the distribution channels, bot filtering, and the competition for attention. The vintage-bias reading used throughout rests on the aggregate flow-and-stock arithmetic above, not on any rescaling of old thresholds. The platform's own year-in-review confirms continued growth: more than 250,000 new papers were posted in 2025 (up 20% on 2024), with over 53 million downloads during the year [145].

3.3 The Mean Is Not the Typical Paper: Direct Evidence

The most recent published distributional evidence for SSRN comes from Jo and Li, who analyze 2,361 papers from SSRN's Marketing Research Network in a single cross-sectional snapshot collected on 15 August 2025 [83]. Table 5 below reproduces their descriptive statistics in full, and the consolidated benchmark map of Section 5 folds them in. Mean downloads were 221.346 against a median of 87: the mean runs approximately 2.54 times the median. The 95th percentile, 726, is more than eight times the median. The mean exceeds even the 75th percentile (206), which puts more than three quarters of the papers in the sample below the "average paper." A researcher told that the average paper in this network has some 221 downloads would conclude that a paper with 100 downloads is substantially below normal, although 100 exceeds the published median by a wide margin. Figure 2 displays this structure directly.

The Mean Is Not the Typical Paper — and the Sample Is Mature: 2,361 Marketing Research Network Papers (Jo & Li, 2026)

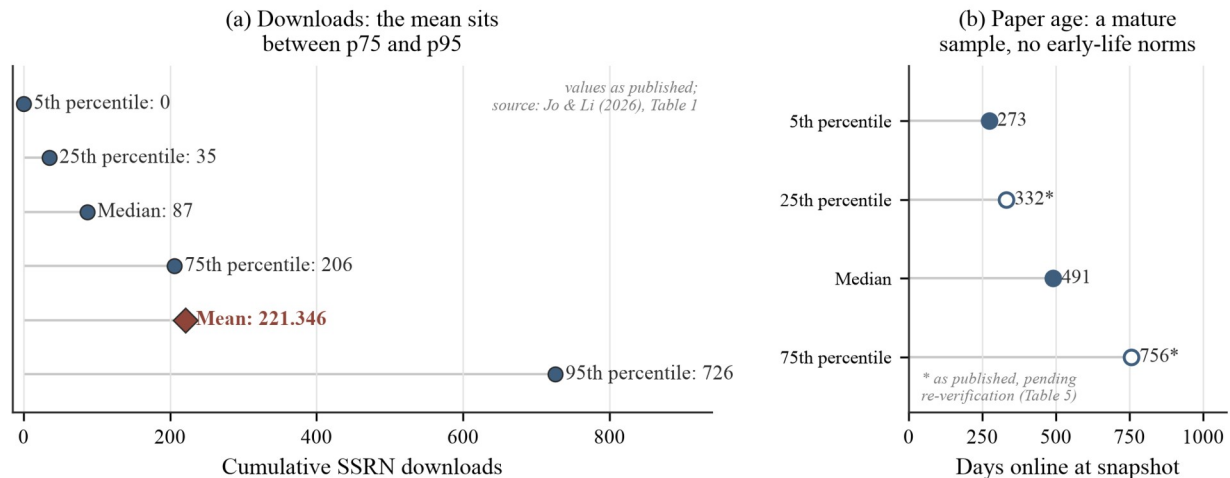


Figure 2. Published descriptive statistics for 2,361 Marketing Research Network papers. Panel (a): downloads — $p_5 = 0$, $p_{25} = 35$, median = 87, $p_{75} = 206$, mean = 221.346, $p_{95} = 726$, plotted in order of value (note that the mean falls between p_{75} and p_{95}). Panel (b): paper age in the same sample — 5th percentile 273 days and median 491 days verified against the source; the 25th (332 days) and 75th (756 days) percentiles reproduced "as published" (Table 5). Two values in the source's Table 1 are flagged and reproduced "as published," but they are anomalies of different kinds. The abstract-views 95th percentile (966) sits only about 0.32 SD above the mean (611.973; SD 1,100.5) — nearly impossible for a right-skewed count whose p_{75} is already 833 — so it is a suspected source error, most plausibly a typo though a transcription or rounding fault cannot be excluded, and the other cells of the abstract-views row are treated as potentially affected by the same correlated error (flagged for an author query to Jo and Li, whose response is not yet available). The conversion 95th percentile (77.7, against 27.2 at p_{75}) is a different phenomenon: a small-denominator ratio artifact of exactly the kind Section 5.2 predicts, not a statistical impossibility. The age panel makes the interpretive limit visible: these are mature-paper statistics, and the sample cannot supply early-life norms. Source: Jo and Li (2026), Table 1 [83].

Table 5 reproduces the study's complete descriptive statistics, its companion variables as well as the download distribution, because this is the only published SSRN sample with a full percentile panel; several of its cells do interpretive work elsewhere in this article. The standard deviations state the heavy tail as a single number: an SD of 940.9 against a mean of 221.3 for downloads. The 5th percentile of downloads is exactly zero, so a zero counter sits inside the normal range even of a mature sample. The conversion percentiles bracket the cross-study band of Section 5.2 from both sides ($p_{25} = 12.1$, median = 18.9, $p_{75} = 27.2$ per 100 views). The age percentiles (interquartile range 332–756 days) fix the sample's maturity, and the authorship and packaging rows supply the covariate norms drawn on in Sections 6.4 and 8.3.

Table 5. Full published descriptive statistics of the Jo and Li (2026) sample, "as published" (N = 2,361 Marketing Research Network papers; snapshot of 15 August 2025). *All values are reproduced from the source's Table 1 [83]. The distributional core — the download and view means, medians, and quartiles; the conversion mean and median; the age median and 5th percentile — has been verified against the published table (see Table 7). The remaining cells (the standard deviations; the conversion and age percentiles beyond the core; the authors, pages, and abstract-length rows) are reproduced "as published" and have not been independently re-derived here; the source's existence, venue, and authorship were confirmed via Crossref on 14 August 2026. Two internal anomalies are carried "as published" and flagged rather than corrected, and they are not the same kind of problem. The abstract-views 95th percentile (966) lies only about 0.32 SD above the reported mean (611.973; SD 1,100.5) — statistically near-impossible for so skewed a count and inconsistent with the row's own p75 of 833 — so it is a suspected typo in the source; the remaining abstract-views cells are flagged as potentially affected by a correlated error and downgraded to re-verification pending an author query to Jo and Li. The conversion 95th percentile (77.704, against 27.187 at p75) is instead a small-denominator ratio artifact predicted by Section 5.2, not a distributional anomaly of the same kind.*

Variable	Mean	SD	p5	p25	Median	p75	p95
Downloads	221.346	940.855	0	35	87	206	726
Abstract views	611.973	1,100.536	151	288	504	833	966 ("as published")
Downloads per 100 abstract views	23.699	20.118	0	12.105	18.944	27.187	77.704 ("as published")
Availability on SSRN, days	550.210	230.557	273	332	491	756	955
Number of authors	2.578	1.705	1	1	2	3	5
Number of pages	32.734	29.186	1	13	27	46	75
Abstract length, words	182.274	71.641	87	137	178	217	295

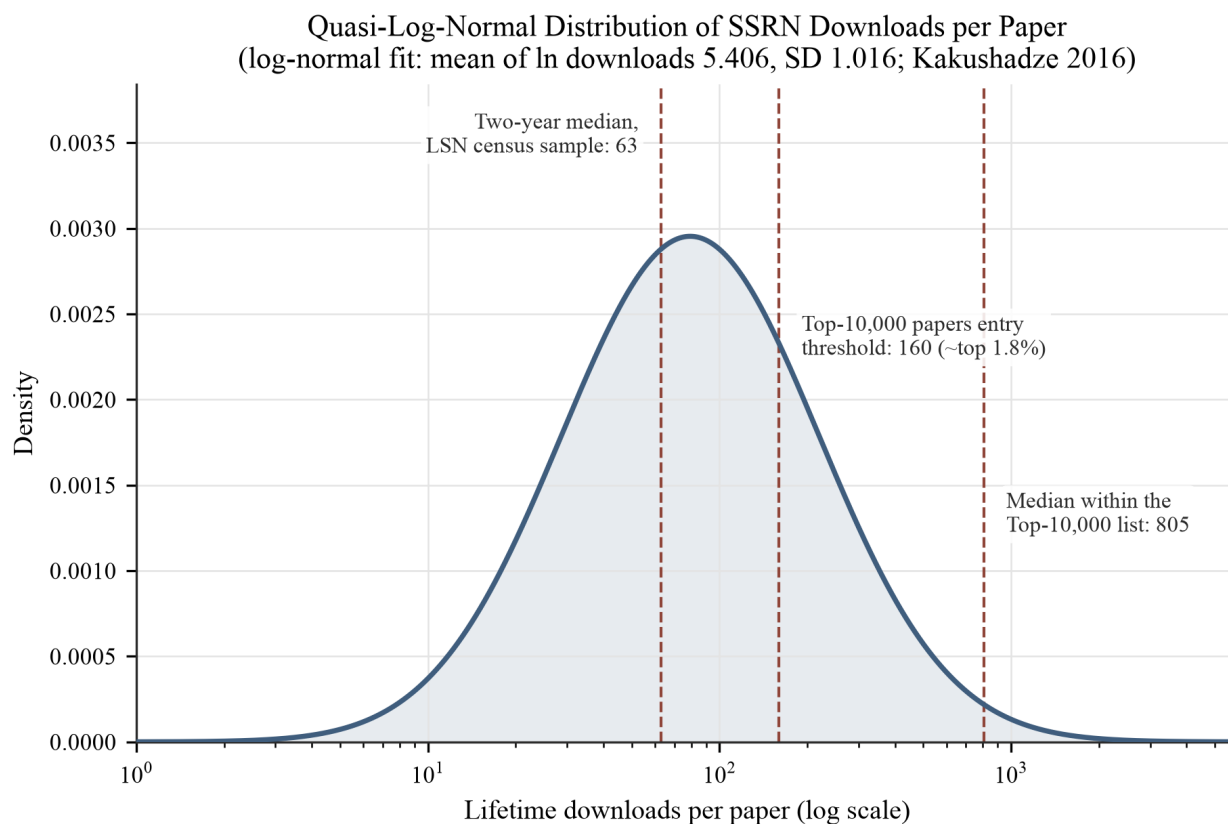
Three further features of the study matter for interpretation. Abstract views show the same asymmetry at a higher level: mean views were 611.973 against a median of 504 (p5 = 151, p25 = 288, p75 = 833; the p95 of 966 is reproduced "as published" but is a suspected source typo, only ≈ 0.32 SD above the mean and below the tail implied by a p75 of 833, and enters no downstream calculation; see the note to Figure 2), with a published conversion of 23.699 downloads per 100 views [83]. The study also separates reach from conversion: in the authors' Poisson exposure model, institutional prestige predicted downloads conditional on abstract views (coefficient 0.192, and for a three-author paper, one top-50-affiliated author corresponded to an estimated 6.6% higher download rate per abstract view), and this prestige premium weakened as AI assistance in abstracts increased (interaction -0.627 , $p < .01$) [83]. SSRN itself now requires disclosure of material generative-AI use at submission [54][56]. So attention tracks status signals even after a reader has landed on the abstract page, which is direct evidence that download counts embed visibility and reputation effects rather than content evaluation alone. The limitation the authors state themselves sets the agenda: a single snapshot of a single disciplinary network cannot reveal time paths of attention, and they explicitly identify panel data as the required next step [83], the step taken up in Section 8.

The sample is one network in one field, observed at one date: these figures are not universal SSRN benchmarks, and their value lies in the structure. The mean/median ratio of 2.54 and the p95/median ratio above 8 are signatures of heavy right skew that recur in every other body of SSRN evidence, including a sample from a different discipline and decade: in the only census-style sample of legal scholarship, papers posted to the Legal Scholarship Network in 2012 and measured two years later ($n = 1,107$) had a median of 63 downloads against a mean of 122.55, a near two-fold gap after only two years of exposure [139].

3.4 The Shape of the Distribution: Quasi-Log-Normality and the Long Tail

The large-scale anatomy of the distribution was established by Kakushadze, who analyzed SSRN's public top lists: 29,985 authors on the Top 30,000 list and the 367,478 papers attached to them (a scrape of 2015-vintage rankings) [85]. The central finding is that downloads are quasi-log-normal, heavily skewed with a long upper tail in levels and approximately Gaussian in the logarithm (fitted parameters for $\ln$ downloads per paper: mean 5.406, standard deviation 1.016) [85]. Kakushadze draws the practical corollary of log-normality explicitly: linear arithmetic fails on such data. A direct application of the h -index to downloads is uninformative: download counts are orders of magnitude larger than citation counts, so the index collapses to the number of papers. He proposes instead an index built on the integer part of the logarithm of downloads, $k = \lfloor \ln D \rfloor$ [85]. For interpretation this means that meaningful differences in downloads are multiplicative, not additive. This shape is not peculiar to downloads: the distribution of individual scientific impact is itself approximately log-normal, as the Q-model of Sinatra, Wang, Deville, Song, and Barabási (2016) established across scientific careers [181], so the multiplicative reading the counters demand is the norm for scholarly attention, not an artifact of SSRN.

The same dataset supplies the only published threshold anchors for the upper distribution. Entry into the Top 10,000 papers list (approximately the top 1.8% of the platform's full-text paper stock, 571,040, July 2016 [1], consistent with rather than above the "top 1–2%" figures sometimes cited) required a minimum of 160 downloads, with a median of 805, a mean of 1,800, and a maximum of 152,242 [85]. Among Top 30,000 authors, minimum total downloads were 129, the median author total was 1,626, the median downloads-per-paper was 210, and the maximum, Michael C. Jensen, was 824,762 [85]. Scarcity persists at the paper level even inside this elite sample: of the 367,478 papers analyzed, only 112,793 (approximately 30.7%) had enough downloads to contribute to their authors' log-index at all [85]. Most individual papers by the platform's most downloaded authors attract modest attention, and outside the top lists the share can only be lower. Figure 3 visualizes the fitted distribution with the published thresholds marked.



Fit estimated from a top-list sample; the body below the lists is inferred from the fitted form. Thresholds are 2015-vintage.

Figure 3. Log-normal density fitted with the parameters reported by Kakushadze (mean of $\ln$ downloads 5.406, SD 1.016) [85]; the fitted median, $\exp(5.406) \approx 223$, sits about 6% above the reported in-group per-paper median of 210, an adequate but imperfect fit to the top-list distribution. Dashed lines mark the two-year median of 63 downloads from the census-style legal-scholarship sample [139], the Top-10,000-papers entry threshold of 160 (approximately the top 1.8% of the platform's full-text paper stock, 571,040 in July 2016 [1]), and the median within the Top 10,000 list (805) [85]. The fit is estimated from a top-list sample — roughly the platform's upper author decile — so the body of the distribution below the lists is inferred from the fitted form rather than observed; the thresholds are 2015-vintage (inflation since then: Section 5.1). A quantitative caution belongs with the figure because the curve invites reading its lower tail: under this fitted density the implied low-tail masses are $P(X \leq 63) \approx 10.7\%$ and $P(X \leq 10) \approx 0.1\%$ — two to three orders of magnitude below the 15–30% ≤ 10 -download share that Study 1 (Section 8.6) predicts for a whole-platform probability sample. The divergence is not a contradiction but a direct measure of the fit's domain: a density estimated on the upper author decile cannot describe the low tail it never observed, so the census-style markers ([139]'s two-year median of 63; the Top-10,000 entry of 160 and its median of 805) are plotted as external reference points laid on the fitted curve, not as claims that the fit reproduces them. The distance between the markers and where the fitted curve places the same values is itself the visual evidence that the body below the lists is unobserved, and no whole-platform low-tail share should be read off this curve.

These anchors need two caveats. The sampling frame is the top of the distribution, so Kakushadze's statistics describe the tail, and the mass of ordinary papers below the lists is characterized only through the fitted log-normal form. The attribution is to a single author: the study is Kakushadze (2016), Journal of Informetrics [85]; a recurring miscitation adds a second author and should be avoided.

An independent archival cross-section reported in Appendix E observes the body of the distribution directly, on a random all-field sample of 1,360 papers whose pre-retirement counters are recovered from public archives. It finds the observed distribution approximately log-normal (skewness of log-downloads ≈ 0.18), corroborating on the body the shape Kakushadze could fit only on the tail, and it independently reproduces the mean-to-median gap, the conversion band, and the field multiplier. Its own upward selection bias leaves the extreme low tail under-observed. It nonetheless reproduces the single modern whole-distribution anchor the synthesis otherwise leans on: the reconstruction's mean (219) coincides with the published marketing-network mean of 221, so the headline shape is corroborated by an entirely separate data source as well. (On recovering historical web data from the Internet Archive as a validated social-science method, and on its coverage and crawl-frequency biases, see Arora, Li, Youtie & Shapira, 2016, *Journal of the Association for Information Science and Technology*, 67(8), 1904–1915, doi:10.1002/asi.23503.)

3.5 Miscalibration by the Visible Tail

The joint structure of Sections 3.2–3.4 explains a systematic psychological error. The most-downloaded author in the 2015 scrape stood at 824,762 downloads, roughly 500 times the median of the top-author group itself (1,626), and four orders of magnitude above the typical paper [85]. Under a quasi-log-normal law, the distances that matter are multiplicative: the step from the ordinary-paper median (tens of downloads) to the top-list entry threshold (160) is smaller than the step from that threshold to the median of the top list (805), which is in turn dwarfed by the distance to the tail maximum [85][139]. Yet the numbers that circulate publicly (milestone announcements, "top author" profiles, institutional press releases [101]) are drawn almost exclusively from the visible tail. Authors therefore calibrate "good" against a reference class located hundreds of multiples above the median, while the mass of the distribution sits near the median. A dashboard reading of 7 downloads is then experienced as failure even where it is unremarkable for the paper's age and field. The mean participates in the same error in milder form: as Section 3.3 showed, it sits above the 75th percentile, so even the official-sounding "average" is a tail-contaminated benchmark. That is the sense in which the mean lies: it is computed correctly and then mistaken for a typical value in a distribution that has no typical value.

Several further bodies of evidence reinforce the conclusion that raw counters reflect visibility structure at least as much as content. Attention responds to position: in arXiv announcement lists, papers appearing in position 1 earned approximately 83% higher median citations in astro-ph, an effect the authors attribute substantially to visibility rather than self-selection alone [72]. It responds to indexing lag as well. SSRN itself notes that external indexing of its content by Google Scholar is controlled by the search engines and can take "6–9 months or longer" [54, accessed 14 August 2026], so a young paper's counters are generated before a major discovery channel has begun operating. Attention also flows backward from citations to downloads: being cited draws readers (among them the cited authors and their own networks) to the citing paper, and Granger-causality analysis finds this citations-to-downloads path active for a substantial share of articles [80][39], so part of a counter indexes a paper's position in the citation and collaboration network rather than any reader's verdict on its content. Each of these mechanisms shapes the counter through visibility structure (position, discovery lag, network location) rather than through a reader's judgment of quality, and none is visible in the dashboard number. The broader attention economy shows the same concentration in extreme form: in social media, attention Gini coefficients reach 0.90–0.94, with the top 1% of accounts receiving more attention than the bottom 99% combined [180], a modern echo of Simon's dictum that a wealth of information creates a poverty of attention [141].

Scholarly attention on SSRN belongs, in a milder form, to the same family of skewed-diffusion processes. The audience most at home on SSRN, economists and legal scholars, will recognize the structure from finance. Bessembinder's analysis of CRSP data for 1926–2016 found that the roughly \$34.8 trillion of net stock-market wealth creation was attributable to fewer than 4% of listed companies, the majority of stocks failing over their lifetimes to outperform Treasury bills [16]. The parallel is offered as an illustration, not proof: in wealth as in attention, aggregate totals are generated by a thin extreme tail, and measuring oneself against an average dominated by giants is a category error. Nor is scarcity confined to usage counters: in legal scholarship at large, 43% of the law review articles in one 2007 census were never cited at all, and only 18% of the 2015–2019 cohort attracted a single citation within five years [171].

3.6 The Field Multiplier: Disciplinary Structure of the Norm

The distribution is not one distribution: it differs from field to field. Whalen's analysis of just under a quarter-million papers classified into the Legal Scholarship Network's 94 eJournals shows that the disciplinary baseline varies by roughly a factor of eight even within a single umbrella discipline: lifetime average downloads per paper ranged from 438.35 in Corporate, Securities & Finance Law to 53.48 in Agricultural Law & Policy [168]. Figure 4 displays the gradient.

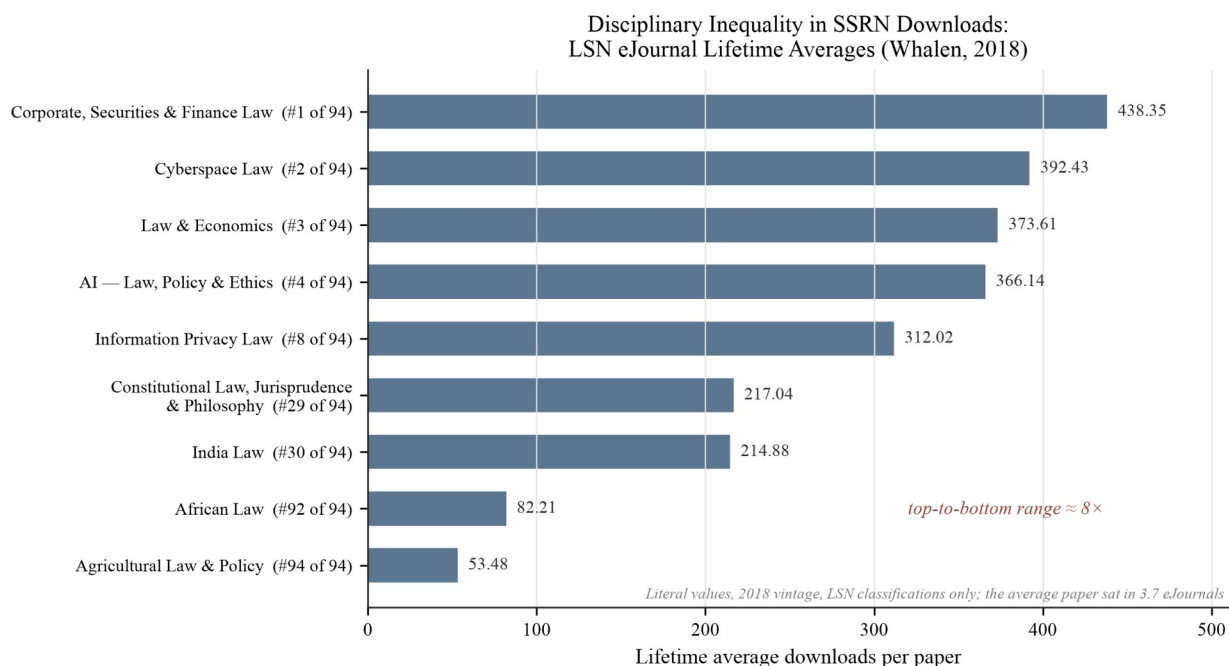


Figure 4. Lifetime average downloads per paper for selected LSN eJournals, literal values from Whalen (2018) [168]: Corporate, Securities & Finance Law 438.35 (rank 1 of 94); Cyberspace Law 392.43 (2); Law & Economics 373.61 (3); AI — Law, Policy & Ethics 366.14 (4); Information Privacy Law 312.02 (8); Constitutional Law, Jurisprudence & Philosophy 217.04 (29); India Law 214.88 (30); African Law 82.21 (92); Agricultural Law & Policy 53.48 (94). The top-to-bottom range is approximately eightfold. Values are 2018-vintage and cover LSN classifications only; the average paper was assigned to 3.7 eJournals.

The mechanism is structural rather than meritocratic: SSRN's discovery is classification-based (subscribers to an eJournal receive alerts for papers classified into it), so a paper's addressable audience is fixed by its field's subscriber base before any quality signal operates. Whalen additionally finds that papers with more classifications earn more downloads and are less likely to record zero downloads [168].

Field structure also shapes the counters through citation culture: Black and Caron observed that downloads-to-citations ratios differ systematically across fields because legal scholarship cites prior work more completely than economics or finance [18], and within law, attention concentrates on particular topics and author characteristics: in the 2012 census sample, corporate and international law topics, US authorship, and top-20 affiliation all predicted significantly higher counts [139], determinants echoed in contemporaneous law-library commentary [175]. So a number that is unremarkable in corporate law may be exceptional in legal history or agricultural law, and any benchmark or scale (including the one constructed in Section 5) must be read as field-relative. There is a well-established precedent for this on the citation side: Radicchi, Fortunato, and Castellano (2008) showed that citation distributions collapse onto a single universal curve only after each paper's count is rescaled by its field's mean [182], the analogue of the reference class that this article argues usage counts equally require.

Three examples, built entirely on the verified anchors above, make the stakes concrete. An author in agricultural law whose paper holds 60 lifetime downloads sits above her eJournal's mean of 53.48 [168], yet, calibrating against the corporate-law counts that circulate most visibly, may read the same number as failure. An author in corporate and securities law with the same 60 sits far below that field's mean of 438.35 [168]. And an author whose marketing paper reaches the network median of 87 [83] may undervalue an outcome that is, by construction, typical. The error of the missing scale runs in both directions: it produces unwarranted despair, and it leaves success unrecognized.

Sections 3.2–3.6 establish the empirical core of this article. The SSRN download distribution is heavily right-skewed and approximately log-normal, and the mean exceeds not only the median but the 75th percentile. The visible tail extends to 500 times elite medians and contaminates informal calibration; the counters embed prestige, position, indexing, and field-structure effects that operate independently of content; and the disciplinary baseline varies eightfold within a single scholarly domain. A raw download count read against the mean, or against the tail, is therefore uninterpretable by construction. A defensible reading requires a reference class: a distribution conditioned on field, paper age, and cohort. The Rankings, for all their defects, were the platform's only public gesture toward such a reference class, and their removal makes the construction of an external one (the benchmark map and interpretation scale of Section 5) the central task of the remainder of this article.

4. The Mechanics of the Counters: A Counter Under a Rule

Before any number on an SSRN dashboard can be interpreted, it is necessary to establish what that number actually records. SSRN displays two primary paper-level usage metrics, abstract views and downloads, alongside citation statistics and PlumX indicators [50]. After the retirement of Rankings on 15 July 2026, these counters remained in place: SSRN stated that "download counts and citation statistics will continue to be recorded and displayed on your individual papers and author profile" [149]. The comparative scale disappeared, but the data stayed. The counters that remain, however, are not raw observations of human readership: they are the joint product of platform architecture, distribution machinery, non-human traffic, and integrity filtering. Each of those layers is taken apart below, and Section 5 then builds an interpretation framework on top of them.

4.1 What a view and a download record, and what they do not

An abstract view registers a request to a paper's landing page. A download records the retrieval of the full-text PDF. Both are engagement events; neither is direct evidence of reading, comprehension,

endorsement, or scientific approval. That distinction bounds every claim made in the remainder of this article.

SSRN's public counters are not simple HTTP-request logs. The platform has operated data-integrity filtering since its early years [69] and has described its approach publicly: automated downloading, repeated self-downloading, and attempts to manipulate counts are screened using "cues for irregular activity," and suspicious downloads (including, in some circumstances, anonymous downloads) are excluded from the author-facing totals [144]. Independent observers documented the same behavior earlier: repeated downloads from a single IP address are not credited to the counter [66]. The download counter is therefore a *curated* statistic, materially more robust than a raw server log. Filtering removes the most detectable non-human and manipulative activity and no more. A counter is always a counter under a rule and not a natural unit, and the rules change (Section 4.3).

4.2 Non-human traffic and the asymmetry between views and downloads

Automated agents account for more than half of overall internet traffic in recent industry measurements [158]. That figure cannot be transferred to SSRN's counters: no published measurement of the bot share in SSRN abstract views or downloads was identified, and any specific percentage asserted for SSRN would be unsupported. What can be established is the *structure* of the problem, approached from two directions: the mechanisms that plausibly inflate views, and the transparent filtering record of a neighboring platform.

Several mechanisms plausibly generate non-human abstract views. Search-engine crawlers routinely index newly posted abstract pages. Institutional email-security systems, most prominently Safe Links in Microsoft Defender for Office 365, scan and pre-fetch URLs contained in incoming mail, which can register server-side requests before any human recipient opens the message [106]. Because SSRN's principal distribution channel is email (Section 4.4), link-scanning of announcement emails is a structurally plausible source of phantom views. No published study measures the contribution of security-proxy pre-fetching to SSRN's view counters, and SSRN itself audits its metrics and excludes irregular activity from download counts [144][66][69], so the platform may well be filtering a substantial share of such traffic, particularly on the download side. The claim that "most SSRN views are not human" appears in informal discussion but has never been measured.

The transparent case is RePEc. Its logging service, LogEc, publishes both its filtering methodology and its filtering history, and the record demonstrates how severe and how *asymmetric* non-human contamination can be. As early as July 2007, robots accounted for approximately 74% of abstract views on monitored RePEc services but only about 7% of downloads [99]. The filter history reads as a chronology of an arms race: heuristics against anti-malware link scanners were added in July 2010; a new type of systematic mass downloading was addressed in January 2017; DDoS-mitigation filters introduced in January 2022 removed 54% of abstract views but only 7% of downloads [99]. From October 2025, RePEc documented a surge of AI-driven bot traffic masquerading as human users [129]. In September 2025, more than 99.5% of raw traffic to IDEAS was discarded as non-human [129]. After cleaning, the RePEc report for June 2026 counted on the order of 2.36 million abstract views and 238 thousand downloads across its services [129].

The RePEc evidence supports two inferences for SSRN. The contamination is asymmetric by construction: automated agents crawl HTML abstract pages far more aggressively than they execute full-text PDF retrievals, so view counters inflate much more than download counters wherever filtering is

imperfect. That asymmetry, robots at ~74% of views but ~7% of downloads in the 2007 measurement [99], is an architectural property of repository traffic and no idiosyncrasy of RePEc. The specific proportions are *not* transferable: SSRN's filtering rules, distribution model, and audience differ, and SSRN does not publish an equivalent filtering log. Section 5.2 develops the practical consequence, which is that ratios built on views (such as the conversion rate) must be read as lower bounds on human engagement quality, because the denominator is the more contaminated quantity. This lower-bound direction is vintage-dependent rather than a law: it holds while view-side filtering lags download-side filtering, which is the historically documented regime [99], but it would weaken or reverse if a platform's view filters caught up with its download filters, or if download-side contamination rose faster than view-side (as the 2025–2026 AI-bot episode raised both [129]). The direction is to be re-checked whenever a platform revises its counting rules.

A further consideration has emerged since these filtering regimes were designed. The filters were built to strip automated requests from a count intended to proxy human attention, and that rule still removes noise. What it can no longer do unexamined is treat the remainder as human, because three classes of automated request now sit where the filter saw one: crawls that harvest text for model training, crawlers that index for a search or answer service, and fetches a system makes in real time on behalf of a person who has asked it something. The first two remain noise for an attention count. The third is a delegated act with a human behind it, since a system that retrieves a paper's full text to answer a user's question or support a literature search has been served the content and used it, although no person ever loaded the page. An automated request is not by itself evidence that anything was read. The proxy has weakened, not the filter: human against machine no longer reliably tracks the distinction this article needs, which is whether an event is evidence of scholarly attention, human or delegated.

Project COUNTER has begun to address this in its reporting vocabulary rather than in its filter. Best-practice guidance on AI usage, issued in April 2026 and updated on 30 June 2026 for Release 5.1.1, makes available an optional `Access_Method` of Agent, defined as content and metadata accessed by an AI system, together with AI-specific metric types, while Release 5.1.1 continues to require that “activity generated by traditional bots and crawlers **MUST** be excluded from all COUNTER reports” [125]. The instrument has three limits that matter here. It is an optional extension, produced only in reports a provider elects to generate and aimed first at AI systems embedded on a provider's own platform rather than at inbound crawlers. Its elements “**MUST NOT** be included in any Standard Views of the COUNTER Reports” [125], so a standard view will not distinguish agent access even where a provider implements the guidance. Whether tagged agent events are then removed from default totals or carried in them unlabeled is a matter of implementation on which the guidance is silent. It also classifies access that was served and not access that was refused: the guidance states that “there are no denial metrics associated specifically with `Access_Method` Agent” [125]. None of this reaches a bare public counter in any case, since such a counter is not a COUNTER report and carries no access method; the standard is an analogue against which it can be judged, not a description of it.

Refusal is not hypothetical either, and it operates upstream of any counting rule, but its instruments differ in kind and do not act as one. A platform's `robots.txt` asks named agents not to crawl; a text-and-data-mining reservation asserts a legal position; terms of use bind contractually; a challenge service withholds the page from clients that cannot complete it. Only the last of these is enforcement. On the platform studied here, as retrieved on 29 August 2026, `robots.txt` named three agents (GPTBot, ChatGPT-User and Google-Extended) and was silent on others, silence being the Robots Exclusion default that well-behaved

crawlers read as permission; responses carried a mining reservation pointing to the publisher's policy, which permits mining for purposes falling under Article 3 of the EU Digital Single Market Directive and requires a license otherwise; terms of use reserve rights over text and data mining and AI training except where an author has chosen an open license, an option available only since 20 July 2026 and defaulting to all rights reserved; and a single unauthenticated request made without a JavaScript engine received a Cloudflare managed challenge rather than the page. One of the named agents is the user-triggered fetcher rather than a training crawler: that is the delegated class of the preceding paragraph. These materials record published restrictions and one configuration observation. They do not establish how many requests, automated or human, were declined. Directives are honored unevenly: in a large-scale measurement of crawler behavior, some agents respected whole-site disallows completely while honoring path-level disallows less than a third of the time, and others ignored the file altogether [87]. Nor is the collateral one-directional, since challenge deployments have disrupted legitimate human users elsewhere (one university library recorded two such incidents in the first four months of use [29]).

The historical asymmetry therefore cannot simply be carried forward. The comfort of the 2007 finding, robots at roughly 74% of abstract views but about 7% of downloads [99], rested on crawlers that browsed landing pages and rarely fetched full text. A system that retrieves documents in order to ground an answer has no reason to stop at the abstract, so the download side is no longer known to be the cleaner tally by construction. Where raw scholarly traffic has been measured in 2025–26 the automated share has not been marginal, and it has been found inside download and visit counts rather than only in page views: one cultural-heritage collection revised a March 2025 visit count from 1,000,000 to 125,000 after suspected bot traffic was removed [167]; a vendor classification of roughly thirty million access requests at one institutional repository labeled only about a quarter of them human [98]; and a repository responding to a survey of open repositories attributed at least a quarter of its own bot downloads, some half a million requests, to eleven named AI agents [135]. None of these figures is a measurement of the platform studied here: they are drawn from different platforms with different, in several cases absent, defences, and two of the three count visits or requests rather than the download events a public counter displays. They establish the prior condition rather than a magnitude. In the current environment the interpretability of any usage count is a function of filter quality, and where a platform does not publish its filtering, that quality is unaudited.

Two identification problems follow, and they compound. Composition first: from a public lifetime counter alone, the shares of human, conventional automated, and agent-mediated events are not estimable, and the platform's own view may be no better, since automated clients have been observed presenting themselves as desktop browsers [167]. On selection, a counter cannot reveal what share of attempted access was served, challenged, or excluded before any counting rule applied. Published policies disclose that such mechanisms exist. They do not disclose how often they operate, how they operated in the past, or their effect on inclusion. Attempted requests, challenged requests, served full-text events, displayed downloads, and answers generated from a retrieved paper are therefore five different populations, and they must not be treated as interchangeable. What remains unlicensed is a magnitude or a net sign for this platform. Unfiltered scholarly usage is now known to be inflatable by automated traffic, downloads included, but whether this platform's unpublished filters, challenge rules and crawler policies leave a more human-weighted remainder, a gap in agent-mediated discovery, both or neither, is unmeasured. Nor is the mixture stable across the years a lifetime total spans, which run from a pre-2024 robot regime through the scrape surge of 2025–26 to a license default introduced in July 2026. The downstream

consequences of the question are themselves documented: sites blocking a major AI crawler are retrieved significantly less often by AI answer engines even where the content remains accessible [70]. No published estimate of the agent share exists for any major preprint server, a research gap rather than evidence of a small effect, and the audit specified in Section 8 is designed to measure it. A public lifetime total is a sum over an unknown mixture, drawn from an unknown share of attempted access, with no external estimator of either and no basis for treating either as stationary.

4.3 Falling counters as measurement events

Counters can also fall, or be restated, when filtering rules change, a further consequence of the filtering record that no dashboard explains. RePEc's January 2022 filter revision removed half of previously counted abstract views in a single step [99]. Its 2025–2026 AI-bot episode forced further methodological changes [129]. Vendor-reported download statistics are, more generally, known to be sensitive to counting conventions: a library-side audit found that download reports from one major publisher ran roughly twice as high as comparable platforms, with "deflated" counts at 43% of the reported figure [176]. A sudden drop in an author's SSRN counters is therefore as likely to be a *measurement event*, a filtering or restatement change on the platform side, as a readership event, and it should never be read as reputational loss without corroboration. The same reasoning governs any dossier use of metrics, where dated records should be preserved (a protocol developed in Section 6). The platform's numbers can be revised, and the author's dated archive cannot be revised.

4.4 Push versus pull: architecture shapes the ratio

Distribution architecture is the last mechanical layer. It differs qualitatively across the major repositories and directly shapes the relationship between views and downloads. SSRN operates predominantly on a *push* model. At submission, authors select between one and seven subject classifications (the average LSN paper carried 3.7 eJournal assignments [168]), and these determine placement in research networks and eligibility for topic-specific eJournal email alerts. Inclusion in an alert is "curated and targeted, not guaranteed," and distribution in high-volume topics can take as long as 45 days [53]. Subscribers receive abstract-level announcements by email [57]. Every recipient who clicks through from an announcement, whether out of focused interest or mild curiosity, registers an abstract view. arXiv, by contrast, operates predominantly as a *pull* platform: a non-profit repository whose traffic is driven by researchers actively visiting listings and feeds [11][65], where even position within a daily announcement list measurably affects readership and citations [72]. RePEc aggregates metadata from decentralized archives and applies the heavy documented filtering described above [99][129].

Table 6 summarizes the comparison across platforms as a qualitative schema of governance and discovery mechanics. Comparable cross-platform traffic data do not exist in the published record.

Table 6. Platform architectures and their mechanical effect on usage metrics (qualitative schema).

Feature	SSRN	arXiv	RePEc / IDEAS
Governance	Commercial (Elsevier, since May 2016) [46]	Non-profit (Cornell University) [11]	Decentralized volunteer network [99]

Feature	SSRN	arXiv	RePEc / IDEAS
Primary discovery channel	<i>Push</i> : curated eJournal email alerts driven by author-selected classifications (1–7; distribution up to 45 days in high-volume topics) [53][57]	<i>Pull</i> : intent-driven visits to listings and feeds; positional effects within announcements documented [11][72]	Aggregation of distributed archives; discovery via IDEAS/EconPapers [99]
Filtering transparency	Policy described qualitatively; no public filtering log [144]	Not centrally published for usage metrics	Fully documented methodology and revision history (LogEc) [99][129]
Expected mechanical effect on the view : download ratio	Elevated views relative to downloads: email click-throughs register views from low-intent exposure	Higher conversion of visits into downloads: traffic is intent-driven	Ratio meaningful only after heavy bot trimming; published counts are post-filter [129]

The architecture bears on author psychology, since a high view-to-download ratio on SSRN is in substantial part a *design outcome*: the push model generates broad, low-intent exposure. Abstract skimming is also legitimate scholarly behavior (researchers routinely monitor fields through abstracts without downloading full texts), so a view that does not convert may mean that the abstract already delivered its informational payload rather than that it repelled the reader. SSRN's own historical rule of thumb, reported by Black and Caron, was that roughly "three abstract views ... result in one actual paper download" [18]. The same authors noted that without integrity filters, raw counters would be inflated "by a factor of five or more" [18]. Section 5.2 places these dated anchors alongside modern measurements.

Two mechanical features of how a paper is presented govern conversion alongside reader interest, which makes the measured conversion rate partly a property of the interface rather than of the research. One is the platform's user interface. Whether the full text is offered as an inline, in-browser HTML rendering or only as a downloadable PDF materially changes whether an interested reader ever triggers the download counter, since a reader who consumes the full text in the browser satisfies the same intent while incrementing views alone. The other is the device composition of the audience. Mobile access depresses full-text PDF retrieval relative to desktop, so a paper whose readers arrive predominantly on phones (from a social-media link or a mobile mail client) will record systematically lower downloads and comparatively higher views for identical scholarly interest. Both effects pull in the same direction as the mechanical inflation of the denominator already established above, and both are further reasons to read the empirical band of Section 5.2 as a platform-conditional range rather than a fixed threshold: a conversion norm estimated on one platform's interface and device mix should be recalibrated rather than transplanted when either changes. The usage-metrics literature raises the same reliability caution about the comparability of download reports across platforms and counting conventions [176].

4.5 Predictable side effects: similarity self-matches and metadata hygiene

Two further mechanical consequences of preprint posting are predictable, and both routinely generate unnecessary alarm, so both belong with the rest of the machinery.

The similarity self-match. When a working paper posted on SSRN is later submitted to a journal, the journal's similarity screening will, in the expected case, flag a near-total match with the author's own

preprint. iThenticate's documentation states both halves of the interpretation explicitly: the software "does not check for plagiarism; it checks for similarity," and "if you see a pre-print match for the same author, this isn't plagiarism" [81]. The mechanism is documented as well. Preprints registered as Crossref Posted Content are included in the similarity-check comparison databases, so a match against one's own deposited preprint is the system operating as designed [36]. Matches can surface even when editors exclude preprint sources, because secondary aggregators duplicating repository metadata are also crawled [81]. The author's response should be administrative rather than defensive: identify the flagged source to the handling editor as the author's own time-stamped preprint. Crossref and iThenticate guidance advises qualitative editorial review rather than mechanical similarity-score thresholds [36][81]. No numeric threshold is specified here, since none is endorsed by the screening services themselves. The composure this advice presupposes has a documented basis in publisher policy: posting a working paper to a preprint server is permitted under the prior-publication policies of most academic journals, though policies vary and the target journal's current policy should be checked; preprint stances are publicly indexed in registries such as Sherpa Romeo. The salient exceptions cluster in medicine and adjacent clinical fields, where some journals still restrict or discourage prior preprint posting. Authors in those fields should verify the target journal's preprint policy before posting rather than after submission.

AI-writing detection: an adjacent false-alarm channel. A newer screening layer carries a structurally similar risk of false alarm. AI-writing detectors produce documented false positives, and the burden falls disproportionately on non-native English writers: in a Stanford evaluation of seven publicly available GPT detectors (the study did not include Turnitin), an average of 61.3% of TOEFL essays written by non-native speakers were misclassified as AI-generated [97]. Institutional practice has registered the concern. In August 2023 Vanderbilt University publicly disabled the AI-detection feature of its similarity-checking software, citing its unreliability: the vendor's own claimed false-positive rate, the absence of any published detection methodology, and the same disproportionate burden on non-native English writers documented above [163]. The defensive posture here parallels the similarity self-match and is administrative rather than adversarial: preserve the manuscript's version history and drafting records as contemporaneous evidence of authorship (a further argument for the versioning hygiene below), and answer any flag with documentation rather than alarm.

Versioning and metadata hygiene. Because counters accumulate per paper page, revision practice directly affects metric integrity. SSRN supports posting revised versions to the same paper page, preserving the accumulated counters and the indexed URL, whereas deleting a paper and re-uploading it resets the metrics and breaks existing links (SSRN FAQ; a support page subject to revision, accessed 14 August 2026 [B-7]). Classification choices likewise have metric consequences, since they determine eJournal routing and hence the paper's principal exposure events [53]. External indexing runs on a longer clock than authors expect: SSRN states that Google Scholar indexing is controlled by the search engine and can take "6–9 months or longer" [55]. A paper's early counters are therefore generated almost entirely by the platform's internal distribution machinery, which anchors the interpretation of early zeros in Section 5.6.

5. An Interpretation Framework: Rebuilding the Missing Scale

The question "is this number good?" is scientifically ill-posed until a comparison group is specified. This section assembles the published empirical record into the missing comparative apparatus: a consolidated

benchmark map with explicit source confidence (5.1), an empirical conversion band (5.2), a graduated interpretation scale (5.3), a two-dimensional reach–conversion diagnostic (5.4), a multidimensional benchmark specification (5.5), and a temporal reading protocol (5.6). Values are stated throughout with their sampling frames and ages. None is a calibrated percentile for the 2026 platform, and all should be read subject to the inflation caveat of Section 5.1.

5.1 A consolidated benchmark map

Table 7 consolidates every usable published benchmark for SSRN usage metrics identified in the Section 2.5 audit, each with its sampling frame and an explicit statement of source confidence. Because a benchmark is misused the moment it is stretched past its frame, each entry also names the inferential role the source can legitimately play and the main limitation that bounds it. The evidence is thin: the map rests on five studies with distinct frames: a census-style sample of one legal network in 2012 [139], a scrape of the platform's own top lists in 2015 [85], a single-network cross-sectional snapshot in 2025 [83], platform-level aggregates [1], and one subfield decomposition [168]. Those frames are not interchangeable: top-list scrapes describe the elite tail, network censuses describe typical papers, and platform aggregates describe neither. Every distribution involved is also strongly right-skewed, so means systematically exceed medians and describe atypical papers.

Table 7. Consolidated benchmark map for SSRN usage metrics, with sampling frames, source confidence, permissible inferential roles, and main limitations.

Benchmark	Value	Source (year; sampling frame)	Source confidence	Permissible inferential role	Main limitation
Median paper downloads after ~2 years on platform	63 (mean 122.55)	Siems 2016 [139]; n = 1,107 LSN papers posted in two 2012 windows, counters read two years later	Medium-high: sample and download figures verified; law only; 2012 cohort	Historical view of the ordinary-paper distribution	A two-year legal-network anchor, not a current platform-wide norm
Median paper abstract views after ~2 years	369 (mean 564.62; views–downloads correlation 0.852)	Siems 2016 [139]	Medium-high: figures verified against the source's Table 1; law only; 2012 cohort	Historical view of ordinary-paper reach	Same aged, single-discipline frame; view counters are the more contaminated quantity (Section 4.2)

Benchmark	Value	Source (year; sampling frame)	Source confidence	Permissible inferential role	Main limitation
Median paper downloads, single-network snapshot	87 (mean 221.3; p25 = 35; p75 = 206; p95 = 726; median paper age 491 days)	Jo & Li 2026 [83]; N = 2,361 Marketing Research Network papers, snapshot of 15 August 2025	High: verified against the published table; note two anomalies of different kinds reported "as published" — views p95 = 966 is a suspected source typo (≈ 0.32 SD above the mean; row downgraded to re-verification), conversion p95 = 77.7 is a small-denominator ratio artifact (Section 5.2); full panel in Table 5	Recent cross-sectional distribution and conversion model	One network, one field, one date; mature papers; no trajectories
Median paper abstract views, single-network snapshot	504 (mean 611.973; p5 = 151; p25 = 288; p75 = 833; p95 = 966 "as published")	Jo & Li 2026 [83]; same sample as above	High: verified against the published table; the p95 value (966) is a suspected source typo — ≈ 0.32 SD above the mean — reproduced "as published" and downgraded to re-verification pending author query	Recent cross-sectional reach distribution	Same single-network, single-date frame
Views→downloads conversion, empirical band	11–19% ($\approx$ one download per 5.3–9 abstract views)	Siems 2016 [139]; Willey & Knapp 2022/23 [172]; Jo & Li 2026 [83]	Medium-high: Siems, Willey–Knapp, and Jo & Li figures verified against the source tables	Conditional-engagement norm (conversion)	Few sampling frames; the denominator is bot-inflatable, so measured conversion is a lower bound while views remain the more inflated counter (Section 4.2)
Historical platform-wide rule of thumb	~ 3 abstract views : 1 download ($\sim 33\%$)	Black & Caron 2006 [18], reporting SSRN's own statistic	Medium: verbatim verified; dated (2006); platform-wide	Metric anatomy and integrity history	Early platform era; an upper anchor, not a current norm
Entry threshold, all-time Top-10,000 papers	≥ 160 lifetime downloads ($\approx$ top 1.8% of the platform's full-text paper stock, 571,040 in July 2016 [1]); median within the group 805; mean 1,800; max 152,242	Kakushadze 2016 [85]; scrape of SSRN top lists, 2015; 367,478 papers	High: verified against the published tables; top-list frame; single-author work	Distribution shape and tail behavior	Top-list selection; unsuitable for estimating the contemporary median paper

Benchmark	Value	Source (year; sampling frame)	Source confidence	Permissible inferential role	Main limitation
Entry threshold, Top-30,000 authors	129 total downloads (minimum); median within the group 1,626 total; median 210 downloads per paper; max 824,762	Kakushadze 2016 [85]; the group SSRN's own letters described as the "top 10%" of authors [161]	High: verified against the published tables	Tail behavior at the author-portfolio level	Same top-list selection; 2015 vintage
Unadjusted platform-wide lifetime mean	~222 downloads per paper (~35 in the trailing 12 months)	Computed from platform totals as of December 2023: 256,678,716 lifetime downloads across 1,157,048 full-text papers; ALL-SIS 2024 [1]	Medium: derived aggregate; mechanically inflated by old papers; not a personal yardstick	Aggregate calibration of platform scale only	An average over two decades of accumulation; describes almost no actual paper (Section 3.3)
Subfield spread within a single discipline	~8× in mean downloads per paper (438.35 Corporate/Securities/Finance Law vs 53.48 Agricultural Law)	Whalen 2018 [168]; 94 LSN eJournals, ~250,000 papers	Medium-high: verified from the source; 2018; means, not medians	Classification and subfield heterogeneity	Means, not age-matched medians; law only; 2018
Early downloads → later citations	$r \approx 0.39\text{--}0.43$ over the first 180 days ($\approx 16\%$ of variance explained)	Brody, Harnad & Carr 2006 [25]; arXiv (HEP), $N = 14,442$	High: verified; different platform, transferability assumed rather than shown	Downstream predictive validation	Different platform and era; paper-level correlations are modest
Calendar structure of access; access→citation lag	Accesses ~11% lower on weekends and ~50% lower in summer; ~unit-elastic citation response ~9 months after an access change	Yuret 2026 [178]; >1,500 articles from seven journals followed ~1 year	High: DOI, venue, and title verified via Crossref (14 August 2026); all four figures — more than 1,500 articles from seven journals, accesses about 11% lower on weekends and about 50% lower in summer, and an approximately unit-elastic citation response (a 10% rise in accesses associated with about a 10% rise in citations) about nine months later — verified verbatim against the published article abstract (Springer, 15 August 2026)	Temporal normalization and predictive validation — the calendar-normalization anchor behind CSAB	Journal-access data, not SSRN-specific

Benchmark	Value	Source (year; sampling frame)	Source confidence	Permissible inferential role	Main limitation
Sensitivity of counts to counting conventions	Vendor download reports can run $\sim 2\times$ those of comparable platforms ("deflated" counts at 43% of the reported figure); article-level metrics require explicit data-integrity screening	Wood-Doughty et al. 2019 [176]; Gordon et al. 2015 [69]	Medium-high: both sources verified and already cited in Sections 4.1–4.3	Integrity and comparability warning — bounds the precision of any cross-platform or cross-vendor comparison	Not a direct SSRN benchmark; no SSRN-specific audit exists

Three structural readings of the map follow.

The mean is not the typical paper. In the only recent published distribution the mean (221.3) is 2.54 times the median (87) and exceeds even the 75th percentile (Section 3.3) [83]. In Siems's earlier census the mean was roughly twice the median [139]. The platform-level mean of ~ 222 lifetime downloads [1] is the least suitable personal yardstick of all: it aggregates two decades of accumulation by papers of all ages.

Every benchmark carries a vintage, and the direction of the vintage bias depends on the comparison. Annual platform downloads grew from roughly 0.5 million in 2000 and 3 million in 2005 [18] to 12.9 million in 2016 and 40.6 million in 2023 [1], a 3.15-fold rise between 2016 and 2023, while the paper stock grew about 2.03-fold over the same window (571,040 to 1,157,048 [1]). Figure 5 plots this trajectory. Because downloads grew faster than the paper stock, the download flow per paper rose by roughly 55% (from about 22.6 to about 35.1 per paper per year [1]), so a fixed 2015 threshold now represents a *smaller* share of what a typical paper earns than it did then. Read as an absolute level band for 2026, an old threshold such as the ~ 160 -download top-list entry is therefore *lenient* rather than conservative: clearing it is less selective today than in 2015, not more. The bias is composition-dependent. For percentile position under continued right skew the old anchors may still understate the tail, but for absolute counts the download inflation makes them soft, so no old threshold should be quoted as a "conservative floor" without this qualification. Multiplying a 2015 threshold by subsequent platform growth is likewise unlicensed: author composition, distribution channels, filtering, and competition for attention all changed alongside the download flow.

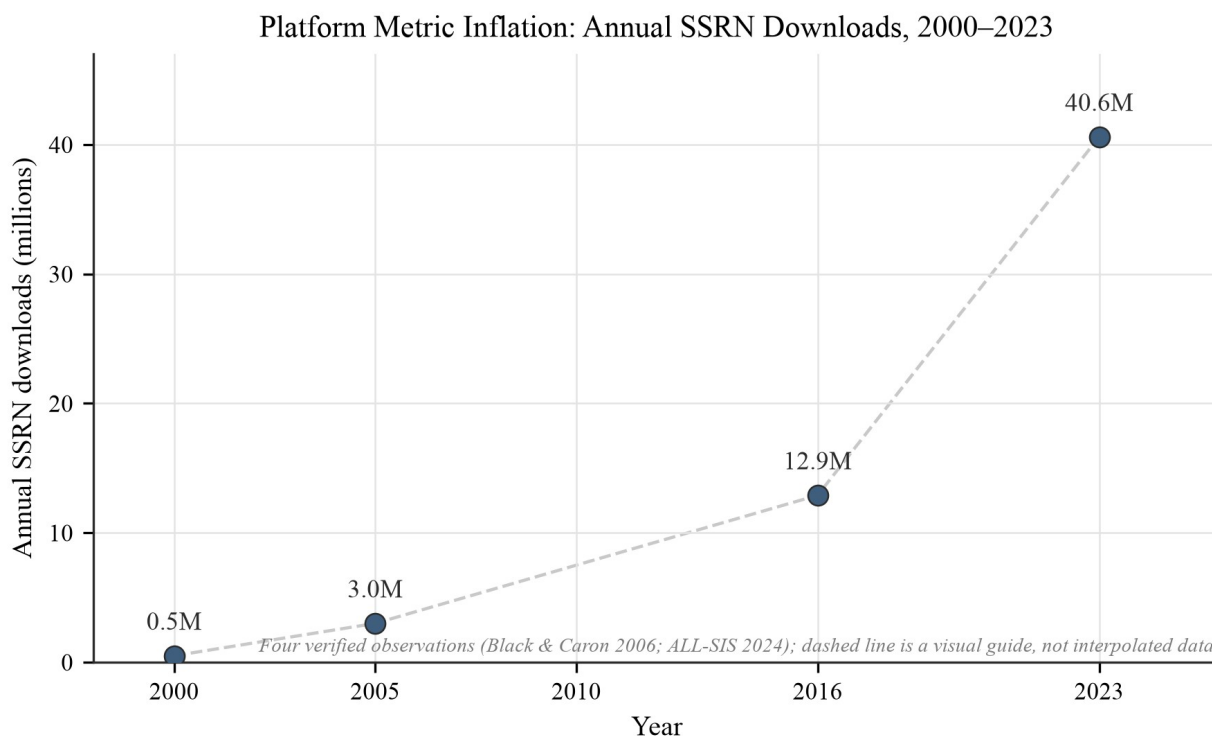


Figure 5. Platform metric inflation, 2000–2023. Annual SSRN downloads: ~0.5 million (2000) and ~3 million (2005) [18]; 12.9 million (2016) and 40.6 million (2023) [1]. Because the download flow has grown faster than the paper stock ($\approx 3.15\times$ versus $\approx 2.03\times$ from 2016 to 2023 [1]), lifting download-per-paper flow by roughly 55%, absolute benchmarks estimated on 2012–2015 samples read as lenient rather than conservative bands for 2026 readings. Only the four verified year-values listed here are plotted; no interpolation between them is implied.

No current percentile source exists. No post-2020 platform-wide distributional study of SSRN downloads was identified in the Section 2.5 audit (the 2025 single-network snapshot [83], published in 2026, is the only recent published external sample), and after 15 July 2026 the platform itself publishes no comparative data from which percentiles could be reconstructed. The gap is one of currency and platform specificity. The distribution and inequality of preprint attention is a mature research area on other repositories (the citation and altmetric profile of bioRxiv deposits has been characterized across the entire server [62]), and the quasi-lognormal shape of SSRN downloads specifically was already established before 2020 by Kakushadze [85]. What is missing is a *current*, SSRN-specific, post-sunset percentile distribution. The map above is assembled from the available frames, which is the argument for the community benchmark dataset proposed in the research agenda (Section 8).

5.2 The conversion band

The ratio of downloads to abstract views is the one dashboard quantity with a reasonably stable cross-study norm, and authors misread it more consistently than any other. Read naively, every view that does not convert counts as a rejection, and an author with 40 views and 7 downloads perceives an "82.5% failure rate." That reading is wrong on mechanical and on empirical grounds.

Section 4 established the mechanics: the denominator (views) is the more bot-contaminated and the more architecture-inflated quantity, since push-model email click-throughs and crawler traffic accumulate views without implying reader intent, while the download counter is actively audited [144][99][106]. So

long as automated traffic inflates views more than downloads, a measured conversion rate is a *lower bound* on the conversion of genuinely interested human visitors; Section 4.2 explains why that asymmetry can no longer be assumed where AI systems retrieve full text, and the bound weakens accordingly. Empirically, the published measurements converge on a narrow band. Siems measured a median downloads-per-view ratio of 0.171 (mean 0.187) in the LSN census [139]; Willey and Knapp report conversion of 15% among their top-article group and 11% among single-citation articles [172]. Jo and Li's snapshot yields a median of 18.9 downloads per 100 abstract views [83]. The resulting empirical band is **11–19%**, that is, approximately **one download per 5.3–9 abstract views**. SSRN's own dated rule of thumb of $\sim 3 : 1$ ($\sim 33\%$) [18] sits above the band and is best read as an early-era, platform-wide upper anchor rather than a current norm. As a cross-domain illustration (the populations and intents differ categorically), a browse-to-commit rate in this range would be considered high in consumer software, where typical freemium free-to-paid conversion benchmarks cluster in the low single digits [10]. The comparison is illustrative only, and it reframes the intuition that "7 out of 40" signals disinterest. Figure 6 plots the estimates.

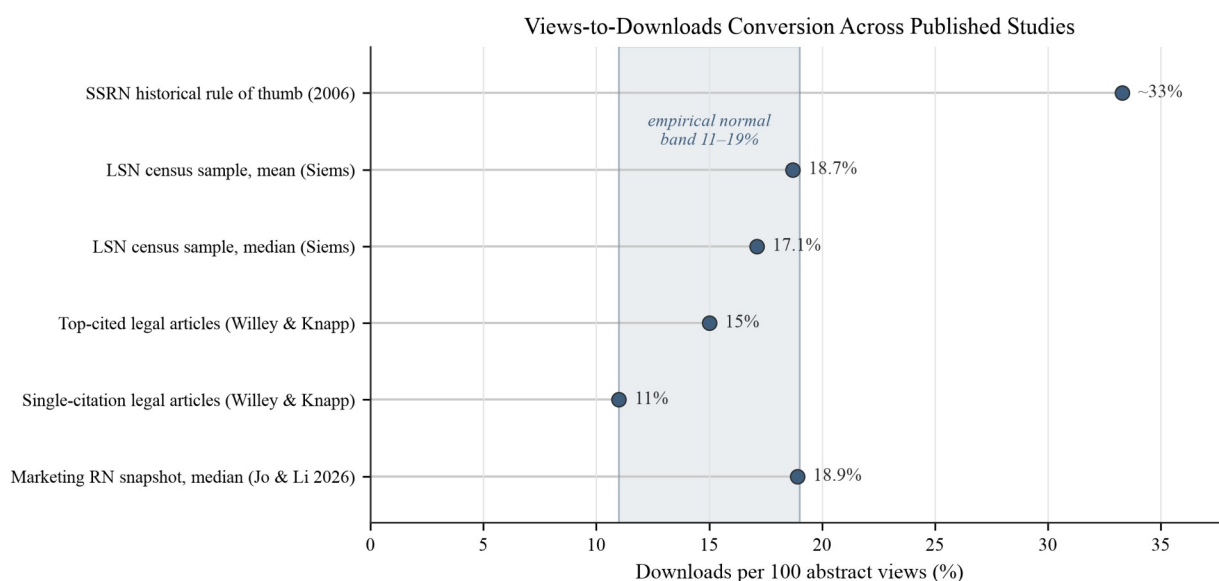


Figure 6. Views-to-downloads conversion across studies. Conversion estimates across the published record: SSRN's historical rule of thumb of $\sim 33\%$ (as reported in 2006) [18]; LSN census mean 18.7% and median 17.1% [139]; 15% among top-cited and 11% among single-citation legal articles [172]; median 18.9 downloads per 100 abstract views in the 2025 Marketing Research Network snapshot [83]. The shaded band marks the empirical normal range of 11–19% — an envelope of heterogeneous estimates (whole-network medians alongside citation-selected subsample figures whose statistic type the source leaves unspecified); the Jo and Li sample mean (23.7 per 100) lies above it. All values are as published in the cited sources; no new estimates are introduced.

The band's members are different statistics of different samples, so it is the envelope of heterogeneous point estimates. The Siems (17.1%) and Jo and Li (18.9%) figures are medians of whole-network samples, whereas the Willey and Knapp 15% and 11% are reported for citation-selected subsamples (top-cited versus single-citation legal articles), whose statistic type the source does not specify and whose selection is the reason the band's lower edge sits where it does. The upper edge is softer still: the mean conversion in the Jo and Li sample, 23.7 downloads per 100 views [83], lies *above* the band, and it diverges sharply from the ratio of that sample's aggregate means ($221.346 \div 611.973 = 36.2$ per 100).

That a ratio-of-means of 36.2 sits so far above a mean-of-ratios of 23.7 is itself diagnostic. Because the ratio-of-means is algebraically the views-weighted average of the per-paper conversion rates, its exceeding the unweighted mean-of-ratios implies a *positive* covariance between reach and conversion: higher-view papers convert at somewhat higher rates, not lower. The reading is conservative: small-denominator inflation would push the unweighted mean-of-ratios *upward*, working against this ordering, so observing the ratio-of-means above it anyway makes the positive-dependence conclusion more robust, not less.

The band supplies this article's central worked example. On the published anchors, a paper with 40 abstract views would be expected to yield roughly 4–8 downloads. The "7 downloads" of the title, on 40 views, is a conversion of 17.5%, a point estimate that lands inside the band but carries wide uncertainty at this denominator. The exact 95% Clopper–Pearson interval on 7 of 40 runs from about 7.3% to about 32.8%, straddling the band on both sides: with only 40 views the estimate cannot be resolved to the band's width, so the reading is a *screening result, not a classification*. Two things further limit the example. The denominator is small (the instability noted below), and the comparison is not age-matched: 40 views at roughly day 20 is an early-life reading, whereas the anchor medians come from mature samples (Siems at two years [139], Jo and Li at a median paper age of 491 days [83]), and Positions 1 and 5 and Section 5.6 argue that early-life exposure is a different regime, so the age axis is not aligned. The defensible reading is therefore that 17.5% is consistent with the cross-study band as a screening result, not a precise population percentile, and not "comfortably inside" it. The distress such a figure causes reflects the absence of a scale rather than the absence of readership. The vintage and scope of the underlying anchors are set out in Section 9.

Small denominators make the ratio unstable. A paper with two views and one download has a nominal conversion of 50%, but that estimate deserves no more confidence than one built on 500 views; for any public-facing benchmarking, empirical-Bayes shrinkage toward the reference-class rate should prevent small denominators from producing extreme scores (Sections 8.3 and 8.8). The band also works as a diagnostic in both directions. Conversion far *below* the band, on a non-trivial denominator, usually indicates a title–abstract–text mismatch: readers arrive, read the abstract, and leave, an actionable communication problem (rewrite the title and abstract), not a verdict on the research. Conversion far *above* the band signals traffic that abstract browsing does not explain (media pickup, course adoption, or viral circulation) and should prompt a check of where the views originate. Neither deviation, by itself, establishes anything about research quality. A further qualification is a hypothesis rather than a finding. Conversion may vary systematically across fields as a property of audience reading style rather than of individual papers: large fields may exhibit "window-shopping" behavior, with many abstract visits and few downloads, while small, highly engaged communities convert at higher rates. The band above is estimated from too few frames to resolve this, so Study 1 (Section 8.6) includes field-stratified conversion among its outputs.

A diagnostic, never a target: the anti-gaming constraint. The conversion band is powerful because it can be read as a norm, and that power carries a hazard. Goodhart's Law holds that a measure adopted as a target ceases to be a good measure [67], a principle paralleled by Campbell's Law on the corruption of social indicators used for decision-making [27]. If the 11–19% band, or any component of the benchmark framework of Section 5.5, were adopted as a promotion *threshold* ("papers should sit inside the conversion band to count"), the ratio would immediately become an object of optimization rather than observation. The attack is concrete and cheap: because the denominator is the more manipulable quantity

(Section 4.2), a coordinated inflation of abstract views in fixed proportion to downloads would hold the ratio inside the band and render an engineered signal statistically indistinguishable from organic conversion. The framework's answer is a matter of definition. The conversion band, the reach–conversion diagnostic (Section 5.4), and the CSAB components (Section 5.5) are a *diagnostic vector*, never a target, and no component may be read as a promotion criterion. The machinery that resists gaming is built into the design: the band is stated as a *lower bound* because its denominator is bot-inflatable, so the reading is conservative against inflation by construction; the anomaly protocol of Section 8.3 is *flag, not delete*, so an engineered spike is preserved and made visible rather than silently absorbed; and the integrity filtering of Section 4.1 already screens the most detectable coordinated activity. Read as diagnosis, the band exposes gaming as an out-of-band or internally inconsistent signature; read as a target, it would invite the very coordination it is meant to detect. This position aligns the framework with the established research-assessment reform consensus: the San Francisco Declaration on Research Assessment (DORA), which rejects single-number and journal-brand proxies in favor of assessment on the research's own content [8], and the Leiden Manifesto, whose principles caution against exactly this substitution of an easy indicator for expert judgment [78].

Downloads have no public audit trail: a further reason to read the band as a floor. A structural asymmetry between citations and downloads sharpens the gaming concern. A citation leaves a public, inspectable trace: the citing document exists, is discoverable, and can be checked, so citation manipulation is in principle auditable by any third party. A download counter leaves no such public trail. It is a private tally on the platform's servers, with no external record of which requests were counted, filtered, or restated, which leaves it uniquely exposed to covert gaming that no outside observer can reconstruct after the fact. That exposure is an independent reason, beyond the bot-contamination of the denominator, to treat the measured conversion rate as a lower bound rather than a point estimate (subject to the Section 4.2 caveat on AI retrieval), and always to pair the counter with the anomaly protocol and the documented integrity filtering the framework relies on (Section 4.1). The absence of an audit trail is not a reason to distrust the counter; it is a reason to read it conservatively and to insist on the flag-not-delete discipline that keeps a manipulated reading detectable (Section 8.3).

The band is a calibration target, not a constant. One consequence follows for how the 11–19% figure, and every CSAB threshold, should be quoted: they are quantities to be calibrated by field, and by platform (Sections 4.4, 5.5). The mechanism for that calibration is already specified: the field-stratified Study 1 (Section 8.6) is designed to deliver the per-field conversion and download distributions the framework requires and is built to consume. Until Study 1 reports, the archival cross-section of Appendix E supplies a field-resolved estimate. Across the fourteen fields holding at least twenty papers, median conversion runs from 0.134 in Energy to 0.195 in Computer Science, a spread of about 1.5 times, and holding paper age fixed it runs from 0.134 to 0.224, about 1.7 times. Set against the three- to fivefold field spread in downloads themselves, conversion is the more portable of the two quantities, which is what makes it usable as a diagnostic. Age is the dimension it does not survive: within every field large enough to test, conversion falls as papers get older, with Spearman coefficients from -0.14 to -0.41 across the eight fields holding at least forty papers. A conversion band quoted without a paper age is doing less work than it appears to.

5.3 The graduated interpretation scale

Post-Rankings authors most conspicuously lack a graduated scale, in the established format of h-index interpretation guides [118], that converts raw counters into plain-language bands. Sections 5.1–5.2 supply

the ingredients for it. Table 8 presents the proposed scale for four quantities: paper-level lifetime downloads, first-month downloads, views-to-downloads conversion, and author-level total downloads.

Table 8. Proposed interpretation scale for SSRN metrics. *The bands synthesize 2012–2015 evidence [139][85][18][172]; because download-per-paper flow has risen since then (Section 5.1), their absolute thresholds are best read as lenient rather than conservative bands for 2026, and in any case as an approximation. The first-month column is not a calibrated rung of the scale but a preregistered prediction carrying wide uncertainty, to be tested by the prospective panel of Study 1 (Section 8.6); its cells must not be read as thresholds or promotion rungs (see text). The Strong/Exceptional boundary at 10,000 is an author's provisional estimate: no published source identified in the Section 2.5 audit ties 10,000 downloads to a percentile. The Nascent band introduces no new numbers: its ceiling is the existing two-year median anchor of 63 [139]. Band boundaries are set at round numbers and kept roughly one full band wide: the precision ceiling established by the integrity literature (Sections 3.1, 4.1–4.3) makes finer gradations spuriously precise, so the coarseness of the bands is by design. The plain-language verdict column describes the downloads scale only and does not transfer to the conversion column: an 18.9% median conversion is ordinary rather than "Solid," and a 27.2% p75 is not "Strong." The lifetime-download boundaries are non-overlapping by construction (Nascent below 63, Modest 63–159, Solid 160 and above); where the table prints a shared endpoint between the higher rows — 800 between Solid and Strong, and the corresponding conversion and author-level cutoffs — the interval is half-open, each shared value belonging to the higher band. The two lifetime boundary values are themselves drawn from incommensurable reference classes: 63 is a two-year-window median for a single-discipline legal-scholarship census [139], whereas 160 is an all-time, lifetime entry threshold for the whole-platform top-list [85]. They therefore sit on different clocks (a fixed two-year paper age versus unbounded lifetime accumulation) and describe different populations (one field versus all fields), so pooling them into a single absolute ladder is a combination of reference classes with different ages, fields, and selection rules — a limitation the field- and age-specific percentiles of Study 1 are designed to remove. A structural limit of these anchors deserves emphasis: because the available reference points cluster at the two-year median (63 [139]) and in the platform's upper tail, four of the five verbal bands — Modest, Solid, Strong, and Exceptional — describe papers at or above the median, so the scale is dense above the median and discriminates only weakly across the broad 50th-to-98th-percentile range in which most papers actually fall. This compression is a direct consequence of anchoring on top-tail evidence; resolving the ordinary from the merely-above-median awaits the field- and age-specific percentiles of Study 1. The scale is illustrative, not normative: no band boundary should enter a tenure, hiring, or promotion decision as a pass/fail threshold until the prospective panel of Study 1 (Section 8.6) has calibrated field- and age-specific percentiles to replace these dated absolute anchors.*

Band	Paper: lifetime downloads	Paper: first-month downloads (extrapolated)	Conversion (downloads / views)	Author: total downloads	Plain-language verdict
Nascent	Below the two-year median of 63 [139]	— (no separate anchor; see text)	— (denominator typically too small; Section 5.2)	—	Too early, or below the field median — the statistical norm for young and niche papers; diagnose exposure first (Sections 5.4, 5.6), not quality
Modest	63–159	<15	<11%	<~130	Typical of the platform; the median band

Band	Paper: lifetime downloads	Paper: first-month downloads (extrapolated)	Conversion (downloads / views)	Author: total downloads	Plain-language verdict
Solid	160–800	15–50	11–19%	~130–1,600	At or beyond historical top-list entry
Strong	800–10,000	50–250	19–30%	1,600–10,000	Typical of the historical top lists
Exceptional	>10,000 (author's provisional estimate)	>250	>30%	>10,000 (author's provisional estimate)	Beyond the documented range of ordinary scholarly circulation

Each interior anchor is tied to a published threshold. The Nascent band covers the region below the two-year LSN median of 63 [139]. That region is populous, not anomalous: the prespecified Study 1 bands (Section 8.6) anticipate that a substantial share of the platform's papers, on the order of 15–30% under the stated prediction band, have ten or fewer lifetime downloads, so a scale whose lowest band began near the median would leave the platform's most numerous papers, this article's protagonist among them, outside the scale entirely. The Nascent verdict is non-evaluative. Below the field median, and especially at a young paper age, the counter licenses no judgment of quality, and the first diagnostic step is exposure (Sections 5.4 and 5.6) rather than content. The Modest band is anchored at the two-year LSN median of 63 [139]. "Modest" is the *median* band: typicality is not failure. The Modest/Solid boundary (~160 lifetime downloads) is the Kakushadze entry point to the all-time Top-10,000 papers, approximately the top 1.8% of the platform's full-text paper stock (571,040, July 2016 [1]) [85]. Operationally, a "Solid" paper is a top-list paper. The Solid/Strong boundary (~800) is the median *within* that elite group (805) [85]: it separates papers that merely enter the historical elite from papers typical of it. On the portfolio scale the author-level column mirrors this logic: the group SSRN's own notifications described as the top 10% of authors [161] had an entry minimum of 129 total downloads and an internal median of 1,626 [85]. Because every author-level anchor is drawn from that ranked list, the author column describes the upper tail of the author distribution rather than the typical author. An author below the ~130 boundary is not merely "Modest": that author falls outside the historically ranked group altogether, so the column's lower bands read as "below the historical elite," not as a calibrated position among all authors. This is the same top-tail anchoring that compresses the paper-level scale; the author column's anchors await the portfolio-level percentiles of Study 1. The conversion column is the cross-study band of Section 5.2. The Strong/Exceptional boundary at 10,000 carries no published percentile anchor and is retained as an order-of-magnitude horizon.

The first-month column has a different status. It is a *preregistered prediction*, a falsifiable forecast with wide uncertainty that Study 1 (Section 8.6) is designed to test, and no first-month cell may be used as a threshold or interpretive rung. No direct first-month benchmark exists in the published literature, so the column is a projection rather than a measurement, extrapolated from the documented early concentration of downloads (attention spikes cluster in the first weeks, particularly under external coverage [18][116] [138]) onto the lifetime thresholds above. The specific boundaries derive from one published anchor: in Kakushadze's top-list sample, approximately the top 1.8% of the platform's full-text paper stock (571,040, July 2016 [1]), the median *twelve-month* download count was 247 [85]. A first month above ~250 would therefore exceed the annual norm of the platform's historical elite, hence the Exceptional boundary, while 15–50 first-month downloads, if sustained, project onto the ~160 lifetime entry threshold, hence Solid.

These cells are projections from a single historical anchor and are the weakest cells of the table; if quoted, they carry this derivation.

Two caveats govern use of the scale. The disciplinary qualifier is not optional: within a single discipline, mean downloads per paper differ by roughly a factor of eight across subfields (438.35 in Corporate/Securities/Finance Law versus 53.48 in Agricultural Law) [168], so every band statement should carry a field label. Those are lifetime means and are not adjusted for paper age, and the archival cross-section of Appendix E shows why that matters: a field's apparent level moves with the age profile of the papers sampled from it, and holding age fixed narrows the spread to roughly three to five times. The second is that alternative single-number heuristics proposed in this space (ratio-based relative-performance indices (a paper's downloads divided by the median of a matched cohort), logarithmic tier heuristics based on $\lfloor \ln D \rfloor$, and verbal scorecard verdicts ("fast start / steady interest / slow start")) are subsumed here as uncalibrated author heuristics, because none has validated numeric bands. The logarithmic intuition itself is sound and worth keeping in words: under a quasi-lognormal distribution [85], moving from 7 to 20 downloads is a shift of one natural-log unit in visibility, and reading it as the addition of 13 clicks misses that.

The scale also explains *why* authors systematically misjudge their own counters. The numbers that circulate in milestones, profiles, and announcements come, as Section 3.5 documented, from a visible tail hundreds of multiples above the median, while the mass of the distribution sits near the LSN median of 63, and under the quasi-lognormal law the distances between rungs are multiplicative, not additive. Figure 7 arranges these values (ordinary-paper median, entry thresholds, elite medians, tail) as a benchmark ladder on a logarithmic scale.

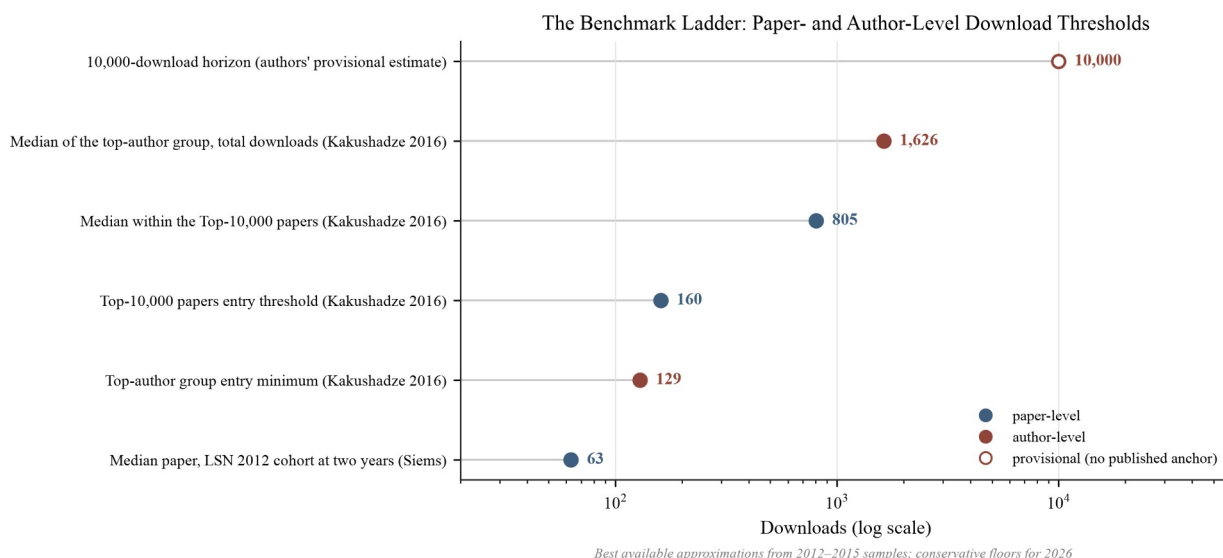


Figure 7. The benchmark ladder. Paper- and author-level download thresholds on a logarithmic scale: LSN two-year median 63 [139]; top-author entry minimum 129 and top-paper entry ~160 [85]; median within the Top-10,000 papers 805 [85]; median of the top-author group 1,626 [85]; the 10,000-download horizon (author's provisional estimate, no published percentile anchor). The region below the two-year median of 63 corresponds to the Nascent band of Table 8. Values are from 2012–2015 samples (inflation since then: Section 5.1).

5.4 Reach versus conversion: a two-dimensional diagnostic

A single count, even a banded one, cannot distinguish between the different mechanisms that produce it. The mechanics of Section 4 have one implication that matters most in practice: visibility failure and content non-response are *different diagnoses*: a paper can be little downloaded because few readers encountered it, or because those who encountered it declined the full text. Table 9 crosses the two normalized dimensions.

Table 9. Reach–conversion diagnostic matrix. *Reach and conversion are evaluated as percentiles within a reference class of comparable papers (same field, posting cohort, paper age, and prior author visibility; Section 5.5), with the class median as the boundary. No absolute numeric cut-offs are defensible on current evidence, and none are given. The 2×2 display dichotomizes each axis at the class median for readability; the underlying reach and conversion percentiles are continuous (Section 5.4), and finer partitions such as terciles are available.*

	Low conversion	High conversion
Low reach	"Cold start" (low reach, low conversion). The paper is neither widely encountered nor frequently downloaded after exposure. The data do not establish whether the underlying research is weak; they establish only that diffusion has not begun.	"Hidden interest" (low reach, high conversion). Readers who encounter the paper respond well, but too few encounter it. Discoverability is the primary observable bottleneck — a distribution problem, not a content problem.

	Low conversion	High conversion
High reach	"Attention without uptake" (high reach, low conversion). The title, topic, or platform placement appears to attract visits, but few visitors obtain the full text. Rule out a contaminated denominator before diagnosing a communication mismatch: high reach may be <i>phantom views</i> rather than genuine attention, because views are the more polluted of the two counters (Section 4.2) — in the transparent RePEc/LogEc record, robots accounted for roughly 74% of abstract views but only about 7% of downloads [99], so unfiltered non-human traffic produces exactly this high-reach, low-conversion signature. Only once bot-inflated views are excluded do the residual data warrant investigating audience fit and title–abstract communication.	"Strong early diffusion" (high reach, high conversion). The paper reaches a comparatively large audience and converts exposure into full-text interest.

The four quadrants are labeled by the diffusion pattern the two percentiles show (cold start, hidden interest, attention without uptake, strong early diffusion). The labels describe a pattern of diffusion, not the research itself, and what matters operationally is the separation of mechanisms; neither low-reach quadrant licenses an inference about research quality, a point that bears directly on any evaluative use of these metrics. The low-diffusion quadrant takes the name of this article's own diagnosis, the cold start. The matrix is practically valuable because the appropriate intervention differs by quadrant: a hidden-interest candidate needs distribution (classification review, announcement timing, external visibility), whereas attention without uptake calls for communication repair (title and abstract) or a traffic audit. Figure 8 renders the plane as a template.

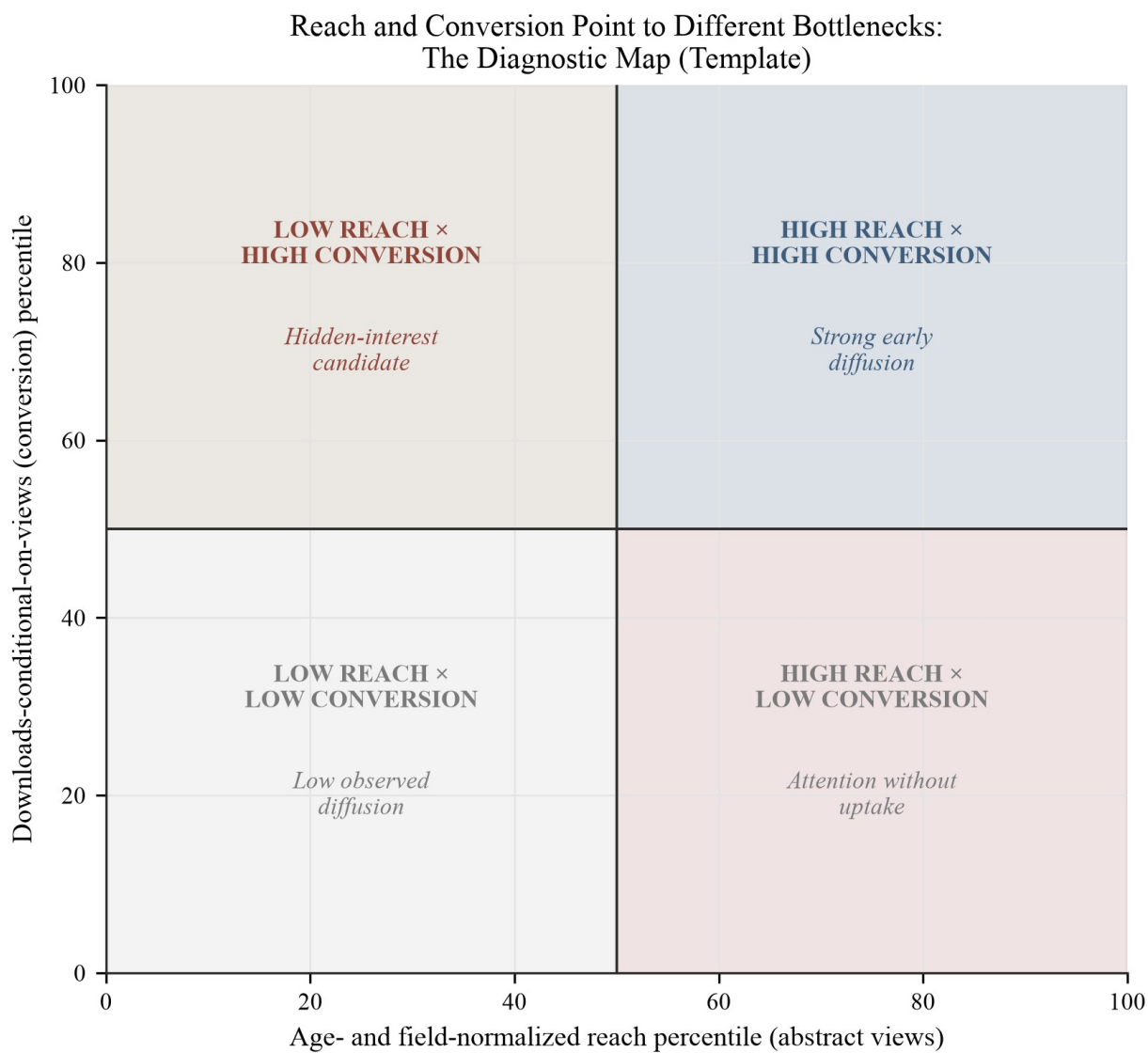


Figure 8. The reach–conversion map. Template for positioning papers by age-, field-, cohort-, and visibility-normalized view percentile (reach, horizontal) and downloads-conditional-on-views percentile (conversion, vertical), with quadrants as in Table 9 (cold start, hidden interest, attention without uptake, strong early diffusion). The upper-left quadrant identifies papers with above-median conversion despite below-median reach, a class that raw download counts systematically undervalue. No empirical values are shown; the map is to be populated from the prospective panel or community benchmark dataset of Section 8.

5.5 From a headline number to a benchmark: the CSAB dimensions

Each of the preceding subsections repairs one failure mode of the raw counter, and the Cold-Start Scholarly Attention Benchmark (CSAB) is that integration: a reporting framework that replaces the single headline number with five normalized components (Table 10). The normalization variables (paper age, field, posting cohort, and prior author visibility) are each empirically motivated: cumulative metrics

mechanically increase with age; field norms differ by up to an order of magnitude in the mean [168]; calendar conditions shift attention measurably (article accesses fall ~11% on weekends and ~50% in summer months [178]); and a platform newcomer does not face the same initial audience as an established author [83]. One population convention has to be fixed before any percentile is computed. A cumulative counter is defined over a paper stock with heterogeneous survival (papers are occasionally withdrawn, merged, or superseded by later versions), so every CSAB percentile is a statement about the surviving stock at the observation date. Study 1 (Section 8.6) prespecifies the inclusion rule (version handling, withdrawals, and duplicate resolution) rather than leaving the reference population implicit.

The two central components admit a compact definition. For paper i at age τ (days since posting), in field f , posting cohort c , and prior-visibility stratum s :

$$B_i^V(\tau) = 100 \cdot \hat{F}_{f,c,s,\tau}^V [V_i(\tau)] \quad \text{and} \quad B_i^D(\tau) = 100 \cdot \hat{F}_{f,c,s,\tau}^D [D_i(\tau)]$$

where $V_i(\tau)$ and $D_i(\tau)$ are cumulative abstract views and downloads and $\hat{F}^V$, $\hat{F}^D$ are the empirical distribution functions of those counts within the reference class defined by f , c , s , and τ . The definition fixes the *estimand*, a percentile within a stated reference class. No $\hat{F}$ can be written down today, because the reference distributions do not exist (Section 5.1). The formula is the specification that Study 1 and the panel of Section 8 are designed to estimate, and the same notation carries through to the predictive-validation design of Section 8.9. Percentile and reference-class normalization is itself an established bibliometric method: percentiles are a standard way to normalize a publication's counts for field, document type, and publication year [22], with dedicated procedures for the field- and time-normalization of indicators dominated by zeros [23], exactly the regime a download counter occupies. Normalizing *usage* metrics specifically (downloads and views rather than citations) against a reference class is likewise a mature line of work: the usage-bibliometrics literature has long treated download-based impact as a distribution to be conditioned on the characteristics of the user community and the sampling frame [92] [19], and journal-usage-factor programs built exactly such normalized usage indicators from COUNTER-compliant data [137]. The standardization of usage counting itself is the subject of a mature practitioner infrastructure, the COUNTER Code of Practice [125]. Field-normalized readership indicators have a comparable established precedent in the altmetrics literature [179]. The method is therefore established. The contribution here is its application to SSRN downloads and the construction of an age-, field-, and cohort-stratified reference class for a post-Rankings environment.

The estimand is clean, but three operationalization questions must be resolved before any $\hat{F}$ is fit. The first is the most consequential, because it concerns what the benchmark is *for*. Conditioning the reference class on prior author visibility (the stratum s) is correct when the question is diagnostic: "how is this paper doing among comparable papers by comparably visible authors?" It is exactly wrong when the question is the equity question of H1–H3, because normalizing *by* prior visibility subtracts out the newcomer-versus-established gap that the equity frame exists to measure. A within- s percentile cannot detect cold-start disadvantage, since the disadvantage has been conditioned away by construction. The framework therefore separates two distinct uses of CSAB that must never be collapsed into one number. As a *descriptive benchmark* it reports an author's position within a fully stratified reference class (including s), which is what an individual author wants at the dashboard. As an *adjusted estimand* for H1–H3 it drops prior visibility from the conditioning set and treats it instead as the *predictor of interest*, measuring the gradient rather than removing it. The two answers differ, they answer different questions, and any report that presents a visibility-normalized percentile must state which of the two it is. A second question is

sample size. Stratifying by field $\times$ cohort $\times$ visibility stratum $\times$ eight age points quickly exhausts a probability sample of 3,000–5,000 papers (Section 8.6): the fully crossed cell count runs into the thousands, so most cells are tiny or empty. Empirical-Bayes shrinkage toward broader reference classes is named for this reason (Sections 8.3, 8.8), but shrinkage requires its own specification: a minimal hierarchical structure (field within broad area, age within coarse age band) and a prespecified minimum-cell rule *for estimation*, below which a cell borrows from its parent rather than reporting a raw percentile. This estimation-side minimum-cell rule is distinct from, and stricter to reason about than, the privacy-driven minimum-cell rule that governs public *release* in Appendix D; the two serve different purposes and are specified separately. The last question is momentum (Table 10), defined as the change in an *estimated* normalized position. A change score in an estimated quantity invites regression to the mean: a paper that lands high at day 7 partly by estimation noise will tend to fall by day 30 for purely statistical reasons. Momentum therefore requires an explicit model of the underlying trajectory (for example, a shrinkage or empirical-Bayes estimate of the latent position with its measurement error propagated) rather than a raw difference of two percentile snapshots, and the raw difference is reported, if at all, only as a descriptive companion to the modeled quantity.

Table 10. The Cold-Start Scholarly Attention Benchmark (CSAB): five dimensions instead of one number. *A copyable sentence-level reporting template for these dimensions is provided in Appendix B.*

Dimension	Metric	Core question	Interpretation
Reach	Age/field/cohort/newcomer-normalized view percentile	Did relevant audiences encounter the paper?	Discoverability
Download performance	Normalized cumulative download percentile	How many full-text acquisitions occurred relative to comparable papers?	Overall early diffusion
Conversion	Downloads conditional on view opportunities (Section 5.2)	Did exposed readers seek the full text?	Conditional engagement
Momentum	Change in normalized position or interval-level attention	Is diffusion accelerating, stable, or fading?	Temporal trajectory
Downstream validation	Later citation or use conditional on early metrics	Did early attention predict subsequent scholarly uptake?	Predictive validity

The framework rejects any universal raw threshold of "success." Success must be defined relative to an analytical purpose: for *discoverability*, a high reach percentile among papers of the same age and field; for *reader uptake*, above-reference-class conversion conditional on exposure; for *diffusion*, the joint distribution of the two; for *scientific influence*, downstream citations or documented use. Early usage alone is insufficient for that last purpose, since early downloads explain only on the order of 16% of later citation variance [25]. That makes the interpretation of the title example exact. Seven downloads is "good" only to the extent that seven downloads places the paper unusually high among otherwise comparable papers at the same stage of their life cycle. A day-7 value and a day-180 value cannot be evaluated on the same benchmark.

Four design choices follow the eighth principle of the Leiden Manifesto: that quantitative indicators must not be allowed to outrun the strength of the assumptions behind them, and that false precision is itself a failure of measurement [78]. The framework reports *bands*, not point scores (Table 8); it fixes CSAB as a diagnostic vector rather than a target (Section 5.2); it draws its output figures, the reach–conversion map

(Figure 8) and the benchmark fan (Figure 14), with no numerical scale, so that a shape cannot be mistaken for a calibrated value; and it attaches a verification tier and vintage to every number it uses (Table 7). Each of these withholds precision the evidence cannot support, as DORA asks of research assessment when it rejects single-number and journal-brand proxies for the substantive evaluation of a paper's content [8]. Collapsing CSAB back into one score would reintroduce exactly the false precision both documents warn against.

To an author, CSAB reporting would replace the bare line "Downloads: 12," with something like this: *12 downloads after 14 days: overall download percentile among comparable papers: 67th; reach percentile: 39th; conversion percentile: 81st; interpretation: high conditional interest, below-median exposure.* Every number in this mock-up is illustrative, since no such percentiles exist until the reference distributions of Section 8 are estimated, and none of these values may be quoted as a benchmark, but the format shows what the decomposition buys. The same raw count that reads as an undifferentiated disappointment resolves into a specific diagnosis: readers who encounter the paper respond unusually well, but too few encounter it, so the indicated intervention is distribution, not revision. The worked example is the CSAB counterpart of the "hidden interest" quadrant of Table 9. Appendix B fixes this format as a copyable reporting template whose percentile fields remain blank until the Section 8 distributions exist.

5.6 Reading trajectories: exposure is a treatment, not background noise

The final component of the framework is temporal: a newly posted paper does not enter a homogeneous attention environment. It passes through distinct exposure regimes (posting → internal discovery → email distribution → external indexing → network diffusion), and a cumulative total collapses these regimes into one number. The regime boundaries are set by the machinery of Section 4: eJournal distribution is curated and can take up to 45 days in high-volume topics [53], and external search indexing can take 6–9 months or longer [55].

The most anxiety-producing dashboard state follows directly from this: **zero downloads in week one is normal pre-distribution, not rejection.** Before the announcement cycle has run, the paper has not yet been exposed to its principal audience. The diagnostic sequence for an early zero is administrative: confirm the paper is publicly posted; confirm eJournal assignment; wait out the announcement cycle; only then read the counters. Conversely, first-month counts, once distribution has occurred, are the earliest and noisiest leading indicator: early downloads correlate with later citations at roughly $r \approx 0.4$ [25], informative but far from a verdict, and early spikes are frequently exposure events (blog or media coverage) rather than durable interest [116][138].

Because retrospective scraping of cumulative counts cannot reconstruct what a paper had on day 3 or day 30, trajectory-based interpretation (and the panel research agenda of Section 8) requires observation from birth on a standardized grid. Table 11 specifies the minimum schedule. It doubles as a self-tracking template for individual authors. A trajectory is classified only provisionally at the 60- and 90-day readings of the Table 11 grid and with confidence at the 180-day reading, never at day 7.

Table 11. Minimum longitudinal observation schedule for SSRN papers.

Paper age	Core purpose of the observation
Day 1	Immediate platform exposure
Day 3	Very-early discovery

Paper age	Core purpose of the observation
Day 7	First-week benchmark (expected value for most papers: at or near zero, pre-distribution)
Day 14	Early stabilization
Day 30	First-month benchmark
Day 60	Distribution and search maturation (eJournal cycles complete [53])
Day 90	Medium-run early diffusion
Day 180	Transition toward downstream impact measurement

Where feasible, daily observation over the first 45 days is preferable, because email distribution can occur at variable times within that window [53]. Interval-level readings on this grid convert the dashboard from a static verdict into a trajectory, the only form in which early SSRN metrics carry defensible information.

Sections 4 and 5 replace the question an author cannot answer ("is 7 downloads good?") with four that the published record can support: where does the count sit against the benchmark map and its vintages (5.1, 5.3); is the conversion inside the empirical band (5.2); which quadrant of the reach–conversion plane does the paper occupy (5.4); and what does the trajectory on the observation grid show (5.6)? None of these questions requires the retired Rankings, but all of them require the comparative data whose collection the final sections of this article propose.

6. A Practical Playbook for Authors

The benchmark map of Section 5 (Table 7) restores the missing scale, and this section puts it to work for the author who opens the statistics panel after 15 July 2026 and sees counts without context [148][149]. It has four tasks in turn: a routine for observing one's own counters (6.1–6.2), substitutes for the comparative frame that the rankings used to supply (6.3), levers that can legitimately improve visibility (6.4), and a protocol for converting counters into defensible career evidence (6.5–6.6). Every numerical anchor used here traces to a published source with an explicit sampling frame and time window; Table 7 consolidates them. Practices circulating in the author community that have no evidentiary support are marked as unverified wherever they appear and are collected, as candidate hypotheses for testing, at the end of Appendix A.

Table 12 compresses the playbook's core into a protocol of seven questions, asked in order, each paired with the evidence worth preserving and the strongest conclusion the answer can legitimately license. Sections 6.1–6.2 unfold this protocol into routines and reading rules, and Table 23 (Appendix A) remains its situational companion, indexed by the dashboard states authors actually encounter.

Table 12. An author protocol for interpreting an SSRN counter: seven questions, the evidence to preserve, and the permissible conclusion. *No step treats a counter as a quality score; the final column states the ceiling of what each answer can support.*

Step	Question	Evidence to preserve	Permissible conclusion
1	How old was the paper when the count was observed?	Dated screenshot or export, plus the posting date	A time-specific observation, not a lifetime verdict

Step	Question	Evidence to preserve	Permissible conclusion
2	Had curated distribution and external indexing occurred?	Distribution notice and date where available; search-index status	Identification of the pre- versus post-exposure regime (Section 5.6)
3	What are views, downloads, and conversion together?	Both counters and the denominator	A reach-versus-conditional-engagement diagnosis (Table 9), not a single verdict
4	What is the narrowest defensible reference class?	Field, cohort, paper age, newcomer status, classification history	Relative standing only within that class (Section 5.5)
5	Is the trajectory changing?	Repeated observations at the fixed ages of Table 11	Momentum or plateau, subject to exposure events — not acceleration of "quality"
6	Is there downstream corroboration?	Citations, substantive contacts, publication or documented use	Layered evidence of uptake, not proof of quality
7	Can the claim be reproduced?	Source, date, counting rule, archived record	An auditable report, not an unsupported milestone

6.1 A measurement routine

Interpretation begins with disciplined recording, because SSRN displays only current cumulative counts and no history under the author's control. Six steps, applied from the day of posting, produce the minimal record that every later inference requires.

1. **Record both counters and their ratio.** Note abstract views, downloads, and the downloads-to-views conversion ratio at fixed intervals. The observation grid of Table 11 (days 1, 3, 7, 14, 30, 60, 90, 180) specifies the schedule and the purpose of each reading. A dated screenshot at each interval is the only version of record the author controls, since platform counters can be restated when filtering rules change [144][69]. The accumulating record is a **versioned evidence packet**: the manuscript version, its metadata and classifications, dated dashboard snapshots, distribution notices, a dated log of the author's own promotion events (a posted thread, a talk, a mailing-list announcement), and later citations and substantive contacts. The promotion log matters because promotions are exposure treatments. Without a dated record of one's own publicity, an author cannot distinguish an organic spike from a self-generated one, and neither can anyone auditing the claim (Section 8.3).

2. **Normalize by time since posting.** Express early performance as downloads per month over the first three to six months rather than as a raw cumulative count. A three-month-old paper compared against the all-time totals of older work yields systematically misleading conclusions, because counters accumulate over a paper's lifetime on a platform whose annual download volume has grown from roughly 12.9 million in 2016 to 40.6 million in 2023 [1].

3. **Never mix observation windows.** A weekly count may not be compared against a twelve-month average, and a 90-day count may not be compared against lifetime totals. Every comparison in this playbook pairs like windows with like. Violations of this rule were among the most common interpretive errors in the pre-sunset use of ranking statistics [161].

4. **Situate the paper in its disciplinary neighborhood, not against the platform.** Mean downloads per paper differed by roughly a factor of eight across Legal Scholarship Network eJournals alone (438.35 in Corporate/Securities/Finance versus 53.48 in Agricultural Law, 2018 data) [168]. A count that is unremarkable in one subfield is exceptional in another. The relevant reference class is the paper's own eJournal or subject network, examined at the same paper age.

5. **Do not over-interpret single small counts.** The integrity literature places a ceiling on the precision of usage counters: publisher-reported downloads can run twice those of comparable platforms [176], platform filters discard a large share of raw traffic [99][129][144], and without such filters SSRN's own counters would be inflated "by a factor of five or more" [18]. Small differences between comparable numbers therefore carry no information.

6. **Monitor long-run accumulation rather than reacting to early figures.** Attention on working-paper platforms concentrates early [116][28] (the latter finding derives from NBER rather than SSRN data), but discovery layers mature slowly: SSRN reports that Google Scholar indexing takes "6–9 months or longer" [55]. Readership, moreover, decays on several timescales at once, and reading and citation need not follow the same temporal process [91], so no single "half-life" reading of one's own curve is safe. A trajectory can first be classified provisionally at the 60–90-day readings of the Table 11 grid and with more confidence at 180 days, never at the day-7 reading.

6.2 Five rules for reading the dashboard

With that record in hand, the distributional evidence of Section 5 comes down to five rules of thumb for reading a dashboard day to day.

(i) **Compare against medians, not means or visible stars.** The only census-style sample of ordinary papers puts the median at 63 downloads and 369 abstract views after two years (LSN, papers posted 2012) [139]. In the one recent published sample the median is 87 downloads, against a mean of 221.3 that sits above even the 75th percentile (Section 3.3) [83]. The milestone numbers that circulate come from the visible tail, whose 2015 maxima ran to hundreds of thousands of downloads (Section 3.5) [85]; an author who calibrates against them is calibrating against the tail of a quasi-lognormal distribution, which guarantees miscalibration. The platform-wide average of ~222 lifetime downloads per paper computed from December 2023 aggregates [1] is just as unusable as a personal yardstick.

(ii) **Check the conversion ratio, not just the count.** Independent measurements converge on a normal band of roughly 11–19% downloads per abstract view: the LSN sample's mean of 18.7% and median of 17.1% [139], and the 15% (top-cited articles) and 11% (single-citation articles) figures of the legal-citation studies [172]. SSRN's long-standing three-views-per-download rule of thumb [18] sits above that band rather than within it. The recent Marketing-network sample sits slightly above that band, with a mean of 23.7 downloads per 100 views [83]. Figure 6 plots the estimates together. Run the central worked example of Section 5.2 through the rule: the title's 40 abstract views and 7 downloads convert to 17.5%, inside the band as a point estimate, but with a wide 95% Clopper–Pearson interval (≈ 7.3 –32.8%). At $n = 40$ the reading is not failure and not a precise percentile: a screening result consistent with ordinary conversion. Conversion far below the band is consistent with a mismatch between the abstract and the paper, where readers arrive, read the abstract, and leave. That is actionable, and the action is to revise the title and abstract. Conversion far above the band signals traffic that abstract browsing does not explain, such as direct links from media or aggregator coverage, and it should prompt a check of where the views originate. Both readings are screening diagnostics: no study has causally validated the inference from low

conversion to abstract quality, and the diagnosis should be recorded as a hypothesis about one's own paper.

(iii) Treat first-month downloads as the earliest, and noisiest, citation signal. Early downloads correlate with later citations at roughly $r \approx 0.4$ in the arXiv data, which the original authors themselves translate into only about 16% of explained variance [25]. The earliest such signal on record is biomedical: online hit counts in the first week after publication in the BMJ correlated at roughly 0.50 with citations accrued over the following years [119]. In legal scholarship, Black and Caron reported paper-level downloads–citations correlations on the order of 0.50, rising to about 0.55 in a sample of one hundred junior law professors [18]. Across fields and horizons the paper-level correlations range from roughly 0.1 to 0.6 [80][25], and arXiv download counts lose statistical significance as predictors once contemporaneous Twitter attention is controlled [138]. At the aggregate level the association tightens: 0.84 between download-based and citation-based rankings of one hundred law schools [77], and $R = 0.724$ between six-month PDF downloads and later citations in one journal's own analysis [114]. Those aggregate figures do not settle where any single paper will end up, and a first month of downloads is an early signal about later citations.

(iv) Zero downloads in week one is normal pre-distribution, not rejection. A newly posted paper earns attention only after it clears moderation, is classified into subject eJournals, and enters the announcement cycle. SSRN's own documentation describes classification as curated rather than guaranteed, with delays of up to 45 days for heavily loaded topics [53], and eJournal distributions and email alerts were explicitly confirmed to continue unchanged after the sunset [147][57]. The diagnostic sequence runs: confirm that the paper is publicly posted, confirm eJournal classification, wait out the announcement cycle, and read the counters only after that. Discovery through search engines arrives later still [55].

(v) A falling or flat counter can reflect filtering, not lost readership. SSRN publicly documents that it discards apparent repeat downloads, robot downloads, and other irregular activity [144][66][69]. The RePEc/LogEc experience shows how large such corrections can be: in July 2007 robots accounted for roughly 74% of abstract views but only about 7% of downloads [99], and in September 2025 more than 99.5% of raw IDEAS traffic was discarded amid AI-bot interference [129]. A sudden drop, or a restatement, is therefore as likely to be a measurement event as a readership event, and it should never be read as reputational loss without corroboration.

6.3 Benchmarking without rankings

The rankings supplied a comparative frame. An author today has three substitutes for it, ordered from least work to most.

Micro-benchmarks. The most direct replacement is a reference group one builds oneself: identify three to five methodologically comparable papers by peers at comparable institutions, posted in the same month, and track their public counters on the same Table 11 schedule as one's own. The result is a localized percentile baseline against exactly the right reference class, matched by construction for field, paper age, and platform era, where the retired rankings compared papers of all ages and fields in a single league. It is a proposed procedure that uses only public per-paper counts, and it should be recorded as such in any dossier use. It also presupposes some familiarity with a peer literature that the most isolated newcomers, the very authors the cold-start condition most affects, may not have. For them the platform-

side remedy of Section 10.4 (anonymous percentile bands) is the structural fix, since it requires no self-assembled network to use.

A relative performance ratio. The same logic can be expressed as a single number: the ratio of a paper's downloads at a given age to the median downloads of a matched cohort, meaning the same discipline, the same paper age, and where possible the same paper type. A ratio of 1.0 means exactly median visibility for the cohort, and values above or below 1.0 are read accordingly. The index orders papers within a matched cohort; percentile bands await the reference distribution. The published medians from which a cohort baseline could be built are aged and narrow (63 downloads at two years for legal scholarship posted in 2012 [139], 87 for marketing papers observed in 2025 [83]), and mapping ratio values onto percentiles requires the full distribution, not the median alone. Assigning percentile labels to ratio bands without that distribution produces false precision. The pre-registered agenda of Section 8 includes calibration of the bands; until then the ratio is an ordering device, not a grading scale. Two further cautions bound the device. It should not be computed at all on small counts, since the small-denominator instability of Section 5.2 applies to the ratio's numerator and to a thin cohort median alike, and empirical-Bayes shrinkage (Sections 8.3, 8.8) is the appropriate repair wherever such a ratio is reported publicly. A given ratio does not translate into a fixed percentile: "twice the cohort median" corresponds to different percentile positions in distributions of different shapes, so the ratio orders papers within one cohort but cannot be compared across cohorts or eras whose distributional shapes differ.

Verdict language for trajectories. Community-facing tools have proposed classifying a paper's early trajectory into three qualitative verdicts (fast start, steady interest, slow start) as a humane alternative to raw numbers. The vocabulary is worth retaining, since it converts a count into a statement about a trajectory, and "slow start" is explicitly not a synonym for failure. The thresholds originally proposed for these verdicts relied on comparisons of mismatched time windows and on a platform-wide "network average" statistic that does not exist as a public metric: SSRN offers no public self-service API or statistics endpoint [51], and the nearest historical metric ("New Downloads Per Paper," an author-level ranking statistic) was retired with the rankings ([B-9], cited with the caveat that the historical ranking pages are no longer accessible). Assign verdicts, then, only from matched comparisons: the micro-benchmark group or the cohort-median ratio above, read no earlier than the 60–90-day mark.

Whichever substitute is used, external triangulation strengthens it. Citation-layer data from Crossref and OpenAlex [35][124] and the PlumX indicators displayed on SSRN pages [50] answer a different question, about downstream use rather than visibility. A policy-citation tool such as Overton extends the same triangulation to policy and practice reach by mapping where a paper is cited in government and institutional documents [187], which suits SSRN's law, economics, and policy base. These layers are examined jointly with the counters in the CSAB framework of Table 10 and the reach–conversion diagnostics of Table 9. Altmetric indicators are a further triangulation layer. Conceived from the outset as early attention signals [123][122], they follow patterns distinct from citations [74][20] and admit field normalization [21]. They are also vulnerable to automated inflation [75] and require theory-guided interpretation rather than face-value reading [76].

6.4 Improving legitimate visibility

Before the sunset, the Rankings supplied crude but real feedback for one recurring author decision: whether further promotion of a paper was worth the effort (Section 1.3). Post-sunset, that decision has to be made from the levers and diagnostics assembled here. The determinants literature identifies a small set

of levers under the author's control and consistent with platform rules. Each acts on visibility, not quality, and none changes what the paper contains. Rapid dissemination and early feedback are among the recognized purposes of preprint posting in the methodological literature on preprints [24], and a playbook for legitimate visibility uses that stated function.

Packaging. In the LSN sample, title characteristics, abstract length, and topic salience were significantly associated with attention, with the author's own caution that such variables may proxy for unobservable quality [139]. In the legal-citation studies, top-decile papers carried more keywords and longer abstracts than bottom-decile papers [172] (the specific counts await verification against the original). The abstract page is the conversion surface on which rule (ii) operates, so when conversion sits persistently below the band the title and abstract are the first things to revise. Any such revision aims at alignment between promise and content, not indiscriminate traffic: rewriting an already clear abstract to maximize clicks can damage its precision, and by attracting mismatched visitors it can push conversion further from the band rather than toward it.

Classification. Papers classified into more eJournals averaged more downloads and fewer zero-download outcomes in the 2018 LSN analysis [168]. SSRN permits one to seven classifications per paper and curates inclusion [53]. A popular community refinement, unverified because no study compares classification strategies, caps the classifications at three, chosen as one core, one broad, and one interdisciplinary eJournal; that advice belongs among the testable hypotheses of Appendix A.

Distribution and external attention. The eJournal email alert remains the platform's principal push channel and continues post-sunset [147][57][152]. Off the platform, blog coverage produced sharp early download spikes in the one empirical test of the question [116], visible list position raised readership and citations in the arXiv positional studies [72], and conference participation was associated with an additional 17–26 downloads per paper in a quasi-experimental study built around a cancelled political-science meeting ([B-4], cited with the caveat that the setting is APSA 2012, not an SSRN-economics sample). Practitioner guides collect further distribution tools for legal scholarship [169]. The case for early open dissemination is independently established: a randomized trial of open access found more downloads though no short-run citation gain [37], while cohort and field studies report citation advantages of roughly 10% to twofold [59][86][17][93], including +53% for American law reviews [40], and an audit of legal citations found that 87% of the most-cited articles were available on SSRN versus 44–46% of little-cited articles, with 95% of the most-cited group posted on SSRN in or before their year of publication [172]. Preprint deposit is associated with higher citations, and lower publisher-site downloads, on arXiv [38], and bioRxiv downloads correlate with journal outcomes [62]. In the census of all bioRxiv preprints, download counts correlate with the impact factor of the journal in which a preprint is later published (Kendall $\tau_b = 0.59$), an association whose direction of causality is unknown [2]. Machine discovery layers (search engines and generative answer engines) are treated separately in the audit framework of Section 8 and the feature-audit agenda of Table 20 [5][164].

Versioning hygiene. Revisions uploaded to the same SSRN page preserve the accumulated counters and the stable URL, whereas deleting and re-posting resets the metrics to zero ([B-7], cited with the caveat that support-page wording changes over time and should be re-verified at submission). The further community claim that a minor revision on day 3 triggers re-announcement and a "second wave" of attention is undocumented, and it sits among the unverified practices in Appendix A. Identity needs the same care. Linking a verified ORCID iD to the SSRN author profile keeps papers, and their counters, attached to the correct person rather than to a namesake, and it is a low-cost guard against the author-

disambiguation problem that the panel design of Section 8.3 must otherwise solve statistically ([B-7], same re-verification caveat).

What not to do. Self-downloading and related manipulation are documented in SSRN's own server logs, intensified near top-list thresholds, and were driven by social comparison [44]. The platform's fraud filters are public [144][69], and evaluation committees aware of this history discount unsupported claims. So the gaming literature carries a practical warning alongside the integrity one: manipulated counts are detectable and reputationally expensive. Figure 9 illustrates the temporal signature at issue on simulated data. Organic attention accumulates as an early peak with a long tail, whereas manipulation produces concentrated spikes of exactly the kind integrity filters are designed to flag [144][69].

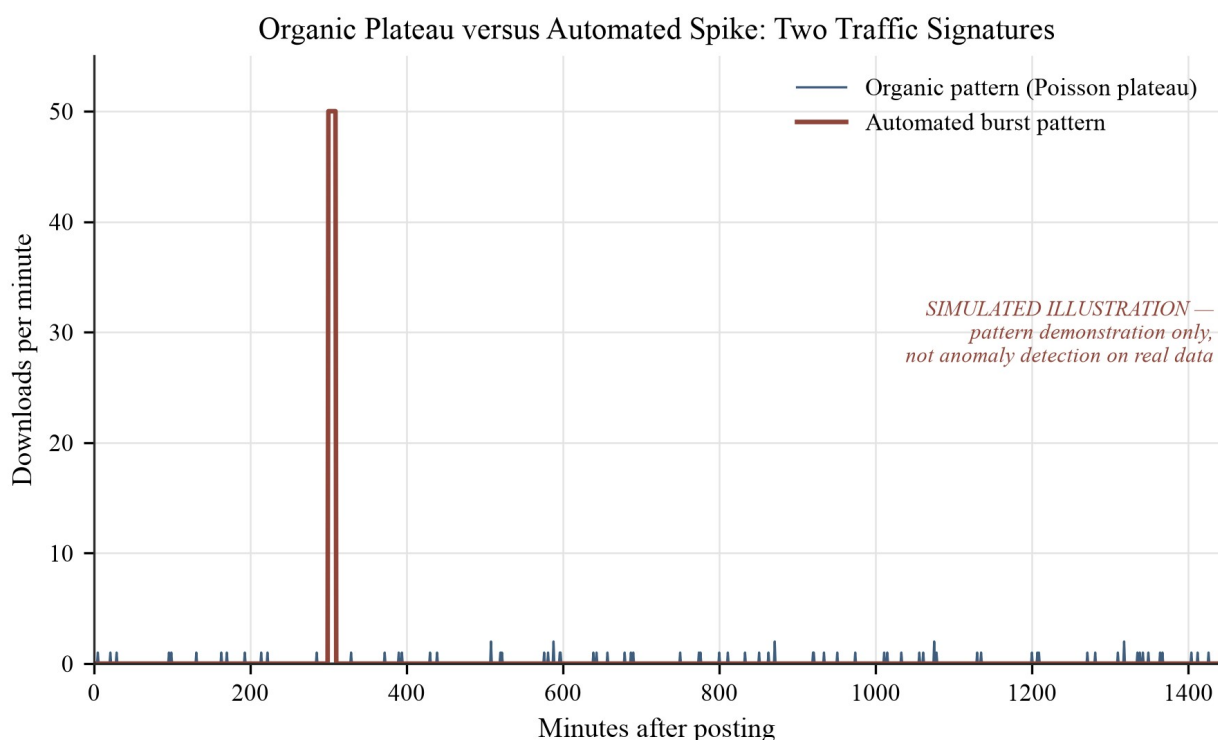


Figure 9. Organic accumulation versus a manipulation spike. A smooth organic download trajectory contrasted with an abrupt single-day spike characteristic of self-downloading, motivating the manipulation warning of Section 6.4 and the integrity discussion of Section 4 [44][69][144]. All data shown are simulated; no empirical anomaly-detection dataset exists for SSRN.

6.5 Converting counters into career evidence: the dossier protocol

Before the sunset, rank-based SSRN evidence was a standardized dossier component. An author's public rank entered promotion files directly: "ranked 13th out of 3,000 legal authors in total number of all-time downloads, which the law school has taken into account in various promotion and related decisions" [12]. Cover-letter guides advised citing download counts and Top Ten list appearances [95], practitioner guidance simultaneously cautioned that self-posted work "probably will not count for law-school promotion and tenure purposes" [33], and the platform's own documentation described its metrics as data "that authors use for promotion and tenure purposes and that institutions use in their hiring and ranking activities" [52]. Every one of these artifacts presupposed a public comparative frame that no longer exists [149]. The evaluative use of the counters was contested from its first proposal, moreover: Black and

Caron's 2006 case for downloads as a "beta" measure of scholarly performance [18] was answered in the same symposium volume by Eisenberg's critique of what SSRN-based rankings can validly measure [45]. A dossier that cites downloads inherits that unresolved debate. Table 13 sets out a seven-step protocol for reconstructing defensible evidence from what remains, and its governing requirement, applied in the third and fourth columns, is that **every number cited in a dossier must travel with its source, its time window, and its limitations**. A translated claim that hardens into unsourced precision is worse than no claim at all.

Table 13. Post-sunset evidence protocol for SSRN metrics in tenure, promotion, and hiring dossiers. *Each step names the action, the evidentiary basis with its time window, the limitation that must be disclosed alongside any number, the anti-pattern to avoid, and the strongest conclusion the step can license.*

Step	Action	Evidentiary basis (source; time window)	Limitations to disclose	Avoid	Permissible conclusion
1. Preserve	Archive dated screenshots of each paper's statistics panel and the author profile at the Table 11 intervals and before any submission; assemble them, with distribution notices and a dated log of the author's own promotion events, into the versioned evidence packet of Section 6.1.	Per-paper counters survive the sunset ([148][149]; 2026). Counters can be restated when filtering rules change ([144][69]; documented since 2018).	A dated personal archive is the only version of record under the author's control; platform numbers are revisable and no public history exists.	"The paper has X downloads" — quoted without date or denominator.	A dated, auditable record of what the platform displayed — a time-specific observation, not a lifetime verdict.
2. Translate	Convert counts into band language using Table 8 (e.g., "lifetime downloads above the ~160 entry threshold of the 2015 Top-10,000 papers list, i.e., approximately the top 1.8% of the platform's full-text paper stock, 571,040 in July 2016 [1], at that time"), always naming source, sample, and year; the CSAB template of Appendix B supplies the sentence-level format, with its percentile fields left blank until a reference panel exists.	Paper-level thresholds: entry ≥ 160 , in-group median 805 (top-list scrape; 2015) [85]. Author-level: entry minimum 129 total downloads, in-group median 1,626 (same scrape) [85]. Typical-paper anchors: median 63 at two years (LSN; 2012→2014) [139]; median 87 (Marketing network; snapshot 2025) [83].	All thresholds derive from 2012–2015 samples (inflation since then [1]: Section 5.1); the 2015 frame covers only top lists; the two median anchors are single-discipline. Band language must not be restated as a calibrated percentile.	Universal labels such as "excellent" hung on an old raw threshold.	Relative standing within the named sample, field, and year — nothing broader.

Step	Action	Evidentiary basis (source; time window)	Limitations to disclose	Avoid	Permissible conclusion
3. Triangulate	Pair SSRN counters with independent citation evidence: Google Scholar, HeinOnline for law, RePEc for economics, Crossref/OpenAlex as open citation layers, PlumX as displayed on SSRN.	School-level SSRN-download and HeinOnline-citation rankings correlate at 0.84 (100 U.S. law schools; 2019) [77]; paper-level download–citation correlations run ~0.1–0.6 depending on field and horizon [80][25]; open citation infrastructure [35] [124]; PlumX [50].	The 0.84 figure is a <i>school-level</i> correlation and does not descend to any individual paper or author (ecological-inference caution, Section 6.2(iii)); downloads measure demand, not reading or endorsement; paper-level correlations are modest, and highly downloaded and highly cited sets overlap only partially [108]. Convergent citation evidence, not the counter alone, carries the impact claim.	Equating a download with endorsement or correctness.	Layered evidence of uptake, not proof of quality.
4. Field-normalize	State the disciplinary frame explicitly (e.g., "within corporate-law scholarship, where eJournal means differ by roughly a factor of eight across legal subfields").	Subfield means 438.35 vs. 53.48 downloads/paper (LSN eJournals; 2018) [168]; downloads-to-citations ratios are themselves field-structured [18].	The spread estimate is law-only and dated 2018; cross-disciplinary transfer of any threshold requires caution, and percentile language without a field qualifier is uninterpretable.	Comparing a new paper with a multi-year cumulative total, or transplanting a threshold across fields.	A field- and age-qualified comparison — never a universal grade.
5. Complement	Cite institutional-repository and other distribution channels alongside SSRN.	SSRN counters say nothing about downloads of the same paper from other services [161]; the 2026 reforms made repository-based series the migration path for institutional distribution [47] [146].	No single platform's counter measures a paper's total circulation; the dossier claim is about the paper, not about one counter.	Presenting one platform's counter as the paper's total circulation.	A claim about the paper's circulation, with each channel's counter named separately.

Step	Action	Evidentiary basis (source; time window)	Limitations to disclose	Avoid	Permissible conclusion
6. Integrity	State the platform's counting rule alongside any cited count, and keep the dated archive that documents the number as displayed.	Filtering and restatement are documented platform behavior ([144][69] [66]; ongoing); vendor-reported download statistics depend on counting conventions ([69] [176]; 2015–2019).	Counters are filtered, revisable estimates; no public restatement history exists, so the author's dated archive is the only stable record.	Treating counters as exact numbers of unique human readers.	A filtered estimate under a stated rule at a stated date — not a census of readers.
7. Expert assessment	Present every metric line as context for qualitative evaluation of the work's contribution and rigor, never as its replacement.	Metrics measure stages of diffusion, not validity (Sections 3.1, 5.5); highly downloaded and highly cited sets overlap only partially [108].	No usage metric, normalized or raw, assesses whether the research is correct or important.	Metric substitution for peer judgment.	Contextual evidence accompanying expert review — the quality claim rests with the review, not with the number.

The Avoid column collects the protocol's anti-patterns. Two warnings bound the protocol as a whole. A download records passive demand rather than reading, endorsement, or quality, and counts respond to publicity dynamics (blog coverage, salient titles, topical waves) as much as to scholarship [116][139] [138]. Committees that know the manipulation record [44] will discount unsupported numbers, and the ingredients that survive such discounting are exactly those the protocol requires: dated screenshots, band language with sources attached, field qualifiers, and triangulation.

6.6 The solicitation corollary

Posting has one documented side effect that no dashboard reports: unsolicited invitations from predatory journals and conferences, produced by the harvesting of names, affiliations, and paper titles from preprint platforms [107][159]. The volume is substantial in adjacent audits, where specialty studies document invitation streams ranging from dozens per academic per fortnight to thousands per specialty group per month (Section 7.2) [9][84][63]. The first contact with a predatory journal is an unsolicited email in about 41% of surveyed cases [32]. Section 7 and Table 14 analyze the phenomenon. For the playbook the consequences reduce to four instructions. Expect solicitations as a default consequence of visibility, and never read their volume as scholarly recognition: the senders target visibility mechanically, not merit selectively [159][63]. Route correspondence through a dedicated address so harvesting does not contaminate a primary inbox; a widely repeated refinement, concealing the profile email address for the first two weeks after posting, is an unverified practice (Appendix A). Never pay, because the business model is the fee, with article-processing charges in audited predatory outlets ranging from \$25 to \$1,800 [136]. The money is the smaller stake: the same evaluation committees that scrutinize dossiers (Section 6.5) also read bibliographies, and a known predatory venue on a publication list imposes a reputational cost that outlasts any fee. Verify any inviting venue against a maintained screening service such as

Cabells Predatory Reports ([B-1]; Beall's list has not been updated since January 2017 and must not be used as a live source), and against the checklists of Think. Check. Submit. [159]. The harvesting mechanism also supplies a qualitative screen: a journal that emails an invitation within days of a preprint's posting, before any human readership could plausibly have formed, is overwhelmingly unlikely to be legitimate. That is a heuristic; no study has measured its incidence [63].

6.7 Scope of the playbook

The playbook's anchors date from 2012 to 2015 and are law- and marketing-heavy [139][85][83][1]. Download-per-paper flow has risen since then (Section 5.1). Fresh, field-stratified percentile data do not yet exist; their collection, by the platform or by the community, is the subject of Sections 8 and 10. The playbook reduces misreading, and it does not measure quality. Four rules make its use mechanical: compare matched cases, hold the observation window constant, attach a source to every number, and label a hypothesis as a hypothesis.

7. The Parasitic Attention Layer: Visibility and Predatory Solicitation

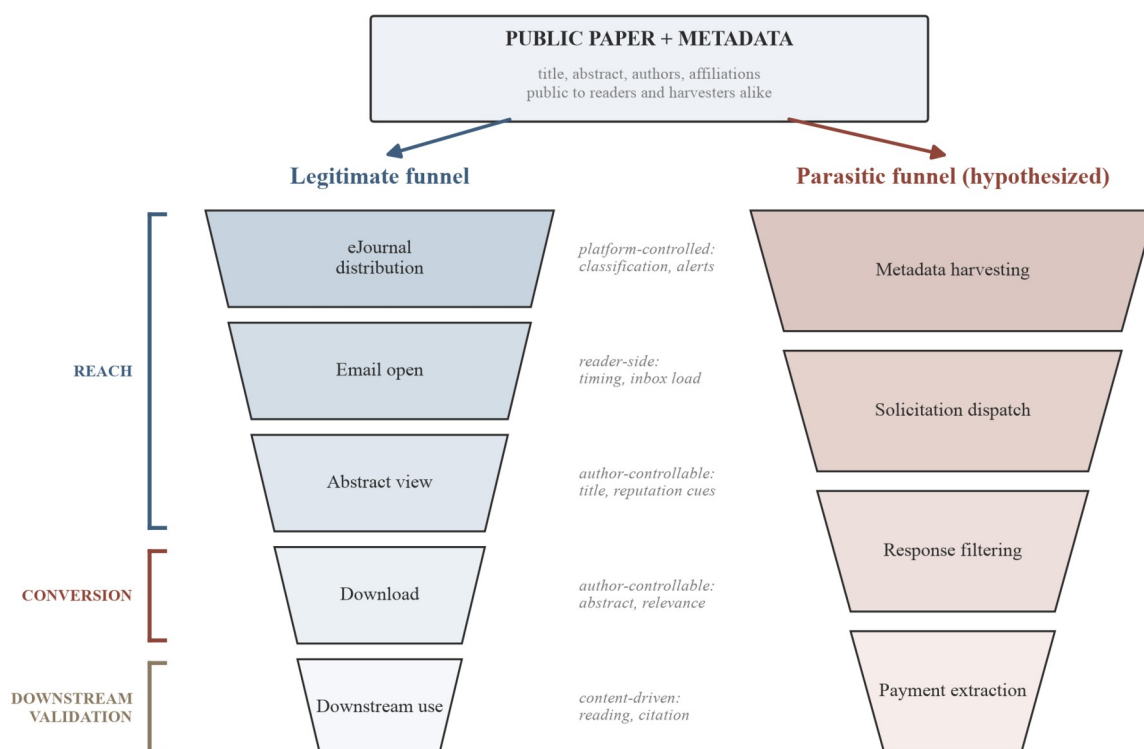
The interpretive framework developed so far treats attention as a single funnel running from exposure to abstract view to download to downstream use. The lived experience of posting a preprint, however, frequently includes a second stream of "attention" that arrives through a different channel entirely: unsolicited email invitations from journals, conferences, and editorial services of questionable legitimacy [14][61][117][142]. A conceptual account of that second stream follows, together with the audit evidence for its documented scale.

7.1 A Two-Layer Model of Attention

The post-posting environment is treated here as a **two-layer attention economy** (Figure 10). Layer 1 is the legitimate funnel described in Sections 4 and 6: eJournal distribution, abstract views, downloads, and eventual scholarly use. Layer 2 is a **parasitic layer**, an ecosystem of predatory publishers and service providers that feeds on the same public infrastructure (titles, abstracts, author names, and affiliations) and aims to convert authors into paying customers, not to read their work. The solicitation-audit literature documents the lifecycle. Public preprint metadata (SSRN titles, abstracts, author names, and affiliations) is harvested [107][159], mimetic invitations that mirror legitimate calls for papers are dispatched [61][63], and authors who engage are converted into paying customers. The economic engine is documented too: article processing charges at presumed predatory journals average roughly \$178, with a reported range of \$25–\$1,800 [136].

The asymmetry between the two layers matters for interpretation. Layer 1 requires a human reader to notice, open, and value the paper, while Layer 2 requires only that the paper's metadata become machine-readable. A preprint can therefore in principle be fully processed by the parasitic layer while remaining effectively unread in the legitimate one. That is the situation of the platform's near-zero-download papers, the "ghost papers," which machines process and humans effectively do not read.

One Public Posting, Two Funnels: The Two-Layer Attention Economy



*A low final count can originate at different stages — and different stages call for different diagnoses.
Conceptual model — no measured stage-to-stage conversion rates exist for SSRN; none are shown. A counter records an event, not an actor;
an unsolicited email is evidence of harvesting, not of scholarly uptake.*

Figure 10. Conceptual model of the two-layer attention economy. Both funnels are shown as qualitative stage sequences; no stage-to-stage conversion rates are displayed because no such rates have been measured for SSRN. The parasitic funnel is a hypothesized reconstruction based on the solicitation-audit literature [61][63][107][159] (Section 7.3).

7.2 The Documented Scale of Predatory Solicitation

That predatory solicitation is a large-scale phenomenon is not in doubt. A 2026 review identified twenty-nine empirical studies of academic spam, converging on a repertoire of tactics: flattering solicitation, false editorial claims, and dubious metrics and indexing assertions [61]. The measured volumes are substantial. Fourteen academic ophthalmologists received 1,813 solicitation emails from 383 unique publishers, spanning 696 unique journals, within a single 30-day window [84]. Academic radiologists averaged 20.7 manuscript invitations and 4.1 speaker invitations over a two-week period [9]. An exploratory audit cataloged 422 solicitation emails received by a single urologist over 2023–24 [63], and a case study in scholarly publishing logged 1,024 spam emails from journals over 391 days [89]. A year-long single-researcher observation following preprint postings recorded 875 invitations from 256 journals, 76% of which appeared on curated blocklists, a conditional source, cited here with the caveat that arrival lags after posting were not measured [B-10]. Solicitation also demonstrably converts. Among authors who had published in presumed predatory journals, 41% reported that their first contact with the journal was an

unsolicited email invitation [32], and content analyses of solicitation corpora document flattering language, fabricated deadlines, and vague topical scope as recurring persuasion devices [96].

Nor is the evidence confined to medicine. In the social sciences, disciplinarily far closer to SSRN's own base, three researchers in educational sciences received 210 unsolicited journal invitations within three months, traced to 139 journals and 37 publishers, most of them classifiable as predatory; roughly half of the invitations fell outside the recipients' fields of expertise, the signature of metadata-driven rather than readership-driven targeting [155]. The burden falls earliest on this article's protagonist, the early-career author. One longitudinal self-audit of an early-career academic logged 1,280 predatory solicitations (990 from journals, 220 from conferences, and 70 of other kinds) across roughly eight years, together with documented and largely unsuccessful attempts to unsubscribe [160].

Two features of this evidence base must be kept in view. The largest audits come from medicine, and the social-science audits above, while disciplinarily closer to SSRN, are smaller; no volume from either literature can be transplanted to SSRN authors. Decisive for this article, none of these studies measured the timing of solicitations relative to a preprint posting event. The audit in [63], for instance, cataloged content and provenance but did not link arrival times to SSRN activity. Every claim connecting the parasitic layer to the act of posting is therefore a mechanism hypothesis, not a finding.

7.3 Hypothesized Targeting Mechanisms

To reach a newly posted preprint, the parasitic layer must first locate it. The mechanisms below are hypotheses with partial support in the literature; no corpus ties solicitations to SSRN postings.

H-S1 (metadata harvesting). Predatory operations scrape public abstract pages and new-paper listings for titles, author names, and affiliations. Practitioner and advisory sources describe this behavior: predatory journals are reported to scan preprint servers and dispatch unsolicited invitations that reference the paper's title [159]. A legal-scholarship commentary states that solicitors "collect emails, names, affiliations etc by trawling existing scholarly papers, e.g. those posted on SSRN" [107], and a law-library guide links preprint posting to "a rise in spam from predatory publishers who solicit researchers to re-publish their articles for a fee" [117]. One measured regularity fits the mechanism. In the radiology audit, the specialty audits report high solicitation volumes among actively publishing authors [9], consistent with harvesting from an author's cumulative bibliographic surface rather than from any single posting event. Study 3 (Section 8.6) accordingly records author-level exposure controls alongside the focal posting.

H-S2 (contact acquisition). Solicitation requires an email address. SSRN abstract pages historically included author contact information [151], and addresses also circulate through the papers themselves, institutional pages, and prior publications. Whether SSRN currently exposes author emails in harvestable form is untested; the instrumented solicitation-logging study is specified to measure it. The acquisition channel for any given solicitation is unobserved.

H-S3 (distribution piggybacking). eJournal digests and RSS-type feeds redistribute new-paper metadata to subscribers [57][51]. A parasitic actor subscribed to these channels would receive structured leads without scraping. No audit has verified this route.

The causal chain from posting to being targeted has not been empirically established; Section 8 specifies a pre-registered test of these mechanisms: logging solicitation arrivals against known posting timestamps for freshly posted papers.

7.4 A Hypothesized Typology of Solicitation Signals

If the targeting mechanisms of Section 7.3 operate, different classes of solicitation should carry different information about a paper's machine visibility. Table 14 sets this out as a hypothesized typology, its strategy categories grounded in documented persuasion tactics: credential mismatch and off-target personalization [9][63], fabricated deadlines and flattery [96], false metric and indexing claims [61], and concealed fees against a documented APC market [136]. The final column is inferential and untested. The table gives no prevalence percentages and no arrival-time estimates.

Table 14. A hypothesized typology of predatory solicitation signals following preprint posting.

Example phrasings are illustrative, invented paraphrases of patterns documented in solicitation audits [9][61][96], not verbatim quotations from any specific message.

Solicitation pattern	Description	Illustrative phrasing	Hypothesized information about visibility
Generic address, no paper reference	Automated invitation with no mention of the specific paper or field	"Dear Esteemed Researcher, submit your valuable work..."	Low: the address may derive from any source; no evidence the paper itself was seen
Credential mismatch	Invitation citing expertise in a field unrelated to the paper [9]	"Given your expertise in quantum physics..." (to a law paper)	Low-to-none: indicates bulk harvesting without metadata parsing
Targeted, title-referencing	Invitation naming the specific preprint and a topically adjacent venue [159]	"Your recent paper '[exact title]' would suit our journal..."	Higher (hypothesized): the paper's metadata reached the sender, implying successful indexing and scraper access
Flattery personalization	Praise of the specific work as "groundbreaking" or "esteemed" [96]	"Your groundbreaking contribution..."	None beyond title access: praise is templated, not read
Fabricated urgency and metrics	Invented deadlines, asterisked "impact factors," indexing claims [61][96]	"Impact Factor 4.2*; submit within 7 days"	None: persuasion device independent of the target paper
Concealed fees	APC disclosed late or only in attachments; market range \$25–\$1,800 [136]	Fee absent from invitation body	None: pricing strategy, not visibility signal
Conference invitation, paper-referencing	Invitation to present at or submit to a questionable conference, often naming the specific paper; speaker invitations are a measured component of academic spam [61][9]	"We invite you to present your paper '[exact title]' as a keynote at..."	As for targeted solicitation (hypothesized): the paper's metadata reached the sender; no evidence of substantive reading
Paid author services	Offers of rapid peer review, translation, or editing services pitched at a freshly posted preprint [61][84]	"Fast-track review and professional language polishing for your new preprint..."	None established: service spam tracks new-posting metadata, not paper content

The flattery rows carry a psychological dimension that the mechanical reading obscures. Templated praise lands on an author who may be staring at a stagnant counter, for whom the solicitation is the only

"validation" the posting has yet produced. The two-layer model thus predicts a cruel inversion in which the first responders to a new preprint are frequently predators rather than readers (Section 7.5), and a newcomer can misread spam as interest because no legitimate signal has yet arrived. An author who expects the hook can discount it.

Table 14 reads incoming messages for what they might reveal about a paper's visibility. How to code those messages at all is the complementary question, and it turns on a distinction left implicit so far, between *legitimate scholarly contact* and higher-risk solicitation. Table 15 sets out that distinction as a six-signal typology whose final column carries the coding cautions: every signal is probabilistic and none is proof on its own. The typology is also the skeleton of the Study 3 codebook (Section 8.6), since a solicitation study without a positive class of legitimate contact cannot measure the asymmetry that Section 7.5 hypothesizes.

Table 15. Separating legitimate scholarly contact from higher-risk solicitation: six signals with coding cautions. *A hypothesized coding scheme, not a validated classifier. Table 14's visibility column asks a different question and is unaffected: the same message can be coded on both grids.*

Signal	Legitimate scholarly contact	Higher-risk solicitation	Coding caution
Specificity	Engages a concrete argument, method, or result	Generic praise with the title or keywords inserted	Specificity can be imitated
Venue identity	Verifiable journal, editor, institution, or conference history	Unclear ownership, copied branding, unverifiable board	New venues are not automatically predatory
Request	Substantive question, collaboration, or transparent invitation	Urgent submission or payment request unrelated to the work	Fees alone are not decisive; transparency matters
Timing	Any timing; often follows actual reading	Rapid templated contact soon after public posting	Temporal proximity does not prove the harvesting source
Language	Topic-specific and professionally accountable	Flattery, false familiarity, artificial urgency, vague scope	Language features are probabilistic, not proof
Verification	Independently verifiable contact details and policies	Broken identifiers, false metrics, misleading indexing claims	Human review under documented criteria is required

In practice the typology yields a screening heuristic rather than a diagnostic: a journal solicitation arriving within the first days after posting, before legitimate distribution channels have plausibly operated, is very likely predatory. Legitimate high-impact journals are not reported to cold-call authors of just-posted manuscripts [174], whereas early unsolicited invitation is the documented first-contact mode of predatory venues [32][174]. We advance the specific cutoff sometimes proposed in practitioner discussions (treating any journal email in roughly the first 96 hours as presumptively predatory) as an author heuristic pending validation. No study has measured the lag distribution that would calibrate such a threshold. The timing caution of Table 15 applies to the heuristic itself: temporal proximity raises suspicion but does not prove that the sender harvested this particular posting, since the address may predate the preprint entirely, so the 96-hour screen sorts inboxes rather than evidence.

7.5 The Spam Paradox and the Speed Asymmetry

Two more propositions about the parasitic layer follow from the two-layer model.

The Spam Paradox. If targeted solicitation requires successful metadata harvesting, then in the post-Rankings environment the *absence* of predatory spam may be a stronger indicator of platform invisibility than its presence is of anything else. A paper that never triggers even the parasitic layer may have failed to clear the minimal bar of machine-readable distribution. The analytic attraction of this inversion, with spam as an unwilling and toxic proxy for indexing, is that it turns an unavoidable nuisance into a free observational instrument. Its diagnostic power is at present only a hypothesis: the correlation between solicitation receipt and any visibility metric has never been measured, and the generic-solicitation rows of Table 14 show why raw spam volume would be a noisy signal even in principle.

Predators arrive before readers. The legitimate funnel operates on institutional lags: eJournal inclusion is curated rather than guaranteed, with delays of up to 45 days reported for high-volume classifications [53], and Google Scholar indexing takes six to nine months or longer [55]. Near-zero downloads in the first week are consistent with normal pre-distribution rather than rejection (Section 6). The parasitic layer, if it harvests metadata programmatically, faces none of these lags. We therefore hypothesize that for an author whose contact information is acquirable, the expected time to first predatory solicitation attributable to a posting is shorter than the expected time to first genuine scholarly contact, so that the unknown author would meet the parasitic layer before the legitimate one. The falsifier is direct: genuine contact reliably precedes solicitation.

Documented solicitation volumes measured in days and weeks [9][84], set against distribution mechanics measured in weeks and months [55][53], make the hypothesis plausible. Only the logging study specified in Section 8 can test it, because the existing corpora were not time-locked to posting events [63].

7.6 What the Parasitic Layer Does Not Tell an Author

Three boundary statements bound the conclusions of this section.

First, **solicitation email does not move the platform's counters.** An invitation sent *to* the author generates no abstract view and no download of the paper. To the extent scrapers do visit abstract pages, SSRN states that it filters irregular and automated activity from its download counts [66][69][144]. Spam volume is therefore never evidence of readership, and a solicitation-filled inbox coexisting with single-digit downloads involves no contradiction: the two layers count different things.

Second, the parasitic layer also falls hardest on the authors least equipped to absorb its costs: early-career and non-native-English researchers are noted as particularly vulnerable to solicitation [142]. The practical protocol of Section 6 therefore treats inbox hygiene (blocklist screening of any soliciting venue, non-response to unsolicited invitations, verification through curated services rather than defunct lists) as part of metric interpretation rather than a separate concern. Automated screening of predatory venues is itself an active research direction [157].

Third, early visibility has one further side effect, which authors encounter as an apparent accusation: on later journal submission, similarity-checking software will flag a near-complete match between the manuscript and its own SSRN precursor, since preprints are indexed in the Crossref Posted Content corpus used by iThenticate [36]. The vendor's documentation is explicit that a same-author preprint match is an expected result and not plagiarism [81]. The appropriate response is to disclose the preprint to the editor, not to withdraw it.

The parasitic layer is where this article's central claim is easiest to see: after the Rankings sunset, raw signals (downloads, views, and now even spam) carry interpretable meaning only when the mechanism that generated them is specified. That specification is partially available for the legitimate layer, and for the parasitic layer it is still a research program.

8. A Preregistration-Ready Research Agenda: Rebuilding the Missing Reference Class

The preceding sections synthesized the verifiable evidence on SSRN attention metrics and proposed an interpretation framework for the post-Rankings environment. This section is a preregistration-ready agenda for answering the framework's open questions empirically. It reports no results, by design: prespecified, falsifiable bands rather than findings. Every design, hypothesis, and prediction band below is ready to be registered on the Open Science Framework before any data collection, with analysis code released alongside. The retirement of SSRN Rankings on 15 July 2026 [149], announced on 12 June 2026 with an initially earlier target date [148] and embedded in a broader platform refocusing [146], makes this agenda unusually time-sensitive: some of the designs below exploit a discontinuity whose pre-period recedes with every month, and one of them, the percentile registry of Section 8.8, addresses a measurement gap that the sunset itself made permanent for external observers.

8.1 Design principles

Four principles govern the agenda.

Observation from birth. Retrospective scraping of cumulative counters cannot reconstruct what a paper had on day 3, day 7, or day 30. Any credible benchmark panel must therefore enroll papers at posting and record their counters at standardized ages: days 1, 3, 7, 14, 30, 60, 90, and 180 (Table 11), with daily observation preferable during the first 45 days, because SSRN states that alert distribution in heavily subscribed topics can take up to 45 days [53]. This principle also protects against leakage: prior-author reputation measures must be frozen at the focal posting date so that future success cannot contaminate baseline covariates. Its companion discipline, carried over from Section 5.6, is that distribution, indexing, and promotion events enter the designs below as dated treatments and never as unmodeled context.

Separation of reach and conversion. Following the diagnostic logic of Section 5 (Table 9), every design below records abstract views and downloads as distinct outcomes and models downloads conditional on views. Low reach and low conversion are different diagnoses: failure to enter the reader's consideration set versus non-acquisition after exposure. Pooling them into a single count discards the information an author needs.

Distributional outcomes, not only means. Because platform attention is heavy-tailed [83][85][180], interventions and platform changes can alter the shape of the distribution while leaving its mean nearly unchanged. Outcomes therefore include the Gini coefficient of attention, top-1% and top-10% download shares, the share of papers at or near zero, and newcomer-to-established attention ratios, in the tradition of attention-inequality measurement in other media [180].

Falsifiability. Each hypothesis is paired with an identification strategy and a stated condition under which it fails. The framework itself is exposed to rejection: if age- and field-normalized early metrics do not outperform raw counts out of sample (RQ3 below), the claimed measurement advantage of normalization

should be rejected or narrowed. Falsifiability is also quantitative: sample sizes throughout the agenda are to be finalized by power simulation against the smallest effect sizes of interest. The multiplicity strategy for a program of five hypotheses, three research questions, and five propositions is committed at preregistration, and confirmatory and exploratory analyses are labeled as such in advance. Attrition (withdrawn, merged, or materially revised papers) is tracked and reported rather than silently dropped, with a cohort flow diagram in every empirical report of the program (Section 8.8).

8.2 Hypotheses and propositions

The framework generates five hypotheses and three research questions concerning the allocation of early attention, and five propositions concerning its structure. Table 16 links the hypotheses, the research questions, and the propositions with dedicated designs (P2, P3b, P4) to an outcome, a contrast, and a preferred design; P1 and P3 are carried by Studies 1 and 2 in Table 17.

H1 (Cold-start disadvantage). Conditional on field, posting cohort, observable topic, document characteristics, and calendar conditions, papers authored entirely by platform newcomers receive fewer initial abstract-page views than papers with at least one previously visible SSRN author. The central outcome is views rather than downloads per view, because H1 concerns audience formation and not conversion.

H2 (Prestige and exposure). Institutional and author prestige positively predicts early abstract-page reach after controlling for prior platform activity and observable paper characteristics, extending the status-attention literature [104][105][111][140][60] to the first stage of preprint diffusion.

H3 (Prestige and conditional conversion). Institutional prestige positively predicts full-text downloading conditional on abstract-page exposure. Jo and Li provide recent field-specific evidence consistent with this pattern in Marketing, a prestige premium of approximately +6.6% in download rate for a three-author paper with one top-50 author [83], but a broader longitudinal test remains necessary. Jointly, H2 and H3 define the panel's headline estimand: a reach-versus-conversion decomposition of the prestige gradient, meaning the fraction of observed prestige-related inequality in downloads that arises *before* the abstract-page visit (exposure) versus *after* it (conditional conversion). A finding that most of the gradient forms before any reader has seen the abstract would locate the inequality in discovery mechanics rather than in content evaluation, in directly measurable form.

H4 (Communication as an equalizer). Greater abstract clarity and informational efficiency are associated with larger improvements in conditional conversion among low-prior-visibility authors than among already visible authors. H4 is the only hypothesis whose treatment lies fully within the author's control. It should be tested with prespecified textual measures and, separately, a randomized abstract experiment designed to online-field-experiment standards [30]. It should not rely solely on automated "AI-written" classifiers, which introduce their own measurement uncertainty.

H5 (Persistence of exposure shocks). Exogenous or quasi-exogenous early exposure shocks predict elevated subsequent attention beyond the immediate exposure window: the scholarly-platform analogue of success-breeds-success dynamics documented experimentally elsewhere [112][131][162] and of positional exposure effects on arXiv [72][73].

RQ1. Did the 15 July 2026 retirement of Rankings change the concentration, distribution, or prestige gradient of attention? Two opposing predictions are both plausible, and the question is left two-sided: removing rankings may dampen popularity-feedback loops [112][131], or it may instead increase

dependence on preexisting author networks by removing a discovery surface through which unknown papers could become visible once they began attracting attention.

RQ2. Does curated email-alert distribution produce larger proportional attention gains for newcomers or for established authors?

RQ3. Do normalized early metrics observed at days 7 and 30 predict later scholarly citation or publication outcomes better than raw cumulative downloads?

The propositions restate the framework's structural claims in falsifiable form:

P1 (Skew and the invisible majority). Over a random stratified sample of SSRN papers, the distribution of lifetime downloads is heavy-tailed and right-skewed such that the median is less than half the mean, and a substantial share of papers have ten or fewer lifetime downloads. *Falsifier*: median $\geq 0.5 \times$ mean, or a negligible ≤ 10 -download share. The published Marketing-network ratio of mean to median (approximately 2.54) [83] and the quasi-lognormal fit of the 2015 top-list data [85] motivate but do not establish P1 for the platform at large.

P2 (Classification governs surfacing). Within the repository layer, the number of eJournal classifications is positively associated with downloads and negatively associated with the probability of near-zero downloads, controlling for age and network, a replication and update of Whalen's 2018 cross-sectional pattern [168] under the platform's curated distribution mechanics [53][57]. *Falsifier*: no significant association after controls. Classification count is itself potentially endogenous: papers of broader relevance may both receive more classifications and attract more readers, so the cross-sectional association overstates any surfacing effect. The preferred identification is therefore within-paper classification changes where available, with the cross-sectional estimate reported explicitly as an upper bound.

P3 (Machine visibility is low and feature-dependent). For queries constructed so that a specific SSRN preprint is the best available source, consumer AI answer engines cite the preprint itself in a minority of cases, and machine visibility rises with features analogous to documented generative-engine-optimization factors [5][164]. *Falsifier*: preprints are cited a majority of the time regardless of features, or no feature predicts visibility.

P3b (A machine-layer Matthew effect). Conditional on expert-validated query relevance, papers with greater prior web and scholarly visibility are cited by AI answer systems more often than otherwise comparable cold-start papers. Under the equity alternative, generative engines retrieve on relevance without a premium on prior visibility and P3b is rejected, an outcome under which the machine layer would be *more* egalitarian than the human discovery layers of H1–H3, not less. *Falsifier*: no visibility gradient at matched query relevance. P3b extends the cumulative-advantage line of H1–H3 to the top of the discovery stack. Testing it requires the prior-visibility covariate and query stratification specified in Study 2 (Section 8.6).

P4 (Predators arrive before readers). For a newly posted preprint whose bibliographic metadata and contact route can be harvested from the posted document or associated public pages, the expected time to the first predatory solicitation attributable to the posting is shorter than the expected time to the first genuine scholarly contact. *Falsifier*: genuine contact reliably precedes predatory solicitation. P4 is stated conditionally on harvestability because the public visibility of author contact details on the current platform interface is not established. The harvesting mechanism itself is documented for preprint

repositories in general [107][159][151], and unsolicited email is the first contact with a predatory journal for a substantial share of solicited authors [32]. Attribution therefore requires the instrumented design of Study 3 rather than inference from inbox timing alone.

Table 16. Hypotheses, research questions, and identification strategies.

Question	Primary outcome	Main contrast	Preferred design
H1: newcomer disadvantage	Incremental views	Newcomer vs. established	Hierarchical longitudinal count model
H2: prestige and reach	Incremental views	Higher vs. lower prior prestige	Covariate-adjusted panel + matched sensitivity analysis
H3: prestige and conversion	Downloads conditional on views	Higher vs. lower prestige	Negative-binomial/Poisson exposure model
H4: communication equalizer	Conditional download probability/rate	Clarity $\times$ prior visibility	Observational text model + randomized replication
H5: cumulative attention	Post-shock views/downloads	Exposure shock vs. pre-shock trajectory	Within-paper event study
RQ1: rankings retirement	Daily attention; Gini; top shares; new-posting volume by field	Formerly ranking-exposed vs. comparison fields	Interrupted/event-study design
RQ2: email distribution	Daily views/downloads	Before vs. after distribution	Within-paper event study (Distribution $\times$ Newcomer)
RQ3: predictive validity	6/12/18-month citations	Normalized vs. raw early metrics	Out-of-sample predictive comparison
P2: classification governs surfacing	Downloads; probability of near-zero downloads	Classification count; within-paper classification changes where available	Panel regression + within-paper change design
P3b: machine-layer Matthew effect	Direct citation by answer system	High vs. low prior visibility at matched query relevance	Repeated multilevel retrieval audit (Study 2)
P4: predators before readers	Time-to-first and rate of coded inbound messages	Pre- vs. post-posting windows; predatory vs. genuine contact	Instrumented prospective cohort + event-history model (Study 3)

Note. All rows are prospective designs. None has been executed. Outcomes, contrasts, exclusion rules, and reference classes are to be preregistered before data collection.

8.3 The core panel

The reference design is a multi-discipline cohort of newly posted SSRN papers followed for at least twelve months from posting, observed on the schedule of Table 11. At the paper level the dataset preserves the earliest available version of title, abstract, author names and order, affiliations, posting date, classifications, keywords, page count, DOI where available, revision history, public abstract views, public downloads, and subsequent citation indicators, with author-identity matching against open bibliographic infrastructure [35][124] manually validated on a random subsample before inferred reputation is used at scale. Newcomer status is defined as no author on the focal paper having an SSRN paper dated before the focal posting, with sensitivity definitions based on prior paper counts and pre-posting citation stock.

Institutional prestige is not part of this definition: because prestige is itself an explanatory variable in H2 and H3, building it into newcomer status would mechanically confound platform novelty with prior author visibility and collapse the hypotheses the design is meant to separate. Prestige therefore enters the models only as a distinct covariate. Normalization by prior visibility must be handled with corresponding care on equity grounds: subtracting a prior-visibility baseline too aggressively would penalize the genuine early success of a junior author and would statistically erase, rather than expose, any systemic under-exposure of marginalized and Global-South authors. The framework's answer is to *measure* where inequality forms rather than to net it out: the reach-versus-conversion decomposition localizes whether a visibility gap arises at exposure or at conversion, and prestige enters as a covariate to be understood, not as noise to be silently differenced away. Prestige's distributional effect on attention is a primary outcome of the quasi-experiments (Sections 8.4–8.5). Calendar controls include day of week, week of year, academic season, and date fixed effects. The documented weekend and seasonal variation in article access (activity approximately 11% lower on weekends and 50% lower in summer in a recent seven-journal panel [178]) makes their omission difficult to justify in any high-frequency usage study. The full panel variable dictionary (on the order of two dozen fields, including interval increments rather than only cumulative counts, the distribution date where consented, citation stocks at 6, 12, and 18 months, and revision number) is specified in Appendix C (Table C1), and the same schema file is deposited with the preregistration materials (Section 8.8; Declarations).

Reach is modeled as a hierarchical negative-binomial process for incremental views, with field-by-age effects, calendar effects, and paper-level heterogeneity. Conditional conversion is modeled as incremental downloads with log incremental views as exposure, so that gaining a landing-page visit and converting that visit are statistically separated. In compact form, for paper i in interval t :

$$\Delta V_{it} \sim \text{NegBin}(\mu_{it}, \theta_V), \quad \log \mu_{it} = \lambda_{f(i) \times \text{age}(it)} + \delta_t + x_i' \beta + u_i$$

$$\Delta D_{it} \sim \text{NegBin}(\nu_{it}, \theta_D), \quad \log \nu_{it} = \log \Delta V_{it} + \gamma_0 + \gamma_1 \text{Newcomer}_i + \gamma_2 \text{Clarity}_i + \gamma_3 (\text{Newcomer}_i \times \text{Clarity}_i) + w_i$$

where $\lambda_{f \times \text{age}}$ are field-by-age effects, δ_t calendar-date effects, u_i and w_i paper-level heterogeneity, and the log-views offset in the second equation makes downloads conditional on landing-page exposure. The Newcomer $\times$ Clarity interaction γ_3 is the H4 test. These equations fix the prespecified estimands and the separation of reach from conversion, and the hierarchical structure performs exactly the empirical-Bayes partial pooling that Sections 5.2 and 8.8 require before any sparse cell's percentile is published. Topic controls in x_i are constructed from title–abstract embeddings frozen and clustered at preregistration, before any outcome is observed, so that topical composition cannot be tuned to results. Because papers occupy up to seven classifications simultaneously, classification effects are estimated with multiple-membership hierarchical structures rather than by forcing each paper into a single field. Because paper fixed effects absorb time-invariant status measures, baseline newcomer and prestige effects are estimated between papers, while paper fixed effects are reserved for within-paper treatments such as distribution events. Overdispersion and excess zeros are formally tested, hurdle and zero-inflated alternatives reported where justified, and the heavy tail handled by quantile and robust estimation rather than by deleting successful papers. Data-cleaning rules are themselves prespecified rather than discretionary, and the anomaly protocol is *flag, not delete*. The prespecified signal list covers: an interval extreme against the paper's own history, a view spike with no movement in downloads, a discontinuity in conversion, repeated identical timestamps, and a platform-side restatement. A flagged interval is never removed from the archive: it is retained with a reason code, and every headline result is

reported on a four-layer sensitivity ladder: (i) unmodified counters; (ii) excluding only platform-confirmed invalid activity; (iii) robust or winsorized estimation; (iv) excluding all flagged intervals, so that any conclusion surviving on only one layer is visible as fragile. Thresholds are calibrated against the observed distribution rather than set a priori. Where partner logs are available, referrer, user-agent, and geography fields distinguish newsletter, blog, prefetch, and crawler traffic. The documented blog-driven download spikes on SSRN [116] illustrate why an extreme value is often organic success rather than contamination. Anomaly detection is a measurement tool. It does not license discarding inconvenient success.

The same discipline applies to temporal structure. The panel distinguishes marginal from cumulative attention: cohort curves of interval downloads ΔD , and the share of a paper's 180-day downloads accumulated by each age, rather than reading only lifetime totals. It imposes no single exponential half-life on decay, since readership decays on several timescales at once [91]. Age curves are therefore fitted with splines or nonparametric estimators by default, with logarithmic transforms used to study structure and raw counts retained for benchmark reporting. Robustness is handled the same way: prestige is operationalized through several alternative measures rather than a single ranking threshold, newcomer status is re-tested under progressively stricter definitions, and results are reported with winsorized-sensitivity checks and full-distribution plots alongside the quantile estimates.

Because there is no public SSRN API [134][51], data collection proceeds by three routes, in order of preference: a formal data request to or partnership with SSRN/Elsevier; author-consented dashboard exports, in which participating authors contribute their own statistics-panel data under informed consent, the intermediate tier on which the registry of Section 8.8 already relies; and rate-limited, robots.txt-respecting per-paper sampling of public pages, whose technical feasibility is established by existing open-source collection tools [134]. The design excludes bulk scraping that would violate the terms of service, and it represents no undocumented endpoint as a stable API. Access constraints are discussed as a limitation in Section 9.

8.4 Quasi-experiment I: email-alert distribution as an attention shock

Sections 8.4 and 8.5 form one analytical module, studying two platform-side interventions as complementary treatments on the same panel cohort. Alert distribution *adds* an exposure surface, the Rankings retirement *removed* one, and both are read against a common set of distributional outcomes (newcomer share, prestige gradient, attention Gini, top shares). SSRN's alert mechanism is analytically valuable because posting and distribution need not coincide. Papers become publicly available first, and eligible work is later distributed through curated eJournal alerts whose timing varies and is not guaranteed [53][57]. For authors consenting to provide their distribution-status dates, a within-paper event study can estimate the attention path around the distribution date, each paper serving as its own control so that stable differences in topic appeal, author identity, and institution are absorbed. The specification is the standard event-study form: for outcome Y_{it} (interval views or downloads),

$$Y_{it} = \alpha_i + \sum_{k \neq -1} \beta_k \cdot \mathbf{1}[t - E_i = k] + \delta_t + \varepsilon_{it}$$

where E_i is paper i 's distribution date, α_i are paper fixed effects, δ_t calendar-date effects, and the β_k , normalized to the $k = -1$ pre-distribution reference period, trace the attention path around the event. The equation specifies the estimand for preregistration. The key heterogeneous effect is the interaction $\text{Distribution} \times \text{Newcomer}$, estimated by interacting the β_k with newcomer status. If alerts

disproportionately improve newcomers' reach, curated distribution functions as a platform-level equalizer; if gains accrue disproportionately to prestigious authors, distribution amplifies existing reputation advantages (RQ2). Prior evidence that discrete exposure events (front-position listings on arXiv [72][73], blog coverage of SSRN papers [116]) produce sharp early attention spikes indicates that the design has a detectable first stage. The prespecified specification also includes alternative event windows around the distribution date and separate reporting for papers distributed quickly versus after a long classification queue, since the up-to-45-day delay [53] otherwise mixes heterogeneous exposure histories into a single event time. Paper fixed effects absorb only *time-invariant* appeal, so three design conditions are load-bearing. Distribution timing is curated, not random: the platform selects which eligible papers it distributes and when [53][57]. The event date therefore carries a time-varying selection on topicality (a paper distributed the week a topic becomes salient differs from the same paper distributed in a quiet week), and this selection survives paper fixed effects, which is the reason the specification conditions on distribution latency and reports distribution-timing strata separately. Distribution-status dates are also observed only for consenting authors, so the estimation sample is a self-selected subset of the panel. Consent status is therefore recorded, and the event-study sample is compared to the full cohort on observable covariates; the selection limits external validity. And authors can time their own publicity to the alert: a coordinated self-promotion spike would masquerade as a distribution effect. A dated author-promotion log (the versioned evidence packet of Section 6.1) is therefore a *requirement* of this design: without it an organic distribution spike cannot be separated from a self-generated one, and the anomaly ladder of Section 8.3 flags rather than deletes any spike that coincides with logged self-promotion.

8.5 Quasi-experiment II: the Rankings retirement as a natural experiment

The 15 July 2026 sunset [149] permits a second analysis if historical ranking exposure can be reconstructed from archived ranking pages. Two features of the design weaken a clean single-date event study. The treatment did not arrive on one date: the per-eJournal ranking pages stopped updating from 11 December 2025 (Section 2.2 [B-2]), roughly seven months before the 15 July 2026 removal, so the withdrawal of the ranking discovery surface is better modeled as a *staggered* process, a first-stage freeze of the live pages followed by their final removal, than as a single common event date. The specification therefore carries at least two event times (the December freeze and the July removal), with the staggered form as the primary model and the single-date fit tested against it. Nor can "formerly ranking-exposed" disciplines be assumed from the platform's structure: the rankings were platform-wide, so differential exposure is endogenous and is not operationalized by naming fields that "felt" ranked. Exposure must be *measured*, not assigned: for example, from archived ranking pages recovered through the Internet Archive, as the share of a field's papers that ever entered the public top lists or the density of a field's presence on the per-eJournal pages. Whether the archive covers each field's ranking pages densely enough to yield a usable exposure measure is not yet established; the design reports that density alongside the exposure measure. With exposure so measured, a paper-level event-study specification compares attention trajectories in high- and low-exposure fields around the freeze and removal dates: the equation of Section 8.4 with the coefficients of interest on $\text{Exposure}_f \times \text{EventTime}_k$ interactions rather than on the event-time terms alone. The most informative outcomes are distributional (Section 8.1). The ranking mechanism could change the concentration of attention while leaving means nearly intact, and because any effect is expected to concentrate in the already-visible tail papers that populated the lists, field-*mean* outcomes are underpowered by construction. Gini coefficients, top-share measures, newcomer/established ratios, and prestige gradients are therefore the primary outcomes, with the two-sided prediction of RQ1,

and the power simulation of Section 8.1 must be targeted on these distributional statistics, not on mean attention, or the study risks being declared null for lack of power on the wrong estimand.

Two named, directional predictions sharpen the two-sided RQ1 without replacing it, both of them prespecified heterogeneity tests. The **visibility-cliff hypothesis** holds that post-sunset posting cohorts in formerly ranking-exposed fields experience a discrete drop in early reach relative to matched pre-sunset cohorts, because a discovery surface was removed on a known date. Under a **compensatory hypothesis**, if the Rankings functioned as a compensatory visibility mechanism for fields with weak discovery ecosystems, the effect of their removal should be largest in fields with small eJournal subscription bases, the lower end of the Whalen gradient [168].

A cheap secondary outcome extends the design from the allocation of attention to the health of the platform itself: the volume of new postings by field, in windows before and after 15 July 2026. If the Rankings supplied part of the incentive to post, their retirement could chill new submissions. That is a behavioral outcome distinct from the attention outcomes above, observable from public new-paper listings against the platform's own baseline of more than 250,000 new papers in 2025 [145], and a direct empirical answer to the contemporaneous "jumped the shark" polemic recorded in Section 2.2 [12][B-11]. The prediction is left two-sided, like RQ1, and the posting-volume outcome connects to item (3) of the Study 4 survey ("avoiding new uploads").

The analysis is quasi-experimental. Five threats are material. One is announcement and anticipation effects: the retirement was announced on 12 June 2026 with an initial target of 1 July and executed on 15 July, so a genuine anticipation window existed and shifted [148][149]. A *cluster* of concurrent 2026 platform changes shares the sunset's calendar and cannot be separated from it by any single event date: the strategic refocusing announced on 13 April 2026, the retirement of the commercial Research Paper Series running through the end of December 2026, and the author-license change of 20 July 2026, five days after the removal [146][150]. Because these interventions overlap the July window, a single common event date would attribute their combined effect to the sunset alone. Identification therefore leans on the staggered timing above (the December freeze precedes every one of them), on placebo dates, and on the distributional signature specific to a ranking withdrawal, and the design reports which co-timed changes remain unseparated. Then there is the summer attention trough: mid-July coincides with the documented seasonal decline in scholarly access, with accesses falling on the order of 50% in summer months [178], so a cross-field contrast run across the removal date confounds the sunset with a field's own summer seasonality, which differs by discipline. Calendar-date fixed effects absorb only the *common* seasonal component. The design therefore requires field-specific seasonal baselines (a prior-year same-window comparison per field). Non-random differences between high- and low-exposure fields are a further threat, addressed by the measured-exposure operationalization above and by matched comparison groups. The remaining threat is prospective: SSRN has stated it may revisit rankings, so a partial reinstatement during the observation window would introduce a second event. The design pre-registers treating any such reinstatement as a censoring point for the post-period and modeling it as a distinct treatment rather than folding it into the sunset contrast. In all reporting, event time is partitioned into three named periods, **announcement** (from 12 June 2026), **transition**, and **retirement** (from 15 July 2026), so that anticipation effects are estimated. The prespecified design therefore includes event-time pre-trend tests, alternative bandwidths, placebo dates, field-specific time trends and field-specific seasonal baselines, and sensitivity analyses excluding the announcement-to-retirement transition period. Figure 11 presents the intended outputs of both event studies as design templates.

Event-Study Templates: Exposure Shocks Are Estimated, Not Assumed

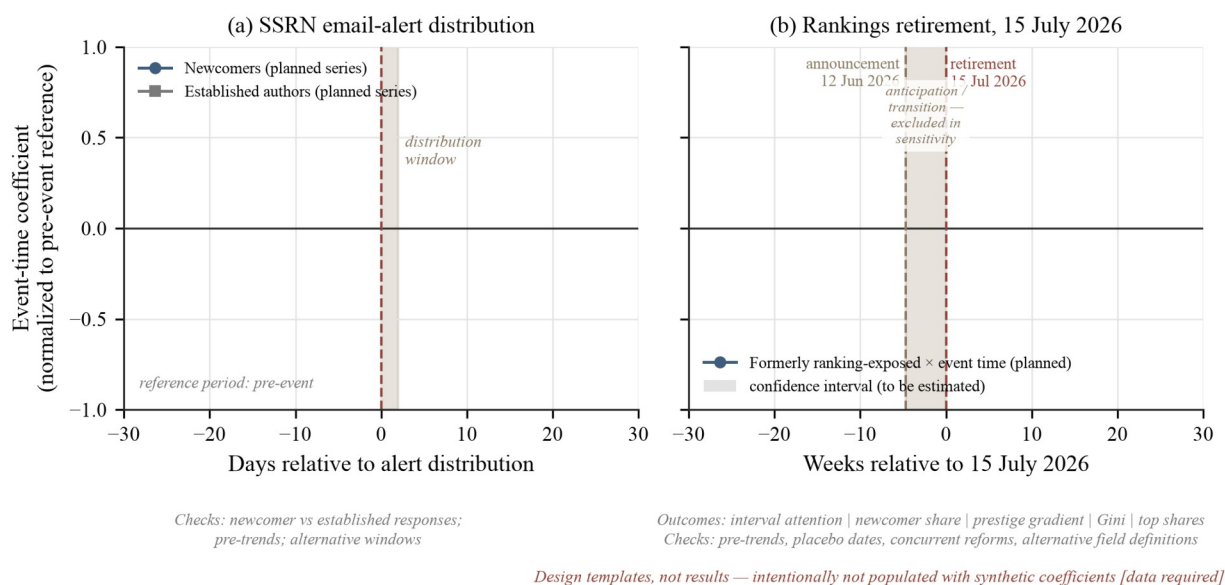


Figure 11. Event-study design templates for the two platform quasi-experiments. Prespecified event-study outputs for the designs of Sections 8.4–8.5: (a) attention around email-alert distribution, newcomers versus established authors, with the distribution window shaded because alerts arrive as a window, not an instant; (b) attention and inequality outcomes around the Rankings retirement, with both the announcement (12 June 2026) and the retirement (15 July 2026) dates marked and the anticipation/transition period shaded — the window excluded in the prespecified sensitivity analyses. The distributional outcomes (interval attention, newcomer share, prestige gradient, Gini, top shares) and required checks are stated on the figure itself. Coefficients are normalized to the pre-event reference period; shaded gray regions denote confidence intervals. These panels are design templates; no empirical values are shown, and generation code will be released with the preregistration materials.

8.6 Four prespecified studies

Alongside the core panel, four studies are specified and summarized in Table 17. The two public-data studies (1 and 2) can be attempted by any researcher without restricted-access permissions, within the access limits noted in Section 9 (rate-limited collection under the platform's terms of service). The two human-subjects studies (3 and 4) require participants and proceed under the ethics approvals stated in the Declarations. All four exist to begin filling the benchmark vacuum immediately rather than after a multi-year panel matures.

Table 17. The four prespecified studies.

	Study 1	Study 2	Study 3	Study 4
Research question	What is the median/percentile download distribution?	Do AI answer engines cite the preprint itself?	Reader or predator first?	How do authors read their own counters?
Data	3,000–5,000 stratified SSRN paper IDs	100–200 niche questions × ~5 engines	2 instrumented postings, 90-day inbound log	Preregistered author survey, four fixed-wording items

	Study 1	Study 2	Study 3	Study 4
Method	Median, Gini, decay curves, classification-count regression	Retrieval coding + logistic regression	Per-channel alias tracking + message coding	Closed-form questionnaire + descriptive shares
Primary outcome	Median; ≤ 10 -download share; attention Gini	Machine visibility V	t_reader vs. t_predator	Share never comparing counters to any benchmark; first-reading distribution for "7 downloads"
Falsifiable prediction	median $< 0.5 \times$ mean; Gini 0.75–0.90	Preprint cited in 10–35% of trials	t_predator $<$ t_reader	Majority report never comparing to any benchmark
Main deliverable	Public, machine-readable age-, field-, and cohort-stratified percentile benchmark tables	Reproducible retrieval-and-attribution audit protocol with a timestamped snapshot	Open instrumented-posting protocol with pre-posting baseline and coding codebook	Documented distribution of author interpretations and reported behavior change
The design fails if	Normalization does not improve stability or out-of-sample prediction over raw counts	Features predict nothing beyond query relevance, engine, and date	No post-posting discontinuity after baseline and exposure controls, or unreliable coding	Recruitment or self-selection leaves even the descriptive shares uninterpretable
Resource profile	Public data; standard scraping and statistics	Public AI-engine queries; manual coding	Two instrumented postings; inbox logging (IRB)	Anonymous author survey (IRB)

Note. Each study is to be registered on OSF before data collection, with code released.

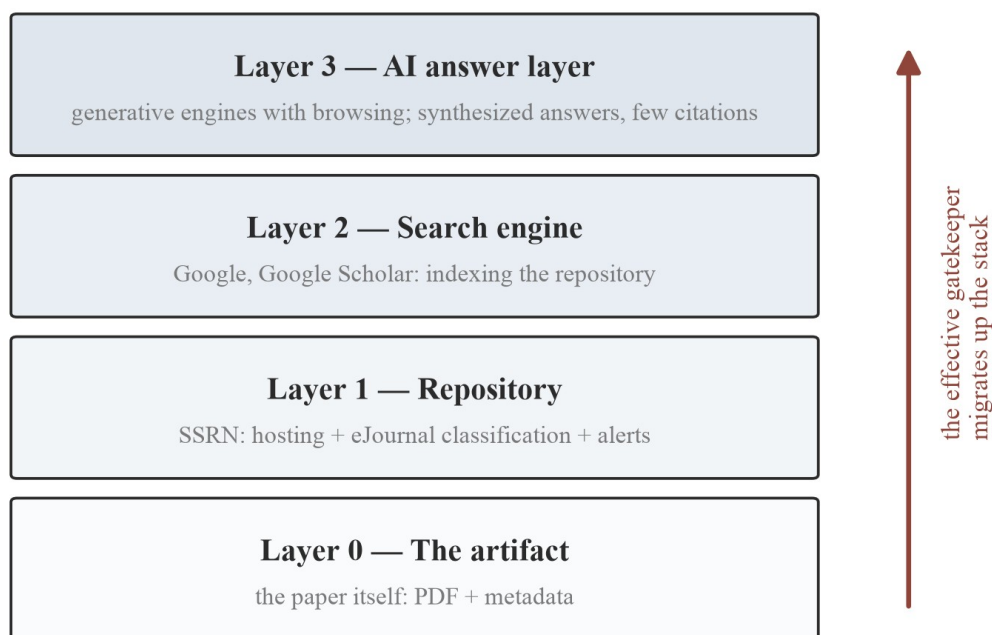
Study 1: "Anatomy of the Median." The study draws a random stratified sample of 3,000–5,000 SSRN abstract IDs, a planning range, with the final N set by power simulation against the prespecified prediction bands and the smallest effects of interest (Section 8.1). Stratification is by posting year and network, and per-paper lifetime downloads, abstract views, posting date, and classification count are collected through the access routes of Section 8.3. Estimands: median lifetime downloads; the shares of papers with zero and with ten or fewer downloads; the attention Gini with bootstrap confidence intervals; the views-to-downloads conversion rate, estimated overall and stratified by network (a test of the audience-reading-style hypothesis of Section 5.2); download decay curves; and a regression of log downloads on classification count with year and network fixed effects (a direct test of P2 [168]). The prespecified prediction bands (median lifetime downloads roughly 30–80, a ≤ 10 -download share of 15–30%, Gini 0.75–0.90) are extrapolations from the dated anchors of Section 5 [83][85][139] and from attention-Gini magnitudes in adjacent media [180], stated in advance so the study can falsify them. Study 1 is the empirical seed of the percentile benchmark discussed in Section 8.8.

Study 2: "Do the Machines Read You?" An increasing share of information seeking now begins at generative interfaces: roughly half of U.S. adults report using AI chatbots, and about a quarter use them daily [121], with a substantial minority already treating a chatbot as their first search engine in convenience samples [4]. The introduction (Section 1.6) conceptualized this as an upward migration of the discovery gatekeeper (journal brand $\rightarrow$ repository $\rightarrow$ search engine $\rightarrow$ AI answer layer) with each step relocating control further from the author and further from any public benchmark (Figure 12). Whether the top layer retrieves preprints at all is unknown: the generative-engine-optimization literature

has so far concentrated on brands and commerce [5][100][164], and audits of AI search have concerned news attribution rather than scholarly sources [82]. Those audits counsel skepticism about the new gatekeeper's reliability: in the Tow Center audit (1,600 queries across eight engines), the engines collectively failed to retrieve correct information in more than 60% of trials, with the best performer (Perplexity) answering incorrectly 37% of the time and the worst (Grok-3) 94%, and they typically erred with unwarranted confidence rather than declining to answer [82], figures from a news-attribution test, not from scholarly search. No controlled study of scholarly preprint visibility in AI answer engines existed as of August 2026 [100]. The contribution sits against two adjacent lines of work. The generative-engine-optimization studies this article builds on measure and manipulate the visibility of brands and commercial pages [5][164], not the retrieval of scholarly artifacts. A separate strand shows that when language models generate references at all, they reproduce and amplify the human tendency to cite already-highly-cited work, extending a Matthew effect into the machine layer [7], which is the mechanism P3b (Section 8.2) tests. Neither strand has examined the discoverability of scientific preprints under controlled query relevance, and specifically their retrievability after a platform ranking surface is removed, which is the object Study 2 isolates.

The Attention-Migration Stack: The Gatekeeper Moves Up

(historically: the journal brand sat above all layers)



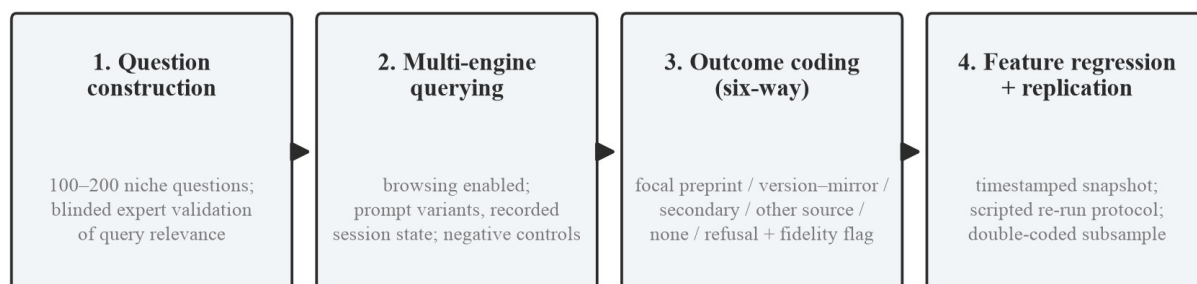
*Conceptual diagram — no data displayed.
Layers overlap; lower layers do not disappear — each new layer adds a gate above the last.*

Figure 12. The attention-migration stack. The layered path from a finished manuscript to a reader's attention. The migration thesis holds that the effective gatekeeper — the layer that decides which artifacts receive attention — has moved upward over time, from journal brand through repository and search engine to the AI answer layer; each upward move relocates control further from the author and further from any public benchmark. The stack is a conceptual model.

Study 2 constructs 100–200 niche questions, each engineered so that one specific SSRN preprint is the best available source, a relevance claim validated by blinded expert raters rather than self-certified by the query authors. It queries major consumer AI systems with browsing enabled and codes the outcome into six mutually exclusive categories: (1) the focal preprint itself cited; (2) the published version or a mirror of the same paper cited, a case that a coarser scheme would misclassify as absence; (3) secondary coverage cited without the paper; (4) a different but defensible source cited; (5) no traceable source; and (6) engine refusal (Figure 13). For every answer that cites the focal paper or a version of it, a binary **attribution-fidelity** flag additionally records whether the claim in the answer is actually supported by the cited source. Visibility and fidelity are reported separately and never merged into a single score, because an engine can cite a paper prominently while misstating it (citation laundering) or reproduce its content with no citation at all (influence without attribution). The confirmatory regression remains logistic, on the contrast "focal paper or its version cited" versus all other outcomes, with standard errors clustered, or preferably a crossed random-effects structure fitted, to respect the non-independence of observations that share an engine, a query, or a focal paper, since the same paper queried across engines and the same engine queried across papers each induce correlated outcomes that naive independent-trial errors would understate. Formally, the audit's estimand is M_{pqe} : the probability that paper p is cited by engine e in

response to query q . M_{pqe} is not a property of the paper alone, since it varies across queries, engines, dates, and model versions, so the notation defines the estimand and each timestamped run is read as an estimate of M_{pqe} at its date. The engine sample should extend beyond consumer chatbots to the emerging scholarly retrieval layer (tools such as Consensus, Elicit, Scite, and Semantic Scholar, which are grounded in academic corpora and cite real papers), since a preprint invisible to consumer engines may still be surfaced by discipline-specific ones, and the divergence between the two layers is itself an outcome of interest. Machine visibility is then regressed on paper features operationalized in Table 18, including the prior-visibility covariate that P3b requires: each paper's pre-audit web and scholarly footprint, frozen at query construction. The query sample is stratified by prior visibility so that cold-start papers are represented by design. Because generative engines are non-stationary, the design pairs a timestamped primary snapshot with a scripted re-run protocol that others can repeat. Three further audit-hygiene requirements are prespecified: negative controls (prestigious but irrelevant papers, title variants, and prompts that should *not* retrieve the focal work) make indiscriminate citation detectable; each query runs in equivalent prompt variants under recorded session state (clean versus logged-in sessions; engine, model version, date, and geography), so that prompt- and session-specific artifacts are separable from stable retrieval behavior; and a random subsample of outcomes is double-coded, with inter-rater agreement reported. The full step-by-step audit checklist is deposited with the preregistration materials. The study's principal input is researcher time (the blinded relevance validation of 100–200 questions and the double-coding) rather than restricted-access data or specialized infrastructure. A study that must pay for expert-rater time at market rates should budget for that labor explicitly. The prespecified prediction is that the preprint itself is cited in a minority of trials (band 10–35%), with structured-abstract and published-version features raising visibility (P3). Study 2 is one protocol among five agenda items. A full audit of scholarly machine visibility, including cross-engine and cross-repository comparisons, is the subject of a planned companion paper.

Machine-Visibility Audit Pipeline (Study 2)



Conceptual design — protocol schematic, not results

Figure 13. The machine-visibility audit pipeline (Study 2). The four-step Study 2 pipeline. The six-way coding separates citation of the focal preprint from citation of its published version or mirror — preventing false classification of absence — and the attribution-fidelity flag keeps "was it cited" distinct from "was it cited truthfully." The snapshot-plus-replication structure addresses model non-stationarity: the primary snapshot is dated, and the re-run protocol allows any researcher to reproduce the audit at a later date.

Table 18. Operationalization of machine-visibility features (Study 2).

Variable	Definition	Coding
DOI present	Paper page exposes a registered DOI	1/0
Title type	Descriptive vs. allusive title	descriptive = 1
Abstract structure	Structured/statistic-rich vs. narrative	structured = 1
Age	Months since first posting	integer
Classification count	Number of eJournals assigned	integer
Secondary coverage	≥1 indexed blog/news mention exists	1/0
Published version	A journal version of record exists	1/0
Prior visibility	Pre-audit web and scholarly footprint of the paper and its authors, frozen at query construction (P3b)	stratified ordinal
Outcome V	Retrieval outcome for the focal paper	six categories: focal preprint / published version or mirror / secondary only / other defensible source / none traceable / refusal; plus binary attribution-fidelity flag for cited cases

Note. Feature choices adapt documented generative-engine-optimization factors [5][164] to scholarly artifacts. The largest such audit, on the order of 252,000 trials across six LLM-based engines, identified

topical relevance and *list position* within the retrieved context as the leading drivers of being cited first [164]. List position is therefore a separate, transferable hypothesis for Study 2, and result-list position should be recorded wherever it is observable.

Table 19 complements Table 18 at a different level of description. Table 18 operationalizes the paper-level features whose effects Study 2 estimates, while Table 19 specifies the audit's measurement constructs: what is measured, under which controls, and against which primary threat to validity. Every construct is a distinct claim about the machine layer, and each threat names the specific way a naive audit would go wrong.

Table 19. Measurement constructs of the machine-visibility audit (Study 2).

Construct	Operational measure	Required controls	Primary threat
Direct machine visibility	Focal preprint cited or linked in the answer	Query relevance, engine, date, session state	Model and index drift
Source substitution	Published version, mirror, or secondary coverage cited instead of the preprint	Version matching, DOI, title variants	False classification as absence
Citation position	Rank and prominence of the focal citation within the answer	Answer length and citation count	Interface-specific presentation
Attribution fidelity	Claim in the answer accurately supported by the cited paper	Blinded claim–source verification	Citation laundering or hallucination
Query sensitivity	Variation of outcomes across equivalent query phrasings	Randomized prompt variants	Prompt-specific artifacts
Replicability	Same outcome on a later date or a clean re-run	Archived protocol and system metadata	Engine non-stationarity

The most defensible transfer from the commercial GEO literature is a feature-audit agenda, not a promotional recipe: Table 20 lists candidate features whose association with machine visibility Study 2 is designed to confirm or refute. Each row is a hypothesis rather than an established tactic, and a feature associated with retrieval may still be a proxy for paper maturity, topic popularity, or prior visibility instead of a lever.

Table 20. Feature-audit agenda for machine visibility (hypotheses for Study 2, not recommendations).

Candidate feature	Hypothesized mechanism	Supporting evidence
Register a DOI	Stable identifier aids retrieval and citation matching	[164]
Structured, statistic-rich abstract	Statistics addition raised generative visibility by up to ~40% in commercial benchmarks	[5]
Descriptive, relevance-dense title	Topical relevance is a top driver of citation in AI answers	[164]
Maximize accurate eJournal classifications	More classifications associated with more downloads, fewer zeros	[168]
Secure ≥1 item of secondary coverage	Engines often cite secondary coverage over primary sources	[82]

Study 3: "Who Arrives First?" An instrumented-posting protocol tests P4 directly. Two real working papers are posted using unique plus-addressed email aliases per exposure channel, with a 90-day inbound log preceded by a pre-posting baseline window of two to four weeks on the same aliases. Without a baseline, an increase in solicitation volume cannot be attributed to the posting at all, since academic spam arrives continuously against a high background rate. Each message is coded as genuine scholarly contact or predatory solicitation using published solicitation signatures [96] and curated blocklist-type criteria, against the high background rates documented in the specialty audits reviewed in Section 7.2 [9][84][63][61]. The outcome is the comparison of time-to-first-genuine-contact and time-to-first-predatory-solicitation attributable to the posting. The prespecified prediction, per P4, is that solicitation precedes genuine contact. Two measurement rules are fixed before launch. The baseline and post-posting windows are unequal in length (roughly two to four weeks against ninety days), so the comparison is of arrival *rates* (events per alias-day) and of survival functions, never of raw event counts, which the window asymmetry would otherwise distort. And time-to-first-contact is right-censored by construction: with two papers, an alias that receives no genuine contact within the 90-day window yields a censored, not a zero, observation, so the censoring rule is prespecified. Observations are right-censored at 90 days, aliases with no event contribute censored time to the survival estimate rather than being dropped or coded as immediate, and with $N = 2$ the resulting curves are reported with their exact small-sample uncertainty rather than as point estimates. The causal estimand of the distributed-replication design is the **change in the hazard of solicitation after posting** (not the identity of any downloader), estimated with survival/event-history methods on matched pre- and post-posting windows. Because an author's visible productivity plausibly predicts solicitation volume independently of any single posting (solicitation is harvested from an author's accumulated public output, so the high background rates documented in the specialty audits of Section 7.2 scale with that output), distributed replications must record author-level exposure controls (prior postings, publication volume, and public contact surface) alongside the focal event. This is an $N = 2$ instrumented case study, a proof of mechanism and an open protocol inviting distributed replication, which would collectively yield a powered estimate. $N = 2$ demonstrates the feasibility of the instrumentation, not prevalence and not causal magnitude. Unlike Studies 1, 2, and 4, which are powered against prespecified effect sizes, Study 3 is a feasibility pilot and open replication protocol.

The coding frame is fixed before launch. Inbound messages are coded into seven types: journal solicitation, conference invitation, peer-review request, editorial-board invitation, book-chapter invitation, commercial service offer, and genuine scholarly correspondence, a wider frame than a binary predatory/genuine split. Each record carries a minimum field set: timestamp, receiving alias (channel), venue name, fee disclosure, specificity to the focal paper, urgency language, verification outcome against maintained screening services, and coder confidence. Borderline cases and legitimate low-prestige venues are governed by prespecified rules rather than ad hoc judgment. A new or obscure venue is not automatically predatory. A random subsample is double-coded with inter-rater agreement reported. Only aggregate counts are released (Section 9; Declarations), and the full codebook is deposited with the preregistration. Solicitation volume is not evidence of readership, and solicitation does not move the public counters: SSRN's integrity system excludes anomalous automated activity from official counts [144], and no inference from inbox to dashboard is licensed in either direction.

Study 4: "How Authors Read Their Counters." A prespecified author survey converts the forum evidence of demand recorded in Table 1 into direct measurement. The instrument comprises four fixed-wording items: (1) "Have you ever compared your SSRN downloads to any benchmark or norm?", with response options distinguishing a specific published number, a discipline-specific table, informal peer comparison, and never; (2) "When you see '7 downloads' on your dashboard, what is your first thought?", with options spanning inferred failure, too early to tell, no idea what is normal, and expected or fine; (3) "Did the July 2026 Rankings sunset change your preprint behavior?", with options for stopped checking statistics, avoiding new uploads, uploading but checking less often, and no effect; and (4) "Do you think hiring and tenure committees understand SSRN metrics?". The prespecified predictions are that a majority of respondents report never having compared their counters to any benchmark, and that first readings of a low count skew toward inferred failure rather than toward the too-early reading that Section 5.6 shows is usually the correct default. Sample-size target, recruitment channels, and analysis code are fixed at registration. Self-selection toward engaged, and plausibly more anxious, authors limits the sample, as it does the Section 8.8 registry. The fixed-choice format also admits a qualitative companion strand: prespecified thematic coding of public author questions (with explicit inclusion criteria and inter-coder reliability, under the quotation ethics stated in Section 9) and semi-structured interviews on how authors read, and feel about, their counters, reaching the human consequences of metric presentation that closed-form items cannot. One design limit bears on the central construct. Four fixed items administered to SSRN authors alone measure a *level* of benchmark-free reading, not a *contrast*, and a level cannot by itself identify the article's claimed path from benchmark opacity to metric blindness. A high never-compared share on SSRN is equally consistent with a general indifference to counters that has nothing to do with platform opacity. Discriminant and causal traction therefore require a comparison condition: the same instrument administered to authors who post on platforms with different benchmark exposure, arXiv (endorsement- and listing-based, with its own visibility surfaces) and OSF Preprints, turns "how SSRN authors read their counters" into "whether authors on more- and less-opaque platforms read their counters differently," which is the contrast the opacity-to-blindness claim actually predicts. The cross-platform arm is prespecified as part of Study 4; without it a single-platform result would remain descriptive of SSRN rather than a test of the mechanism. Study 4 then tests the interpretive claims of Sections 1 and 5 directly, measuring what authors believe their counters mean rather than what the counters are.

8.7 The article as its own experiment

This article's own download trajectory is prespecified too. Its counters will be recorded at 30, 90, and 180 days after posting and reported against the Study 1 median once that median is known. The record is descriptive and carries no interpretation rules, because both directional readings would flatter the thesis, and a device on which every outcome confirms the argument tests nothing; the numbers will be published as they arrive. One reading would cut against the article's framing: a trajectory sitting at the field-and-age median, neither elevated nor suppressed, would indicate that self-referential design confers no measurable discoverability advantage in this single case. The reflexivity cannot be avoided, since this paper studies a metric it will itself accrue; the cataloged levers are evaluated by the direct tests in Section 8, not by this record.

8.8 The percentile benchmark as a public good

The single most useful artifact this agenda can produce is a percentile table for preprint downloads, the number every author searches for and cannot find. The platform has historically published means and top-N lists but not medians or percentiles, and since 15 July 2026 it displays only per-paper and per-author

counts with no comparative context at all [149]. For the h-index and similar indicators there is a mass-audience genre of graduated "is X good?" guides [118][128]; the Section 2.5 audit found no equivalent for SSRN downloads. The raw materials for such a table have existed in print for years [18][83][85][139][168], and what is missing is the translation of those distributions into an author-facing scale.

After the sunset, external collection of full-platform percentiles is no longer possible: the top lists that leaked distributional thresholds were the only public window onto the tail, and they are gone [149]. Two complementary routes remain. One is Study 1's stratified sample, which yields platform percentiles at the paper level. The other route is a crowdsourced registry of self-reported counters: a voluntary community dataset in which authors contribute dated dashboard snapshots (paper age, field, abstract views, downloads), assembled before institutional memory of the pre-sunset distributions decays further. The preservation step of the dossier protocol (Table 13), if adopted at scale, is the seed of exactly such a registry. There is also a working precedent to extend rather than invent. IRUS (Jisc) already aggregates standardized, COUNTER-conformant item-level usage statistics, `Total_Item_Requests` among them, across participating institutional repositories [185], which shows that multi-institution usage aggregation is operationally tractable. A preprint registry is that established model widened to carry the reference-class variables that matter here (field, posting cohort, and career stage). A North American counterpart, CC-PLUS, supplies the same capability as open-source infrastructure, with automated SUSHI harvesting of COUNTER 5 usage across consortia [188], so a registry can extend a working platform instead of being commissioned from scratch. A registry of this kind would constitute the first fresh, multi-discipline download distribution since the mid-2010s samples [85][139]. The known weaknesses are self-selection toward engaged and likely above-median authors, unverifiable self-reports unless dated screenshots are required, and coverage gaps in small fields. Percentiles derived from it must therefore be published with explicit selection caveats and, where possible, calibrated against the Study 1 probability sample, and sparse field-by-age cells must be reported under empirical-Bayes shrinkage toward broader reference classes rather than as unstable raw percentiles (Sections 5.2, 8.3). The registry protocol and verification rules will be published with the preregistration materials, and the contribution format is fixed here: a registry record is a defined subset of the panel schema of Appendix C (the identity, timing, usage, exposure, and audit domains), so that self-reported records remain mergeable with the Study 1 panel rather than accumulating as an incompatible second dataset.

The registry's output formats are prespecified along with its protocol, and two artifacts are committed to. One is a translation table that takes a single number (say 10, 40, or 100 downloads) and returns, for each covered field, its percentile position and the plain-language band of Table 8; with it, the answer to "is 7 downloads good?" can finally be looked up rather than guessed. The other is a percentile "fan" figure plotting the 25th, 50th, 75th, and 95th percentile download curves as functions of paper age, by field and newcomer stratum: the format in which a benchmark answers "how much is normal on day X" (Figure 14; [data required]: neither artifact may be populated with synthetic values). Every percentile release from either route (the probability sample or the registry) is published under the transparency specification of Appendix D, including its minimum-cell privacy rules, so that the benchmark this agenda proposes does not reproduce the opacity it was designed to cure. On top of the data belongs a thin self-service layer, a public online calculator into which an author enters field, paper age, and counters and receives a percentile band and a verbal verdict. It would turn the uncalibrated scorecard-style heuristics of Section 6.3 into a calibrated instrument once, and only once, the reference distributions exist. Section 8.10 proposes a multi-model architecture through which such a service could be operated and kept current.

That same self-service layer keeps the reporting protocol from imposing a new workload on hiring and tenure committees, since the per-number normalization the dossier format (Table 13) asks for is discharged by the calculator (this section) and, at scale, by the orchestration layer (Section 8.10). Its realistic deployment path runs through existing research-information plumbing, namely current research information systems (CRIS) and ORCID-linked profiles, so that a benchmarked reading can be produced from a paper's identifiers rather than assembled by hand. One natural cross-platform extension is a parallel-deposit protocol comparing the same paper's trajectory on SSRN, arXiv (where field-appropriate), and OSF Preprints, a direct test of whether the bands of Table 8 transfer between push- and pull-architectures (Table 6).

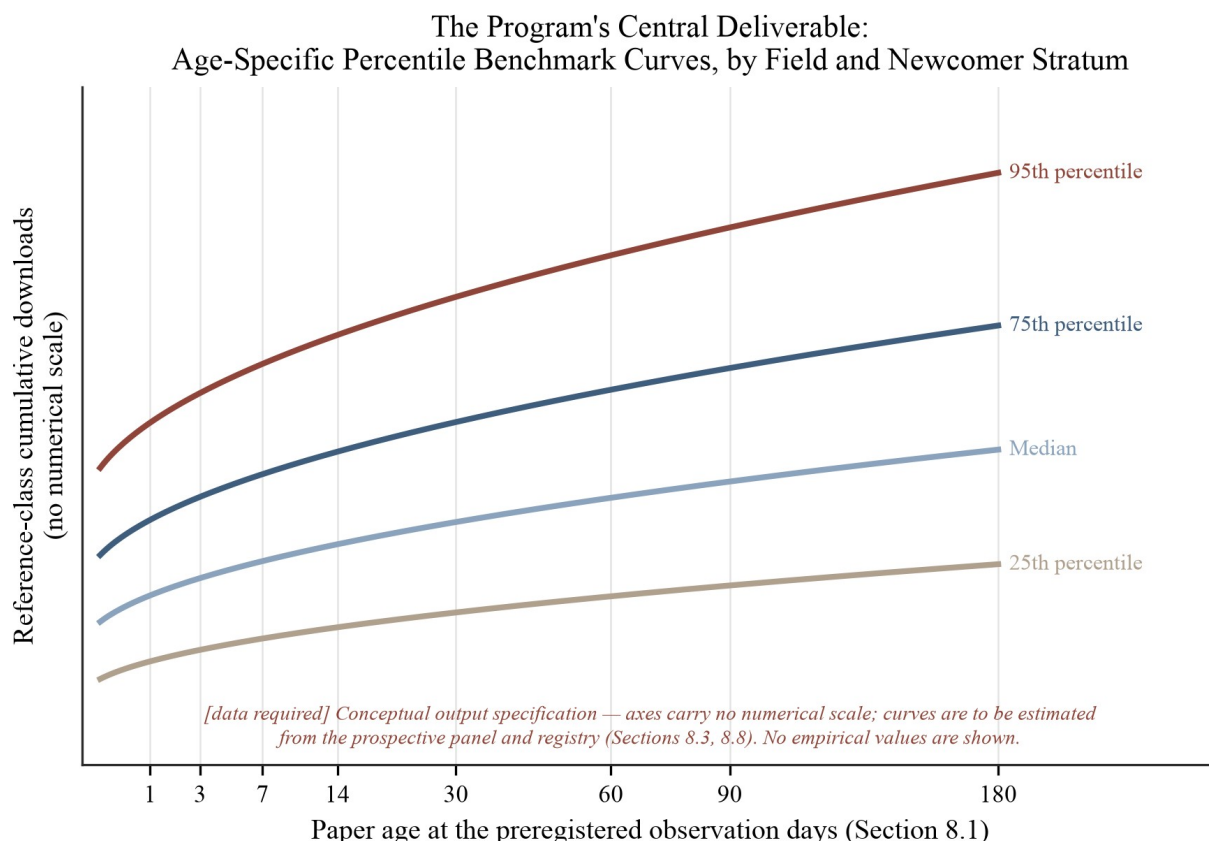


Figure 14. The age-specific benchmark fan: the program's central deliverable. Percentile download curves (25th, 50th, 75th, and 95th) as functions of paper age, evaluated at the prespecified observation days of Section 8.1 and to be estimated by field and newcomer stratum from the Study 1 sample and the community registry. The axes carry no numerical scale; the figure is a conceptual output specification, and it is the display format behind both the platform recommendation of Section 10.4 and the public calculator described above.

Every empirical article of the program commits to a **minimum reporting package**: a cohort flow diagram; raw and normalized distributions side by side; reach and conversion reported separately; exposure histories; anomaly-sensitivity results on the four-layer ladder of Section 8.3; null and contrary results reported with the same prominence as confirmations; and a machine-readable benchmark table, the condition under which Study 1's output can actually feed the public calculator described above. Abstracts of program articles distinguish verified estimates from prespecified expectations.

8.9 Predictive validation

The agenda's summative test (RQ3) asks whether normalization earns its complexity. Early usage and later citations are positively but imperfectly associated across settings: correlations on the order of 0.11–0.35 in journal-level Granger analyses [80], approximately 0.4 in early arXiv cohorts, where $r \approx 0.4$ explains only about 16% of variance [25], with comparable associations reported for chemistry preprints two decades ago [26], for downloads at journal and paper levels generally [71], in RePEc usage data [31], and in bioRxiv cohorts [62], and moderate rank correlations in mega-journal usage [102]. The highly downloaded and the highly cited sets overlap only partially [108][39]. Recent panel evidence links access to citations with a lag: in one seven-journal panel the association was approximately unit-elastic, with a 10% increase in accesses associated with roughly a 10% increase in citations about nine months later [178]. In the notation of Section 5.5, the summative comparison estimates

$$\text{Citations}_{i,12m} = g(B_i^V(7), B_i^D(7), \text{Conversion}_i(7), \text{Momentum}_i(30), X_i)$$

with g fitted flexibly and evaluated only out of sample. The equation fixes which normalized inputs compete against raw counts. Models predicting 6-, 12-, and 18-month citations from age-, field-, and cohort-normalized day-7 and day-30 metrics should be compared out of sample (by cross-validated likelihood, calibration, rank correlation, and predictive error) with models using raw cumulative downloads alone, a prediction policy problem in the sense of Kleinberg et al. [88]. A particularly informative outcome would be raw counts favoring established authors while normalized cold-start measures better predict later citation uptake. That pattern would imply that part of what raw early counters measure is initial opportunity rather than subsequent intellectual demand. If normalization fails to add predictive information, the framework's measurement claim is rejected or narrowed accordingly.

Figure 15 draws the agenda together. One prospective cohort feeds the four studies and the two platform quasi-experiments in parallel, and every module deposits into the same set of versioned public outputs. That shared structure makes the agenda a research program.

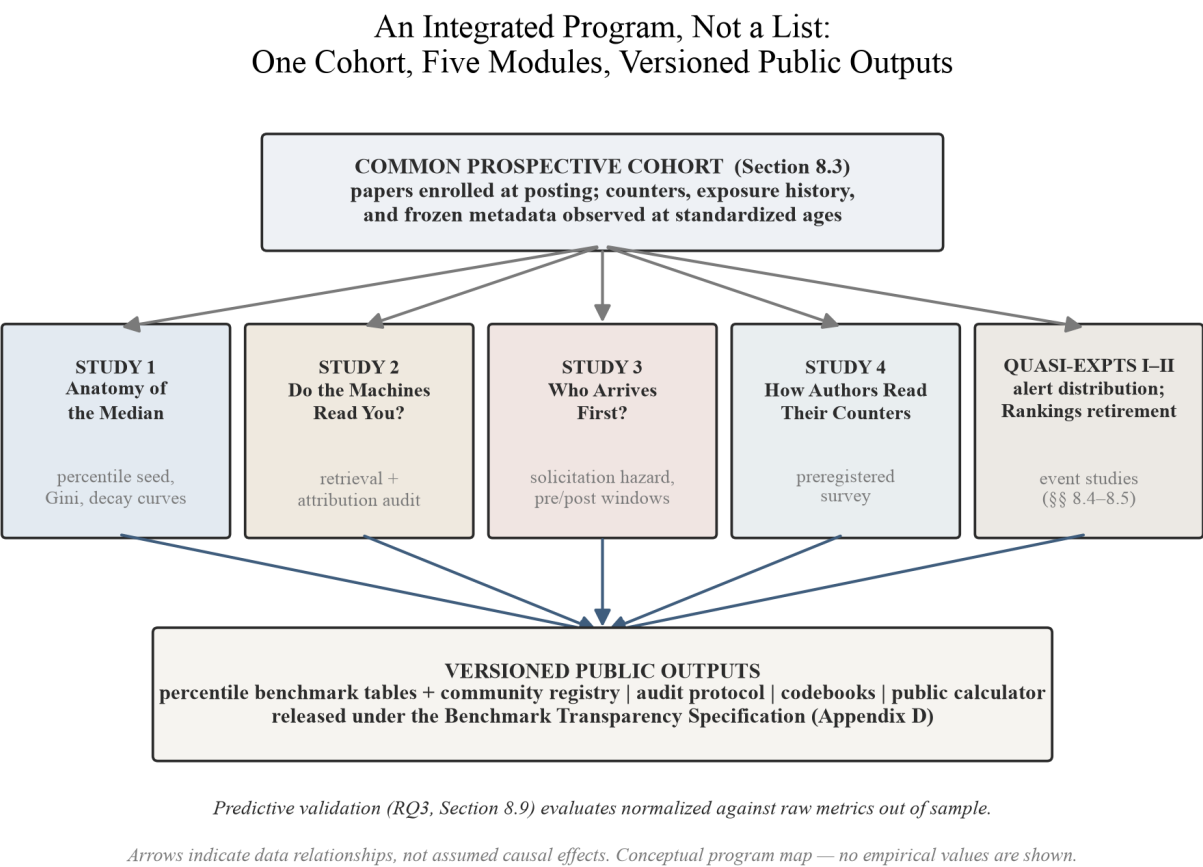


Figure 15. The integrated research program. One prospective cohort (Section 8.3) feeds the four studies (Section 8.6) and the two platform quasi-experiments (Sections 8.4–8.5); all modules deposit into versioned public outputs — the percentile benchmark tables and community registry (Section 8.8), the audit protocol, the coding codebooks, and the public calculator — released under the Benchmark Transparency Specification of Appendix D, while predictive validation (Section 8.9) evaluates the resulting benchmark out of sample. Arrows indicate data relationships, not assumed causal effects.

8.10 A design proposal: orchestrated frontier models as the benchmark's delivery layer

One-off publication cannot deliver what the agenda's outputs require. The percentile tables, the community registry, the CSAB reporting format (Section 5.5; Appendix B), and the public calculator (Section 8.8) must all be assembled from scattered and decaying sources, verified claim by claim, kept current as counting rules and platform features change, and delivered to individual authors in calibrated language. A candidate architecture for that delivery-and-maintenance layer is an orchestrated pipeline of specialized frontier AI models (a class of system the author develops; see Competing interests). Everything in this subsection is a design proposal. No implementation exists, no empirical performance is claimed, and no figure for accuracy, cost, or time saved is asserted anywhere below.

One alternative to an orchestrated system is a single generalist model. An author who asks a consumer chatbot "is 7 downloads good?" receives a fluent, unverifiable assertion, the failure mode the answer-engine audits document, in which systems err with unwarranted confidence instead of declining to answer [82]. Underneath that sits a structural problem. Metric interpretation decomposes into subtasks with different, partly adversarial objectives. Acquiring and dating the evidence, screening a counter for non-human or anomalous traffic, normalizing the screened count against a reference class of age, field, cohort, and prior visibility, translating the normalized position into banded language, and challenging that translation before release are not the same task, and a component optimized for one is a poor judge of another. The model that drafts a reassuring reading should not be the sole authority on whether the underlying number is contaminated, and the model that drafts dossier language should not certify that the language avoids overstatement. The organizing principle is therefore a division of epistemic labor: generation and verification are assigned to separate components with independent objectives, and disagreement is preserved and reported rather than averaged away. Convergence of independently prompted or architecturally distinct evaluators is not truth, though it is stronger ground for a reading than a single model's fluency, a design intuition consistent with experimental evidence that multiagent debate among language-model instances measurably improves factual validity and reduces hallucination relative to a single model answering alone [43]. Where evaluators diverge, the divergence is itself diagnostic information, just as conflicting sources were treated as signal rather than noise in this article's own synthesis rules (Scope and Evidence Discipline). Two design cautions temper that intuition. Agreement among components is not itself evidence: it can arise from shared training data, shared retrieval results, common prompt assumptions, or contamination through the orchestrator itself. Model families drawn from different vendors are not scientifically independent, and independence must be engineered, through separate instructions, restricted information flow between components, preserved intermediate outputs, and evaluation against evidence that no generating component is permitted to edit. Verification also does not become reliable merely because a second model performs it: claims that are disputed or consequential require inspection of the underlying source by the human user.

Figure 16 shows the proposed pipeline, and Table 21 decomposes it by function. Processing begins from a frozen case file, not a free-form prompt: the counters with their observation dates, paper age, declared field and candidate reference cohorts, exposure history, and the author's actual question. Missing inputs stay missing, and the system may not manufacture an absent percentile distribution, impute a denominator, or accept a model-generated number in place of an observed one. The information flow then reproduces the interpretive sequence the article defends throughout: **raw counters** → **integrity screening** → **reference-class normalization** → **interpretation under constraint**. A monitoring component retrieves counters, platform documents, and registry records with dated provenance, so that every downstream value is traceable to an archived source. Every factual output travels between components as an entry in a claim ledger. Each entry records the claim, its supporting source or data field, a stable locator, retrieval time, any transformation applied, its reference class, and its verification state, so that no unsupported output can silently become a premise of a later stage. Integrity screening applies conventional time-series anomaly detection (a task for statistical methods, not free-text generation) against the documented background of bot filtering and AI-crawler contamination [129][144]. It attributes trajectory breaks to dated exposure events where the exposure record permits (Section 5.6). Normalization is performed by deterministic statistical code invoked as a tool: percentile positions are computed against the versioned Study 1 and registry tables under the shrinkage rules of Section 8.3. Percentile arithmetic is never delegated to free-text generation. Only then does an interpreter model draft the author-facing reading, confined to the band language of Table 8, the reach-versus-conversion diagnostic of Table 9, and the CSAB template of Appendix B, with source, vintage, and confidence tier attached to every number. A separate verifier and critic check each claim against its recorded evidence before release, reporting each as verified, contested, unsupported, or not assessable, four distinct states that are never compressed into a single synthetic confidence score. Two constraints carry over from the framework unchanged. Any recommendation is conditional on the diagnosed bottleneck, since reach-limited and conversion-limited papers warrant different actions (Table 9), and a system that cannot tell them apart should say nothing instead of emitting generic advice. The other constraint is that a released recommendation is a hypothesis: it predicts which mechanism should move if the diagnosis is right, and subsequent observations test that prediction. The system's objective function is correspondingly narrow: remove identifiable interpretive friction between a paper and its appropriate reference class, and never the maximization of any counter, a boundary fixed by the platform's integrity and generative-AI policies [54] [144] and by the gaming history of Section 3 [44].

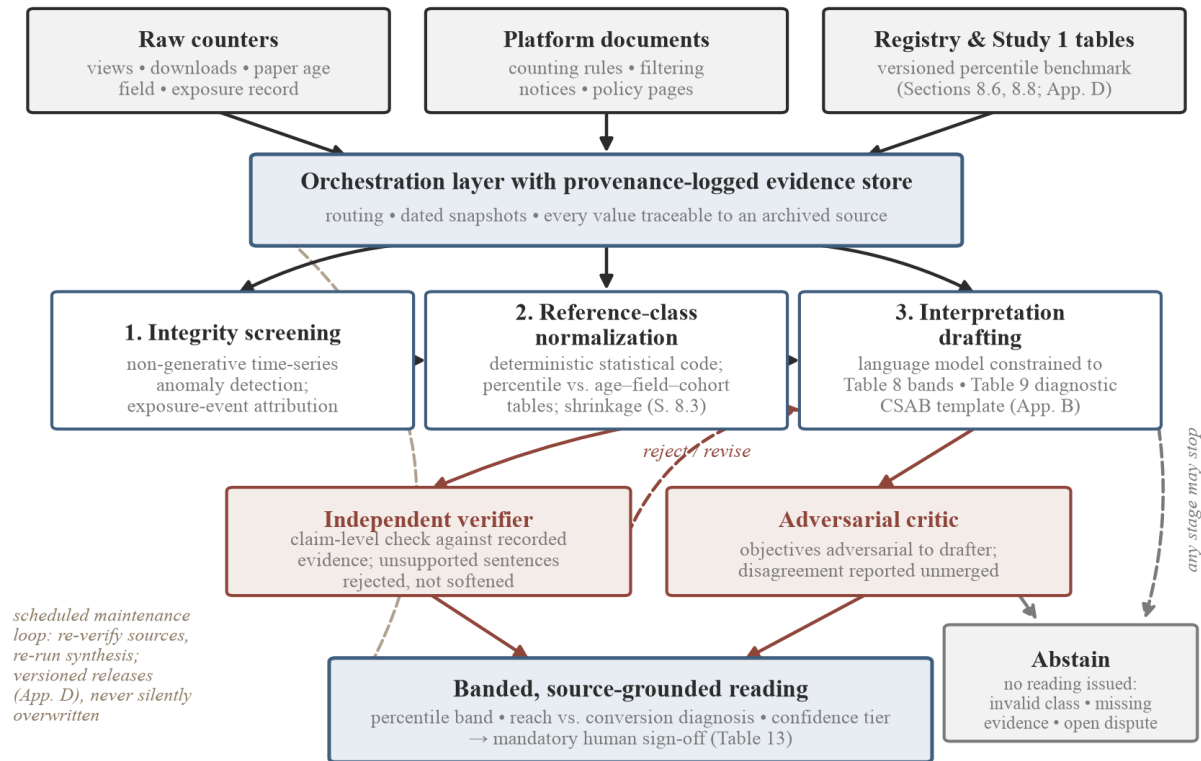
Table 21. Functional decomposition of the proposed orchestration layer: function, component type, principal failure mode, and the mechanism that independently verifies it. Every row is a hypothesis about a workable division of labor.

Function	Model or component type	Principal failure mode	Independent verification mechanism
Orchestration and provenance	Task router maintaining the frozen case file and the claim ledger	An early unsupported assumption propagates as a premise of every later stage	No unverified output may become a trusted premise; source, transformation, retrieval time, and uncertainty preserved at every handoff

Function	Model or component type	Principal failure mode	Independent verification mechanism
Evidence acquisition and monitoring	Tool-using retrieval model on a re-run schedule; bibliographic matching against open citation infrastructure [35][124]	Fabricated, mismatched, or selectively retrieved sources	Dated provenance log; every value traceable to an archived snapshot; random subsample independently re-fetched by a second component; failed searches and contrary results recorded, not discarded
Integrity screening	Non-generative time-series anomaly detection [129] [144]	A spike labeled "bot" — or "reader" — without evidence for either identity	Anomaly flags published alongside aggregates (Section 8.8); sensitivity reported on the four-layer robustness ladder of Section 8.3; anomaly treated as a flag, never as an identity claim
Reference-class normalization	Deterministic statistical code invoked as a tool	Convenient peers mistaken for a valid comparison population	Recomputation against the versioned public benchmark tables (Appendix D); field, age, cohort, prior visibility, and source vintage stated or the module abstains; no free-text arithmetic accepted into the record
Interpretation drafting	Generalist language model constrained to Table 8 band language and the CSAB template (Appendix B)	Fluent overstatement — certainty, scope, or causal status beyond what the evidence licenses	Claim-level check by a separate verifier against the recorded evidence; each claim reported as verified, contested, unsupported, or not assessable; unsupported sentences rejected, not softened
Adversarial review	Critic model with objectives adversarial to the drafter	Ritual critique, or a synthetic consensus that hides material dissent	Structured disagreement report — evidence status, evaluator pattern, decision sensitivity (does the dispute change the diagnosis or only its wording), residual uncertainty — delivered unmerged, never collapsed into one confidence number
Solicitation triage	Specialized classifier over the coding frame of Tables 14–15	A categorical "predatory" verdict that outruns the evidence and risks reputational harm	Flag-and-review only, never auto-delete and never a published accusation; error rates reported on labeled sets; documented limits and biases of AI text classification acknowledged [157][97]
Benchmark maintenance	Scheduled orchestration of all components	Silent drift: stale sources, or releases overwritten without a trace	Versioned, dated releases under Appendix D; revisions logged, never silently overwritten — a living benchmark map rather than a static snapshot
Final interpretation and any career use	Human author	Automation bias and diffusion of responsibility	Mandatory human sign-off before any external use; dossier language follows the Table 13 protocol, with source, time window, and limitations attached

Two boundaries complete the specification. One is a permission boundary: read-only access is the default, and the layer drafts but never executes. It may not autonomously edit a manuscript, alter platform metadata, contact scholars or venues, post publicly, or delete a solicitation message, and any implemented workflow must preserve the human responsibility and disclosure obligations of the applicable platform and submission policies [54][56][144]. The other is an abstention boundary: the null output is a legitimate (often the correct) result, and every stage may stop rather than pass a defective input forward. Mandatory abstention triggers include a missing or indefensible reference class; an unverifiable source central to the reading; unresolved disagreement about whether a proposed formulation preserves the paper's claims; any request for a causal conclusion from a purely descriptive pattern or for a prediction of citations or career outcomes without a validated prediction dataset (Section 8.9); and any request to manipulate, purchase, simulate, or otherwise inflate attention. That last trigger keeps the layer from becoming the metric-gaming instrument whose history Section 3 documents [44][69][144]. A last separation governs the agenda's two evidence streams. The versioned percentile tables of Sections 8.6 and 8.8 are the layer's only admissible normalization source, while the machine-visibility audit of Study 2 feeds a distinct retrieval-and-attribution check. Machine retrieval is not a percentile of human attention, and the layer may not blend the two.

Proposed orchestration architecture for the benchmark's delivery layer



Design proposal — no empirical performance is claimed.

Figure 16. Proposed orchestration architecture for the benchmark's delivery layer (Section 8.10; Table 21). Specialized components separate evidence acquisition, integrity screening, normalization, and interpretation, with independent verification interposed between generation and release, an abstention path available at every stage, and a scheduled loop maintaining the versioned benchmark of Section 8.8.

The same discipline bounds what the layer may do once it has a diagnosis. Table 22 fixes the action policy: each observed pattern (reach and conversion read jointly with momentum, in the vocabulary of Tables 9–10) maps to a provisional diagnosis, a bounded form of permissible assistance, the mechanism test that could corroborate or weaken the diagnosis, and the inference or action that remains prohibited. The mechanism-test column operationalizes the rule stated above that a released recommendation is a hypothesis. If a classification correction is supposed to remove a reach bottleneck, the outcome that counts is age-adjusted reach, and a retrospective narrative about total downloads will not stand in for it. A failed prediction is recorded as a failed intervention and returned to diagnosis, never re-scored against a more flattering success measure. The full sequence thereby closes into a loop: **measure** → **validate** → **normalize** → **diagnose** → **propose** → **criticize and verify** → **decide** → **observe** → **update or stop** (Figure 17). Within that loop the human decision stands between verification and any action, post-decision observations return to the case file, and the abstention path terminates the cycle whenever the evidence is insufficient. The loop controls complexity too: when the evidence supports only preservation and monitoring, or when a deterministic rule of the author playbook (Section 6) already answers the question, orchestration adds nothing and should not run.

Table 22. Diagnosis-to-action policy for the proposed orchestration layer. *Patterns are read within a defensible reference class under the percentile conventions of Table 9; every row is a policy hypothesis, and no cell asserts an empirical frequency.*

Observed pattern (reach, conversion, momentum)	Provisional diagnosis	Permissible assistance	Mechanism test	Prohibited inference or action
Low reach, high conversion, stable or rising	A discoverability bottleneck is more consistent with the pattern than a manuscript-uptake bottleneck ("hidden interest," Table 9)	Audit identifiers, classifications, metadata, indexing, and title discoverability; bounded, genuinely relevant distribution	After a metadata or classification change, does reach rise while conversion holds?	Rewriting a well-converting abstract merely to make it more promotional
High reach, low conversion, stable or falling	Communication mismatch, audience mismatch, and low-intent exposure are competing explanations ("attention without uptake")	Test title–abstract alignment, scope clarity, and audience fit; claim-preserving alternatives only	Does conversion change under comparable exposure, with a fidelity check on every revision?	Increasing indiscriminate promotion; concluding that readers rejected the science
Low reach, low conversion	Underidentified: weak exposure, wrong reference class, metadata failure, communication friction, or niche demand all fit ("cold start")	Validate the reference class and exposure record first; then audit metadata and communication separately	Change one measurable mechanism at a time; preserve the null result	Inferring low quality from low usage; returning a single-cause verdict
High reach, high conversion, rising	Strong early diffusion relative to the stated class	Preserve provenance; monitor at fixed ages; change nothing without cause	Does the pattern persist after the exposure event and under anomaly sensitivity?	Manufacturing further amplification; treating attention as proof of quality
High reach, high conversion, sharply falling increments	Event-driven attention, ordinary decay, or saturation (Section 5.6)	Identify the exposure event; compare interval increments, not cumulative totals	Does the decay differ from age-matched trajectories conditional on the event?	Calling declining increments scientific failure — cumulative counts almost always flatten in rate terms
Low reach, high conversion, falling increments	Conditional interest may be real while continuing exposure is weak	Audit whether distribution or indexing ended; identify bounded, relevant channels	Does restored exposure raise reach without degrading conversion?	Broadening claims or removing technical specificity to chase a wider audience
High reach, low conversion, rising exposure	Arrivals are growing without proportional full-text uptake	Test low-intent traffic (Section 4.2), audience mismatch, and presentation clarity	Does conversion respond to a claim-faithful presentation change at comparable exposure?	Optimizing only the top of the funnel; counting every view as a human reader

layer becomes a new gatekeeper in the stack this article maps (Figure 12), and it must not inherit the opacity it was built to repair: reasoning traces and provenance logs must be open to inspection, and the layer's own retrieval and interpretation behavior should be audited with the discipline of Study 2 (timestamped snapshots plus a scripted replication protocol), because model non-stationarity makes every output a dated estimate, never a durable value. **Cost and equity.** Operating such a pipeline is not free, and a paid interpretive layer would exclude the low-resource and unaffiliated authors this framework centers. Known biases of AI text classifiers against non-native English writers [97] would, unmitigated, do the same. **False precision is the residual risk.** A layer that emits banded verdicts without confidence tiers would automate the very miscalibration diagnosed in Section 3: a count with no comparative scale replaced by a verdict with no uncertainty attached. The architecture is itself a falsifiable object and belongs in the agenda on those terms, evaluated against simpler assistance and not inferred from the sophistication of its diagram. The appropriate design is graded. Realistic case packets with planted ambiguities (invalid citations, counter restatements, indefensible cohorts, genuinely borderline solicitations) are randomized across conditions that add one ingredient at a time: the manual protocol of this article's tables (the correct baseline, since a non-AI method already exists), a single generalist model, the same model with retrieval and deterministic calculation, generation plus independent criticism, and the full governed pipeline. Ablations that remove the verifier, the deterministic calculator, the visibility of disagreement, or the human gate then isolate which governance ingredient earns its cost. Outcomes are scored separately, never as one composite: claim and citation fidelity; reference-class validity; calibration of author beliefs (Study 4's instrument) and appropriate abstention on ambiguous cases; researcher time and oversight burden; governance failures (autonomous external actions, metric-gaming advice, categorical venue accusations); and, central to the article's cold-start argument, heterogeneity of benefit by prior author visibility and newcomer status, since a layer that helps mainly already-visible authors would reproduce the inequality it was designed to reduce. Null and adverse equity results are reportable findings. Model versions, prompts, tool permissions, and source snapshots are archived, because in a non-stationary system every result is a dated estimate. Two estimands stay separate throughout: the effect on interpretation quality and the effect on subsequent diffusion. A layer that improves diagnosis without moving downloads is a success, while one that raises downloads while distorting claims is a governance failure. The design fails if the governed pipeline does not outperform a strong single model on evidence accuracy and calibration: "several models" is not the contribution, and the hypothesis under test is the division of epistemic labor. The proposal specifies a testable role for AI-assisted interpretation. AI can help an author ask a better question of a number; only verified evidence can answer it.

9. Limitations

The claims of this article are bounded by the limitations below, each attached to the artifact it constrains.

Aging and composition of the evidence base. The older distributional anchors in the benchmark map (Table 7) rest on samples collected between 2012 and 2015 and dominated by legal scholarship [85][139]; the one recent entry is the single-network snapshot discussed below [83]. The interpretation scale's bands (Table 8) are best read as lenient rather than conservative absolute bands for 2026 (Section 5.1). Platform-wide annual download volume grew severalfold between the mid-2010s and the 2020s [1][143], lifting download-per-paper flow and making a fixed old threshold a softer bar today than when it was estimated. No multiplicative correction can repair this, since growth changes numerator, denominator, and composition simultaneously (Section 3.2). The bands of Table 8 are approximations derived from

aggregate arithmetic rather than rescaled norms. No post-2020 platform-wide distributional study of SSRN downloads was identified.

Single-source contemporary evidence. The only recent distributional evidence used here is a single published snapshot of one disciplinary network: 2,361 Marketing papers observed on 15 August 2025 [83]. It demonstrates skew and a conditional prestige premium, but it does not establish platform-wide benchmarks, and two internal values of its published table are reproduced "as published" with a cautionary note. Generalizations from this source to other SSRN networks are extrapolations. Appendix E reduces this dependence by adding an independent, all-field archival cross-section ($n = 1,360$) that reproduces the skew, the mean-to-median gap, the conversion band, the field multiplier, and log-normality on the observed body. It carries its own upward selection bias and corroborates the distribution's shape across disciplines. The agenda of Section 8 still targets the platform-representative benchmark.

Top-list bias. The 2015 large-sample fit covers only the top lists (roughly the upper decile of authors and their papers [85]), so the body of the distribution below the elite lists is inferred from a fitted lognormal rather than observed. Percentile statements derived from that fit are model-based.

Counter semantics are unstable. What views and downloads record has changed over time and will continue to change with interface redesigns and evolving bot filtering. The platform's integrity system excludes some suspicious activity from official counts [144][66][69], and the adjacent RePEc/LogEc infrastructure documents both the historical scale of robot traffic and recent AI-crawler waves that required discarding the overwhelming majority of raw traffic [99][129]. Conversion norms cited in Section 5 are therefore vintage-specific quantities, and comparisons of counters across years conflate behavioral and filtering changes. General web-bot statistics [158] cannot be extrapolated to SSRN's filtered counters in either direction.

Platform volatility and source drift. The Rankings retirement is embedded in a broader 2026 product restructuring [146][149][150], so no analysis should treat 15 July 2026 as an isolated intervention. This threat is built into the design of Section 8.5 but cannot be fully eliminated. Platform support pages are revised without notice; all such pages are cited here with access dates, and quoted platform language should be re-verified against archived copies before formal submission.

No percentile data. No public full-platform percentile data have ever existed, and after 15 July 2026 no platform-representative percentile frame can be collected externally, because the rankings that previously leaked top-list thresholds have been removed [149]. A sample-based reference distribution can still be assembled from archived captures, and Appendix E reports one; it fixes the shape of the distribution and does not deliver platform percentiles. Some quantities of interest (for example, the fate of the "top 10%" author notification emails [161]) are officially undocumented, and we note their disappearance from the public frame without asserting their termination.

Usage is not readership, and neither is quality. A download may precede careful reading, superficial inspection, archival storage, or no use at all. Downloads are not a direct indicator of research quality and not a guaranteed causal source of citations: the downloads–citations association is positive but heterogeneous (roughly 0.1 to 0.6 depending on field, level of aggregation, and horizon [18][25][39][80][102][108]) and the highly downloaded and highly cited populations overlap only partially [108]. Citations themselves can be positive, negative, perfunctory, or strategic, and the citation process is

systematically biased against novel work [166]. Nothing in this article licenses metric substitution for expert judgment; the framework's defensible use is diagnostic and descriptive.

Unobservable exposure and non-random prestige. Not every exposure is observable: readers encounter titles in email, social media, search results, and correspondence before any abstract-page view is logged, and well-networked authors generate off-platform exposure correlated with both prestige and quality. Observational prestige coefficients therefore mix status-based attention with genuine differences in resources, topics, and networks [111][140]. Experimental manipulation of author-identity signals [154] [64] complements but does not replace platform observation. Field classification is partly author-selected, and papers occupy several categories, complicating single-field assignments.

Limitations of the agenda itself. The quasi-experimental designs of Sections 8.4–8.5 face announcement effects, concurrent platform changes, and non-random field differences. They are labeled quasi-experimental because a platform-wide event offers no untreated control group. Study 2's snapshot will age with the engines it audits, and its replication protocol mitigates but does not remove non-stationarity. Study 3 is an $N = 2$ case study whose value is mechanism and protocol, not a population estimate, and its attribution logic depends on per-channel aliases rather than on any claim about public email visibility. The crowdsourced registry inherits self-selection and verification problems stated in Section 8.8. Data access for all designs depends on either platform cooperation or rate-limited collection within terms of service [134][51]; neither is guaranteed.

Breadth of the bibliography. The bibliography is broad because the article is a synthesis; many of the cited studies share frames, fields, and design limits, and community posts document experience without establishing population parameters. A claim's evidence is its verification tier and sampling frame, not the number of references behind it.

Bands are calibration targets, not constants. The 11–19% conversion band and the CSAB thresholds are quantities to be calibrated by field and by platform, not universal constants of scholarly behavior (Sections 4.4, 5.2). Their present form is a single cross-field band, to be replaced by the per-field distributions the framework requires when the field-stratified Study 1 (Section 8.6) reports.

Generalization beyond SSRN requires recalibration, not transplantation. The interpretability problem diagnosed here (usage counters displayed without a public comparative scale) and the reach-versus-conversion logic that repairs it are not specific to SSRN; they apply to any preprint server that shows usage counts, including bioRxiv, medRxiv, arXiv, and OSF Preprints [11][62][2][24]. What does not transfer is the specific 11–19% band or the CSAB thresholds. Curation, distribution architecture, and interface differ across servers: arXiv's intent-driven pull traffic and its email announcement lists against SSRN's curated push alerts (Table 6). The numeric anchors must therefore be re-estimated on each platform before they are quoted there. The framework is portable; its calibration is local. The parallel-deposit protocol of Section 8.8 is the direct empirical test of how far the SSRN bands travel. The scale of that recalibration is easy to underestimate: bioRxiv's all-corpus median download (about 279, and roughly 496 in genomics [2]) is several times the median of 60 in the archival cross-section of Appendix E: roughly five times overall and about eight times in genomics. The comparison is indicative only. The bioRxiv figure comes from a whole-corpus census, while the 60 comes from an archival sample that over-represents captured pages, mixes ages and fields whose composition shifts with age, and resolves age to whole years. The direction of the gap is safe to read; its size is not.

Ethics. The community-discourse strand of this article consists of self-selected public posts used only as evidence of demand, never as representative sentiment, and quoted only where a live, verifiable source exists. User posts are cited with attention to the ethics of quoting semi-public forum content, including paraphrase in place of verbatim quotation where authors could be identified against their expectations, consistent with transparency norms for qualitative evidence [6]. Study 3 logs only messages sent unsolicited to the instrumented author and collects no third-party personal data beyond what senders volunteer. Study 4 collects only voluntary, anonymized questionnaire responses under informed consent. The panel design uses publicly displayed counters and bibliographic metadata, and any collection of non-public dashboard information (such as distribution-status dates) requires informed author consent and, where applicable, institutional review. The two studies that involve human participants and participant-generated data are Study 4 (the author survey) and Study 3 (instrumented posting with inbox logging). Because both are human-subjects research, both will be submitted for approval by the relevant institutional review board or research-ethics committee and conducted under participants' informed consent *before* any data collection begins, and personal data contained in the inbox logs are aggregated and de-identified prior to any analysis or reporting. This protocol is a precondition of execution and is reproduced in the Declarations.

10. Conclusion and Policy Recommendations

10.1 What the analysis established

The retirement of SSRN Rankings on 15 July 2026 removed the only public comparative scale [149]. Counters remained on every paper and profile, which left metrics that were unreadable rather than meaningless. Against that loss this article assembles every verifiable empirical anchor into a benchmark map with explicit vintage and confidence caveats (Table 7), builds an interpretation framework, the five-band scale, the reach–conversion decomposition, and the dossier protocol (Tables 8, 9, 13), that translates counters into diagnoses, and sets out the preregistration-ready agenda of Section 8, which specifies how the missing distributions can be rebuilt.

The mechanism at the core of the argument comes down to this:

Nobody knows what seven downloads means without knowing seven downloads compared with what, and after 15 July 2026 the platform no longer answers "with what."

The contributions differ in kind. Naming the reference-class problem is the diagnostic one: a raw count answers no question until age, field, cohort, and exposure define its comparison group. The instrumental contribution is the benchmark map and CSAB, which normalize an emotionally salient but scientifically ambiguous number into an interpretable diagnostic. The playbook and the dossier protocol are the practical contribution, letting authors and committees act on counters today without misreading them. In policy terms the contribution is the anonymized percentile band: comparative context restored in a form that identifies no one and revives no league table.

On the assembled evidence, the question in this article's title is ill-posed until a reference class is specified. Against every available anchor, small counts are the statistical norm of the platform rather than a verdict on the work. A paper with 40 abstract views and 7 downloads is consistent, as a screening result, with the empirical conversion band of Section 5 [139][172][83], while the older three-views-per-

download rule of thumb [18] sits above that band rather than within it. The median paper in the only census-style sample accumulated 63 downloads in two years [139], and entry into the 2015 top-10,000 list, roughly the top 1.8% of the platform's full-text paper stock (571,040, July 2016 [1]), required 160 lifetime downloads [85]. Under a quasi-lognormal attention regime, most authors stand on the ladder's lower rungs by construction. The median preprint lives in the long tail of a log-normal distribution its author cannot see, and it is surfaced by classification mechanics that favor the affiliated. It is solicited by predators before it is read by peers, and it has now lost the ranking that once located it. The scale places a paper somewhere on that ladder, and the distress the question carries is produced by the missing scale rather than by the number.

Nothing in this is peculiar to one platform, and the event is not unprecedented. ResearchGate retired its own heavily criticized composite RG Score in August 2022 [130], after sustained methodological critique [90], and in 2022–2023 dozens of U.S. law schools, led by Yale and Harvard, withdrew from the U.S. News law-school rankings [13]. Withdrawing a contested measure is not a cosmetic act, because the measure was never a passive description: the reactivity literature that grew out of the law-school rankings shows that a public ranking reshapes the behavior of those it ranks [58][132], so both its imposition and its removal change conduct rather than merely change what is reported. A counter and its comparative frame are jointly a metric, and retiring the frame while keeping the counter leaves every user holding a count with nothing to read it against. Interpretability, the availability of a public distribution for a public counter, should be treated as a first-class quality dimension of any indicator, alongside validity and reliability. Metrics outlive their interpretive infrastructure, and a metric without context is noise, an instance of the general discipline of skeptical reading of decontextualized numbers [15]. The same metric *with* context is a calibration that empowers authors, improves scholarly communication, and sustains trust in the repositories that display it.

The framework asks a few ordinary habits of its user. **Compare like with like, and keep reach separate from conversion. Read trajectories rather than totals, treat every number as dated, and let no count travel without its source, window, and limitations.**

10.2 Recommendations for authors

Authors can act now, without waiting for new data. Take dated screenshots of paper-level counters, which preserve context: after the sunset they are the only record of standing that remains under the author's control, and they seed the community registry of Section 8.8. Report counts only with their frame (age, field, cohort, and source), following the dossier protocol of Table 13, in which every number carries its source, time window, and limitations. Read the dashboard diagnostically rather than emotionally: separate reach from conversion (Table 9), treat a zero-download first week as normal pre-distribution behavior [53], and benchmark against three to five papers by colleagues posted in the same month rather than against the visible tail, whose circulating milestone numbers guarantee miscalibration under a heavy-tailed regime [85][180]. Treat solicitation email as a scraping phenomenon, not as recognition, and not as evidence about the counters [61][107][144].

10.3 Recommendations for evaluators

Institutions that imported SSRN downloads as an evaluation currency [12][18][161][52] face a governance question now that the currency's public scale is gone. Raw download counts should not enter hiring, promotion, or tenure files: they are gamed at exactly the thresholds that matter [44], they vary up to eightfold across subfields within a single network [168], and they are now benchmark-opaque. The

minimally defensible unit is a percentile within a defined cohort (same field, same age band, same posting period) accompanied by its source and vintage, and usage evidence should complement rather than substitute for citations, peer assessment, and substantive judgment [108][109][110]. The established research-assessment reform consensus works on that principle. DORA's core injunction is not to use a single number or a venue proxy as a substitute for assessing the research itself [8], and the Leiden Manifesto insists that quantitative evaluation support, not supplant, qualitative expert judgment [78]. The Metric Tide review reinforces that insistence, calling for indicators that are robust, humble, transparent, diverse, and reflexive and used in support of expert assessment rather than as a substitute for it [173], as does the more recent CoARA Agreement on Reforming Research Assessment, whose several hundred signatory organizations commit to qualitative evaluation supported by, not supplanted by, responsible quantitative use [183]. The whole consensus applies with full force to a download counter whose public scale has been removed. Two of the Metric Tide's five dimensions bear directly on the method proposed here: diversity is the field-normalization this article insists on, and reflexivity is the stance of Section 8.7, in which the article measures its own attention. The review's REF-facing successor carries the same five principles into current UK practice [184]. This consensus is itself one expression of the broader critique of metric overreach in institutional life [113]. Evaluation committees should expect and accept the protocol format of Table 13 (the contextualized, narrative form of evidence that emerging assessment formats such as the UK's Résumé for Research and Innovation (R4RI) are built to carry [186]) in place of the retired "top 10%" ranking language. In the United States the same dossier-language lever is the focus of the Higher Education Leadership Initiative for Open Scholarship (HELIOS Open), which works with universities to rewrite promotion and tenure criteria [189]; the contextualized usage line specified here is exactly the concrete rubric such reforms often lack.

10.4 Recommendations for the platform

Retiring the rankings without replacement context trades one flawed signal for no signal. The consequences reach past author comfort. Whether the sunset chills new submissions, the posting-volume outcome specified in Section 8.5, is an open empirical question about the platform's own health. The framework implies a concrete product direction that does not recreate the league-table apparatus SSRN chose to retire [149]:

1. **Anonymized percentile bands by discipline and paper age.** Display each paper's download and view standing as a coarse percentile band within its discipline-by-age cohort: information every author needs, published in a form that identifies no one, creates no leaderboard, and reduces the winner-take-all dynamic of top-N lists. Percentile context preserves privacy while restoring readability. The target display format is exactly the age-specific benchmark fan of Figure 14.
2. **Time-normalized display.** Show age-matched comparisons rather than lifetime counts alone, so that a 14-day-old paper is never implicitly compared with work that has accumulated attention for years.
3. **Separate reach from conversion.** Report abstract views and downloads as distinct signals with their ratio, rather than inviting a single-number reading that conflates invisibility with rejection.
4. **Transparent reference classes.** Any comparative indicator should state the population, field, period, and age band it compares against; an unlabeled "top X%" is scientifically uninformative. Appendix D specifies the ten disclosure fields that a versioned percentile release should carry, from population and cohort definitions to minimum-cell privacy rules and a citable version history.

5. Audit design for cumulative advantage. Classification, alert distribution, search ranking, and visible popularity signals allocate scarce attention and should be evaluated for their distributional effects. The empirical designs of Sections 8.4–8.5 are directly usable for such an audit.

6. Plan for interpretive residue when retiring features. When comparative features are removed, the interpretive loss should be assessed and mitigated with the same care as data migration, through archived distributions, published summary percentiles, or an explicit handover to the community, since a metric that cannot be read cannot inform.

7. Harvest-resistant contact mechanics, reducing the solicitation externality of visibility documented in Section 7 [107][159]. One concrete candidate: delay or obfuscate the public exposure of author contact routes for roughly the first 48–72 hours after posting, the window in which legitimate distribution has typically not yet begun but programmatic harvesting is most likely. The effect of any such measure should be evaluated empirically, not assumed. The trade-off is real, since hiding contact information could also disadvantage exactly the newcomers who depend on direct feedback.

8. A public filtering log and documented distribution mechanics. Publish a LogEc-style record of what the integrity filters discard, and document how eJournal distribution and digest selection operate. The adjacent RePEc/LogEc infrastructure demonstrates that such transparency is feasible at scale [99][129], whereas SSRN currently describes its filters without publishing their history [144] (Table 6). When the filtering is documented, restatements of a counter can be read as measurement events rather than as reputational ones.

9. Make exposure history legible. A paper page or author dashboard should distinguish public posting, classification assignment, curated distribution, indexing confirmation, and major counter restatements, so that an author can separate "not yet exposed" from "exposed but not converted." This is the platform-side counterpart of the exposure-regime reading of Section 5.6, and the cheapest treatment for the most anxiety-producing dashboard state, the early zero.

10.5 Recommendations for the research community

The community need not wait for the platform. The registry of Section 8.8 (dated, self-reported dashboard snapshots contributed under an open protocol) can produce the first fresh download distribution in a decade. The two public-data studies of Table 17 (1 and 2) are executable by any researcher within the access limits noted in Section 9, while the two participant studies (3 and 4) require volunteers and the ethics approvals stated in the Declarations. Scientometrics should treat metric interpretability as part of its subject matter: the present case, in which a decade of counters outlived their only public scale, is unlikely to be the last. Distributed replication of the instrumented-posting protocol, community maintenance of percentile tables (and of the public author-facing calculator built on them, Sections 8.8 and 8.10), and interpretive-impact assessment of platform changes are all research contributions in their own right, and each of them converts this article's repair work into prevention.

10.6 Closing

The question "what do my seven downloads mean?" has run through this article because it is the question the platform's nearly two million authors were left holding on 15 July 2026 [49][149]. The answer assembled here is neither reassurance nor alarm but calibration. The number is consistent with the available conversion evidence but remains a screening result; the anxiety is manufactured by the missing

scale, which can be rebuilt, partly from the literature, as Sections 3–6 did, and partly from new data, as Section 8 specifies.

The article opened with the question every author is left holding:

"What do my 7 downloads mean?"

The framework replaces it with this:

"Where does this paper stand among comparable papers of the same age, field, cohort, and distribution history? Is its constraint reach or conversion? And is it being retrieved and credited by the systems through which readers now discover research?"

That is a question scientometrics can answer, and this article has specified how.

What the counters record is events, and only events. Views are not downloads; downloads are not reading; a mean is not a norm; a counter is not a verdict; solicitation is not recognition; AI retrieval is not faithful attribution; and a number without a scale is not information.

As long as platforms show a counter apart from its context, authors, evaluators, and platforms will keep reading distributional facts as personal verdicts. Open access won the battle it fought (the right to post and to download without price barriers) and, in winning it, exposed a battle it never fought: the allocation of attention. Authors can read counters with field, age, and conversion. SSRN's metrics still hold value, but only for an author who understands what they mean.

Appendix A. Situation → Action Guide

Table 23 condenses Sections 4–6 into the situations authors most frequently report [161][127][139]. Its middle column states only what the assembled evidence licenses, with sources and windows. The caveats column carries the limitation that must accompany each reading. Numbers appearing here are those of Tables 6–7 and Figure 6 of the main text.

Table 23. Situation → what the evidence says → what to do.

#	Your situation	What the evidence says	What to do	Sources and caveats
1	"7 downloads and 40 abstract views in the first month."	Conversion is 17.5%, inside the cross-study band of 11–19% as a point estimate, though the 95% Clopper–Pearson interval on 7/40 (≈ 7.3 –32.8%) is wide and straddles the band; the median LSN paper reached 63 downloads only after two years, and the median marketing paper stood at 87 at a median age of 16 months — mature-sample anchors, not age-matched to a day-20 reading. A screening result with a wide interval, not evidence of failure.	Continue the Table 11 observation schedule; classify the trajectory no earlier than day 60–90.	Band: [18][139][172]; medians: [139] (2012 LSN), [83] (2025 Marketing). All anchors are single-discipline and aged; the 7/40 conversion is a wide-interval screening estimate, not age-matched to the mature anchors.

#	Your situation	What the evidence says	What to do	Sources and caveats
2	"Zero downloads in week one."	Normal pre-distribution: classification is curated and can take up to 45 days for loaded topics; eJournal alerts drive early views and continue post-sunset; Google Scholar indexing takes 6–9 months or longer.	Confirm public posting; confirm eJournal classification; check whether the paper is indexed under its exact title — the cheapest available indexing test; wait out the announcement cycle; only then read the counters.	[53][147][57][55]. Support-page wording changes; re-verify with access dates.
3	"Many views, few downloads — conversion well below 11%."	Consistent with an abstract–text mismatch: readers arrive and leave at the abstract. Packaging variables (title, abstract, keywords) are the attention levers most robustly associated with counters.	Revise title and abstract; re-observe conversion over the next window before further changes.	[139][168][172]. The causal reading "low conversion = weak abstract" is a screening hypothesis, not a validated diagnosis.
4	"Conversion far above the band, or a sudden view spike."	Traffic that abstract browsing does not explain — media or aggregator links, or non-human traffic that filters have not yet removed.	Identify the external source (coverage, mailing lists); do not cite the spike as readership without corroboration.	[116][138][144][102 — Safe Links pre-fetch mechanism, with the caveat that SSRN's own filtering may already exclude such traffic].
5	"My counter dropped, or the numbers were restated."	Platforms revise counters when filtering rules change; RePEc discarded the large majority of raw traffic in documented episodes (robots $\approx 74\%$ of views in 2007; $>99.5\%$ of IDEAS traffic in September 2025).	Treat it as a measurement event; check your dated screenshots; never report it as lost readership without corroboration.	[144][69][99][129][66].
6	"The rankings are gone and I do not know where I stand."	No public percentile frame remains, and no platform-representative one can be collected externally. The archival cross-section of Appendix E supplies a sample-based reference distribution, and matched comparison remains the substitute for platform percentiles.	Build a micro-benchmark of 3–5 peer papers posted the same month and track them on your own schedule; optionally compute the ratio of your downloads to a matched cohort median.	[149][51]; anchors [139][83][85]. The ratio has no calibrated interpretation bands (Section 6.3).
7	"A committee will see my counts."	Rank lines in dossiers are documented practice, and their comparative frame is gone; band language with sources, field qualifiers, and citation triangulation survives scrutiny where raw counts do not.	Apply the seven-step protocol of Table 13: preserve, translate, triangulate, field-normalize, complement, integrity, expert assessment.	[12][95][33][52][77][108]. Downloads measure demand, not quality; the manipulation history [44] is known to evaluators.

#	Your situation	What the evidence says	What to do	Sources and caveats
8	"Predatory invitations began arriving days after I posted."	A documented byproduct of harvesting from preprint platforms; adjacent audits report tens of invitations per fortnight per academic. Volume signals visibility to scrapers, not recognition by readers.	Never pay; verify venues via Cabells and Think. Check. Submit.; route correspondence through a dedicated address; see Table 14 for the typology.	[107][159][9][84][63][32][136]; [B-1] (Beall's list defunct since 2017). Timing regularities are unmeasured; the "days-after-posting" screen is qualitative.
9	"A journal's similarity check flagged my submission as matching my own SSRN preprint."	The screening system is operating as designed: preprints registered as Crossref Posted Content are part of the comparison corpus, and iThenticate's guidance states that a same-author preprint match is not plagiarism. Posting a working paper is permitted by most journals' prior-publication policies, but policies vary and the target journal should be checked; the main exceptions are in medicine and clinical fields.	Identify the flagged source to the handling editor as your own time-stamped preprint; do not withdraw the preprint; for medical and clinical venues, check the target journal's preprint policy (e.g., via the Sherpa Romeo registry) before submission.	[81][36]; Section 4.5. Journal policies vary and change; verify the target journal's current policy rather than relying on the general rule.
10	"The Rankings are gone — should I move my papers to ResearchGate or Academia.edu?"	No post-sunset author migration is documented; download counters are not transferable across platforms; and the comparator platforms carry the same interpretability problem — ResearchGate retired its own composite RG Score in August 2022 after methodological critique.	Keep the SSRN record and complement rather than relocate: deposit the same paper in an institutional or subject repository (Table 13, step 5), preserve dated counter snapshots, and do not delete-and-repost elsewhere, which forfeits accumulated counters.	[130][90][149]; [47][146] (repository complement); [161] (one platform's counter says nothing about other channels); [B-7] (delete-and-repost resets metrics). Absence of documented migration is absence of evidence, not proof of stability.

Unverified community practices (candidate hypotheses for the Section 8 agenda). The following recommendations circulate in author communities and appear in earlier drafts of practical guides; no study supports them, and they are recorded here as testable hypotheses and must not be cited as findings: (a) capping eJournal classifications at three (one core, one broad, one interdisciplinary) rather than using more of the permitted seven [53][168]; (b) citing SSRN working papers so that their authors are notified through Google Scholar alerts and reciprocate with attention; (c) concealing the profile email address for the first fourteen days after posting to reduce harvesting; (d) uploading a minor revision on day 3 to trigger re-announcement (SSRN documents only that in-place revisions preserve counters and URL [B-7]); (e) any numerical incidence attached to the "invitation within 96 hours is predatory" screen; (f) direct outreach to three to five authors cited in the paper (a short message noting the citation of their SSRN work) as a channel to a first legitimate audience, framed as ordinary scholarly communication rather than promotion. Each is falsifiable with the observational designs of Section 8; until tested, their proper citation form is "community practice, unverified."

Table 24 re-sorts the playbook chronologically. Every cell points back to a section and its sources. The timeline view answers the question authors actually ask ("what do I do *now*?") in the order in which the situations arise, and complements the situation-indexed Table 23.

Table 24. The playbook as a timeline: three phases × three axes. *All cells summarize evidence and instructions already stated in Sections 4–7; unverified community practices are excluded (see the list above).*

Phase	Platform mechanics and measurement	Audience and distribution	Solicitation defense
Before posting	Treat the title and abstract as the conversion surface and draft them accordingly (§6.4; [139][172]); plan eJournal classifications — up to seven are permitted, and more classifications are associated with more downloads and fewer zeros [53][168]	Identify the 3–5-paper micro-benchmark cohort to track alongside your own paper (§6.3)	Route correspondence through a dedicated address before the metadata becomes public (§6.6)
At posting	Confirm public posting and eJournal classification; record the day-1 baseline of the Table 11 grid (§6.1)	eJournal email alerts are the platform's principal push channel and continue post-sunset [147][57]	Expect solicitations as a default consequence of visibility, not a signal of merit (§6.6; [159][63])
First weeks (days 1–45)	Read an early zero as pre-distribution, not rejection (§6.2 rule iv; [53]); log counters on the Table 11 schedule; classify the trajectory no earlier than day 60–90 (§6.3)	Apply legitimate external levers where appropriate — blog and list coverage produced measurable early spikes (§6.4; [116] [72])	Never pay; screen any inviting venue ([B-1]; Think. Check. Submit. [159]); check invitations against the Table 14 typology; never read spam volume as readership (§7.6)

Appendix B. CSAB Reporting Template

A minimal author- or evaluator-facing CSAB statement (Section 5.5; Table 10) reports:

At paper age [X days], the paper had [V] abstract views and [D] validated downloads. Relative to [defined field / cohort / newcomer reference class], its reach percentile was [RV], its download percentile [RD], and its shrinkage-adjusted conversion percentile [RC]. The paper [had / had not] received documented curated distribution by that observation. Between [age 1] and [age 2], its normalized position [rose / remained stable / fell]. These indicators describe diffusion and do not establish research quality or correctness.

Until a contemporary reference panel has been estimated (Sections 5.5, 8.3), **the percentile fields must remain blank**: no value may be entered from intuition, from platform-wide means, or from aged samples. Historical anchors may be reported alongside the template: the 2012 LSN median of 63 downloads at two years [139] and the 2025 Marketing-network median of 87 [83], each with its sample, field, and age stated, but never substituted for the missing reference class. The template is the sentence-level counterpart of the dossier protocol's Translate step (Table 13, step 2).

Appendix C. Panel Variable Dictionary

Table C1 states the minimum paper-interval record for the core panel of Section 8.3: the smallest schema sufficient to reproduce the CSAB estimates of Section 5.5 and the platform-event analyses of Sections 8.4–8.5. Public releases replace direct identifiers with stable study keys wherever linkage is not analytically necessary, and registry contributions (Section 8.8) use the subset of this schema indicated there.

Table C1. Minimum paper-interval data schema for the core panel.

Domain	Variable	Definition or format
Identity	paper_id	Stable study identifier
Identity	ssrn_abstract_id	Platform identifier; hashed or restricted in public releases
Timing	posted_at	First public posting date
Timing	observed_at	Metric observation timestamp (Table 11 grid)
Timing	paper_age_days	Days since first posting
Usage	abstract_views_cum	Public cumulative abstract-page views
Usage	downloads_cum	Public cumulative validated downloads
Usage	views_interval, downloads_interval	Increments from the preceding valid observation
Exposure	classification_ids	All eJournal classifications active at observation
Exposure	distributed_status, distributed_at	Curated-distribution indicator and date, where consented
Exposure	search_indexed	Protocol-defined external index status (exact-title test)
Paper	title_frozen, abstract_frozen, keywords, page_count, doi	Earliest observed manuscript metadata, frozen at first observation
Paper	authors_frozen, affiliations_frozen	Author names, order, and affiliations at first observation
Versioning	revision_id, revision_at, material_change_flag	Version history with a material-change indicator
Authors	author_count, platform_newcomer, prior_ssrm_papers	Team structure and pre-posting platform history (Section 8.3 definitions)
Reputation	prior_citations, prior_works, institution_measure	Frozen before the focal posting and field-normalized
Calendar	day_of_week, week, season, holiday_flag	Predefined temporal controls [178]
Downstream	citations_6m, citations_12m, citations_18m, publication_status	Later validation outcomes with source and date (Section 8.9 horizons)
Audit	collection_method, source_url, quality_flag	Provenance, access route, and anomaly documentation (flag, not delete; Section 8.3)

Appendix D. Benchmark Transparency Specification

To keep the percentile benchmark of Section 8.8, and any successor maintained by the platform or by the community (Section 10.4, recommendations 1 and 4), from reproducing the benchmark opacity this article criticizes, every percentile release is published with a versioned companion specifying:

1. the eligible paper population, networks, document types, and exclusions;
2. posting-cohort dates and observation ages;
3. the field-assignment rule, including the treatment of multiple classifications;
4. the definitions of views, downloads, unique actors, repeats, and validated traffic in force at release;
5. the median, interquartile range, P10/P25/P75/P90/P95, zero share, Gini coefficient, and top shares;
6. sample sizes and uncertainty intervals for every reported reference class;
7. the handling of missingness, withdrawals, revisions, merges, and counter restatements;
8. the distribution and external-promotion variables included in any adjustment;
9. privacy and minimum-cell rules: for the crowdsourced registry of Section 8.8, where small field-by-age cells could otherwise identify individual contributors, this item is a binding constraint governed by data-protection law (in the UK, UK GDPR and the Data Protection Act 2018), not a formality. Disclosure control should extend beyond simple anonymization to k-anonymity or noise-added (differentially private) aggregates for thin cells, and a permitted-use rule should bind the registry to interpretation rather than evaluation, forbidding its use for hiring, promotion, resource allocation, or reassembly into any public leaderboard;
10. a version history and a permanent, citable reference for the release.

No composite score may replace these components. Simple ratios to a cohort median (the uncalibrated fallback of Section 6.3) can be unstable when the denominator cell is small and can falsely suggest that "twice the median" maps to a fixed percentile across differently shaped fields. Author-facing reporting therefore uses percentile estimates under the empirical-Bayes shrinkage already specified for sparse cells in Section 8.3, while raw distributions remain available for audit.

Appendix E. An Independent Archival Cross-Section of the Download Distribution

For recent whole-distribution evidence the synthesis of Sections 3 and 5 leans on a single published snapshot [83]. This appendix reports an independent, reproducible cross-section assembled for this article that corroborates the distributional structure across fields, using only public data and no privileged platform access. It is a one-time archival reconstruction with a known upward bias, offered as corroboration of shape, distinct from the preregistered panel of Section 8.6, which observes papers prospectively from posting.

Method. A random sample of SSRN papers was drawn from OpenAlex (source S4210172589, "SSRN Electronic Journal") across twenty-three independent random seeds. For each sampled paper the pre-retirement download and abstract-view counters were read from the paper's most recent capture in the Internet Archive, located through the Wayback Machine's public availability endpoint, rather than from the live site, which is access-gated (Section 4) and no longer publishes comparative frames. Reading the

counters from archived copies rather than by querying the platform keeps the procedure rate-limited, respectful of the platform, and reproducible from the released identifiers. Because the per-paper counters survived the Rankings sunset unchanged, and only the comparative Rankings were withdrawn, a capture dated after 15 July 2026 remains a valid reading of a paper's cumulative counts. Snapshots were therefore not restricted to the pre-sunset window, and the capture date is retained only to compute paper age. Captures that returned an interstitial bot-challenge rather than page content were discarded. Of 2,600 sampled papers, 64% had a usable archived capture and 52% yielded a machine-readable download count, giving $n = 1,360$ papers spanning twenty-five fields, fourteen of them with at least twenty papers, and paper ages from the posting year to over two decades. The dataset is released as appendixE-wayback-downloads.csv and deposited at <https://doi.org/10.5281/zenodo.22285871>, where the same file carries the name study1-wayback-downloads.csv. Two properties of the sample constrain what can be read from it. Posting dates carry a year and no month: 1,358 of the 1,360 records are stamped 1 January, because OpenAlex reports SSRN postings at year resolution. Age is therefore computed as capture year minus posting year, reported in whole years, and no shorter band is defensible. Captures are also concentrated in time: 688 of them, half the sample, fall in a single crawl between 25 April and 7 May 2025, and 81.5 per cent fall in 2025, so a paper's age here runs close to a restatement of its posting year and the two cannot be separated. The first of these is removable. The posting line “Posted: DD Mon YYYY” appears in the same archived captures, so full dates can be recovered by re-parsing them, which the access gate on the live site does not prevent. Table E1 makes this attrition explicit.

Table E1. Construction of the archival cross-section and attrition at each stage.

Stage	Papers	% of initial	% of previous
Randomly sampled OpenAlex records (SSRN Electronic Journal)	2,600	100.0	—
Usable Internet Archive capture located	≈1,660	63.8	63.8
Capture carried a machine-readable download count	1,360	52.3	81.9
Final analytical sample	1,360	52.3	100.0

Results. The cross-section reproduces every qualitative feature the published anchors assert (Table E4). Downloads are heavily right-skewed: mean 219, median 60 (95% bootstrap CI 55–69), mean-to-median ratio 3.65, with the mean exceeding the 75th percentile (158). They also carry a heavy upper tail (95th percentile 816, more than thirteen times the median; Gini 0.75, 95% CI 0.71–0.78). In the logarithm the distribution is close to symmetric (skewness of $\ln$ downloads ≈ 0.18), consistent with approximate log-normality on the observed body of the distribution, the property Kakushadze could fit only on the top lists (Section 3.4). The download-to-view conversion has a median of 16.2% and a mean of 17.3%, inside the 11–19% cross-study band (Section 5.2). Median downloads differ across fields at the same paper age. Holding the age band fixed, the spread between the highest and the lowest field median runs from about three to about five times, the exact figure depending on the band and on the treatment of thinly-sampled

fields. Pooled across all ages the spread reaches about twelvefold, and that larger figure carries the sample's changing composition: Engineering supplies 21 to 27 per cent of papers aged three years and under and at most 5 per cent of those nine years and older, while Economics supplies 9 to 12 per cent and 37 to 54 per cent of the same groups, and the median age of an Economics paper here is nine years against one year for Engineering. The direction agrees with the field multiplier of Section 3.6. The magnitudes are not comparable, because that multiplier rests on lifetime means across 94 subfields of a single discipline and this spread on medians across 25 disciplines, and neither is adjusted for age. Median downloads are higher in the older age bands, with non-overlapping 95% bootstrap confidence intervals between most adjacent bands (median 12 in the posting year, 172 beyond ten years), which is the age dependence that motivates the reference-class argument. Within a single field the difference is concentrated in the first three or four years: between the five-to-ten band and the beyond-ten band it does not reach significance (Mann-Whitney $p = 0.57$ for Economics, 0.82 for Social Sciences, 0.23 for Business). About a quarter of the pooled gradient is field composition rather than age. Age in this sample also runs close to a restatement of posting year, since half the captures were taken in one crawl between 25 April and 7 May 2025 and 81.5 per cent fall in 2025, so age and cohort cannot be separated here and the bands report levels rather than the path of any single paper. The sample mean (219) coincides with the published marketing-network mean (221) [83], while the sample median (60) sits at the level of the census-style legal median (63) [139]. Table E2 gives the download distribution by paper age and Table E3 by field, both with 95% bootstrap confidence intervals and sample sizes. Confidence intervals are widest for the smallest fields, which is why the field spread is described as roughly an order of magnitude rather than a single precise multiple. The article's opening case can be read against the sample directly. Among papers under two years old carrying 30 to 50 abstract views, the band that case falls in, the median download count is 7, with an interquartile range of 4 to 11 ($n = 42$). The case in the title is the typical case for its view band and age. In the posting year 35 per cent of the sample records 7 downloads or fewer, the median abstract-view count is 66, and the case's own conversion of 17.5 per cent sits at the 57th percentile of the sample. The 11–19% conversion band of Section 5.2 is a band of medians rather than of papers: 41 per cent of the sample falls inside it, 23 per cent below and 36 per cent above. Half the sample, 51 per cent, sits below the census-style legal median of 63 [139], which places the nascent range in the populated part of the distribution.

Table E2. Download distribution by paper age in the archival cross-section.

Paper age	n	Median [95% CI]	Mean	IQR (P25–P75)
posting year	203	12 [11–14]	33	5–21
1 year	253	25 [23–29]	94	15–44
2–4 years	375	52 [47–57]	110	28–100
5–9 years	188	130 [110–151]	384	72–374
10 years and over	341	172 [157–189]	453	96–336
All papers	1,360	60 [55–69]	219	23–158

Table E3. Download distribution by field (fields with $n \geq 20$) in the archival cross-section.

Field ($n \geq 20$)	n	Median [95% CI]	Mean	IQR
Economics	285	150 [134–169]	443	76–302
Business	161	150 [126–188]	335	72–330
Psychology	24	107 [35–192]	179	32–229
Social Sciences	241	85 [78–117]	219	45–189
Decision Sciences	32	74 [56–158]	615	49–260
Computer Science	63	42 [36–73]	167	26–104
Medicine	40	42 [24–59]	81	22–66
Agric. & Biological Sci.	36	31 [24–57]	65	23–67
Biochemistry	21	30 [23–37]	38	20–38
Environmental Science	72	26 [18–32]	44	14–50
Engineering	196	24 [20–26]	46	13–46
Physics & Astronomy	21	23 [15–38]	112	13–45
Materials Science	55	18 [15–25]	26	11–36
Energy	21	12 [11–22]	25	11–22

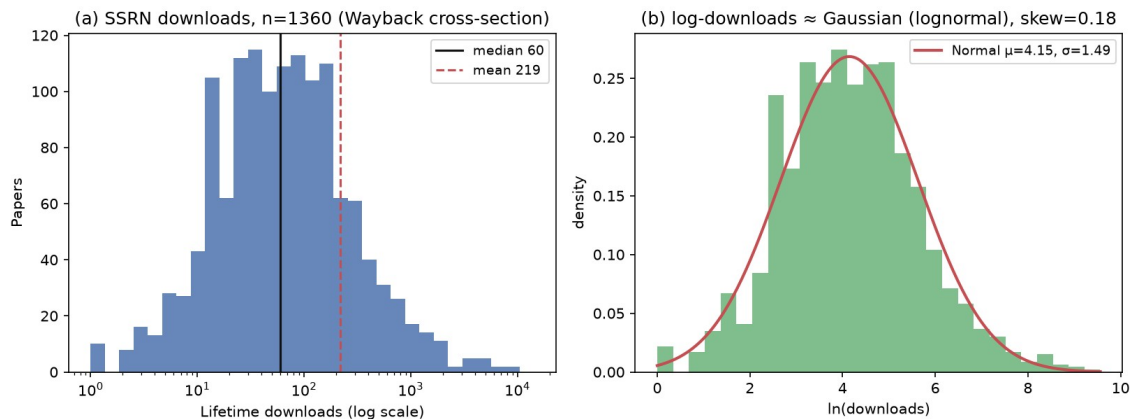


Figure E1. (a) The download distribution on a logarithmic scale with the mean and median marked; (b) a histogram of $\ln(\text{downloads})$ with a fitted normal density. Skewness of log-downloads ≈ 0.18 . $n = 1,360$.

Table E4. Published anchors versus the independent archival cross-section ($n = 1,360$).

Feature	Published anchor	Archival cross-section
Mean downloads	221 (marketing) [83]	219
Median downloads	87 (marketing) [83]; 63 (law) [139]	60
Mean / median	2.54 [83]	3.65
95th percentile / median	> 8 [83]	13.6
Download-to-view conversion	11–19% band (§5.2)	16.2% median
Field spread (medians)	$\approx 8\times$ field means, not age-adjusted [168]	$\approx 3\text{--}5\times$ at equal age; $12.5\times$ pooled
Distributional shape	log-normal, tail-only [85]	log-skew 0.18 (observed body)

What the cross-section establishes. The cross-section corroborates the shape of the distribution and its central and upper structure across disciplines, on an independent all-field sample far broader in field coverage than any single published SSRN snapshot. It does not estimate platform-representative percentiles, and two features of its construction bias it upward. Archived captures over-represent visible pages, so download-heavy papers are the more likely to be captured. A paper also enters the sample only if it had an archived page carrying a readable counter, which censors the extreme low tail: the sample's share of papers with ten or fewer downloads (about 10%) is accordingly far below the 15–30% that Sections 3.4 and 8.6 anticipate for an unbiased whole-platform sample, and must not be read as an estimate of that share. That the sample median (60) nonetheless lands at or below the published legal (63) and marketing (87) medians indicates the upward bias is modest for the central mass and that the skew and heavy-tail conclusions are, if anything, conservative. For calibrated benchmarks the appropriate instruments remain the platform-representative percentile bands recommended in Section 10 and the preregistered panel of Section 8.6; this cross-section corroborates that the structure they would measure is the one this article describes. Three of the prediction bands that Study 1 fixes in advance (Section 8.6) can be read against this sample, with its upward bias in mind. The sample median of 60 downloads sits inside the predicted 30 to 80. The attention Gini of 0.75 sits on the lower edge of the predicted 0.75 to 0.90, which is the direction the sampling bias predicts, since the censored low tail compresses measured

inequality. The median-to-mean ratio of 0.27 is below the predicted ceiling of 0.5. None of this validates the prospective design, whose value is that it observes papers from posting on a probability sample; it does record that the bands were not set implausibly.

Declarations

Data availability. Every synthesized anchor reproduced here traces to the published sources cited in the text and tables, with sampling frames and verification status stated at the point of use, and no new population-representative dataset underlies the article's benchmark claims. The one original dataset generated for this article is the independent archival cross-section of Appendix E, a random all-field sample of 1,360 papers whose pre-retirement counters were recovered from public archives, released as `appendixE-wayback-downloads.csv` (an Appendix E dataset, not the prospective Study 1) with the analysis code, so the procedure can be re-run and extended; it corroborates the distribution's shape and is not a platform-representative benchmark. The dataset is deposited at <https://doi.org/10.5281/zenodo.22285871> under CC BY 4.0, where the same file carries the name `study1-wayback-downloads.csv`, and it ships with an open tool that reads a download count against it while conditioning on both paper age and field, at <https://github.com/vadimchernets/ssrn-benchmark> (MIT licence, Python standard library only). Readers who want to place their own counter in this reference class can run that tool rather than reconstruct the analysis. Materials for the preregistration-ready agenda of Section 8, including the panel variable dictionary (Appendix C, Table C1) and the search-audit query inventory, are prepared for deposit at preregistration (Section 8.8); registry contributions follow the schema subset defined in Section 8.8, and any percentile release from the benchmark program is published under the transparency specification of Appendix D. The CSAB reporting template of Appendix B may be reused freely, subject to its blank-percentile rule. On 24 August 2026 a formal request was filed with SSRN Support (ticket 260824-003051) for the aggregate reference class the retired Rankings implied — median, interquartile range, and decile bands of downloads stratified by field and paper age, with no per-paper identifiers; its outcome (grant, refusal, or non-response) will be reported in a subsequent revision.

Open access to this article. An openly licensed copy of this manuscript is deposited at <https://doi.org/10.5281/zenodo.22166896> (CC BY 4.0), where the full text is retrievable directly and without an access challenge. The deposit is provided because the canonical record on SSRN is served behind bot-mitigation that declines clients without a JavaScript engine, a condition documented in Section 4.2; readers, reference managers and automated clients that cannot reach the canonical copy can obtain the same text from the deposit. The companion papers are deposited on the same terms at <https://doi.org/10.5281/zenodo.22166995> and <https://doi.org/10.5281/zenodo.22168724>.

Code availability. The harvesting and analysis code for the independent archival cross-section of Appendix E is released alongside its dataset at <https://github.com/vadimchernets/ssrn-benchmark>, so the procedure can be reproduced and the reference class re-derived from the raw counters. The article otherwise reports no results requiring analysis code; figure-generation templates for the empirical designs of Section 8 are to be released as an online appendix accompanying the preregistration materials.

Ethics. The empirical designs of Section 8 rely on publicly displayed paper-level counters and bibliographic metadata; any collection of non-public dashboard information, distribution-status dates, or email messages requires informed consent and, where applicable, institutional review (see Section 9). The

two designs that constitute human-subjects research, Study 4 (the author survey) and Study 3 (instrumented posting with inbox logging), will be submitted for institutional review board or research-ethics-committee approval and conducted under participants' informed consent before any data collection begins. Message datasets are aggregated and de-identified before any reporting, and community posts are quoted only where a live, verifiable source exists, with paraphrase where authors could be identified against their expectations. Data collection is confined to means permitted by the platforms' terms of service: the designs use publicly displayed counters and bibliographic metadata and employ no scraping that would breach a platform's terms of use, and any design requiring bulk or non-public data proceeds through a formal data request and legal review rather than automated extraction. Quotation of semi-public forum content attends to contributor consent and re-identification risk, which is the reason such material is paraphrased rather than reproduced verbatim wherever an author could be identified against their expectations. The article labels no specific journal or publisher "predatory"; venue screening applies validated, transparent criteria, which keeps any classification defensible and avoids defamation.

Acknowledgements. The author thanks Stefanie Haustein for critical comments, sent by email in August 2026, that led to revisions of Section 4.2, and Vincent Larivière, whose question in correspondence — which single normalisation an author should insist on first, and his insistence on seeing data rather than argument — prompted the archival cross-section reported in Appendix E.

Competing interests. The author is the developer of the multi-model orchestration methods discussed in Section 8.10, which are open and reproducible, an interest readers should weigh. The author is a named applicant on pending patent applications, not granted patents, covering AI orchestration and multi-model validation systems of the kind discussed in that section. The orchestration-layer delivery design proposed in Section 8.10 is a conceptual specification on which none of the article's findings rests. Section 8.10 could be deleted in full without altering any diagnosis, result, or recommendation elsewhere in the article.

Tools and verification. AI assistants were used as instruments under the author's direction; the ideas, research and conclusions are the author's. All quantitative claims and citations were verified by the author against their primary sources or archived copies; the author bears sole responsibility for the content.

References

Sources cited in the text as [B-n] (platform pages, forum threads, blog posts, and other supporting materials) are listed separately below under "Additional sources." Access-dated platform pages are unstable and should be re-verified against archived copies.

- [1] AALL / ALL-SIS (2024, June 14). Citation Methods White Paper. https://www.aallnet.org/allsis/wp-content/uploads/sites/4/2024/07/ALL-SIS_CitationMethodsWhitePaper_6_14_24_final.pdf
- [2] Abdill, R. J., & Blekhman, R. (2019). Tracking the popularity and outcomes of all bioRxiv preprints. *eLife*, 8, e45133. DOI 10.7554/eLife.45133
- [3] Academia StackExchange (2016, May 28). "Importance of paper download statistics". <https://academia.stackexchange.com/questions/69421/>
- [4] Adobe / Adobe Express Blog (2025; survey fielded May 2025). "Using ChatGPT as a search engine". <https://www.adobe.com/express/learn/blog/chatgpt-as-a-search-engine>
- [5] Aggarwal, P., Murahari, V., Rajpurohit, T., Kalyan, A., Narasimhan, K., & Deshpande, A. (2024). GEO: Generative Engine Optimization. KDD '24. DOI 10.1145/3637528.3671900; arXiv:2311.09735
- [6] Aguinis, H., & Solarino, A. M. (2019). Transparency and replicability in qualitative research. *Strategic Management Journal*, 40(8), 1291–1315. DOI 10.1002/smj.3015
- [7] Algaba, A., Mazijn, C., Holst, V., Tori, F., Wenmackers, S., & Ginis, V. (2025). Large Language Models Reflect Human Citation Patterns with a Heightened Citation Bias. Findings of the Association for Computational Linguistics: NAACL 2025, 6844–6879. DOI 10.18653/v1/2025.findings-naacl.381; arXiv:2405.15739
- [8] American Society for Cell Biology and co-signing organizations and individuals (2012/2013). San Francisco Declaration on Research Assessment (DORA). Drafted 16 December 2012 at the ASCB Annual Meeting, San Francisco; released 2013. <https://sfdora.org/read/>
- [9] Ansari, S., Khan, M., Weinstein, D., Conant, E. F., Mirza-Aghazadeh-Attari, M., & Yousem, D. M. (2023). Pervasiveness of open journal invitations across radiology specialties. *Current Problems in Diagnostic Radiology*, 52(6), 534–539. DOI 10.1067/j.cpradiol.2023.04.002; PubMed 37150715
- [10] Artisan Growth Strategies; Lenny's Newsletter; First Page Sage — freemium-conversion benchmarks. <https://www.artisangrowthstrategies.com/blog/freemium-conversion-rate-benchmarks> ; <https://www.lennysnewsletter.com/p/what-is-a-good-free-to-paid-conversion>
- [11] arXiv (Cornell University). Help pages on submission formats. https://info.arxiv.org/help/submit_pdf.html ; https://info.arxiv.org/help/policies/format_requirements.html
- [12] Bainbridge, S. (2026). "The Social Science Research Network Has Jumped the Shark" (03.06.2026), <https://www.stephenbainbridge.com/p/the-social-science-research-network> ; "SSRN Responds to My Complaints With Changes and/or Clarifications" (04.06.2026), <https://www.stephenbainbridge.com/p/ssrn-responds-to-my-complaints-with>
- [13] Bauer-Wolf, J. (2022, November 16). Yale, Harvard law schools drop out of U.S. News rankings, saying they undermine legal profession's tenets. Higher Ed Dive.

<https://www.highereddive.com/news/yale-harvard-law-schools-drop-out-of-us-news-rankings-saying-they-under/636742/>

- [14] Beall, J. (2012). Predatory publishers are corrupting open access. *Nature*, 489, 179. DOI 10.1038/489179a
- [15] Bergstrom, C. T., & West, J. D. (2020). *Calling Bullshit: The Art of Skepticism in a Data-Driven World*. Random House.
- [16] Bessembinder, H. (2018). Do stocks outperform Treasury bills? *Journal of Financial Economics*, 129(3), 440–457. DOI 10.1016/j.jfineco.2018.06.004; SSRN 2900447
- [17] Björk, B.-C., & Solomon, D. (2012). Open access versus subscription journals: a comparison of scientific impact. *BMC Medicine*, 10:73. DOI 10.1186/1741-7015-10-73
- [18] Black, B. S., & Caron, P. L. (2006). Ranking Law Schools: Using SSRN to Measure Scholarly Performance. *Indiana Law Journal*, 81(1), 83 ff. SSRN abstract_id=784764; <https://www.repository.law.indiana.edu/ilj/vol81/iss1/7/> ; PDF: https://ilj.law.indiana.edu/articles/81/81_1_Black.pdf
- [19] Bollen, J., & Van de Sompel, H. (2008). Usage impact factor: The effects of sample characteristics on usage-based impact metrics. *Journal of the American Society for Information Science and Technology*, 59(1), 136–149. DOI 10.1002/asi.20746
- [20] Bornmann, L. (2014). Do altmetrics point to the broader impact of research? *Journal of Informetrics*, 8(4), 895–903. DOI 10.1016/j.joi.2014.09.005; arXiv:1406.7091
- [21] Bornmann, L., & Haunschild, R. (2016). How to normalize Twitter counts? A first attempt based on journals in the Twitter Index. *Scientometrics*, 107(3), 1405–1422. DOI 10.1007/s11192-016-1893-6
- [22] Bornmann, L., Leydesdorff, L., & Mutz, R. (2013). The use of percentiles and percentile rank classes in the analysis of bibliometric data: Opportunities and limits. *Journal of Informetrics*, 7(1), 158–165. DOI 10.1016/j.joi.2012.10.001; arXiv:1211.0381
- [23] Bornmann, L., & Haunschild, R. (2018). Normalization of zero-inflated data: An empirical analysis of a new indicator family and its use with altmetrics data. *Journal of Informetrics*, 12(3), 998–1011. DOI 10.1016/j.joi.2018.01.010
- [24] Bourne, P. E., Polka, J. K., Vale, R. D., & Kiley, R. (2017). Ten simple rules to consider regarding preprint submission. *PLOS Computational Biology*, 13(5), e1005473. DOI 10.1371/journal.pcbi.1005473
- [25] Brody, T., Harnad, S., & Carr, L. (2006). Earlier Web usage statistics as predictors of later citation impact. *JASIST*, 57(8), 1060–1072. DOI 10.1002/asi.20373; arXiv cs/0503020
- [26] Brown, C. (2003). The role of electronic preprints in chemical communication: Analysis of citation, usage, and acceptance in the journal literature. *JASIST*, 54(5), 362–371. DOI 10.1002/asi.10223
- [27] Campbell, D. T. (1979). Assessing the impact of planned social change. *Evaluation and Program Planning*, 2(1), 67–90. DOI 10.1016/0149-7189(79)90048-X
- [28] Carrell, S., Figlio, D., & Lusher, L. (2023). Congestion on the information superhighway: Inefficiencies in economics working papers. *Journal of Public Economics*, 225, 104978 (also NBER WP 29153)

- [29] Casden, J., Romani, D., Shearer, T., & Campbell, J. (2025). Mitigating Aggressive Crawler Traffic in the Age of Generative AI: A Collaborative Approach from the University of North Carolina at Chapel Hill Libraries. *The Code4Lib Journal*, 61. <https://journal.code4lib.org/articles/18489>
- [30] Chen, Y., & Konstan, J. A. (2015). Online field experiments. *Journal of the Economic Science Association*, 1(1), 29–42. DOI 10.1007/s40881-015-0005-3
- [31] Chu, H., & Krichel, T. (2007). Downloads vs. citations in economics: relationships, contributing factors and beyond. In *Proceedings of ISSI 2007 (11th International Conference of the International Society for Scientometrics and Informetrics)*, Madrid.
- [32] Cobey, K. D., et al. (2019). Knowledge and motivations of researchers publishing in presumed predatory journals. *BMJ Open*, 9(3), e026516. DOI 10.1136/bmjopen-2018-026516
- [33] Cooley Law School Blog (2021, December 7). "Showcase Your Scholarship, Part Three: Self-Publishing". <https://cooley.edu/blog/showcase-your-scholarship-part-three-self-publishing>
- [34] CRIV Connection / CRIV Blog (2026, April 20). "Social Science Research Network (SSRN) Ending Commercial Products". <https://crivblog.com/2026/04/20/social-science-research-network-ssrn-ending-commercial-products/>
- [35] Crossref. REST API documentation. <https://www.crossref.org/documentation/retrieve-metadata/rest-api/>
- [36] Crossref. Similarity Check documentation. <https://www.crossref.org/documentation/similarity-check/similarity-report-understand/>
- [37] Davis, P. M., et al. (2008). Open access publishing, article downloads, and citations: randomised controlled trial. *BMJ*, 337, a568. DOI 10.1136/bmj.a568
- [38] Davis, P. M., & Fromerth, M. J. (2007). Does the arXiv lead to higher citations and reduced publisher downloads for mathematics articles? *Scientometrics*, 71(2), 203–215. DOI 10.1007/s11192-007-1661-8
- [39] Ding, Y., Dong, X., Bu, Y., Zhang, B., Lin, K., & Hu, B. (2021). Revisiting the relationship between downloads and citations: a perspective from papers with different citation patterns in the case of the Lancet. *Scientometrics*, 126(9), 7609–7621. DOI 10.1007/s11192-021-04099-3
- [40] Donovan, J. M., Watson, C. A., & Osborne, C. (2014). The Open Access Advantage for American Law Reviews. SSRN abstract_id=2506913. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=2506913
- [41] Dorta-González, P., & Dorta-González, M. I. (2023). The funding effect on citation and social attention: the UN Sustainable Development Goals (SDGs) as a case study. *Online Information Review*, 47(7), 1358–1376. DOI 10.1108/OIR-05-2022-0300; arXiv:2304.00862
- [42] Douma, M. J. (2021, January 12). "How many reads does a typical academic article get?" <https://michaeljdouma.com/2021/01/12/how-many-reads-does-a-typical-academic-article-get/>
- [43] Du, Y., Li, S., Torralba, A., Tenenbaum, J. B., & Mordatch, I. (2023). Improving Factuality and Reasoning in Language Models through Multiagent Debate. arXiv:2305.14325. <https://arxiv.org/abs/2305.14325>

- [44] Edelman, B., & Larkin, I. (2015). Social Comparisons and Deception Across Workplace Hierarchies: Field and Experimental Evidence. *Organization Science*, 26(1), 78–98. DOI 10.1287/orsc.2014.0938; SSRN abstract_id=1346397; HBS WP 09-096 (circulated since 2009)
- [45] Eisenberg, T. (2006). Assessing the SSRN-Based Law School Rankings. *Indiana Law Journal*, 81(1), 285–291. <https://www.repository.law.indiana.edu/ilj/vol81/iss1/13/>
- [46] Elsevier (2016, May). "Elsevier Acquires the Social Science Research Network (SSRN)...". PR Newswire. <https://www.prnewswire.com/news-releases/elsevier-acquires-the-social-science-research-network-ssrn-the-leading-social-science-and-humanities-repository-and-online-community-579752311.html> (also LSE Impact Blog, 18 May 2016)
- [47] Elsevier Digital Commons. "SSRN Research Paper Series Sunset — DC Options". https://digitalcommons.elsevier.com/en_US/integration-preservation/ssrn-research-paper-series-sunset-dc-options
- [48] Elsevier. "SSRN First Look" (product page). <https://www.elsevier.com/products/ssrn-preprint-services/first-look>
- [49] Elsevier. "SSRN for Authors and Researchers" (product page). <https://www.elsevier.com/products/ssrn-preprint-services/author-researchers>
- [50] Elsevier/SSRN Support. "How can researchers gauge a paper's impact?" <https://www.elsevier.support/ssrn/answer/gauge-impact> (accessed 14.08.2026)
- [51] Elsevier/SSRN Support. "What are RSS and Data Feeds on SSRN". <https://www.elsevier.support/ssrn/answer/what-are-rss-and-data-feeds-on-ssrn>
- [52] Elsevier/SSRN Support. "What is SSRN?" <https://www.elsevier.support/ssrn/answer/what-is-ssrn> (accessed 14.08.2026)
- [53] Elsevier/SSRN Support. Classification and distribution. <https://www.elsevier.support/ssrn/answer/classification-distribution> (accessed 14.08.2026)
- [54] Elsevier/SSRN Support. Generative AI policy. <https://www.elsevier.support/ssrn/answer/AI> (accessed 14.08.2026)
- [55] Elsevier/SSRN Support. Google / Google Scholar indexing. <https://www.elsevier.support/ssrn/answer/google-scholar> (accessed 14.08.2026)
- [56] Elsevier/SSRN Support. Submission guidelines / Get started. <https://www.elsevier.support/ssrn/answer/get-started> (accessed 14.08.2026)
- [57] Elsevier/SSRN Support. Subscriptions (eJournal mailings). <https://www.elsevier.support/ssrn/answer/subscriptions> ; also <https://www.ssrn.com/index.cfm/en/subscribe/>
- [58] Espeland, W. N., & Sauder, M. (2007). Rankings and reactivity: How public measures recreate social worlds. *American Journal of Sociology*, 113(1), 1–40. DOI 10.1086/517897
- [59] Eysenbach, G. (2006). Citation Advantage of Open Access Articles. *PLoS Biology*, 4(5), e157. DOI 10.1371/journal.pbio.0040157

- [60] Farys, R., & Wolbring, T. (2021). Matthew effects in science and the serial diffusion of ideas. *Quantitative Science Studies*, 2(2), 505–526. DOI 10.1162/qss_a_00129
- [61] Fernández-Ramos, A., Comas-Forgas, R., & Rodríguez-Bravo, B. (2026). Identifying academic spam and recommendations for dealing with it: a literature review. *European Science Editing*, 52, e174423. DOI 10.3897/ese.2026.e174423 (published 8 May 2026)
- [62] Fraser, N., et al. (2020). The relationship between bioRxiv preprints, citations and altmetrics. *Quantitative Science Studies*, 1(2), 618–638. DOI 10.1162/qss_a_00043
- [63] Fung, A., Razi, B., Basto, C., Bariol, S., Hossack, T., Baskaranathan, S., Ende, D., & Woo, H. (2026). Trust, but Verify: An Exploratory Audit of Predatory Journal and Publisher Solicitations. *ANZ Journal of Surgery*. DOI 10.1111/ans.70568
- [64] Giannakakos, V., et al. (2025). Impact of author characteristics on outcomes of single- versus double-blind peer review: a systematic review of comparative studies in scientific abstracts and publications. *Scientometrics*, 130, 399–421. DOI 10.1007/s11192-024-05213-x
- [65] Ginsparg, P. (2011). ArXiv at 20. *Nature*, 476(7359), 145–147. <https://www.nature.com/articles/476145a>
- [66] Goldman, E. "SSRN and Download Statistics" (2005), https://personal.ericgoldman.org/ssrn_and_downlo/ ; "Update on SSRN", https://personal.ericgoldman.org/update_on_ssrn/
- [67] Goodhart, C. A. E. (1975). Problems of monetary management: The U.K. experience. In *Papers in Monetary Economics*, Vol. I. Sydney: Reserve Bank of Australia. Reprinted in C. A. E. Goodhart (1984), *Monetary Theory and Practice: The U.K. Experience* (ch. 4). London: Macmillan. DOI 10.1007/978-1-349-17295-5_4
- [68] Gordon, G. J. (2016). The Digital Lawyer's Evolving Education in Scholarly Research. SSRN abstract_id=2754647
- [69] Gordon, G., Lin, J., Cave, R., & Dandrea, R. (2015). The Question of Data Integrity in Article-Level Metrics. *PLOS Biology*. <https://pmc.ncbi.nlm.nih.gov/articles/PMC4546647/>
- [70] Grossman, R., Liu, S., Chen, M. K., Smith, M., Borcea, C., & Chen, Y. (2026). How Generative AI Disrupts Search: An Empirical Study of Google Search, Gemini, and AI Overviews. arXiv:2604.27790. <https://arxiv.org/abs/2604.27790>
- [71] Guerrero-Bote, V. P., & Moya-Anegón, F. (2014). Relationship between downloads and citations at journal and paper levels, and the influence of language. *Scientometrics*. DOI 10.1007/s11192-014-1243-5
- [72] Haque, A., & Ginsparg, P. (2009). Positional effects on citation and readership in arXiv. *JASIST*, 60(11), 2203–2218. DOI 10.1002/asi.21166
- [73] Haque, A., & Ginsparg, P. (2010). Last but not least: Additional positional effects on citation and readership in arXiv. *JASIST*, 61(12), 2381–2388. DOI 10.1002/asi.21428; arXiv:1010.2757
- [74] Haustein, S., Costas, R., & Larivière, V. (2015). Characterizing social media metrics of scholarly papers. *PLOS ONE*, 10(3), e0120495. DOI 10.1371/journal.pone.0120495

- [75] Haustein, S., Bowman, T. D., Holmberg, K., Tsou, A., Sugimoto, C. R., & Larivière, V. (2016). Tweets as impact indicators: Examining the implications of automated "bot" accounts on Twitter. *JASIST*, 67(1), 232–238. DOI 10.1002/asi.23456
- [76] Haustein, S., Bowman, T. D., & Costas, R. (2016). Interpreting "altmetrics": viewing acts on social media through the lens of citation and social theories. In: *Theories of Informetrics and Scholarly Communication* (De Gruyter Mouton), 372–405; arXiv:1502.05701
- [77] Heald, P. J., & Sichelman, T. (2019). Ranking the Academic Impact of 100 American Law Schools. *Jurimetrics*, 60(1), 1–39 (Fall 2019). SSRN abstract_id=3483325
- [78] Hicks, D., Wouters, P., Waltman, L., de Rijcke, S., & Rafols, I. (2015). Bibliometrics: The Leiden Manifesto for research metrics. *Nature*, 520(7548), 429–431. DOI 10.1038/520429a. <https://www.nature.com/articles/520429a>
- [79] Hirsch, J. E. (2005). An index to quantify an individual's scientific research output. *PNAS*, 102(46), 16569–16572.
- [80] Hu, B., et al. (2021). On the relationship between download and citation counts: An introduction of Granger-causality inference. *Journal of Informetrics*, 15(2), 101125. DOI 10.1016/j.joi.2020.101125
- [81] iThenticate/Turnitin. Guides: "Understanding the Similarity Report"; "Preprints". <https://guides.ithenticate.com/hc/en-us/articles/27842305774605> ; <https://guides.ithenticate.com/hc/en-us/articles/27848481223693-Preprints>
- [82] Jaźwińska, K., & Chandrasekar, A. (2025, March 6). AI search has a citation problem. *Columbia Journalism Review* / Tow Center. https://www.cjr.org/tow_center/we-compared-eight-ai-search-engines-theyre-all-bad-at-citing-news.php
- [83] Jo, W., & Li, J. (2026). Can AI reduce the reputation premium in knowledge diffusion? Evidence from SSRN marketing working papers. *Marketing Letters*, 37, Article 32. DOI 10.1007/s11002-026-09827-4
- [84] Justin, G. A., et al. (2024). An Analysis of Solicitations From Predatory Journals in Ophthalmology. *American Journal of Ophthalmology*. PMC11257792; [https://www.ajo.com/article/S0002-9394\(24\)00087-4/abstract](https://www.ajo.com/article/S0002-9394(24)00087-4/abstract) — source of the "1813 emails" figure
- [85] Kakushadze, Z. (2016). An index for SSRN downloads. *Journal of Informetrics*, 10(1), 9–28. DOI 10.1016/j.joi.2015.11.005; arXiv:1511.04275; SSRN 2656600
- [86] Kim, Heekyung Hellen (2014). The Effect of Open Access (AEA Conference 2014). <https://www.aeaweb.org/conference/2014/retrieve.php?pdfid=58>
- [87] Kim, T., Bock, K., Luo, C., Liswood, A., Poroslay, C., & Wenger, E. (2025). Scrapers selectively respect robots.txt directives: evidence from a large-scale empirical study. *Proceedings of the ACM Internet Measurement Conference (IMC '25)*. DOI 10.1145/3730567.3764471
- [88] Kleinberg, J., Ludwig, J., Mullainathan, S., & Obermeyer, Z. (2015). Prediction policy problems. *American Economic Review*, 105(5), 491–495. DOI 10.1257/aer.p20151023
- [89] Kozak, M., Iefremova, O., & Hartley, J. (2016). Spamming in Scholarly Publishing: A Case Study. *JASIST*, 67(8), 2009–2015. DOI 10.1002/asi.23521

- [90] Kraker, P., & Lex, E. (2015). A critical look at the ResearchGate Score as a measure of scientific reputation. Presented at ASCW'15 (Quantifying and Analysing Scholarly Communication on the Web), Web Science Conference 2015, Oxford. <http://ascw.know-center.tugraz.at/2015/05/26/kraker-lex-a-critical-look-at-the-researchgate-score/> (popular version: "The ResearchGate Score: a good example of a bad metric," LSE Impact Blog, 9 Dec 2015)
- [91] Kurtz, M. J., et al. (2005). The bibliometric properties of article readership information. *JASIST*, 56(2), 111–128. DOI 10.1002/asi.20096
- [92] Kurtz, M. J., & Bollen, J. (2010). Usage bibliometrics. *Annual Review of Information Science and Technology*, 44(1), 3–64. DOI 10.1002/aris.2010.1440440108
- [93] Langham-Putrow, A., et al. (2021). Is the open access citation advantage real? A systematic review of the citation of open access and subscription-based articles. *PLOS ONE*, 16(6), e0253129. DOI 10.1371/journal.pone.0253129
- [94] Larivière, V., Haustein, S., & Mongeon, P. (2015). The oligopoly of academic publishers in the digital era. *PLOS ONE*, 10(6), e0127502. DOI 10.1371/journal.pone.0127502
- [95] Legal Writing Institute. "Publishing Tips" (PDF). <https://www.lwionline.org/sites/default/files/Publishing-Tips.pdf>
- [96] Lewinski, A. A., & Oermann, M. H. (2018). Characteristics of e-mail solicitations from predatory nursing journals and publishers. *The Journal of Continuing Education in Nursing*, 49(4), 171–177. DOI 10.3928/00220124-20180320-07
- [97] Liang, W., Yuksekogonul, M., Mao, Y., Wu, E., & Zou, J. (2023). GPT detectors are biased against non-native English writers. *Patterns*, 4(7), 100779. DOI 10.1016/j.patter.2023.100779
- [98] Lloyd, T. (2026, May 7). Guest Post — Love, Death & Robots: Scholarly Edition. The Scholarly Kitchen. <https://scholarlykitchen.sspnet.org/2026/05/07/guest-post-love-death-robots-scholarly-edition/>
- [99] LogEc / RePEc. About LogEc (traffic-filtering methodology). <https://logec.repec.org/about.htm> ; <https://logec.repec.org/>
- [100] Martinez, O. (2026). Optimizing Visibility in Generative Engines: A Critical Survey of Generative Engine Optimization (2023–2026). arXiv:2607.14035 (submitted 15 July 2026). <https://arxiv.org/abs/2607.14035>
- [101] Maynooth University News (2026, January). "Maynooth University's School of Law and Criminology ranks first...". <https://www.maynoothuniversity.ie/news-events/maynooth-university-school-law-and-criminology-ranks-first-social-science-research-networks-2>
- [102] McGillivray, B., & Astell, M. (2019). The relationship between usage and citations in an open access mega-journal. *Scientometrics*, 121, 817–838. DOI 10.1007/s11192-019-03228-3
- [103] Medoff, M. H. (2006). Evidence of a Harvard and Chicago Matthew Effect. *Journal of Economic Methodology*, 13(4), 485–506. DOI 10.1080/13501780601049079
- [104] Merton, R. K. (1968). The Matthew effect in science. *Science*, 159(3810), 56–63. DOI 10.1126/science.159.3810.56

- [105] Merton, R. K. (1988). The Matthew effect in science, II: Cumulative advantage and the symbolism of intellectual property. *Isis*, 79(4), 606–623. DOI 10.1086/354848
- [106] Microsoft Learn. "Safe Links in Microsoft Defender for Office 365". <https://learn.microsoft.com/en-us/defender-office-365/safe-links-about> (+ Microsoft Q&A on clicks in UrlClickEvents; + Suped on false opens)
- [107] Milanovic, M. EJIL:Talk! "Academic Spam [UPDATED]". <https://www.ejiltalk.org/academic-spam/>
- [108] Moed, H. F., & Halevi, G. (2016). On full text download and citation distributions in scientific-scholarly journals. *JASIST*, 67(2), 412–431. DOI 10.1002/asi.23405
- [109] Moed, H. F. (2017). *Applied Evaluative Informetrics*. Springer. DOI 10.1007/978-3-319-60522-7
- [110] Moher, D., et al. (2018). Assessing scientists for hiring, promotion, and tenure. *PLoS Biology*, 16(3), e2004089. DOI 10.1371/journal.pbio.2004089
- [111] Morgan, A. C., et al. (2018). Prestige drives epistemic inequality in the diffusion of scientific ideas. *EPJ Data Science*, 7:40. DOI 10.1140/epjds/s13688-018-0166-4
- [112] Muchnik, L., Aral, S., & Taylor, S. J. (2013). Social influence bias: A randomized experiment. *Science*, 341(6146), 647–651. DOI 10.1126/science.1240466
- [113] Muller, J. Z. (2018). *The Tyranny of Metrics*. Princeton, NJ: Princeton University Press. ISBN 978-0-691-17495-2.
- [114] *Nature Neuroscience*, Editorial (2008). <https://www.nature.com/articles/nn0608-619>
- [115] Northwestern University LibGuides (research assessment/metrics). <https://libguides.northwestern.edu/c.php?g=1372178&p=10205721>
- [116] Ohm, P. (2007). Do Blogs Influence SSRN Downloads? Empirically Testing the Volokh and Slashdot Effects. U. Colorado Law Legal Studies Research Paper No. 07-15. <https://ssrn.com/abstract=980484>
- [117] Pace Law Library. "Get Published — Guide for Law Review Members". <https://libraryguides.law.pace.edu/lawreviews/publishing>
- [118] Paperpile. "What is a good h-index?" <https://paperpile.com/g/what-is-a-good-h-index/> (similar guides: Elsevier, Researcher.Life, Academia Insider)
- [119] Perneger, T. V. (2004). Relation between online "hit counts" and subsequent citations: prospective study of research papers in the BMJ. *BMJ*, 329(7465), 546–547. DOI 10.1136/bmj.329.7465.546
- [120] Petersen, A. M., et al. (2014). Reputation and impact in academic careers. *PNAS*, 111(43), 15316–15321. DOI 10.1073/pnas.1323111111
- [121] Pew Research Center (2026, June 17). Americans and AI 2026: Chatbots, Smart Devices and AI's Impact. <https://www.pewresearch.org/internet/2026/06/17/americans-and-ai-2026-chatbots-smart-devices-and-views-on-impact/>
- [122] Piwowar, H. (2013). Altmetrics: Value all research products. *Nature*, 493(7431), 159. <https://www.nature.com/articles/493159a>

- [123] Priem, J., Taraborelli, D., Groth, P., & Neylon, C. (2010). Altmetrics: A manifesto. <http://altmetrics.org/manifesto/> (archived copy at web.archive.org)
- [124] Priem, J., Piwowar, H., & Orr, R. (2022). OpenAlex: A fully-open index of scholarly works, authors, venues, institutions, and concepts. arXiv:2205.01833. DOI 10.48550/arXiv.2205.01833
- [125] Project COUNTER (2026). COUNTER Code of Practice, Release 5.1.1. Published 30 June 2026 (Release 5.1 published 5 May 2023). <https://www.countermetrics.org/> ; “Best Practice on AI Usage” (guidance for R5.1.1, updated 30 June 2026), <https://www.countermetrics.org/code-of-practice/best-practice/bp-ai/>
- [126] Quora. "To what extent would you recommend or discourage publishing an academic working paper on SSRN...". <https://www.quora.com/To-what-extent-would-you-recommend-or-discourage-publishing-an-academic-working-paper-on-SSRN-to-claim-and-put-a-stake-in-the-ground-e-g-before-any-journal-submission>
- [127] Reddit r/AskAcademia (2023, August 10). Thread 15newc2 "Endless spam since I published a paper a few months ago". <https://www.reddit.com/r/AskAcademia/comments/15newc2/> (confirmed via the pullpush.io archive)
- [128] Remler, D. (2014, April 23). Are 90% of academic papers really never cited? Reviewing the literature on academic citations. LSE Impact Blog. <https://blogs.lse.ac.uk/impactofsocialsciences/2014/04/23/academic-papers-citation-rates-remler/>
- [129] RePEc Blog: "How RePEc counts views and downloads" (06.11.2021), <https://blog.repec.org/2021/11/06/how-repec-counts-views-and-downloads/> ; "AI issues in RePEc" (10.11.2025), <https://blog.repec.org/2025/11/10/ai-issues-in-repec/> ; "RePEc in June 2026" (13.07.2026), <https://blog.repec.org/2026/07/13/repec-in-june-2026/>
- [130] ResearchGate (2022, August). "Removing the RG Score". <https://www.researchgate.net/researchgate-updates/removing-the-rg-score>
- [131] Salganik, M. J., Dodds, P. S., & Watts, D. J. (2006). Experimental study of inequality and unpredictability in an artificial cultural market. *Science*, 311, 854–856. DOI 10.1126/science.1121066
- [132] Sauder, M., & Espeland, W. N. (2009). The discipline of rankings: Tight coupling and organizational change. *American Sociological Review*, 74(1), 63–82. DOI 10.1177/000312240907400104
- [133] Schein, A. I., Popescul, A., Ungar, L. H., & Pennock, D. M. (2002). Methods and metrics for cold-start recommendations. In *Proceedings of the 25th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR '02)*, 253–260. DOI 10.1145/564376.564421
- [134] Schumann, E. SSRN R package. <https://github.com/enricoschumann/SSRN> (also scrapers github.com/talsan/ssrn, github.com/karthiktadepalli1/ssrn-scraper)
- [135] Shearer, K., & Walk, P. (2025, June 3). The impact of AI bots and crawlers on open repositories: Results of a COAR survey, April 2025. Confederation of Open Access Repositories. <https://coar-repositories.org/wp-content/uploads/2025/06/Report-of-the-COAR-Survey-on-AI-Bots-June-2025-1.pdf>
- [136] Shen, C., & Björk, B.-C. (2015). 'Predatory' open access: a longitudinal study of article volumes and market characteristics. *BMC Medicine*, 13, 230. DOI 10.1186/s12916-015-0469-2

- [137] Shepherd, P. T. (2007). The feasibility of developing and implementing journal usage factors: a research project sponsored by UKSG. *Serials: The Journal for the Serials Community*, 20(2), 117–123. DOI 10.1629/20117
- [138] Shuai, X., Pepe, A., & Bollen, J. (2012). How the scientific community reacts to newly submitted preprints. *PLOS ONE*, 7(11), e47523. DOI 10.1371/journal.pone.0047523
- [139] Siems, M. M. (2016). Legal Research in Search of Attention: A Quantitative Assessment. *King's Law Journal*, 27(2), 170–187. DOI 10.1080/09615768.2015.1105560. Working-paper version: "What Determines SSRN Downloads", SSRN abstract_id=2836950; Durham: <https://durham-repository.worktribe.com/OutputFile/1387250>
- [140] Simcoe, T. S., & Waguespack, D. M. (2011). Status, Quality, and Attention: What's in a (Missing) Name? *Management Science*, 57(2), 274–290. DOI 10.1287/mnsc.1100.1270
- [141] Simon, H. A. (1971). Designing organizations for an information-rich world. In M. Greenberger (Ed.), *Computers, Communications, and the Public Interest* (pp. 37–72). Johns Hopkins Press.
- [142] Soler, J., & Cooper, A. (2019). Unexpected Emails to Submit Your Work: Spam or Legitimate Offers? *Publications*, 7(1), 7. DOI 10.3390/publications7010007
- [143] SSRN (2018, April 13). eLibrary Stats announcement. <https://www.ssrn.com/index.cfm/en/fen/ads/04132018ann014/>
- [144] SSRN Blog (2018, June 5). What SSRN Means When We Talk About Data Integrity. <https://blog.ssrn.com/2018/06/05/what-ssrn-means-when-we-talk-about-data-integrity/>
- [145] SSRN Blog (2025, December 16). 2025 Year in Review: How SSRN Enhanced Access to Preprints and Early-Stage Research. <https://blog.ssrn.com/2025/12/16/2025-year-in-review-how-ssrn-enhanced-access-to-preprints-and-early-stage-research/>
- [146] SSRN Blog (2026, April 13). SSRN Strategic Update: Renewed Focus on Core Research Sharing Mission (post author: S. Decker-Lucke). <https://blog.ssrn.com/2026/04/13/ssrn-strategic-update-renewed-focus-on-core-research-sharing-mission/>
- [147] SSRN Blog (2026, June 4). SSRN's Ongoing Commitment to Legal Scholarship. <https://blog.ssrn.com/2026/06/04/ssrns-ongoing-commitment-to-legal-scholarship/>
- [148] SSRN Blog (2026, June 12). SSRN to Sunset Rankings from July 1st 2026. <https://blog.ssrn.com/2026/06/12/ssrn-to-sunset-rankings-from-july-1st-2026/>
- [149] SSRN Blog (2026, July 15). SSRN Ranking has sunsetted as of 15 July 2026. <https://blog.ssrn.com/2026/07/15/ssrn-to-sunset-rankings-from-15th-july-2026/>
- [150] SSRN Blog (2026, July 20). SSRN will now allow you to choose the licence that's right for your work. <https://blog.ssrn.com/2026/07/20/ssrn-will-now-allow-you-to-choose-the-licence-thats-right-for-your-work/>
- [151] SSRN. Historical FAQ (PDF copy, Scuola Superiore Sant'Anna). https://www.sssup.it/UploadDocs/4461_SSRN_FAQ.pdf
- [152] SSRN. Subscription page. <https://www.ssrn.com/index.cfm/en/subscribe/>

- [153] Sugimoto, C. R., & Larivière, V. (2018). *Measuring Research: What Everyone Needs to Know*. Oxford University Press. ISBN 9780190640118
- [154] Sun, M., Barry Danfa, J., & Teplitskiy, M. (2022). Does double-blind peer review reduce bias? Evidence from a top computer science conference. *JASIST*, 73(6), 811–819. DOI 10.1002/asi.24582 (second author's surname per Crossref: "Barry Danfa")
- [155] Sureda-Negre, J., Calvo-Sastre, A., & Comas-Forgas, R. (2022). Predatory journals and publishers: Characteristics and impact of academic spam to researchers in educational sciences. *Learned Publishing*, 35(4), 441–447. DOI 10.1002/leap.1450
- [156] TaxProf Blog (Caron, P.): "SSRN e-Journal Rankings By Average Downloads Per Paper" (2018, May), https://taxprof.typepad.com/taxprof_blog/2018/05/ssrn-e-journal-rankings-by-average-downloads-per-paper.html ; "SSRN Has Not Jumped the Shark" (07.06.2026), <https://taxprofblog.aals.org/2026/06/07/>
- [157] Teixeira da Silva, J. A., Tsigaris, P., & Moussa, S. (2023). Can AI detect predatory journals? The case of FT50 journals. SSRN abstract_id=4391108. DOI 10.2139/ssrn.4391108
- [158] Thales / Imperva (2025). 2025 Imperva Bad Bot Report. <https://cpl.thalesgroup.com/about-us/newsroom/2025-imperva-bad-bot-report-ai-internet-traffic>
- [159] Think. Check. Submit. (2022, February 4). "Is a journal interested in your preprint? Separating the wheat from the chaff in journal solicitations to submit". <https://thinkchecksubmit.org/2022/02/04/is-a-journal-interested-in-your-preprint-separating-the-wheat-from-the-chaff-in-journal-solicitations-to-submit/>
- [160] Tomlinson, O. W. (2023). Analysis of predatory emails in early career academia and attempts at prevention. *Learned Publishing*, 36(2), 156–163. DOI 10.1002/leap.1500
- [161] University of Washington Gallagher Law Library. LibGuide "SSRN Metrics" (updated 29 Sep 2025). <https://lib.law.uw.edu/c.php?g=1238218&p=9061272> (+ page p=9061269 on eJournals)
- [162] van de Rijt, A., et al. (2014). Field experiments of success-breeds-success dynamics. *PNAS*, 111(19), 6934–6939. DOI 10.1073/pnas.1316836111
- [163] Vanderbilt University, Brightspace (Center for Teaching / Office of Innovative Technologies) (2023). Guidance on AI Detection and Why We're Disabling Turnitin's AI Detector. Published 16 August 2023. <https://www.vanderbilt.edu/brightspace/2023/08/16/guidance-on-ai-detection-and-why-were-disabling-turnitins-ai-detector/> (accessed 15 August 2026).
- [164] Vishwakarma, R., Kumar, A., & Jamidar, S. (2026). What gets cited: Competitive GEO in AI answer engines. *SIGIR '26*, 4950–4954. DOI 10.1145/3805712.3808445; arXiv:2605.25517
- [165] Wallace, K. L., Hartley, M., & Lepird, C. (2023). Expanding the Drake Journal of Agricultural Law: The Green Issue. *Drake Journal of Agricultural Law*, 28, 114–127. <https://aglawjournal.wp.drake.edu/wp-content/uploads/sites/66/2023/07/Wallace-Ready-for-PRODUCTION.pdf>
- [166] Wang, J., Veugelers, R., & Stephan, P. (2017). Bias against novelty in science. *Research Policy*, 46(8), 1416–1436. DOI 10.1016/j.respol.2017.06.006

- [167] Weinberg, M. (2025, June). Are AI Bots Knocking Cultural Heritage Offline? GLAM-E Lab (NYU Engelberg Center / University of Exeter).
https://www.glamelab.org/files/Are_AI_Bots_Knocking_Cultural_Heritage_Offline.pdf
- [168] Whalen, R. (2018, May 3). SSRN e-Journals and downloads.
<https://ryanwhalen.com/2018/05/03/ssrn-e-journals-and-downloads/>
- [169] Wherry, J. L. (2022). Tools for Improving Distribution, Discussion, and Downloads... SSRN abstract_id=4191909. DOI 10.2139/ssrn.4191909 (published in U. Detroit Mercy Law Review)
- [170] Wikipedia: "Social Science Research Network"; "List of academic journals by preprint policy"; "Preprint". https://en.wikipedia.org/wiki/Social_Science_Research_Network and related pages.
- [171] Willey, R., & Knapp, M. (2021). How to Increase Citations to Legal Scholarship. Ohio State Technology Law Journal, 18. SSRN 3801190; PDF: <https://moritzlaw.osu.edu/sites/default/files/2022-01/HOW%20TO%20INCREASE%20CITATIONS%20TO%20LEGAL.pdf>
- [172] Willey, R., & Knapp, M. (2022/2023). SSRN's Impact on Citations to Legal Scholarship and How to Maximize It. University of Arkansas at Little Rock Law Review, 45(3), 475–507. SSRN abstract_id=4037744; https://research.ualr.edu/bowen_lawreview/vol45/iss3/2/
- [173] Wilsdon, J., et al. (2015). The Metric Tide: Report of the Independent Review of the Role of Metrics in Research Assessment and Management. HEFCE. DOI 10.13140/RG.2.1.4929.1363.
<https://doi.org/10.13140/RG.2.1.4929.1363>
- [174] Wilson, P. (2022). Unsolicited solicitations: identifying characteristics of unsolicited emails from potentially predatory journals and the role of librarians. Journal of the Medical Library Association, 110(4), 520–524. DOI 10.5195/jmla.2022.1554; PMC10124592.
<https://pmc.ncbi.nlm.nih.gov/articles/PMC10124592/>
- [175] WisBlawg, UW–Madison Law Library (2016, September 21). "Which Legal Articles Generate the Most Attention on SSRN". <https://wisblawg.law.wisc.edu/2016/09/21/which-legal-articles-generate-the-most-attention-on-ssrn>
- [176] Wood-Doughty, A., Bergstrom, T., & Steigerwald, D. G. (2019). Do Download Reports Reliably Measure Journal Usage? College & Research Libraries, 80(5).
<https://crl.acrl.org/index.php/crl/article/view/17824/19653>
- [177] Wouters, P., & Costas, R. (2012). Users, narcissism and control — tracking the impact of scholarly publications in the 21st century. SURF Foundation, Utrecht. <https://www.cni.org/news/surf-report-users-narcissism-and-control>
- [178] Yuret, T. (2026). Article accesses: links to citations, weekend effects and seasonality. Scientometrics (published 23.07.2026). DOI 10.1007/s11192-026-05710-1
- [179] Zahedi, Z., Costas, R., & Wouters, P. (2017). Mendeley readership as a filtering tool to identify highly cited publications. Journal of the Association for Information Science and Technology, 68(10), 2511–2521. DOI 10.1002/asi.23883
- [180] Zhu, L., & Lerman, K. (2016). Attention inequality in social media. arXiv:1601.07200
- [181] Sinatra, R., Wang, D., Deville, P., Song, C., & Barabási, A.-L. (2016). Quantifying the evolution of individual scientific impact. Science, 354(6312), aaf5239. DOI 10.1126/science.aaf5239

- [182] Radicchi, F., Fortunato, S., & Castellano, C. (2008). Universality of citation distributions: Toward an objective measure of scientific impact. *PNAS*, 105(45), 17268–17272. DOI 10.1073/pnas.0806977105
- [183] CoARA (2022). The Agreement on Reforming Research Assessment. Coalition for Advancing Research Assessment. <https://coara.eu/agreement/the-agreement-full-text/>
- [184] Curry, S., Gadd, E., & Wilsdon, J. (2022). Harnessing the Metric Tide: Indicators, infrastructures & priorities for UK responsible research assessment. Research on Research Institute. <https://doi.org/10.6084/m9.figshare.21701624>
- [185] IRUS (Jisc). Standardised, COUNTER-conformant usage statistics for repositories. <https://irus.jisc.ac.uk>
- [186] UK Research and Innovation (2024). Résumé for Research and Innovation (R4RI): guidance. <https://www.ukri.org/apply-for-funding/develop-your-application/resume-for-research-and-innovation-r4ri-guidance/>
- [187] Overton. The policy intelligence tool. <https://www.overton.io>
- [188] CC-PLUS. Consortia Collaborating on a Platform for Library Usage Statistics (open-source SUSHI/COUNTER-5 harvesting; IMLS-funded, PALCI et al.). <https://www.cc-plus.org/>
- [189] HELIOS Open. Higher Education Leadership Initiative for Open Scholarship. <https://www.heliosopen.org/>

Additional sources

Additional sources cited in the text as [B-n]: platform pages, forum threads, blog posts, and other materials that support specific in-text points. Platform pages are unstable and should be re-verified against archived copies.

- [B-1] Cabells Predatory Reports (active predatory-journal screening service); Beall's List (archived only; not updated since January 2017).
- [B-2] Caron, P. TaxProf Blog (2026, January 25). "The SSRN Top Five New Tax Papers Ranking Is Broken". <https://taxprofblog.aals.org/2026/01/25/the-ssrn-top-five-new-tax-papers-ranking-is-broken/>
- [B-3] EconJobRumors (EJMR), discussion threads on SSRN downloads and author rankings (e.g., /topic/ssrn-downloads-and-views; /topic/top-10-of-authors-on-ssrn-by-downloads). Anonymous forum; cited only as evidence that such discussions exist.
- [B-4] de Leon, F. L. L., & McQuillin, B. (2020). The Role of Conferences on the Pathway to Academic Impact. *Journal of Human Resources*, 55(1), 164. <https://jhr.uwpress.org/content/55/1/164>
- [B-5] Reddit r/AskAcademia, thread 1seeq3 (6 Apr 2026, "I've uploaded a working paper to SSRN... honest advice"); Reddit r/academia, thread 13w6zjw (31 May 2023, "SSRN Approval Time is Absurd"). Existence confirmed via third-party Reddit archives.
- [B-6] SSRN — statistical boilerplate of announcements (e.g. <https://www.ssrn.com/index.cfm/en/ern/ads/06042025ann001/>) and the April 2018 announcement (04132018ann002). Platform pages; verify against archived copies. Archived copy: <https://web.archive.org/web/20251229194317/https://www.ssrn.com/index.cfm/en/ern/ads/06042025ann001/> (snapshot 29 December 2025; accessed 15 August 2026).

[B-7] SSRN FAQ — versioning and ORCID. <https://www.ssrn.com/index.cfm/en/faq/> (accessed 2026; verify against an archived copy).

[B-8] SSRN — "SSRN has Sunset Rankings" landing page. <https://www.ssrn.com/ssrn/rankings-sunset>. Archived copy: <https://web.archive.org/web/20260719233259/https://www.ssrn.com/ssrn/rankings-sunset> (snapshot 19 July 2026; accessed 15 August 2026).

[B-9] SSRN — historical rankings pages at hq.ssrn.com / papers.ssrn.com/topten: Top 10,000 Papers (TRN_gID=10), Top 30,000 Authors (TRN_gID=7), Top 3,000 Law Authors, Top 350 U.S. Law Schools, Top 500 International Law Schools, Eigenfactor, Recent Top Papers (60 days), "Ranking Data Explained" ([ranking_data_explain.cfm?id=6](https://hq.ssrn.com/rankings/Ranking_data_explain.cfm?id=6) / [id=10](https://hq.ssrn.com/rankings/Ranking_data_explain.cfm?id=10)). Historical pages; since 15 July 2026 these URLs return 403/placeholder responses — cite via archived copies. Archived copy of the Top 10,000 Papers page: https://web.archive.org/web/20260714112806/https://hq.ssrn.com/rankings/Ranking_display.cfm?TRN_gID=10 (snapshot 14 July 2026, the day before the sunset; accessed 15 August 2026).

[B-10] Teixeira da Silva, J. A. (2024/2025). Spam, unwanted or unsolicited emails following the publication of preprints (year-long observation: 875 invitations from 256 journals, 76% from blocklists). DOI 10.1080/10875301.2025.2550600; ResearchGate publication 384286174.

[B-11] Caron, P. TaxProf Blog (2026, June 7). "SSRN Has Not Jumped the Shark". <https://taxprofblog.aals.org/2026/06/07/ssrn-has-not-jumped-the-shark/> (reply to Bainbridge [12]; accessed September 2026).