

NBER WORKING PAPER SERIES

AI AND THE COLLAPSE OF THE WWW

Alex Chan

Working Paper 35344

<http://www.nber.org/papers/w35344>

NATIONAL BUREAU OF ECONOMIC RESEARCH

1050 Massachusetts Avenue

Cambridge, MA 02138

June 2026

I thank various seminar and workshop participants for their comments. Funding was provided by Harvard Business School AI Institute. Valuable research assistance was provided by Daniel Kornbluth and Ce Li. The views expressed herein are those of the author and do not necessarily reflect the views of the National Bureau of Economic Research.

NBER working papers are circulated for discussion and comment purposes. They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2026 by Alex Chan. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

AI and the Collapse of the www

Alex Chan

NBER Working Paper No. 35344

June 2026

JEL No. D4, D43, D47, D49, D62, D8, D82, D83, L82, L86, O3, O33

ABSTRACT

This paper studies market design for generative AI intermediation. AI answer systems can improve user experience while diverting visits that finance publisher content and generate source-level quality signals. I show that an AI platform that underinternalizes future content reproduction retains too little referral traffic and can make costly open-web information subcritical, even with truthful content, accurate answers, and rational users. The mechanism can be self-reinforcing: less source-level measurement weakens conventional search, inducing further AI reliance. Sustainable repair requires replacing displaced revenue and deleted measurement through visitor-replacement royalties, audited provenance, human-information audits, and keystone-topic compensation.

Alex Chan

Harvard University

Harvard Business School

and NBER

achan@hbs.edu

(Please click [here](#) for the most current version)

1 Introduction

The open web has relied on a simple bargain.¹ Publishers produce content. Search and social discovery send users to that content. Visits generate revenue (“revenue event”). Visits also generate information about quality: users click, return, subscribe, cite, correct, bookmark, link, and leave. These signals help future users and search systems find good sources (“measurement event”).

Generative AI changes this bargain.² An AI answer can use publisher content while keeping the user in the AI interface. Users may be better served in the short run. But the source may lose the visit, the revenue from the visit, and the source-level signal that the visit would have produced. In this paper, I study whether costly human information can still be reproduced when the revenue event and the measurement event move away from publishers.

Publishers and media executives increasingly describe AI summaries and chatbots as a threat to the traffic model of the open web. Recent reporting based on Reuters Institute evidence describes fears of an “end of traffic era,” with media leaders expecting large declines in search referrals. Clickstream evidence points in the same direction: users are less likely to click source links when AI summaries appear in search results. At the same time, infrastructure firms and publishers are experimenting with crawler blocking, pay-per-crawl, and machine-readable licensing standards. These developments motivate the market design question in this paper: what has to be replaced when user attention moves from the publisher page to the AI answer interface?³

The paper focuses on costly human information: reporting, testing, expert judgment, measurement, verification, local observation, maintenance, and other scarce inputs that cannot be produced for free. I do not assume misinformation, irrational users, or bad AI answers.⁴ All publishers may be truthful, all AI answers may be accurate, and users may rationally prefer the AI answer. The problem is instead a missing market.⁵ The AI interface can appropriate current

¹I define the “open web” as the decentralized, publicly accessible web built on free, shared standards such as HTML, HTTP, and DNS. The open web allows anyone to publish content and communicate without requiring permission from a central corporate gatekeeper.

²Generative AI can refer broadly to artificial intelligence that creates new content, such as text, images, audio, video, and code, in response to prompts. I focus on answer engines and large-language-model interfaces such as ChatGPT or Gemini that locate, organize, and synthesize existing information.

³For present discussion of publisher concerns, see [The Guardian \(2026\)](#). For evidence on AI summaries and click behavior, see [Pew Research Center \(2025\)](#). For market responses such as crawler blocking and pay-per-crawl, see [Cloudflare \(2025\)](#).

⁴AI-generated content has in fact increased misinformation online, making false information harder to detect and easier to spread. I do not focus on that channel here. For a model of AI, digital platforms, and the information ecosystem with truthful and untruthful producers, see [Stiglitz and Ventura-Bolet \(2025\)](#).

⁵This is not only the familiar creator-incentive problem in which an intermediary reduces traffic to producers. The additional point is that AI answer intermediation can remove both the revenue event that finances future costly information and the source-level measurement event that reveals quality. The resulting market-design object is therefore the reproduction constraint of costly information, not current traffic or citation counts alone.

answer value without fully paying for the future reproduction value and source-level information value of publisher traffic.

I first isolate the basic participation constraint. If the share of monetizable exposure reaching a publisher goes to zero, then every approximately optimal positive-cost content plan produces arbitrarily little quality-adjusted content. Formally, the norm of the content plan converges to zero. This result uses only bounded visit revenue, the zero option, approximate optimality, and the assumption that positive content has positive avoidable cost. It does not require concavity, differentiability, a representative publisher, or equilibrium uniqueness. I state this as a lemma because several forces in the paper work through the same economic object: retained monetizable exposure.

The central result endogenizes the loss of exposure. The AI platform chooses how much referral traffic to leave for publishers. A social planner chooses the same object. The social planner internalizes more of the continuation value of the content stock and the source-level measurement system. Using monotone comparative statics, I show that the AI platform chooses weakly less retained referral traffic than the social planner. If the AI platform's choice puts the no-transfer reproduction system below replacement, then market-supported costly human information collapses. If the social planner can implement a positive-stock policy that preserves costly human information and yields higher social value, the AI platform's choice is inefficient.

The dynamic object behind this theorem is a reproduction system. Today's content generates immediate traffic and durable attention capital: subscribers, repeat readers, backlinks, bookmarks, reputation, and search authority. These assets persist but depreciate. In the formal model, I derive a next-generation matrix from these primitives. The spectral radius of the resulting matrix gives a one-sided sufficient condition for collapse of the maximal no-transfer reproduction system. The result is deliberately one-sided. If the system is below replacement, collapse follows for every path dominated by the no-transfer upper bound. If the condition fails, the paper does not claim survival; survival also requires reinvestment, entry, financing, allocation across topics, and ranking systems that direct attention to valuable sources.

I then treat visits and links as experiments rather than only payments. AI diversion deletes source-level observations. In the language of Blackwell's comparison of experiments, higher diversion makes the old web's quality experiment less informative. This point is separate from revenue. Even if an AI answer is useful, answer-level feedback need not identify which source created the useful information, which source was original, which source was stale, or which source corrected an error.

Raw content volume is not a one-to-one to costly information. If generative AI makes it cheap to produce pages, the number of URLs can rise while costly human information falls. The web can become larger in page count and thinner in facts, testing, reporting, and maintenance. This can also weaken conventional search through this channel when ranking systems depend on source-level signals generated by real human web activity. These search results are conditional.

If conventional search has hard trust anchors, provenance rules, manual verification, or protected exposure for verified sources, complete exposure collapse need not occur through this channel.

I also offer insights relevant to market design. A visitor-replacement payment can restore the old private publisher problem, but this is only a benchmark. The harder problem is informational. If low cost synthetic content can mimic high cost human information, paying for “quality” is not incentive compatible unless audits or provenance signals are sufficiently informative. I derive a likelihood-ratio condition for such audits. I also show that scarce compensation should not be targeted only by current clicks. At the margin, compensation should go to keystone topics: topics whose additional monetization most relaxes the reproduction constraint of the open-web content system. A practical experimental clearinghouse with sponsored challenges is offered as an example of a market design.

The paper therefore makes a disciplined version of the web-collapse claim. It does not say that every website disappears, that all AI reduces diversity, or that all modern search must fail.⁶ It says that when an AI platform diverts the revenue and measurement events without replacing them, costly human information may fall below replacement. The market-design problem is to move revenue and measurement to the new point of attention while screening costly human information from cheap synthetic imitation.

2 Related literature and contribution

This paper contributes to the economics of media and digital intermediaries by providing a minimal participation-constraint theorem for AI-mediated content markets and by showing how that theorem becomes an endogenous over-diversion and missing-market result once AI platform choice and social value are introduced. The classic media-economics literature studies how advertising, subscriptions, ownership, and market structure shape variety and quality; examples include [Anderson and Coate \(2005\)](#), [George \(2001\)](#), and the survey by [Gentzkow, Shapiro, and Stone \(2015\)](#). The two-sided-market literature studies pricing, indirect network effects, and intermediary design, as in [Rochet and Tirole \(2003\)](#). This paper borrows the idea that media production must be financed, but it makes a different contribution. The result is not a welfare comparison across particular market structures. It is a reproduction result. When the monetizable visit is removed and not replaced, bounded visit revenue and positive avoidable production cost are enough to imply collapse of market-supported content.

The paper is also related to work on aggregators and online news consumption. [Athey, Mobius, and Pal \(2021\)](#) study the Google News shutdown in Spain and show that aggregation affects news consumption and small publishers. [Chiou and Tucker \(2017\)](#) study content aggregation by digital

⁶The paper is deliberately one-sided where the mathematics is one-sided. A subcritical maximal reproduction system gives a sufficient condition for collapse. A supercritical upper bound does not by itself prove survival. Survival requires additional assumptions about reinvestment, entry, topic allocation, and ranking. I keep this distinction because it identifies what a successful market institution must provide.

intermediaries, and [Yan, Miller, and Skiera \(2020\)](#) examine how ad blocking changes online news consumption. This paper contributes a theoretical counterpart to that empirical and applied literature. Aggregators can either substitute for or complement publisher visits. Generative answer engines create a sharper limiting case because the user may obtain the information without visiting any source. The theorem characterizes what must happen in that limiting case under no transfers, without assuming a particular demand curve, click model, or AI platform objective.

The closest economics paper is [Stiglitz and Ventura-Bolet \(2025\)](#), who model an AI-shaped information ecosystem with truthful and untruthful producers, visit revenue, and consumers who differ in their ability to distinguish truth from untruth. Their mechanism explains how AI-mediated digital intermediation can reduce truthful information, increase misinformation, and intensify polarization. The present paper is complementary but distinct. It turns off misinformation and user error. All content may be truthful, all answers accurate, and all users rational. The collapse theorem still holds because it is driven by the failure to finance positive-cost content and by the deletion of source-level experiments.

Recent work studies AI search and publisher incentives more directly. [Wu et al. \(2026\)](#) put forth a game-theoretic model of AI Overviews⁷ in which the search engine improves user experience, diverts traffic from creators, discourages high quality creation, and can use citation or compensation mechanisms. [Zhao and Berman \(2026\)](#) provide early empirical evidence on LLMs, online news consumption and production, bot blocking, and publisher responses. The present paper contributes a sharper theorem that does not rely on a particular AI platform game, equilibrium uniqueness, calibrated click behavior, creator competition, smooth effort choices, or functional forms. This matters because many important web sources are local monopolies, niche databases, expert forums, public agencies, open-source maintainers, or lumpy projects that require a fixed investment before useful content exists. The collapse result holds for every approximately optimal no-transfer content plan under the primitive conditions stated below, and the strategic result requires only that the social planner internalizes more continuation value from future content reproduction than the AI platform. The paper also adds two objects that are not just traffic or effort: topic-support diversity and source-level statistical information.

A related economics and management literature studies generative AI as a supply shock to content markets and publisher ecosystems. [Zhang and Zhang \(2026\)](#) models supply, demand, traffic distribution, and welfare in AI-affected content markets. [Zhang and Zhang \(2026\)](#) studies generative AI as a non-convex supply shock with market bifurcation and information pollution. [Baumann, Akcay, and Plotkin \(2026\)](#) develop a social-welfare model in which individually rational use of generative AI can reduce collective welfare through model collapse. This paper contributes to that line by distinguishing raw content volume from costly human-information content. The extension shows that cheap AI production can increase the number of pages while decreasing the costly information stock and diluting PageRank. The result therefore reconciles two observations

⁷This is a feature of Google search, available (at least) circa 2026.

that may otherwise seem inconsistent: the web can become more abundant in pages and more scarce in costly information.

The computer science literature on model collapse and synthetic data is another close neighbor. [Shumailov et al. \(2023\)](#) and [Shumailov et al. \(2024\)](#) show that recursive training on generated data can make models forget distributional tails; [Seddik et al. \(2024\)](#), [Dohmatob et al. \(2024\)](#), and [Gerstgrasser et al. \(2024\)](#) study statistical mechanisms and qualifications of that phenomenon. The present paper is not a model collapse paper in the machine-learning sense. It studies market collapse of human content production and the collapse of source-level experiments. The two channels are connected: if the market supplies less fresh human content, future training and retrieval systems may become more dependent on synthetic material. But the theorem here does not require recursive training or synthetic data; it follows from no-transfer diversion alone.

A separate computer science literature studies generative engines, optimization for AI answers, and retrieval pollution. [Aggarwal et al. \(2023\)](#) introduce generative engine optimization and show that content creators can optimize visibility in generative engine responses. [Yu, Kim, and Kim \(2026\)](#) study retrieval collapse when AI-generated content pollutes the web, emphasizing source diversity and the infiltration of synthetic evidence into retrieval pipelines. This paper contributes an economic complement. Visibility in an AI answer is not enough if it does not generate visits or transfers. Retrieval pollution can make the economic collapse harder to observe because raw content appears abundant even as costly human information disappears.

The paper also contributes to the literature on ranking, feedback, and link analysis. [Brin and Page \(1998\)](#) and [Kleinberg \(1999\)](#) show how hyperlinks can be used as decentralized authority signals. [Barabasi and Albert \(1999\)](#) provide a canonical model of concentration in growing networks, and [Gyongyi, Garcia-Molina, and Pedersen \(2004\)](#) study trust-based defenses against web spam. PageRank is therefore the natural mathematical benchmark, but it is not a complete description of contemporary web search. Google’s public documentation describes Search as using automated ranking systems that look at many factors and signals, including language understanding systems such as BERT and neural matching, relevance, usability, freshness, original-content systems, source expertise, interaction data, and link analysis; the same documentation says that PageRank has evolved and continues to be part of Google’s core ranking systems ([Google Search Central, 2025](#); [Google Search, 2026](#)). The learning-to-rank and recommendation literatures, including [Joachims, Swaminathan, and Schnabel \(2017\)](#) and [Chaney, Stewart, and Engelhardt \(2018\)](#), show how biased feedback and algorithmic confounding can distort learning. The contribution here is a more primitive signal result. AI diversion deletes source-level observations. The old source experiment and the current-link experiment are Blackwell garbled when higher diversion is represented by state-independent thinning. Modern ranking inputs generated by genuine human web activity satisfy the same conclusion only when their loss is modeled as such a garbling; otherwise the paper’s modern-ranking result below gives a conditional exposure-bound theorem rather than a universal Blackwell comparison. This is a theorem about the loss of information in

the input to ranking, not a proposal for a new ranking algorithm.

I derive the strategic over-diversion result using monotone comparative statics (Milgrom and Shannon, 1994; Topkis, 1998), side-stepping the need to solve a fully differentiable dynamic game.⁸ Formally, the platform and planner are compared state by state. The objective is indexed by a continuation-value weight θ and has increasing differences in retained referrals and θ . Since the planner places a larger weight on continuation value than the platform, Topkis’s theorem implies that the planner’s optimal retained-referral choice is weakly higher under monotone tie-breaking. Economically, retained referrals may reduce current AI-interface rents, but they raise the future value of content reproduction and source-level measurement; the wedge is that the platform internalizes only part of that continuation value.

The reproduction number section also draws on a different mathematical tradition. In finite Markov chains, a fundamental matrix adds up expected future time spent in transient states. I use the same idea for web attention capital. Costly content can create immediate traffic and durable assets such as subscribers, repeat readers, backlinks, bookmarks, reputation, and search authority. These assets persist, but they depreciate. Later I define the matrices that map content into attention capital, attention capital into future traffic, and traffic into future costly content. Combining those objects gives the next-generation matrix used in the reproduction theorem. This is not a new behavioral assumption. It is a reduced-form accounting object built from standard pieces: content creates visits and durable attention capital; attention capital persists but decays; visits become revenue; and revenue finances future costly content. This construction follows the fundamental-matrix logic in Kemeny and Snell (1976) and the standard nonnegative-matrix treatments in Berman and Plemmons (1994). The web interpretation combines that mathematics with the computer-science view that links and random walks encode discovery and authority, as in Brin and Page (1998), Kleinberg (1999), and Langville and Meyer (2006); with marketing work that treats subscribers and repeat users as assets, as in Gupta, Lehmann, and Stuart (2004), Reinartz and Kumar (2003), and Fader, Hardie, and Lee (2005); and with empirical evidence that aggregation affects consumption, page views, and publisher monetization, as in Athey, Mobius, and Pal (2021) and Laub, Miller, and Skiera (2024).

Finally, the mathematical structure is closest to Blackwell’s comparison of experiments (Blackwell, 1953). The paper uses Blackwell garbling to formalize a simple idea: if an AI platform deletes source-level visits or links and replaces them with answer-level feedback, then source-level inference generally becomes less informative. This lets the paper state signal collapse without committing to a Gaussian model, a particular click-through equation, a specific PageRank implementation, or a proprietary modern search ranker.

⁸The choice of lattice-theoretic methods is motivated by generality (and doctoral education at Stanford): the argument requires only monotone order structure, not a fully differentiable dynamic game. The relevant AI and search choices are often discrete, threshold-based, and proprietary: whether to show an answer, how many links to expose, which sources to cite, how much crawl access to grant, and how much traffic to return are not naturally modeled as smooth first-order conditions. Publisher production is also lumpy and nonconvex.

3 Content Production and AI Diversion

Let $X = \{1, \dots, n\}$ be a finite set of topics. A topic may be a local-news beat, a technical vertical, a review category, a medical explanation, a legal guide, a recipe class, a database, a forum category, a software tutorial, a train enthusiast hub, or a cat content category. This section works topic by topic and uses $x \in X$ for a generic topic. Later sections use $i = 1, \dots, n$ for the same topics when the analysis uses vectors and matrices.

For each topic x , Q_x is an arbitrary feasible set of content plans. A plan may encode quantity, quality, updating frequency, staffing, verification effort, maintenance, or other production choices. No vector-space, convexity, or differentiability structure is imposed on Q_x unless explicitly stated. A publisher chooses a content plan $q \in Q_x$. The plan $0_x \in Q_x$ denotes no commercial production. The scalar index

$$\|q\|_x : Q_x \rightarrow \mathbb{R}_+$$

measures quality-adjusted commercial content, with $\|0_x\|_x = 0$. The notation $\|q\|_x$ is not assumed to be a norm generated by a vector space; it is only a topic-specific scalar content index.

Let $r \in [0, 1]$ denote the retained referral share. The complementary diversion rate is $z = 1 - r$. At $r = 1$ (equivalently, $z = 0$), the relevant interaction sends users to the publisher. At $r = 0$ (equivalently, $z = 1$), the AI answer satisfies the relevant interaction without a publisher visit. In a one-topic environment, with conventional search visibility normalized to one, the scalar effective-monetization multiplier is $m = r = 1 - z$.

Let $R_x : Q_x \rightarrow \mathbb{R}_+$ be direct-visit revenue when monetizable exposure is one. This may include advertising, subscription conversion, affiliate revenue, donations, data, leads, or other revenue that requires the relevant human interaction to occur at the publisher. Let $C_x : Q_x \rightarrow \mathbb{R}$ be avoidable net production and maintenance cost after any nonpecuniary benefit or exogenous subsidy has been deducted. The collapse results use lower bounds on this net cost: positive content must have positive avoidable cost on the commercial margin considered below. A hobby site, university website, government service, or philanthropically funded outlet can be represented by a lower C_x or by a transfer outside the no-transfer commercial margin.

The no-transfer payoff is

$$\pi_x(q; m) = mR_x(q) - C_x(q), \quad m \in [0, 1]. \quad (3.1)$$

For pure AI diversion, substitute $m = 1 - z$. Thus writing the payoff as $(1 - z)R_x(q) - C_x(q)$ and writing it as $mR_x(q) - C_x(q)$ refer to the same object. In the finite reproduction system, the corresponding topic-level object is $m_i = r_i \xi_i$, where r_i is retained referral traffic and ξ_i is genuine-source visibility in conventional search. The latter object is introduced formally in the reproduction section.

The model does not require a demand curve, utility function, search technology, or AI platform

objective. The theorem is about the publisher participation constraint conditional on effective monetization. AI diversion is one reason exposure falls; later sections add search degradation, synthetic dilution, and user migration into AI. If AI diversion selects low-ad-value or high-ad-value users, that selection is absorbed into the effective monetization term or into the counterfactual revenue function R_x . The result does not require diverted and nondiverted users to have the same per-person value.

Assumption 1 (Zero option). For every topic x , the zero plan satisfies $0_x \in Q_x$, $R_x(0_x) = 0$, $C_x(0_x) = 0$, and $\|0_x\|_x = 0$.

A publisher can always choose not to produce the commercial content for this topic. This outside option pins down the participation constraint: any active plan must beat zero. I call a topic active under a chosen plan if the chosen plan is not 0_x .

Assumption 2 (Nonnegative bounded direct-visit revenue). For every topic x , direct-visit revenue is nonnegative and bounded:

$$0 \leq R_x(q) \leq \bar{R}_x < \infty \quad \text{for all } q \in Q_x. \quad (3.2)$$

The result does not require a demand curve. It only requires that, for a fixed topic, the amount a publisher can earn from direct visits is not infinite. If AI removes the visits, this bounded revenue base cannot cover a positive avoidable cost.

Assumption 3 (No free positive commercial content). For every topic x and every $\delta > 0$,

$$\underline{C}_x(\delta) \equiv \inf\{C_x(q) : q \in Q_x, \|q\|_x \geq \delta\} > 0. \quad (3.3)$$

Assumption 3 is the key economic assumption. It says that content bounded away from zero has positive avoidable net cost. It is weaker than fixed cost. It permits discontinuities, nonconvexities, multiple qualities, multiple publishers, arbitrary cost curves, and arbitrary revenue functions. It also allows marginal cost to approach zero as $\|q\|_x$ approaches zero.

For results about whether a topic is active, rather than only how much content it produces, the paper adds only the following assumption.

Assumption 4 (Avoidable activation cost). For every topic x to which the active-topic result is applied,

$$F_x \equiv \inf\{C_x(q) : q \in Q_x, q \neq 0_x\} > 0. \quad (3.4)$$

Assumption 4 is not needed for quality-adjusted content collapse. It is needed only to prove collapse of the active topic set. Without it, a topic can remain active with arbitrarily small positive output. Assumption 4 implies Assumption 3, because any plan with positive norm is nonzero and therefore costs at least F_x .

For maximum generality, the theorem uses approximate maximizers rather than assuming exact maximizers exist. Let

$$\Pi_x(m) = \sup_{q \in Q_x} \pi_x(q; m). \quad (3.5)$$

A plan q is ε -optimal at exposure m if

$$\pi_x(q; m) \geq \Pi_x(m) - \varepsilon. \quad (3.6)$$

Because the zero plan is feasible, $\Pi_x(m) \geq 0$ for all $m \in [0, 1]$. When a statement is written in terms of a diversion rate z , ε -optimal at z means ε -optimal at the exposure multiplier $m = 1 - z$.

4 Effective monetization and no-transfer corollaries

The same mathematical force appears repeatedly in the paper. AI diversion, search degradation, PageRank dilution, and migration into AI all reduce the effective share of old-web monetization that reaches the publisher. It is useful to isolate that force once.

Lemma 1 (Effective-monetization collapse). *Fix a topic x . Suppose Assumptions 1–3 hold. Let $m_\ell \in [0, 1]$ satisfy $m_\ell \rightarrow 0$, let $\varepsilon_\ell \rightarrow 0$, and let $q_\ell \in Q_x$ be ε_ℓ -optimal for the payoff*

$$\pi_{x,\ell}(q) = m_\ell R_x(q) - C_x(q). \quad (4.1)$$

Then $\|q_\ell\|_x \rightarrow 0$. If Assumption 4 also holds, then $q_\ell = 0_x$ for all sufficiently large ℓ .

The lemma is the paper’s basic participation-constraint result. It makes no reference to AI, search, PageRank, or equilibrium selection. A bounded monetization base multiplied by a vanishing exposure term cannot cover a positive avoidable cost.

Corollary 1 (Market-supported content is not reproducible under uncompensated full diversion). *In a no-transfer economy, as AI diversion approaches full diversion, every approximately optimal commercial content plan for a fixed topic becomes arbitrarily small. Formally, fix a topic x . Under Assumptions 1–3, let $z_\ell \rightarrow 1$, let $\varepsilon_\ell \rightarrow 0$, and let $q_\ell \in Q_x$ be ε_ℓ -optimal at z_ℓ . Then*

$$\|q_\ell\|_x \rightarrow 0. \quad (4.2)$$

If Assumption 4 also holds, then $q_\ell = 0_x$ for all sufficiently large ℓ .

The corollary is Lemma 1 with $m_\ell = 1 - z_\ell$. The AI-specific content is not the proof technique but the economic interpretation: answer intermediation removes the monetizable visit unless some transfer or other revenue source replaces it.

Corollary 2 (Aggregate quality-adjusted content collapse). *For a finite topic set, aggregate quality-adjusted commercial content converges to zero whenever the topicwise assumptions hold*

for every topic. Formally, suppose Assumptions 1–3 hold for every $x \in X$. Let $z_\ell \rightarrow 1$, let $\varepsilon_\ell \rightarrow 0$, and for each sequence index ℓ and each topic x let $q_\ell(x)$ be an ε_ℓ -optimal selection at z_ℓ . Then

$$Q_\ell \equiv \sum_{x \in X} \|q_\ell(x)\|_x \rightarrow 0. \quad (4.3)$$

Topic-by-topic content collapse aggregates by finite summation. No domination condition is needed because the topic set is finite.

Proposition 1 (Active-topic collapse). *If each topic to which the active-topic result is applied has a positive avoidable activation cost, then approximately optimal production eventually chooses the zero plan. Formally, fix a topic x . Under Assumptions 1, 2, and 4, let $z_\ell \rightarrow 1$, $\varepsilon_\ell \rightarrow 0$, and let q_ℓ be ε_ℓ -optimal at z_ℓ . Then $q_\ell = 0_x$ for all sufficiently large ℓ . If the assumptions hold for every topic in finite X , and for each sequence index ℓ and each topic x the plan $q_\ell(x)$ is ε_ℓ -optimal at z_ℓ with $z_\ell \rightarrow 1$ and $\varepsilon_\ell \rightarrow 0$, then*

$$N_\ell \equiv \sum_{x \in X} \mathbf{1}\{q_\ell(x) \neq 0_x\} \rightarrow 0. \quad (4.4)$$

The distinction between Corollary 1 and Proposition 1 is important. Positive content cannot remain bounded away from zero under the weaker no-free-content assumption. To force the active indicator itself to vanish, one needs a positive cost of any nonzero activity.

Corollary 3 (Finite exit before full diversion). *If active topics have a uniform fixed-cost-to-revenue lower bound, then full commercial exit occurs before full diversion. Formally, suppose Assumptions 1, 2, and 4 hold and there exists $\underline{r} > 0$ such that $F_x/\bar{R}_x \geq \underline{r}$ for every topic with $\bar{R}_x > 0$, with the ratio interpreted as infinite when $\bar{R}_x = 0$ and $F_x > 0$. Then every exact optimum is inactive for every topic whenever*

$$z > 1 - \underline{r}. \quad (4.5)$$

More generally, no active plan can be ε -optimal when $(1 - z)\bar{R}_x - F_x < -\varepsilon$.

With a uniform fixed-cost-to-revenue margin, exit happens before the limiting case of full diversion. A publisher does not need to lose every visit; it exits once the remaining visit revenue cannot cover the activation cost.

Corollary 4 (Revenue-insufficiency lower bound). *Under Assumptions 1–3, any positive content bounded away from zero requires non-vanishing non-visit revenue as diversion approaches one. Formally, suppose a transfer $T_x(q, z)$ is added to the payoff and satisfies $T_x(0_x, z) = 0$:*

$$\pi_x^T(q; z) = (1 - z)R_x(q) + T_x(q, z) - C_x(q). \quad (4.6)$$

If a sequence q_ℓ with $\|q_\ell\|_x \geq \delta > 0$ is ε_ℓ -optimal at $z_\ell \rightarrow 1$, with $\varepsilon_\ell \rightarrow 0$, then necessarily

$$\liminf_{\ell \rightarrow \infty} T_x(q_\ell, z_\ell) \geq \underline{C}_x(\delta). \quad (4.7)$$

Any mechanism that preserves positive costly content near full diversion must replace the missing visit revenue with non-visit compensation. The lower bound identifies the minimum object the mechanism must cover: the avoidable cost of positive content.

5 Sharpness and economic interpretation

The assumptions are close to necessary. The next proposition records the exact margins on which the collapse conclusions depend.

Proposition 2 (Sharpness). *The collapse conclusions can fail when the corresponding primitive assumption is removed. If positive content can be produced at zero net avoidable cost, content collapse need not occur. If publisher value does not vanish with visits because some non-visit revenue or benefit replaces the visit, content collapse need not occur. If there is no avoidable activation cost, active-topic collapse need not occur even though quality-adjusted output collapses.*

The proposition clarifies the scope of the result. The collapse claim is not a claim that all online expression or production disappears. It is a claim about market-supported commercial production whose revenue event is diverted. If content has zero net cost, if it is publicly funded, if it is produced for ideological reasons, or if another revenue source replaces the visit, then the collapse corollary need not apply. The result identifies the object that must be replaced: the monetizable visit.

6 The open-web reproduction number

The central dynamic force in the paper is summarized by the open-web reproduction number. The object is not meant to be primitive. This section derives it from a minimal attention-capital process. The starting observation is that a web page, article, review, database, or tutorial can generate two kinds of traffic. Some traffic is immediate. Other traffic comes later because the content created durable attention capital: subscribers, repeat readers, backlinks, bookmarks, reputation, search authority, or other persistent discovery assets. This is the same logic as the fundamental matrix of a transient Markov chain, which sums all future visits generated by a persistent transition process, and it is also the logic behind customer-lifetime-value models that value a retained customer by the expected stream of future margins.

The one-topic variable q is a publisher-level content plan. The dynamic variables below are aggregate accounting objects. Time is discrete, $t = 0, 1, 2, \dots$. Let $h_t \in \mathbb{R}_+^n$ be the stock of costly human-information content across topics at date t . Let $g_t \in \mathbb{R}_+^n$ denote newly produced costly

human-information content at date t . The component h_{it} is the stock of costly human information in topic i , while g_{it} is new costly human information produced in that topic. The reproduction system does not require a one-to-one mapping from individual plans q to h_t or g_t ; it uses only the financing upper bound implied by visit revenue and unit production costs.

Let $a_t \in \mathbb{R}_+^{d_a}$ be attention capital. The dimension d_a need not equal the number of topics n : attention capital may include several types of assets, such as subscriber relationships, persistent backlinks, bookmarks, reputation, search authority, and other durable objects that convert past content into future visits. The attention-capital law is

$$a_{t+1} \leq Bh_t + P^A a_t, \quad (6.1)$$

where $P^A \in \mathbb{R}_+^{d_a \times d_a}$, $B \in \mathbb{R}_+^{d_a \times n}$, and $\rho(P^A) < 1$.⁹ The matrix B maps current costly content into durable attention capital. The matrix P^A maps existing attention capital into next-period attention capital. The restriction $\rho(P^A) < 1$ says that attention capital is durable but not self-sustaining: subscribers churn, links become stale, bookmarks decay in use, and search authority loses current informativeness if costly content is not replenished.

Let $N_t \in \mathbb{R}_+^n$ denote potential human visits before AI diversion and search degradation. These visits are bounded by

$$N_t \leq Qh_t + Ua_t, \quad (6.2)$$

where $Q \in \mathbb{R}_+^{n \times n}$ captures immediate visits from current content and $U \in \mathbb{R}_+^{n \times d_a}$ converts attention capital into visits. The matrix Q includes immediate direct traffic, referrals from current pages, and cross-topic traffic within sessions. The matrix U includes future traffic generated by subscribers, links, bookmarks, and search authority.

The notation distinguishes one-topic, topic-date, and stationary-vector objects. In the one-topic model, $m = 1 - z$. In the dynamic topic model,

$$m_{it} = (1 - z_{it})\xi_{it} \in [0, 1] \quad (6.3)$$

is topic i 's effective monetizable exposure rate at date t . AI diversion lowers $1 - z_{it}$. Search degradation, PageRank dilution, or synthetic pollution lowers ξ_{it} , the genuine-source visibility of topic i in conventional search at date t . Let $\mathbf{m}_t = (m_{1t}, \dots, m_{nt})$. In stationary vector notation, $\mathbf{m} = r \circ \xi$, where r is retained referral traffic and ξ is genuine-source visibility in conventional search.

Let $\eta_i > 0$ be an upper bound on publisher revenue per monetized visit in topic i . The bridge from the nonconvex publisher problem to the linear reproduction system is the following unit-cost

⁹Here and throughout, $\rho(A)$ denotes the spectral radius of a square matrix A , meaning the largest absolute value of its eigenvalues. This is not the matrix rank. The condition $\rho(P^A) < 1$ ensures that the powers $(P^A)^k$ decay sufficiently for the Neumann series $(I - P^A)^{-1} = \sum_{k=0}^{\infty} (P^A)^k$ to be well defined. Economically, the restriction says that attention capital is durable but not self-sustaining: subscribers churn, links become stale, bookmarks decay in use, and search authority loses current informativeness if costly content is not replenished.

lower bound. It is an average-cost lower bound, not a convexity assumption.

Assumption 5 (Unit-cost lower bound for costly information). For each topic i , let $\text{Cost}_i : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ be the avoidable private cost of producing $g_i \geq 0$ units of new costly human-information content in topic i . There exists $c_i > 0$ such that

$$\text{Cost}_i(g_i) \geq c_i g_i \quad \text{for all } g_i \geq 0. \quad (6.4)$$

Equivalently, one may take $c_i \leq \inf_{g_i > 0} \text{Cost}_i(g_i)/g_i$.

Define the diagonal matrix of maximum visit-financed content units by

$$\Theta = \text{diag}(\eta_1/c_1, \dots, \eta_n/c_n), \quad E(\mathbf{m}) = \text{diag}(\mathbf{m})\Theta. \quad (6.5)$$

If g_{t+1} is newly produced costly content, then a no-transfer market path satisfies the financing bound

$$g_{t+1} \leq E(\mathbf{m}_t)N_t. \quad (6.6)$$

This is an upper bound, not a behavioral equality. It says that market-supported production cannot cost more than the monetized visits available to finance it: topic i has at most $\eta_i m_{it} N_{it}$ in monetized visit revenue, and each unit of new costly human information in that topic costs at least c_i . If revenues can be pooled across topics, the same financing logic gives a nonnegative allocation matrix $A(\mathbf{m}_t)$ in place of the diagonal matrix $E(\mathbf{m}_t)$, and the augmented upper-bound matrix becomes

$$\begin{pmatrix} D + A(\mathbf{m}_t)Q & A(\mathbf{m}_t)U \\ B & P^A \end{pmatrix}.$$

The Schur-complement subcriticality argument extends to the corresponding reduced operator

$$D + A(\mathbf{m}_t) [Q + U(I - P^A)^{-1}B].$$

The closed-form uniform-exposure condition and the keystone derivatives below use the diagonal no-pooling specification $A(\mathbf{m}_t) = E(\mathbf{m}_t)$; with general pooling they must be interpreted after rederiving the relevant reduced operator. External finance, cash reserves, public funding, philanthropy, ownership integration, or cross-subsidies should be modeled as transfers, additional state variables, or reductions in avoidable cost. The reproduction bound here is a no-transfer, visit-financed market-production bound.

Finally, existing costly content depreciates or becomes stale according to

$$h_{t+1} \leq Dh_t + g_{t+1}, \quad (6.7)$$

where $D = \text{diag}(1 - \delta_1, \dots, 1 - \delta_n) \in \mathbb{R}_+^{n \times n}$ and $\delta_i \in (0, 1]$.

The financing inequality is where the nonconvex one-topic model connects to the linear network

system. Section 3 allows individual publisher technologies to be nonconvex and discontinuous. The linear system here is therefore not a hidden convexity assumption. It is a maximal-reproduction upper bound obtained from Assumption 5 and the revenue upper bound η_i . The bound may be loose when publishers face indivisibilities, fixed costs, credit constraints, or imperfect reinvestment. Looseness makes the subcriticality result stronger: if even the maximal financing upper bound is subcritical, every feasible no-transfer path collapses.

The same display becomes an equality only in a benchmark with exact linear reinvestment: all monetized visit revenue is reinvested into costly human information, the relevant unit costs are constant at c_i , the visit bound (6.2) is tight, and there are no additional financing, entry, or allocation frictions. I use that benchmark only to clarify when the spectral condition is a two-sided threshold. The main theorem needs only the upper bound.

Combining (6.1)–(6.7) gives the augmented nonnegative system

$$\begin{pmatrix} h_{t+1} \\ a_{t+1} \end{pmatrix} \leq \mathcal{M}(\mathbf{m}_t) \begin{pmatrix} h_t \\ a_t \end{pmatrix}, \quad \mathcal{M}(\mathbf{m}) = \begin{pmatrix} D + E(\mathbf{m})Q & E(\mathbf{m})U \\ B & P^A \end{pmatrix}. \quad (6.8)$$

Because $\rho(P^A) < 1$, the fundamental matrix

$$(I - P^A)^{-1} = I + P^A + (P^A)^2 + \dots \quad (6.9)$$

exists and is nonnegative. It sums all future attention-capital states generated by one unit of current attention capital. Therefore the lifetime potential-visit operator generated by one unit of current costly content is

$$Q + U(I - P^A)^{-1}B. \quad (6.10)$$

Multiplying by revenue per visit divided by content cost gives the derived next-generation matrix

$$\boxed{K = \Theta [Q + U(I - P^A)^{-1}B]}. \quad (6.11)$$

The entry K_{ij} is the amount of future market-supported costly content in topic i that can be financed by one unit of current costly content in topic j , counting immediate referrals, cross-topic traffic, subscriber retention, link persistence, reputation persistence, visit monetization, and reinvestment into future costly content.

Theorem 1 (Microfounded open-web reproduction-number collapse). *The open-web reproduction number is the spectral radius of the attention-capital next-generation operator. Formally, suppose (6.8) holds, $P^A \geq 0$, $B \geq 0$, $Q \geq 0$, $U \geq 0$, $D \geq 0$, $\Theta \geq 0$, and $\rho(P^A) < 1$. For a stationary exposure vector $\mathbf{m}$, define*

$$\mathcal{R}^{web}(\mathbf{m}) \equiv \rho(D + \text{diag}(\mathbf{m})K), \quad K = \Theta [Q + U(I - P^A)^{-1}B]. \quad (6.12)$$

Then

$$\rho(\mathcal{M}(\mathbf{m})) < 1 \iff \mathcal{R}^{web}(\mathbf{m}) < 1. \quad (6.13)$$

Moreover, if there is $T < \infty$ and a vector $\bar{\mathbf{m}} \in [0, 1]^n$ such that $\mathbf{m}_t \leq \bar{\mathbf{m}}$ componentwise for all $t \geq T$ and

$$\mathcal{R}^{web}(\bar{\mathbf{m}}) < 1, \quad (6.14)$$

then $h_t \rightarrow 0$ and $a_t \rightarrow 0$ for every initial stock (h_0, a_0) .

Under the no-transfer upper bound, costly human information cannot be sustained when today's content generates too little monetized attention to finance tomorrow's content. The matrix K is the lifetime multiplier from current content into future monetizable attention and then back into future content. AI diversion lowers $1 - z_{it}$, search degradation lowers ξ_{it} , and both reduce $m_{it} = (1 - z_{it})\xi_{it}$. When the spectral radius falls below one, the system is subcritical: each generation of costly human information finances less than one replacement generation.

Proposition 3 (Exact linear reinvestment benchmark). *Suppose the inequalities (6.1), (6.2), (6.6), and (6.7) hold as equalities at a stationary exposure vector $\mathbf{m}$. Then the augmented state $y_t = (h_t, a_t)$ satisfies $y_{t+1} = \mathcal{M}(\mathbf{m})y_t$. In this benchmark, $\mathcal{R}^{web}(\mathbf{m}) < 1$ if and only if $y_t \rightarrow 0$ for every initial state. If $\mathcal{R}^{web}(\mathbf{m}) > 1$, then some nonnegative initial state does not converge to zero; if $\mathcal{M}(\mathbf{m})$ is irreducible,¹⁰ every strictly positive initial state fails to converge to zero. Thus the spectral radius is a true critical boundary only under exact linear reinvestment or an equivalent local linearization. Outside that benchmark, $\mathcal{R}^{web} < 1$ remains a sufficient condition for collapse, while $\mathcal{R}^{web} \geq 1$ does not by itself prove survival.*

The Perron-Frobenius theorem gives a shadow-price test for subcriticality. Let $\mathcal{T}(\mathbf{m}) = D + \text{diag}(\mathbf{m})K$ denote the reduced next-generation operator for costly human information. The test below applies to any nonnegative matrix A . When $A = \mathcal{T}(\mathbf{m})$, the prices weight the reduced generational reproduction process. When one works directly with the calendar-time attention-capital system, the same test should instead be applied to the augmented matrix $\mathcal{M}(\mathbf{m})$ and to augmented shadow prices on (h, a) .

Lemma 2 (Standard shadow-price characterization of subcriticality). *For any nonnegative matrix A , the following are equivalent:*

1. $\rho(A) < 1$;
2. there exist $p \gg 0$ and $\beta < 1$ such that $p^\top A \leq \beta p^\top$.

Consequently, if $h_{t+1} \leq Ah_t$ and there exist $p \gg 0$ and $\beta < 1$ such that $p^\top A \leq \beta p^\top$, then

$$p^\top h_t \leq \beta^t p^\top h_0 \rightarrow 0. \quad (6.15)$$

¹⁰Irreducibility means that the topic-reproduction network is connected in the sense that every topic can eventually affect every other topic through the reproduction matrix.

A vector $p \gg 0$ assigns weights to topics according to their role in reproducing costly human information. If, under these weights, one period of market-financed reproduction reduces the value of the content stock by a fixed factor $\beta < 1$, then visit-financed human information cannot be sustained. The test avoids solving the full publisher dynamic; it only uses the no-transfer financing bound.

Corollary 5 (Uniform-exposure below-replacement condition). *Suppose $m_{it} = (1 - z)\xi$ for every topic and date, and suppose $D = (1 - \delta)I$, with $\delta \in (0, 1]$. Then the no-transfer maximal reproduction system is subcritical if and only if*

$$(1 - z)\xi\rho(K) < \delta. \quad (6.16)$$

Whenever this condition holds, every content path dominated by the no-transfer upper bound converges to zero. If $\rho(K) > 0$, the condition is equivalent to $(1 - z)\xi < \delta/\rho(K)$. If $\rho(K) = 0$, the system is subcritical for every $\delta > 0$.

Full diversion is not necessary. The relevant object is the product of residual non-AI exposure, conventional-search usefulness, and the networked reproduction capacity of the web. Search degradation lowers the same reproduction number as AI diversion. A moderate level of answer diversion can therefore be unsustainable if conventional search is polluted or if the underlying reproduction network is weak.

Proposition 4 (Keystone topics and marginal fragility). *Suppose $\mathcal{T}(\mathbf{m}) = D + \text{diag}(\mathbf{m})K$ is irreducible. Let $u \gg 0$ and $v \gg 0$ be the left and right Perron eigenvectors of $\mathcal{T}(\mathbf{m})$, normalized so that $u^\top v = 1$. Then the derivative of the open-web reproduction number with respect to topic i 's effective exposure is*

$$\frac{\partial \mathcal{R}^{web}(\mathbf{m})}{\partial m_i} = u_i(Kv)_i. \quad (6.17)$$

Since $m_i = (1 - z_i)\xi_i$, whenever the derivative exists,

$$\frac{\partial \mathcal{R}^{web}}{\partial z_i} = -\xi_i u_i(Kv)_i, \quad \frac{\partial \mathcal{R}^{web}}{\partial \xi_i} = (1 - z_i)u_i(Kv)_i. \quad (6.18)$$

Aggregate traffic is not the right sufficient statistic. Some topics are keystone topics because they sit in central positions in the reproduction network. An equal marginal reduction in the effective exposure rate m_i does more damage in a topic with high $u_i(Kv)_i$ than in a peripheral topic. A per-visit comparison would require normalizing by the topic's baseline visit volume. If the matrix is reducible or the Perron root is not simple, the derivative statement should be replaced by the corresponding directional derivative or subgradient of the spectral radius.

The ranking mechanisms that affect the visibility component ξ_i are developed in the search and synthetic-entry sections below.

7 Endogenous diversion and the central strategic theorem

The preceding collapse result is conditional on exposure. The central economic question is why exposure becomes too low. This section makes the retained-referral share a choice of the AI platform and compares the private choice with the social choice using monotone comparative statics rather than a fully parametric dynamic game.

There are n topics. Let $r \in [0, 1]^n$ be the vector of retained referral shares, so $z_i = 1 - r_i$ is AI diversion in topic i . Let s be a payoff-relevant state, including the inherited stock of costly human information and attention capital. The AI platform's objective is indexed by a scalar continuation-weight parameter $\theta \in [0, 1]$:

$$J(r; \theta, s) = \Phi(r, s) + \theta \Psi(r, s). \quad (7.1)$$

This is a normalized ordered family. In the primitive interpretation below, Φ includes current interface payoff and the platform continuation value that the platform already internalizes, while Ψ is the incremental continuation value internalized by the social planner but not by the platform. Thus the notation does not require the platform to be myopic. A positive private continuation weight is allowed; the wedge is that the planner additionally values publisher surplus, public information, diversity, source-level measurement, and rival-system or public-option benefits that the platform does not fully capture.¹¹ The wedge is not that the platform is necessarily short-runist. A platform may care about future answer quality and therefore have $\theta_P > 0$. The wedge is that its continuation value is generally narrower than the social continuation value. It need not internalize publisher surplus, option value from future entry, public-good value of verified information, diversity across topics, or the source-level measurements that also discipline conventional search and rival AI systems. Thus a forward-looking platform can still choose a referral policy that is privately optimal and socially subcritical.

Assumption 6 (Increasing differences in retained referrals). For every state s , the feasible referral set $\mathcal{R}_s \subset [0, 1]^n$ is a nonempty compact sublattice. For every θ , $J(\cdot; \theta, s)$ is upper semicontinuous and supermodular on $\mathcal{R}_s$. Moreover J has increasing differences in (r, θ) : whenever $r' \geq r$ and $\theta' \geq \theta$,

$$J(r'; \theta', s) - J(r; \theta', s) \geq J(r'; \theta, s) - J(r; \theta, s). \quad (7.2)$$

¹¹A smooth one-dimensional benchmark gives the intuition. Suppose, for a fixed state s , retained referrals are summarized by a scalar $r \in [0, 1]$ and

$$J(r; \theta, s) = \Phi(r, s) + \theta \Psi(r, s),$$

where $\Phi_r(r, s) < 0$ captures the current AI-interface benefit of keeping users inside the AI product and $\Psi_r(r, s) > 0$ captures the continuation value of future content, measurement, and search quality. If the objective is single-peaked and the optimum is interior, the first-order condition is

$$\Phi_r(r, s) + \theta \Psi_r(r, s) = 0.$$

A decision maker with a larger θ chooses a weakly larger retained-referral share. The lattice result below is the same economic comparison without differentiability, concavity, scalar choices, or unique optima.

This assumption is the lattice-theoretic expression of a simple force: retained referrals are more attractive to a decision maker that places a higher value on future content reproduction. It does not require concavity, differentiability, unique optima, or a separated maximum at full diversion.

The increasing-differences assumption can be derived from the reproduction system itself. Let $\mathcal{S} \subseteq \mathbb{R}_+^d$ be the ordered topological state space for content and attention capital, with the product order and the usual topology. For each fixed state s , write $\Gamma_s(r)$ for the next-period content-and-attention state induced by retained referrals $r \in \mathcal{R}_s$.

Lemma 3 (Primitive source of increasing differences). *Fix s . Suppose $\mathcal{S} \subseteq \mathbb{R}_+^d$ is ordered by the product order and has the usual topology. Suppose $\mathcal{R}_s$ is a compact sublattice, and for each fixed state s let $\Gamma_s : \mathcal{R}_s \rightarrow \mathcal{S}$ be increasing and continuous. Suppose the baseline continuation term $\Phi(\cdot, s) + V^P(\Gamma_s(\cdot))$ is upper semicontinuous and supermodular on $\mathcal{R}_s$. Suppose also that $\Delta V(h) \equiv V^S(h) - V^P(h)$ is increasing and upper semicontinuous on $\mathcal{S}$, and that $\Delta V(\Gamma_s(\cdot))$ is supermodular on $\mathcal{R}_s$. Define, for $\lambda \in [0, 1]$,*

$$J^\lambda(r; s) = \Phi(r, s) + V^P(\Gamma_s(r)) + \lambda \Delta V(\Gamma_s(r)). \quad (7.3)$$

Then $J^\lambda(\cdot; s)$ is upper semicontinuous and supermodular on $\mathcal{R}_s$ for every λ , and J^λ has increasing differences in (r, λ) . Hence Assumption 6 holds with $\theta = \lambda$.

The lemma makes explicit what is needed topologically and directionally. It also explains the normalization in (7.1): take

$$\Phi(r, s) = \Phi^0(r, s) + V^P(\Gamma_s(r)), \quad \Psi(r, s) = \Delta V(\Gamma_s(r)),$$

where Φ^0 is current interface payoff and V^P is the platform's own continuation value. Upper semicontinuity is imposed directly on the baseline continuation term. Supermodularity after composition is maintained as a lattice condition rather than derived from monotonicity alone. The lemma derives the increasing-differences part from the monotonicity of Γ_s and of the incremental social-planner continuation value. For example, if ΔV is a nonnegative weighted sum of future costly-information stocks, source-level measurement precision, topic-support diversity, and public-information values, and if the referral-induced transition Γ_s makes each of these objects increasing and supermodular in retained referrals, then the high-level Topkis assumption follows from the reproduction and measurement primitives. Thus the AI platform need not be myopic; over-diversion follows from the AI platform valuing the future content-and-measurement state less than the social planner does.

The next definition selects among multiple optimal referral choices for a fixed state and continuation weight.

Definition 1 (Monotone tie-breaking). Given an argmax correspondence $M(\theta, s)$, a tie-breaking rule is a selection $\tau(\theta, s) \in M(\theta, s)$. It is monotone if, for fixed s and any $\theta' \geq \theta$, it satisfies

$\tau(\theta', s) \geq \tau(\theta, s)$ componentwise. Under Assumption 6, the least and greatest optimal selections are monotone tie-breaking rules by Topkis's theorem.

Proposition 5 (Endogenous over-diversion by underinternalized continuation value). *Under Assumption 6, the set of optimal referral choices*

$$M(\theta, s) = \arg \max_{r \in \mathcal{R}_s} J(r; \theta, s) \quad (7.4)$$

is nonempty and has least and greatest elements. These extremal selections are nondecreasing in θ . Hence if $\theta_S > \theta_P$, then

$$\min M(\theta_S, s) \geq \min M(\theta_P, s), \quad \max M(\theta_S, s) \geq \max M(\theta_P, s) \quad (7.5)$$

componentwise. Under any monotone tie-breaking rule, the AI platform's retained referrals satisfy $r^P \leq r^S$, and its diversion vector $1 - r^P$ is weakly larger than the social planner's diversion vector $1 - r^S$.

The proposition is the endogenous-diversion result. It says that the AI platform chooses too much answer substitution whenever it underweights the continuation value of the content-reproduction system. It does not assert that every AI platform chooses full diversion.

Corollary 6 (Current-rent benchmark). *Suppose $0 \in \mathcal{R}_s$, $\theta_P = 0$, and $\Phi(\cdot, s)$ has a separated maximum at the full-diversion point $r = 0$: for every $a > 0$,*

$$\Phi(0, s) > \sup\{\Phi(r, s) : r \in \mathcal{R}_s, \|r\|_\infty \geq a\}. \quad (7.6)$$

Then every approximately optimal private AI platform sequence satisfies $r_n \rightarrow 0$, or equivalently $z_n \rightarrow \mathbf{1}$. If Assumptions 1–3 hold for the induced publisher problem, publisher choices converge to zero content by Corollary 1. If Assumption 4 holds as well, they are eventually inactive.

The benchmark recovers the simple full-diversion case when the AI platform is evaluated only on current answer-interface rents. The proposition above is more general: even a forward-looking AI platform over-diverts relative to the social planner if the social planner internalizes additional continuation value.

The social planner side of the central theorem also needs a primitive viability condition. Let $\xi \in [0, 1]^n$ denote genuine-source visibility in conventional search, and define

$$\mathcal{T}(r, \xi) = D + \text{diag}(r \circ \xi)K. \quad (7.7)$$

The inequality $\rho(\mathcal{T}(r, \xi)) < 1$ is a sufficient condition for collapse for paths dominated by $\mathcal{T}(r, \xi)$. The reverse inequality is not a survival theorem. A sufficient preservation condition is instead the existence of a positive superfixed stock that the social planner can actually implement.

Proposition 6 (Planner viability from a superfixed stock). *Fix r and ξ . Suppose there exists $h^* \gg 0$ such that*

$$\mathcal{T}(r, \xi)h^* \geq h^*. \quad (7.8)$$

Suppose further that the social planner can implement a reinvestment policy whose induced content stock satisfies

$$h_{t+1} \geq \mathcal{T}(r, \xi)h_t \quad (7.9)$$

whenever $h_t \geq h^$. Then every path starting from $h_0 \geq h^*$ satisfies $h_t \geq h^*$ for all t . Hence the social planner can preserve a positive costly human-information stock.*

Corollary 7 (Augmented-state planner viability). *Let $y_t = (h_t, a_t)$ and fix r and ξ . Suppose there exists $y^* \gg 0$ such that*

$$\mathcal{M}(r \circ \xi)y^* \geq y^*. \quad (7.10)$$

Suppose further that the social planner can implement a reinvestment and attention-capital policy satisfying

$$y_{t+1} \geq \mathcal{M}(r \circ \xi)y_t \quad (7.11)$$

whenever $y_t \geq y^$. Then every path starting from $y_0 \geq y^*$ satisfies $y_t \geq y^*$ for all t . Hence the social planner can preserve both a positive costly human-information stock and a positive attention-capital stock.*

This proposition and corollary are deliberately sufficient conditions. It does not assert that a supercritical upper bound alone implies survival. It states exactly what must be added on the social side: an implementable reinvestment policy and a positive stock that is self-sustaining under that policy.

Effective monetization is $\mathbf{m} = r \circ \xi$. Section 6 derives the open-web reproduction number

$$\mathcal{R}^{web}(r, \xi) \equiv \rho(D + \text{diag}(r \circ \xi)K). \quad (7.12)$$

The inequality $\mathcal{R}^{web}(r, \xi) < 1$ is a sufficient condition for collapse of the maximal no-transfer reproduction system. The reverse inequality is not, by itself, a survival theorem. Given a private referral choice r^P , define the private reduced maximal reproduction operator

$$\mathcal{T}^P = D + \text{diag}(r^P \circ \xi)K. \quad (7.13)$$

Theorem 2 (Strategic over-diversion and inefficient subcriticality). *Suppose Assumption 6 holds, and let r^P and r^S be private and social referral choices selected by the same monotone tie-breaking rule from $M(\theta_P, s)$ and $M(\theta_S, s)$, with $\theta_S > \theta_P$. Fix a genuine-source visibility vector ξ and define $\mathcal{T}^P = D + \text{diag}(r^P \circ \xi)K$. Then:*

1. **Over-diversion.** *The AI platform retains weakly less referral traffic than the social planner: $r^P \leq r^S$.*

2. **Private collapse.** If $\mathcal{R}^{web}(r^P, \xi) < 1$, then every reduced no-transfer content path satisfying $h_{t+1} \leq \mathcal{T}^P h_t$ converges to zero. Equivalently, in the microfounded attention-capital system, every augmented no-transfer path dominated by $\mathcal{M}(r^P \circ \xi)$ converges to zero.
3. **Inefficiency.** If, at r^S , the sufficient viability condition in Proposition 6 holds for some $h^* \gg 0$, or the augmented-state condition in Corollary 7 holds for some $y^* \gg 0$, and the resulting positive-stock path yields social-planner value strictly above the private collapsing path, then the private AI platform choice is inefficient.

This is the paper’s main theorem. It turns a conditional collapse statement into a strategic market-design result. The private AI platform chooses a lower retained-referral vector than the social planner. When that choice moves the maximal web-reproduction upper bound below one, costly human information collapses. The social side is not inferred from $\mathcal{R}^{web} \geq 1$ alone; it is supplied by the superfixed-stock viability condition in Proposition 6 or, in the augmented attention-capital system, by Corollary 7. The theorem therefore remains one-sided where the mathematics is one-sided.

The theorem holds ξ fixed in order to isolate the referral-choice distortion. The measurement event enters the same framework through ξ and through the continuation value in J . Source-level visits, links, corrections, and repeat use are inputs into search visibility and into the information system that future AI answers rely on. Section 9 shows that AI diversion Blackwell-garbles these source-level experiments; the search and migration sections then study reduced-form cases in which loss of measurement lowers ξ . Thus the revenue channel and the measurement channel are separated for clarity, but both operate through the same effective-monetization and continuation-value objects.

For fixed-diversion welfare, define value components as mutually exclusive marginal values relative to the same counterfactual. Let $\mathcal{A}_x(q)$ be AI-answer value not already counted in direct visits or AI platform current payoff. Let $I_x(q)$ be source-level information value from visits, links, corrections, and other experiments. Normalize $\mathcal{A}_x(0_x) = I_x(0_x) = 0$. The social planner’s fixed- z payoff for topic x is

$$W_x(q; z) = (1 - z)R_x(q) + \mathcal{A}_x(q) + I_x(q) - C_x(q). \quad (7.14)$$

The publisher internalizes only residual visit revenue net of its own production cost.

Proposition 7 (Inefficient collapse from nonappropriable AI-use and signal value). *At a fixed diversion rate, no-transfer equilibrium can select zero content even when positive content is socially valuable. Formally, fix $z \in [0, 1]$. Suppose*

$$\sup_{q: \|q\|_x > 0} \{(1 - z)R_x(q) - C_x(q)\} < 0. \quad (7.15)$$

Then every exact private optimum has zero quality-adjusted commercial content. If there exists

$q^s \in Q_x$ with $\|q^s\|_x > 0$ such that

$$(1 - z)R_x(q^s) + \mathcal{A}_x(q^s) + I_x(q^s) - C_x(q^s) > 0, \quad (7.16)$$

then zero content is not socially optimal for the fixed- z social-planner problem. Thus the no-transfer outcome is inefficient even with truthful content, accurate AI answers, and rational users.

Corollary 8 (High-diversion inefficient collapse). *Suppose Assumptions 1, 2, and 4 hold for topic x . If there exists $q^s \in Q_x$ with $\|q^s\|_x > 0$ such that*

$$\mathcal{A}_x(q^s) + I_x(q^s) - C_x(q^s) > 0, \quad (7.17)$$

then for every z sufficiently close to one, every exact no-transfer private optimum has zero quality-adjusted commercial content, but zero content is not socially optimal.

The inequality in (7.15) is a private participation condition. The inequality in (7.16) is a social surplus condition. Their coexistence is the core missing market: AI answers and source-level measurements can be valuable, but the producer may be paid only for residual visits.

8 Diversity collapse

There are two diversity margins.¹² The first is topic support: which topics are commercially produced at all. The second is source diversity conditional on a topic: how attention is distributed among sources that cover the topic.

Define the active support at diversion z by

$$S_z = \{x \in X : q_z(x) \neq 0_x\}. \quad (8.1)$$

Under Proposition 1, $|S_z| \rightarrow 0$ in the finite topic economy, meaning every topic is eventually inactive as full diversion is approached. This is the minimal diversity result. If active topics disappear, topic diversity disappears. More generally, for any bounded social weight $\omega : X \rightarrow [0, \bar{\omega}]$, define weighted topic diversity by

$$D_\omega(z) = \sum_{x \in X} \omega(x) \mathbf{1}\{x \in S_z\}. \quad (8.2)$$

Corollary 9 (Weighted topic diversity collapse). *When active-topic collapse holds, every bounded weighted measure of active topic diversity converges to zero. Formally, under the assumptions of Proposition 1,*

$$D_\omega(z) \rightarrow 0 \quad (8.3)$$

¹²In this paper, “diversity” refers to a broad, unweighted notion of variety or heterogeneity. It does not denote demographic differences or the specialized uses of the term found in the author’s other work.

for every bounded ω .

This result is stronger than a statement about concentration among surviving sources. It says that commercially covered topics themselves vanish under no transfers.

The same logic gives a sharper social-composition result. Define the private exposure-cost ratio

$$\tau_x = \frac{F_x}{\bar{R}_x}, \quad (8.4)$$

with $\tau_x = \infty$ when $\bar{R}_x = 0 < F_x$. A topic with high τ_x has high avoidable activation cost relative to its maximum direct-visit revenue. These are the topics that exit first under diversion.

Proposition 8 (Diversion selects against high-cost, low-click topics). *Under Assumptions 1, 2, and 4, any active exact optimum at diversion z must satisfy*

$$\tau_x \leq 1 - z. \quad (8.5)$$

Hence the active set is contained in $\{x : \tau_x \leq 1 - z\}$. For any $a > 0$, all topics with $\tau_x > a$ are inactive whenever $z > 1 - a$.

If $V : X \rightarrow [0, \bar{V}]$ is any bounded social-value index and $H_a = \{x : \tau_x > a\}$, then the active social-planner value of the high-cost, low-click set is zero for all $z > 1 - a$:

$$\sum_{x \in H_a} V_x \mathbf{1}\{q_z(x) \neq 0_x\} = 0. \quad (8.6)$$

Thus diversion removes high- τ_x topics regardless of their social-planner value unless that social value is monetized through transfers or other non-visit revenue.

This proposition is the topic-by-topic analogue of Corollary 3: instead of imposing a uniform lower bound on $F_x/\bar{R}_x$, it identifies the topics that exit at each diversion level. It gives a more economic diversity statement. AI diversion does not merely reduce the number of active topics. It selects against topics with high fixed costs and low direct-click monetization, precisely the profile of many local, investigative, minority-language, technical-maintenance, and public-interest topics.

Source-concentration results require an additional assumption about how AI agents compress heterogeneous human search. Let $w = (w_1, \dots, w_L)$ be the old human attention distribution across the L sources within a topic, with $w_\ell > 0$ and $\sum_{\ell=1}^L w_\ell = 1$. A concentrated AI retrieval policy is represented by the power transform

$$T_\ell^\gamma(w) = \frac{w_\ell^\gamma}{\sum_{j=1}^L w_j^\gamma}, \quad \gamma \geq 1. \quad (8.7)$$

Let Shannon entropy be $\text{Ent}(p) = -\sum_{\ell=1}^L p_\ell \log p_\ell$ for $p \in \Delta^L$.

Lemma 4 (Source concentration lowers entropy). *Fix a topic x and let J_x be the number of sources covering that topic. Let $w = (w_1, \dots, w_{J_x})$ be the old human attention distribution across*

these sources, with $w_j > 0$ and $\sum_{j=1}^{J_x} w_j = 1$. A concentrated AI retrieval policy is represented by the power transform

$$T_j^\gamma(w) = \frac{w_j^\gamma}{\sum_{\ell=1}^{J_x} w_\ell^\gamma}, \quad \gamma \geq 1. \quad (8.8)$$

Let Shannon entropy be $\text{Ent}(p) = -\sum_{j=1}^{J_x} p_j \log p_j$ for $p \in \Delta^{J_x}$. Then

$$\frac{d}{d\gamma} \text{Ent}(T^\gamma(w)) = -\gamma \text{Var}_{T^\gamma(w)}(\log w_j) \leq 0. \quad (8.9)$$

The inequality is strict unless the components of w are all equal. Hence $\text{Ent}(T^\gamma(w)) < \text{Ent}(w)$ for $\gamma > 1$ unless w is uniform.

The lemma is not needed for the support-collapse theorem. It formalizes a pre-collapse force: when many heterogeneous human search paths are replaced by a more concentrated agentic retrieval rule, source attention becomes less entropic even before fixed costs induce exit.

The diversity result is therefore conditional on how AI gathers information. Let μ be the source-attention distribution generated by heterogeneous human search in a topic, and let ν be the distribution generated by an AI retrieval policy or by a mixture of AI agents. If ν is more concentrated than μ in the sense of majorization, entropy falls. If instead AI agents represent diverse interests so that their mixture preserves μ , or if the AI mixture is more equal than μ , this source-concentration channel does not imply diversity loss.

Proposition 9 (Diversity loss is a concentration result). *Let $\mu, \nu \in \Delta^L$ be source-attention distributions. If ν majorizes μ , then*

$$-\sum_{\ell=1}^L \nu_\ell \log \nu_\ell \leq -\sum_{\ell=1}^L \mu_\ell \log \mu_\ell. \quad (8.10)$$

The inequality is strict unless μ and ν differ only by a permutation. Consequently, AI reduces entropy diversity through this channel only when its information-gathering policy is more concentrated than the heterogeneous human benchmark. A sufficiently diverse mixture of AI agents need not reduce source diversity.

This proposition matters for interpretation and policy. Diversity collapse follows from concentration in retrieval, from topic-specific participation constraints, or from both. It is not implied by the word “AI” alone. If AI systems preserve heterogeneous interests, assign exploration probability to tail sources, or protect underrepresented topics through hard constraints, the diversity channel can be weakened or reversed.

9 Source-level revealed-preference signal collapse

The old web generated measurements. Each visit could reveal source quality or user preference through clicks, dwell time, repeat visits, subscriptions, comments, corrections, shares, and bookmarks. AI diversion can preserve answer-level satisfaction while deleting source-level observations.

Let Ω_x be the state space for source quality or demand in topic x . At $z = 0$, human visits generate an experiment $\mathcal{E}_x^0$ about $\vartheta \in \Omega_x$. Formally, $\mathcal{E}_x^0$ is a Markov kernel from Ω_x to finite observation multisets, and the number of observations has finite expectation under every state. At diversion z , each observation that would have been produced under $z = 0$ is retained with probability $1 - z$ and otherwise deleted, independently conditional on the old observation multiset. Let the resulting experiment be $\mathcal{E}_x^z$.

Recall that experiment $\mathcal{E}$ Blackwell dominates experiment $\mathcal{F}$, written $\mathcal{E} \succeq_B \mathcal{F}$, if $\mathcal{F}$ can be generated from $\mathcal{E}$ by a stochastic transformation independent of the state.

Theorem 3 (Blackwell signal collapse). *Higher AI diversion makes the source-level revealed-preference experiment less informative in the Blackwell order, and full diversion eliminates the old source-level experiment. Formally, if $0 \leq z_1 < z_2 \leq 1$, then*

$$\mathcal{E}_x^{z_1} \succeq_B \mathcal{E}_x^{z_2}. \quad (9.1)$$

Therefore the value of source-level revealed-preference data is weakly decreasing in z for every decision problem with bounded payoffs. As $z \rightarrow 1$, $\mathcal{E}_x^z$ converges pointwise in total variation under every state to the null experiment¹³ that contains no old-web visit observations. If the old observation count has uniformly bounded expectation over states, the convergence is uniform over states.

For search-ranking applications, the Blackwell implication is a statement about statistical decision problems with bounded measurable payoffs. If a ranker has an unbounded objective or unbounded loss, the garbling order still holds, but welfare monotonicity requires the usual integrability or uniform-integrability conditions. The total-variation convergence statement is independent of the ranker’s objective.

A parametric version makes the implication transparent. If source i receives n_i^0 old-web observations and each observation is $s_{i\ell} = \vartheta_i + \varepsilon_{i\ell}$ with variance σ_i^2 , then under diversion z the expected contribution to posterior precision is

$$\frac{(1 - z)n_i^0}{\sigma_i^2}, \quad (9.2)$$

which converges to zero. Independent Gaussian observations add precision, so this expression is

¹³Total variation distance is the maximum difference in probability assigned to any event by two distributions. Convergence to the null experiment means that, as diversion approaches one, the probability of observing any old-web source-level interaction converges to zero.

the expected source-level information contribution of retained visits.

AI answer feedback need not solve the identification problem. Let Z^{ans} denote answer-level feedback, such as a thumbs-up, a reformulation, session continuation, or subscription retention. Suppose Z^{ans} has distribution $P_{ans}(\cdot \mid \vartheta)$.

Proposition 10 (Answer-level feedback need not identify source quality). *Answer-level satisfaction is not generally a substitute for source-level revealed preference. This is a simple nonidentification observation. If there exist $\vartheta, \vartheta' \in \Omega_x$ such that*

$$P_{ans}(\cdot \mid \vartheta) = P_{ans}(\cdot \mid \vartheta') \quad (9.3)$$

but $\vartheta_i \neq \vartheta'_i$ for some source i , then source i 's quality is not identified from answer-level feedback when direct source-level observations vanish.

The condition is mild. When an answer blends sources, answer satisfaction can identify that the answer was useful without identifying which source was original, stale, copied, decisive, or wrong.

10 Link-based and modern search ranking

Link-based ranking is another revealed-preference system. A human or institution reads, trusts, cites, links, bookmarks, or references a source. The resulting graph is then used by search engines like Google Search. The original PageRank insight was that hyperlinks contain decentralized information about authority or quality (Brin and Page, 1998). But PageRank is not the whole of modern Google Search or other mainstream web search. Public Google documentation describes a many-signal ranking architecture: language-understanding systems such as BERT and neural matching, relevance, quality, usability, freshness, original-content systems, interaction data, source expertise, context, and link analysis. PageRank remains part of Google's core ranking systems, but it is one component of a broader system rather than the single key algorithm (Google Search Central, 2025; Google Search, 2026).

The result below uses PageRank as the mathematical benchmark because it is the canonical graph-ranking operator (and perhaps best known). My result's economic content is broader. Any search system that relies on source-level signals generated by genuine web use becomes less informative when those signals are thinned by AI diversion or diluted by cheap synthetic entry.

The ranking results in this section are unanchored-ranking benchmarks. They are not claims that every modern search engine must lose genuine-source exposure. They identify what happens when the ranking system depends on contemporaneous source-level signals and does not reserve sufficient protected exposure for verified sources. Trust anchors, provenance rules, manual verification, or protected exploration can break the limiting collapse result.

Let $\mathcal{L}_x^0$ be the old-web experiment generated by current human links in topic x . As in the

previous section, assume the old current-link experiment produces a finite observation multiset with finite expected size under every state. Under AI diversion, suppose current link opportunities are thinned with the human visits that would have generated them. Let $\mathcal{L}_x^z$ be the retained current-link experiment.

Corollary 10 (Current-link signal collapse). *AI diversion Blackwell-garbles the current-link experiment used by link-based ranking systems. Formally, if $0 \leq z_1 < z_2 \leq 1$, then*

$$\mathcal{L}_x^{z_1} \succeq_B \mathcal{L}_x^{z_2}. \quad (10.1)$$

As $z \rightarrow 1$, the current-link experiment converges pointwise in total variation under every state to the empty new-link experiment. If the old current-link count has uniformly bounded expectation over states, the convergence is uniform over states. For any bounded ranking decision problem whose payoff depends on the current-link experiment, Blackwell monotonicity implies weakly lower value as diversion rises. The total-variation convergence statement itself does not require bounded payoffs. Hence any ranking rule whose information about current quality comes only through new human links loses that information as diversion approaches one.

A common parametric case is Poisson link formation. Suppose new links to source i satisfy

$$L_i(z) \sim \text{Poisson}((1-z)\lambda_i(q_i)), \quad (10.2)$$

where $\lambda_i(q_i)$ is increasing in current quality. If λ_i is differentiable, the Fisher information about q_i is

$$\mathcal{I}_i^{\text{link}}(z) = (1-z) \frac{[\lambda_i'(q_i)]^2}{\lambda_i(q_i)}, \quad (10.3)$$

which converges linearly to zero. Thus PageRank-like scores can remain informative about inherited authority while becoming stale with respect to current quality.

A PageRank statement is not enough for current search. Modern search systems use language models, quality classifiers, freshness systems, interaction data, site-wide signals, and link analysis. The next result states the same idea at this level of abstraction. Let $G_t \in \mathbb{R}_+^k$ denote contemporaneous genuine-source signals generated by human web activity: new links, citations, clicks, repeat visits, subscriptions, freshness updates, corrections, and other source-level interactions. Let $Y_t \in \mathcal{Y}$ denote all remaining signals, including query text, page text, static metadata, and AI platform-side features, where $\mathcal{Y}$ is an arbitrary measurable signal space. Let

$$\Xi_G : \mathbb{R}_+^k \times \mathcal{Y} \rightarrow [0, 1]$$

be the exposure assigned to genuine sources. Let $\|\cdot\|$ be any fixed norm on $\mathbb{R}^k$. Say that a ranking system is locally source-signal-dependent if there exist $L < \infty$ and a nonnegative anchor function

$\alpha_0 : \mathcal{Y} \rightarrow [0, 1]$ such that, in a neighborhood of $G = 0$,

$$0 \leq \Xi_G(G, Y) \leq \alpha_0(Y) + L\|G\|. \quad (10.4)$$

The term $\alpha_0(Y)$ is exposure protected by information that does not require contemporaneous genuine web activity, such as a verified-source prior, a fixed trust anchor, or an explicit provenance rule.

Proposition 11 (Modern ranking signal collapse). *If a conventional search system is locally source-signal-dependent and AI diversion, or synthetic dilution that also drives absolute contemporaneous genuine-source signals to zero, implies $\|G_t\| \rightarrow 0$, then the ranking system’s anchor is an asymptotic upper bound on genuine-source exposure:*

$$\limsup_{t \rightarrow \infty} \Xi_G(G_t, Y_t) \leq \limsup_{t \rightarrow \infty} \alpha_0(Y_t). \quad (10.5)$$

In particular, in an unanchored environment with $\alpha_0(Y_t) \rightarrow 0$, genuine-source exposure converges to zero. If a fixed trust anchor keeps $\alpha_0(Y_t)$ bounded away from zero, complete exposure collapse need not occur.

This proposition is not a claim that semantic matching systems such as BERT or neural matching fail. Those systems can improve the matching of queries to text. The proposition concerns the source-exposure problem: when cheap synthetic content can mimic text and genuine source-level signals vanish, a ranker without anchors loses the input that distinguishes costly human information from low-cost imitations. If synthetic entry leaves $\|G_t\|$ bounded away from zero but swamps it with a much larger synthetic signal mass, that is a relative signal-to-noise or ranking-dilution problem rather than an implication of Proposition 11; the PageRank dilution benchmark below is one way to model that different channel. Anchors, audits, provenance, or verified-source priors can prevent the limiting zero-exposure conclusion.

The unanchored environment above is most relevant when three conditions hold: (1) synthetic pages can mimic the textual and graph features used by the ranker; (2) contemporaneous source-level signals such as new links, repeat visits, corrections, and subscriptions are important inputs into exposure; and (3) the ranker does not reserve sufficient exposure for verified sources through hard provenance, editorial rules, or exploration. If any of these conditions fails, the limiting zero-exposure conclusion need not follow. The theorem is therefore a diagnostic for the institutional role of anchors, not a prediction that all modern search must collapse.

11 Cheap AI content, synthetic entry, and ranking dilution

The baseline theorem studies the demand-side effect of AI: users obtain answers without visiting sources. Generative AI also creates a supply-side effect: it lowers the cost of producing websites,

pages, summaries, images, reviews, and link structures. This section shows that the two effects do not offset one another. Cheap AI content can increase raw page counts while accelerating the collapse of costly human-information content and weakening search-ranking exposure for genuine sources. PageRank is again the canonical graph case; the modern-ranking result above explains why the same logic applies to broader search systems that rely on genuine source-level signals.

The key distinction is between raw content and costly human-information content. A synthetic page can be cheap without containing new reporting, testing, measurement, verification, maintenance, or local knowledge. Thus the relevant collapse object is not the number of URLs. It is the stock of positive-cost information that must be financed.

Definition 2 (Human-information content and synthetic page volume). A plan $q \in Q_x$ has two nonnegative indices,

$$\mathcal{H}_x(q) \geq 0, \quad Y_x(q) \geq 0. \quad (11.1)$$

The index $\mathcal{H}_x(q)$ is costly human-information content: the component of content that requires scarce non-AI inputs such as original reporting, experiments, expert judgment, local observation, product testing, editorial verification, or fieldwork. The index $Y_x(q)$ is raw synthetic page volume: low-cost transformation, duplication, templating, paraphrase, summarization, or synthetic recombination. The theorem does not assume synthetic content has zero value. It assumes only that abundance of Y_x is not the same object as reproduction of costly information $\mathcal{H}_x$.

The distinction is between creating more pages and creating new costly facts. The proposition below permits synthetic content to be useful; it only says that synthetic volume is not a substitute for information that requires scarce human or institutional inputs.

Let $\kappa \geq 0$ be the synthetic-cost shifter. Let $\xi \in [0, 1]$ be an exposure multiplier: conditional on not being replaced by the AI answer, ξ is the share of old direct-visit monetization that remains available to costly sources after synthetic-entry dilution. The payoff is

$$\pi_x(q; z, \xi, \kappa) = (1 - z)\xi R_x(q) - C_x^H(q) - \kappa C_x^S(q), \quad (11.2)$$

where $0 \leq R_x(q) \leq \bar{R}_x < \infty$, $C_x^H(q) \geq 0$, $C_x^S(q) \geq 0$, and the zero plan gives zero payoff.

The only substantive new assumption is the same no-free-positive-content condition, now applied to the human-information component.

Assumption 7 (No free positive human information). For every x and every $\delta > 0$,

$$\underline{C}_x^H(\delta) \equiv \inf\{C_x^H(q) : q \in Q_x, \mathcal{H}_x(q) \geq \delta\} > 0. \quad (11.3)$$

Human-information content cannot be made positive for free. AI may lower the cost of drafting, paraphrasing, or templating, but original measurement, testing, verification, reporting, and maintenance still require scarce inputs.

Proposition 12 (Cheap synthetic content does not prevent human-information collapse). *If the combined exposure and monetization share for costly sources vanishes, then costly human-information content collapses even when synthetic content is cheap. Formally, suppose Assumption 7 holds. Let (z_n, ξ_n, κ_n) be any sequence such that*

$$m_n \equiv (1 - z_n)\xi_n \rightarrow 0. \quad (11.4)$$

Let $\varepsilon_n \rightarrow 0$, and let q_n be ε_n -optimal for (11.2). Then

$$\mathcal{H}_x(q_n) \rightarrow 0. \quad (11.5)$$

This proposition extends the baseline result. If $\xi_n = 1$ and the synthetic margin is absent, it reduces to the same participation-constraint force as Lemma 1; with arbitrary κ_n , it also permits cheap synthetic content. It shows why cheap AI publishing can make collapse happen earlier. The relevant scalar is not only the nondiverted visit share, but the product of nondiverted visits and remaining exposure after synthetic dilution. If synthetic entry drives ξ to zero, human-information collapse follows even when AI answer diversion is bounded away from one.

Corollary 11 (Raw pages can rise while human information collapses). *Under Assumption 7, the number of cheap pages can diverge even while every approximately optimal plan contains vanishing costly human information. Formally, suppose $m_n \rightarrow 0$ and there exist synthetic plans $s_n \in Q_x$ such that*

$$\mathcal{H}_x(s_n) = 0, \quad C_x^H(s_n) = 0, \quad Y_x(s_n) \rightarrow \infty, \quad \kappa_n C_x^S(s_n) \rightarrow 0. \quad (11.6)$$

Then the high-volume synthetic sequence s_n is $o(1)$ -optimal, while Proposition 12 implies that every $o(1)$ -optimal sequence has $\mathcal{H}_x(q_n) \rightarrow 0$.

The corollary explains why collapse may be hard to diagnose from page counts. AI can produce a large web in the literal URL-count sense and a collapsed web in the economically relevant sense. The observable web becomes bigger and thinner: more pages, fewer costly facts.

Proposition 13 (Cheap synthetic free entry and the abundance paradox). *Cheap synthetic production can make the indexed web grow without preserving costly human information. Let $Y \geq 0$ be the mass of synthetic pages. Suppose a representative synthetic page earns gross visibility payoff $\varphi(Y)$, where φ is continuous, strictly decreasing, $\varphi(Y) > 0$ for every finite Y , and $\varphi(Y) \rightarrow 0$ as $Y \rightarrow \infty$. If the per-page synthetic cost is $\kappa_n \downarrow 0$ and the free-entry synthetic mass Y_n solves*

$$\varphi(Y_n) = \kappa_n, \quad (11.7)$$

then $Y_n \rightarrow \infty$. If, in addition, along the same sequence the hypotheses of Proposition 12 hold for a publisher sequence q_n —in particular $m_n = (1 - z_n)\xi_n \rightarrow 0$, $\varepsilon_n \rightarrow 0$, and q_n is ε_n -optimal for (11.2)—then raw synthetic volume diverges while costly human-information content converges to zero.

This is the abundance paradox. Making synthetic pages cheap increases the mass of pages until their per-page visibility payoff is competed down to the tiny synthetic cost. But that entry margin is not the costly human-information margin. The web can therefore become larger in page count and smaller in reproduced knowledge.

Cheap pages also change the graph on which link-analysis algorithms operate. The next result is stated for PageRank because PageRank has an exact Markov-chain representation. The theorem should be read as a canonical case of a broader search-ranking problem: unanchored ranking systems can assign exposure according to raw indexed volume and graph structure unless they deliberately reserve weight for verified genuine sources. Let G_n be the set of genuine or human-information pages and Syn_n the set of synthetic pages. Let P_n be the row-stochastic link-transition matrix on $\text{Syn}_n \cup G_n$, meaning that each row gives the probability of moving from one page to another and sums to one. For damping parameter $\alpha \in (0, 1)$ and teleportation distribution v_n , PageRank is the unique probability vector π_n satisfying¹⁴

$$\pi_n = (1 - \alpha)v_n + \alpha\pi_n P_n. \quad (11.8)$$

The first term is baseline teleportation attention; the second term is link-following attention. Define

$$\delta_n \equiv \sup_{i \in \text{Syn}_n} P_n(i, G_n), \quad (11.9)$$

where $P_n(i, G_n)$ is the probability that a random-walk step from synthetic page i moves to a genuine page. Thus small δ_n means that random walks starting from synthetic pages rarely move to genuine pages.

The PageRank results should be read as a transparent graph-ranking benchmark. The strong conclusions require two institutional/technological conditions: the teleportation prior does not reserve mass for verified genuine sources, and synthetic subgraphs are sufficiently closed that random walks starting in synthetic clusters rarely reach genuine pages. These conditions are useful because they identify the failure mode; they are not asserted to hold in every search ecosystem.

Remark 1 (Lipschitz continuity of PageRank in the transition matrix). Fix damping parameter $\alpha \in (0, 1)$ and teleportation distribution v . For any two row-stochastic transition matrices P and Q on the same finite state space, their PageRank vectors satisfy

$$\|\pi(P) - \pi(Q)\|_1 \leq \frac{\alpha}{1 - \alpha} \max_i \|P(i, \cdot) - Q(i, \cdot)\|_1. \quad (11.10)$$

¹⁴The PageRank literature calls v_n the teleportation distribution. In the random-surfer interpretation, a user follows links according to P_n with probability α and, with probability $1 - \alpha$, “teleports” to a page drawn from v_n . Economically, v_n is an exogenous baseline attention prior over indexed pages: it determines where attention restarts when the process stops following links. If v_n is approximately uniform over indexed pages, then indexed volume directly affects baseline PageRank mass. If v_n instead places protected mass on verified sources, that is a trust anchor.

Small perturbations of the link-transition matrix create bounded perturbations in PageRank. This lets the paper state that when synthetic links dominate the transition matrix, PageRank becomes close to the synthetic-link PageRank and insensitive to current human quality.

A second formulation uses mixture contamination. Suppose the observed transition matrix is

$$P_n = (1 - \omega_n)P_n^H + \omega_n P_n^S, \quad (11.11)$$

where P_n^H is the human link matrix, P_n^S is a synthetic link matrix independent of current human quality, and $\omega_n \rightarrow 1$. Since

$$\max_i \|P_n(i, \cdot) - P_n^S(i, \cdot)\|_1 \leq 2(1 - \omega_n) \rightarrow 0, \quad (11.12)$$

Remark 1 implies

$$\|\pi(P_n) - \pi(P_n^S)\|_1 \rightarrow 0. \quad (11.13)$$

If P_n^S is independent of current human quality, PageRank becomes asymptotically insensitive to current human quality. Baseline AI diversion starves new human links. Cheap synthetic entry floods the graph with low-cost links and nodes.

The implication is not that PageRank is useless, nor that modern search is literally PageRank. It is that the original self-scaling intuition breaks. When publication and linking are cheap enough, link analysis and modern multi-signal ranking must be supplemented by costly trust anchors, provenance, audits, human-source priors, or payments for verified information. Otherwise ranking measures the equilibrium supply of cheap pages, cheap links, and cheap signals rather than the equilibrium supply of costly knowledge.

Lemma 5 (Volume-based teleportation and synthetic dominance). *The assumption that synthetic pages dominate the teleportation prior follows from raw URL dominance under weak per-page weight bounds. Suppose*

$$v_n(i) = \frac{w_{ni}}{\sum_{j \in \text{Syn}_n \cup G_n} w_{nj}}, \quad (11.14)$$

where $0 < w_{ni} < \infty$. If there exist constants $\bar{w}_G < \infty$ and $\underline{w}_S > 0$ such that $w_{ni} \leq \bar{w}_G$ for $i \in G_n$ and $w_{ni} \geq \underline{w}_S$ for $i \in \text{Syn}_n$, then

$$v_n(G_n) \leq \frac{\bar{w}_G |G_n|}{\bar{w}_G |G_n| + \underline{w}_S |\text{Syn}_n|}. \quad (11.15)$$

Hence $v_n(G_n) \rightarrow 0$ whenever $|\text{Syn}_n|/|G_n| \rightarrow \infty$.

The teleportation prior is not assumed to favor synthetic pages by fiat. It does so in unanchored search environments when the prior is approximately volume-based over indexed URLs, or more generally when per-page teleportation weights are not allowed to offset an exploding synthetic-to-genuine page ratio. The result can fail if search deliberately anchors a fixed teleportation mass on verified genuine sources.

Proposition 14 (Synthetic PageRank dilution). *If cheap synthetic pages dominate the teleportation prior and synthetic subgraphs are nearly closed, then PageRank mass on genuine pages vanishes. Formally, for every n ,*

$$\pi_n(G_n) \leq v_n(G_n) + \frac{\alpha}{1-\alpha} \delta_n. \quad (11.16)$$

Consequently, if $v_n(G_n) \rightarrow 0$ and $\delta_n \rightarrow 0$, then $\pi_n(G_n) \rightarrow 0$. Every genuine page $j \in G_n$ then has $\pi_n(j) \leq \pi_n(G_n) \rightarrow 0$.

The result does not require synthetic pages to be false or adversarial. It requires them to be numerous in the teleportation prior and graph-closed enough that random walks started in synthetic clusters rarely reach genuine pages. If teleportation is approximately uniform over indexed pages, then

$$v_n(G_n) = \frac{|G_n|}{|\text{Syn}_n| + |G_n|}, \quad (11.17)$$

so $v_n(G_n) \rightarrow 0$ whenever $|\text{Syn}_n|/|G_n| \rightarrow \infty$.

Proposition 15 (Trust anchors prevent complete PageRank mass loss). *If the teleportation prior assigns uniformly positive mass to genuine pages, $v_n(G_n) \geq \underline{v} > 0$, then*

$$\pi_n(G_n) \geq (1-\alpha)\underline{v}. \quad (11.18)$$

Thus synthetic entry alone does not force PageRank mass on genuine pages to vanish when the ranking system deliberately anchors attention on verified genuine sources.

The PageRank dilution result is an unanchored-ranking result, not an impossibility theorem for all ranking systems. Fixed trust priors, verified-source teleportation mass, provenance, or other costly anchors can preserve some ranking mass for genuine pages. The same caveat applies to modern ranking systems. Without anchors, raw-volume synthetic entry can dominate the graph or the source-level signal environment; with anchors, the ranker can reserve exposure for verified genuine sources.

Proposition 16 (Ranking-feedback collapse). *PageRank dilution can turn synthetic abundance into content collapse. Consider a sequence of economies in which publisher payoff for a genuine topic is*

$$U_n(q) = (1 - z_n)\xi_n R_x(q) - C_x(q), \quad (11.19)$$

where $\xi_n \in [0, 1]$ is the exposure share delivered to genuine pages by a PageRank-based discovery system. Suppose $\xi_n \leq L_\xi \pi_n(G_n)$ for some finite L_ξ , where $\pi_n(G_n)$ is PageRank mass on genuine pages. If $v_n(G_n) \rightarrow 0$ and $\delta_n \rightarrow 0$, then every $o(1)$ -optimal genuine-publisher plan has vanishing quality-adjusted content under Assumptions 1–3. If, in addition, the expected number of new human links generated by the topic is bounded above by a constant times $(1 - z_n)\xi_n$, then the current human-link experiment also converges pointwise in total variation to the null experiment.

The theorem closes the feedback loop. Synthetic pages dilute the ranking mass of genuine pages. Lower ranking mass reduces genuine-source exposure. Lower exposure reduces publisher revenue and content investment. Lower nondiverted exposure also reduces the new human links that would otherwise replenish the ranking system. The result can occur even when answer diversion z_n is bounded away from one, because ranking exposure ξ_n itself converges to zero.

12 Search degradation and AI migration

This section studies a tipping benchmark. The contraction condition below is not a universal law of user behavior or search technology. It describes environments in which conventional search quality is maintained by current genuine-source activity and, once that activity becomes too thin, user migration toward AI further reduces the activity needed to rebuild search quality.

The previous result maps a given search-exposure sequence into publisher collapse. There is a further demand-side feedback. If conventional search becomes less useful because genuine sources are harder to find, users have a stronger reason to rely on the AI answer interface. That user migration raises diversion, which removes more visits, links, and source-level feedback from conventional search.

Let $\xi_t \in [0, 1]$ denote the usefulness of conventional search for genuine costly-information pages at date t . It can be interpreted as genuine-source exposure, source-level ranking quality, or the probability that a conventional search session reaches a useful genuine source. Let $z_t \in [0, 1]$ be the AI diversion share chosen by users or by the AI platform in response to that search usefulness. The effective monetization term in the publisher problem is then $m_t = (1 - z_t)\xi_t$. The migration function below maps search usefulness into AI diversion; the updating function maps today's remaining nondiverted genuine search activity into tomorrow's search usefulness.

Assumption 8 (AI migration and search-quality updating). There exist functions $\zeta : [0, 1] \rightarrow [0, 1]$ and $\phi : [0, 1] \rightarrow [0, 1]$ such that

$$z_t = \zeta(\xi_t), \quad \xi_{t+1} \leq \phi((1 - z_t)\xi_t). \quad (12.1)$$

The migration function ζ is weakly decreasing and continuous at zero. The updating function satisfies $\phi(0) = 0$. Finally, there exist $\bar{\xi} > 0$ and $\lambda \in (0, 1)$ such that, for every $s \in [0, \bar{\xi}]$,

$$\phi((1 - \zeta(s))s) \leq \lambda s. \quad (12.2)$$

The first part says that users move toward AI when conventional search is less useful. The second part says that tomorrow's search quality is maintained by today's nondiverted genuine search activity. The last part is a local tipping condition: once search quality is low enough, the remaining non-AI activity is insufficient to rebuild the ranking and feedback system.

Proposition 17 (Search degradation and AI migration create a tipping loop). *If conventional search usefulness falls below the tipping region, AI migration and search degradation reinforce each other. Formally, suppose Assumption 8 holds and $\xi_0 \leq \bar{\xi}$. Then*

$$\xi_t \leq \lambda^t \xi_0 \rightarrow 0, \quad (12.3)$$

and $z_t \rightarrow \zeta(0)$, the maximum-diversion value induced by vanishing conventional-search usefulness. If publisher payoff at date t is

$$U_t(q) = (1 - z_t)\xi_t R_x(q) - C_x(q), \quad (12.4)$$

then every $o(1)$ -optimal publisher sequence satisfies $\|q_t\|_x \rightarrow 0$ under Assumptions 1–3. If Assumption 4 also holds, every sufficiently accurate publisher optimum is inactive for all large t . If the expected number of new human links is bounded above by a constant times $(1 - z_t)\xi_t$, then the current-link experiment converges pointwise in total variation to the null experiment.

This is the positive feedback loop. Worse conventional search sends more users to AI. More AI use removes the visits and links that would have improved conventional search. Once the system enters the contraction region, effective publisher monetization and source-level link measurement both converge to zero.

The next proposition collects the modular channels into one compact benchmark. It is deliberately a benchmark: it links the search-quality contraction to the reproduction system without claiming that all search ecosystems must satisfy the contraction condition.

Proposition 18 (Integrated feedback benchmark). *Let $y_t = (h_t, a_t)$ be the augmented content-and-attention state. Suppose Assumption 8 holds, $\xi_0 \leq \bar{\xi}$, $z_t = \zeta(\xi_t)$, and $\mathbf{m}_t = ((1 - z_t)\xi_t)\mathbf{1}$. Suppose the no-transfer augmented state is dominated by*

$$y_{t+1} \leq \mathcal{M}(\mathbf{m}_t)y_t. \quad (12.5)$$

Then $\xi_t \rightarrow 0$, $z_t \rightarrow \zeta(0)$, and $\mathbf{m}_t \rightarrow 0$. Moreover, because $\rho(D) < 1$, there exists $T < \infty$ and a vector $\bar{\mathbf{m}} \gg 0$ with $\mathbf{m}_t \leq \bar{\mathbf{m}}$ for all $t \geq T$ and $\mathcal{R}^{web}(\bar{\mathbf{m}}) < 1$. Hence $h_t \rightarrow 0$ and $a_t \rightarrow 0$ for every initial augmented state. If contemporaneous genuine-source signals satisfy $\|G_t\| \leq L_G\|y_t\|$ for some finite L_G , then $\|G_t\| \rightarrow 0$ as well.

The integrated benchmark shows how the revenue and measurement channels can reinforce one another along a single path: reduced search usefulness raises AI reliance, higher AI reliance lowers effective monetization, the reproduction operator becomes subcritical, and the source-level signals generated by the surviving content stock vanish. The result remains conditional on the maintained search-quality contraction and on the absence of hard trust anchors.

13 Market design: moving the revenue and measurement events

If AI answers replace visits, a sustainable institution must move revenue from the replaced publisher visit to the AI-mediated interaction. It must also recreate measurement at the new point of attention.¹⁵ If AI answers satisfy users while deleting source-level observations, a payment rule based only on answer-level outcomes cannot by itself distinguish costly human information from low-cost synthetic imitation.

In practice, this mechanism would likely not be implementable by a single firm or government acting alone. It would likely require a shared institution—a clearinghouse, standards body, or public-private arrangement involving AI platforms, publishers, search intermediaries, independent auditors, and possibly regulators that can measure displaced visits, verify provenance, run audits and exploration, and settle payments.

I outline key market design constraints rather than a complete institutional mechanism. I identify what a feasible institution must measure and what incentive constraints it must satisfy. The paper’s “negative” results say that ordinary visits and links cease to provide source-level experiments, while the repair requires new source-level experiments. Thus, the relevant question is not whether answer-level feedback alone suffices—it generally does not—but whether auditable experiments with enough informativeness to support transfers can be organized.¹⁶ A full implementation could require platform participation or legal authority, and dispute resolution over source attribution. I separate those institutional constraints from the benchmark formulas because the formulas identify the objects that the institution must be designed to recover but not sufficient conditions in practice.

The answer is not straightforward. For instance, a citation pool is not enough. Citations are chosen by the AI system, are endogenous to retrieval and presentation rules, and need not measure either displaced visits or marginal contribution to answer quality. A use-based pool is also not enough if use is concentrated on already visible sources or on synthetic content optimized to be cited. The institution must replace the displaced revenue event and restore a source-level measurement event.¹⁷

Let $d_x \in [0, 1]$ be the effective displaced-revenue share in topic x : the fraction of counterfactual monetizable publisher exposure, measured in the same units as $R_x(q)$, that is displaced by AI-mediated interaction. Thus d_x is not necessarily the raw probability that an arbitrary query

¹⁵Thus the mechanism should use AI-side feedback rather than ignore it. Randomized source holdouts, answer-level A/B tests, retrieval logs, user feedback, correction data, and citation/provenance audits can all help estimate source contribution. The market design problem is to convert these platform-side signals into auditable source-level measurements that can support transfers and screen costly human information from low-cost imitation.

¹⁶e.g., using inputs like platform telemetry, retrieval logs, randomized AI-answer holdouts, source-withholding tests, provenance audits, or exploration.

¹⁷In other words, revenue sufficient to finance new costly information and measurement sufficient to distinguish genuine sources from synthetic imitation.

receives an AI answer. Let $T_x(q, d_x)$ be an AI-use transfer. The compensated payoff is

$$\pi_x^T(q; d_x) = (1 - d_x)R_x(q) + T_x(q, d_x) - C_x(q). \quad (13.1)$$

13.1 A revenue-and-measurement replacement mechanism

Below, I will call the relevant institution a revenue-and-measurement replacement mechanism. Its purpose is not to recreate the old advertising web. Its purpose is to move the revenue event and measurement event to the point where user attention now occurs. In the model, such a mechanism has five elements.

- It measures displacement. The mechanism must estimate d_x , the probability that an AI answer displaced a publisher visit in topic x . In data, this can be estimated using randomized AI-answer holdouts, query-class variation, browser or device variation, platform logs, or changes in crawler and licensing policies.
- It pays a visitor-replacement royalty. The benchmark payment is $d_x R_x(q)$. This restores the private publisher problem that would have obtained if the displaced visit had still occurred. Proposition 20 formalizes this benchmark.
- It uses audited provenance. For each AI answer y , let $\mathcal{I}(y)$ be the finite set of candidate sources used or audited for that answer. The mechanism constructs a source-level provenance vector

$$s(y) = \{s_j(y)\}_{j \in \mathcal{I}(y)}, \quad s_j(y) \geq 0, \quad \sum_{j \in \mathcal{I}(y)} s_j(y) = 1, \quad (13.2)$$

where $s_j(y)$ is source j 's audited contribution to answer y . This is not merely a citation indicator. A source can be cited without being marginal, and a source can be marginal without being displayed prominently. If source j belongs to $\text{topic}(j) \in X$, source-level provenance can be aggregated to topic-level provenance by

$$s_x^{\text{topic}}(y) = \sum_{j \in \mathcal{I}(y): \text{topic}(j)=x} s_j(y). \quad (13.3)$$

- It audits human-information production. The mechanism must distinguish costly human information from low-cost synthetic mimicry. Provenance, randomized exploration, source audits, costly verification, legal warranties, and trust anchors can all serve this role. The likelihood-ratio condition in Theorem 4 gives the formal discipline: the audit must be informative enough to overcome the mimic's cost advantage.
- It uses keystone-topic weights when compensation is scarce. The marginal dollar should not necessarily go to the topic with the most current clicks. It should go to the topic whose

additional monetization most relaxes the reproduction constraint. In the reproduction-number model, this marginal value is proportional to

$$\frac{u_i(Kv)_i}{\kappa_i}, \quad (13.4)$$

where u and v are the Perron eigenvectors of the reproduction operator and κ_i is the cost of increasing topic i 's effective monetization. Theorem 5 formalizes this rule.

The rest of this section formalizes these elements. Visitor replacement gives the revenue benchmark. Audits and provenance solve the incentive problem created by cheap synthetic mimicry. Keystone-topic weights identify where scarce compensation does the most to keep costly human information above replacement. I also offer one market design example (of experimental clearing-house with sponsored challenges).

13.2 Adverse selection and proof of human information

The externality-payment theorem is an implementation benchmark. It assumes that the mechanism can measure the value terms it pays for. The signal-collapse theorem explains why this is not automatic. If revealed-preference and link experiments collapse, an AI platform that pays for “quality” using only residual answer-level observables can attract low-cost synthetic mimics that are observationally equivalent to high-cost human-information sources.

Proposition 19 (No quality payments from a collapsed experiment). *Suppose a mechanism observes a signal Y and pays a transfer $T(Y)$ with finite expectation. There are two source types, G and M . Type G produces genuine costly human information at cost c_G , type M produces a low-cost synthetic mimic at cost c_M , where $0 < c_M < c_G$, and the signal has the same distribution under both types. Then the expected transfer is the same for both types. If the mechanism induces participation of type G , so that $\mathbb{E}[T(Y)|G] \geq c_G$, then type M earns rent at least $c_G - c_M$. Conversely, any mechanism that gives type M nonpositive rent, $\mathbb{E}[T(Y)|M] \leq c_M$, cannot induce participation of type G .*

If the available signal cannot distinguish high-cost human information from low-cost synthetic imitation, a payment for quality becomes a payment for “appearance.” To finance human information without inducing low-cost synthetic mimicry, the mechanism needs an additional separating instrument. This is why the paper’s mechanism emphasizes restoring measurement as well as replacing revenue.

Theorem 4 (Audit likelihood-ratio threshold). *Suppose an audit or provenance experiment produces a “pass” signal with probability p_G for a genuine human-information source and p_M for a synthetic mimic, where $p_G > p_M \geq 0$ and $0 < c_M < c_G$. In the class of nonnegative pass-contingent mechanisms that pay $T \geq 0$ only after a pass and pay zero after a failure, the mechanism can both induce participation by the genuine source and give the synthetic mimic non-positive*

rent if and only if

$$\frac{p_G}{p_M} \geq \frac{c_G}{c_M}, \quad (13.5)$$

with the convention that the left side is infinite when $p_M = 0 < p_G$. More generally, the same condition is the likelihood-ratio requirement for a one-signal audit to overcome the cost advantage of synthetic imitation.

The theorem gives a quantitative version of the information problem for this simple pass-only payment class. A mechanism does not merely need provenance labels; it needs a signal whose likelihood ratio is high enough to offset the lower cost of imitation.¹⁸ Randomized exploration and audits are useful because they create precisely the source-level information that answer diversion otherwise deletes.

The theorem treats the audit technology as given. In practice, lowering p_M or raising p_G is costly, especially when synthetic mimics are cheap and scalable. Thus the likelihood-ratio condition is not a claim that separation is easy. It is an implementability test: any payment rule that purports to reward costly human information must either satisfy this informativeness condition, pay the rents induced by mimicry, or use additional instruments such as randomized audits.

With this incentive constraint in mind, the simplest policy is visitor replacement. In data, an observed AI-answer interaction may include queries that would not have produced a publisher click. If z_x is an AI-answer rate measured on the same counterfactual monetizable-exposure base as $R_x(q)$, and $\alpha_x \in [0, 1]$ is the probability that such an answer displaces a publisher visit, then $d_x = \alpha_x z_x$. If instead z_x is measured per raw query, let b_x be the counterfactual probability that a raw query would have generated a monetizable publisher visit absent AI. Then the corresponding displaced-revenue share is $d_x = \alpha_x z_x / b_x$, truncated at one when necessary; equivalently, one may define α_x to include this normalization. If z_x is already measured as effective displacement, then $d_x = z_x$. Residual direct-visit revenue is $(1 - d_x)R_x(q)$, and the economically relevant displaced revenue is $d_x R_x(q)$.

Proposition 20 (Counterfactual visitor-replacement royalty). *A transfer equal to expected displaced visit revenue restores the counterfactual old-web publisher problem. Formally, suppose the residual visit payoff under effective displacement d_x is*

$$(1 - d_x)R_x(q) - C_x(q), \quad (13.6)$$

where d_x is the effective displacement rate. If

$$T_x(q, d_x) = d_x R_x(q) \quad (13.7)$$

¹⁸If mechanisms can use deposits, penalties, failure-contingent clawbacks, or legal warranties, the exact threshold changes; those instruments work by adding enforcement power beyond the nonnegative pass-contingent payment considered here.

for every q and d_x , then

$$(1 - d_x)R_x(q) + T_x(q, d_x) - C_x(q) = R_x(q) - C_x(q). \quad (13.8)$$

Therefore the set of exact maximizers, the set of ε -maximizers, and the entry decision are identical to the old-web benchmark. The earlier formula $T_x(q, z) = zR_x(q)$ is the special case in which z is already measured as effective diversion, or equivalently $d_x = z$.

The proposition is intentionally simple and nonparametric. It does not say the old web was socially optimal. It says that if the policy objective is to reproduce the old publisher problem, expected displaced visit revenue is the exact object to replace. Three implementation issues remain. First, $R_x(q)$ and d_x are counterfactual once AI diverts visits. Second, AI answers often combine many sources, so attribution must be audited rather than assumed. And third, restoring the old private problem may reproduce old distortions, such as overproduction of clickbait and underproduction of public goods. These issues motivate the externality-payment result below.

However, a social planner may want more than restoration. The old web may underprovide certain low-demand high-social-value material (e.g., minority-language content). To avoid double counting, define each term as a marginal value relative to the same counterfactual and net of the other terms: $\mathcal{A}_x(q)$ is AI-answer value, $I_x(q)$ is source-level signal value, $\Lambda_x(q)$ is diversity value, and $\text{Expl}_x(q)$ is exploration value. A combined restoration-and-externality royalty can be written as

$$T_x(q, d_x) = d_x R_x(q) + \mathcal{A}_x(q) + I_x(q) + \Lambda_x(q) + \text{Expl}_x(q). \quad (13.9)$$

The first term restores displaced visit revenue. The additional terms correct answer-quality, signal, diversity, and exploration externalities. The proposition below separates the exact Pigouvian benchmark from the visitor-replacement benchmark because the old web need not have been socially optimal.

Proposition 21 (Exact externality payments implement the planner problem). *Fix an effective displacement rate d_x and suppose the social planner's topicwise objective is*

$$\mathcal{W}_x(q; d_x) = (1 - d_x)R_x(q) + \mathcal{A}_x(q) + I_x(q) + \Lambda_x(q) + \text{Expl}_x(q) - C_x(q), \quad (13.10)$$

where $\mathcal{A}_x$ is AI-answer value, I_x is source-level signal value, Λ_x is diversity value, and Expl_x is exploration value. If the transfer satisfies

$$T_x(q, d_x) = \mathcal{A}_x(q) + I_x(q) + \Lambda_x(q) + \text{Expl}_x(q), \quad (13.11)$$

then the publisher's compensated payoff equals $\mathcal{W}_x(q; d_x)$. Hence the exact and approximate publisher maximizers coincide with the exact and approximate maximizers of the social planner's topicwise objective. If the objective is instead to restore the old-web private problem and add these externalities, the implementing transfer is $d_x R_x(q) + \mathcal{A}_x(q) + I_x(q) + \Lambda_x(q) + \text{Expl}_x(q)$.

If publishers are paid exactly for the AI-answer value, signal value, diversity value, and exploration value that they create, their private objective becomes the social planner’s objective. The proposition is an implementation benchmark; the hard empirical work is estimating those terms with audited provenance and counterfactual quality measurements.

13.3 Experimental clearinghouses and sponsored challenges

One way to implement these benchmark objects is an experimental clearinghouse with sponsored challenges. The clearinghouse is not assumed to know the correct source payment. It posts provisional source payments, runs certified experiments when those payments are challenged, and uses the resulting settlements to update payments. Let $\mathfrak{U}$ be a finite set of source-settlement units. A unit $j \in \mathfrak{U}$ may be a source, or a source-query-bundle pair, over a settlement period. Let Pay_j^* denote the benchmark payment for unit j : either the visitor-replacement payment or the full replacement-plus-externality payment. The clearinghouse first posts a provisional value $\widehat{\text{Pay}}_j$. The AI platform may submit a proposed value based on its retrieval logs, answer logs, citations, and user-feedback data, but that proposal is only an input. The clearinghouse’s posted value is subject to source registration, provenance checks, random audits, and sponsored challenges. If source access is technologically excludable, participating sources can agree to block off-clearinghouse AI access; if access is not excludable, the same logic requires legal authority, such as reporting mandates, compulsory licensing, or an AI-platform levy.

For each j , a certified settlement experiment is specified by four objects: a bounded real-valued settlement random variable Settle_j , an experiment cost $\text{TestCost}_j \geq 0$, a statistical tolerance $\text{Tol}_j \geq 0$, and an error-probability bound $\text{ErrProb}_j \in [0, 1]$. The experiment may be a source-withholding test, cluster-withholding test, randomized source-panel experiment, provenance audit, freshness audit, originality audit, or expert review. The settlement rule is fixed in advance. Assume

$$\mathbb{E}[\text{Settle}_j \mid \text{Pay}_j^*] = \text{Pay}_j^*, \quad \Pr\{|\text{Settle}_j - \text{Pay}_j^*| \leq \text{Tol}_j\} \geq 1 - \text{ErrProb}_j. \quad (13.12)$$

Thus the experiment defines the mechanism’s payoff-relevant settlement object. A challenger may sponsor an upward challenge to $\widehat{\text{Pay}}_j$ by paying the experiment cost and taking a linear claim with payoff

$$\text{RewardRate}_j(\text{Settle}_j - \widehat{\text{Pay}}_j) - \text{TestCost}_j, \quad \text{RewardRate}_j > 0, \quad (13.13)$$

or may sponsor a downward challenge with payoff

$$\text{RewardRate}_j(\widehat{\text{Pay}}_j - \text{Settle}_j) - \text{TestCost}_j. \quad (13.14)$$

A wrong challenge can lose money. If the challenge is run, the clearinghouse updates the settlement value for unit j to Settle_j ; any amount already paid is credited or debited in the next settlement

cycle, with source reserves or bonds used for downward corrections when necessary. Challenges can be repeated only through a pre-specified audit ladder, for example by funding a more precise and more costly experiment. A unit is settled for the period once no admissible sponsored challenge has positive expected payoff.

Proposition 22 (Sponsored challenges bound payment errors). *Fix a settlement period and a finite set $\mathfrak{U}$. Suppose the certified experiment for each $j \in \mathfrak{U}$ satisfies (13.12). Suppose also that whenever*

$$|\text{Pay}_j^* - \widehat{\text{Pay}}_j| > \frac{\text{TestCost}_j}{\text{RewardRate}_j},$$

there exists at least one risk-neutral challenger with sufficient collateral whose information makes the corresponding upward or downward challenge have strictly positive expected payoff. Then any provisional payment vector with no profitable sponsored challenge satisfies

$$|\text{Pay}_j^* - \widehat{\text{Pay}}_j| \leq \frac{\text{TestCost}_j}{\text{RewardRate}_j} \quad \text{for every } j \in \mathfrak{U}. \quad (13.15)$$

If a challenge is sponsored and the clearinghouse updates $\widehat{\text{Pay}}_j$ to the settlement outcome Settle_j , then

$$\Pr\{|\widehat{\text{Pay}}_j - \text{Pay}_j^*| \leq \text{Tol}_j\} \geq 1 - \text{ErrProb}_j. \quad (13.16)$$

Consequently, after all sponsored challenges in a finite settlement period are resolved,

$$\Pr\left\{|\widehat{\text{Pay}}_j - \text{Pay}_j^*| \leq \max\left\{\text{Tol}_j, \frac{\text{TestCost}_j}{\text{RewardRate}_j}\right\} \text{ for all } j \in \mathfrak{U}\right\} \geq 1 - \sum_{j \in \mathfrak{U}} \text{ErrProb}_j. \quad (13.17)$$

The intuition is that sponsored challenges make large payment errors privately worth testing: if a provisional payment is wrong by more than the experiment cost divided by the challenger reward rate, an informed challenger can profitably fund the experiment, so unchallenged payments must lie inside that no-challenge band, while challenged payments inherit the experiment's statistical error bound.

The second issue is financing and participation. Let Bargain_j be source-settlement unit j 's best outside payment from bilateral licensing or other off-clearinghouse use, net of any cost of obtaining that outside payment. Let $\text{Outside}(\mathfrak{H})$ be the outside value of a coalition $\mathfrak{H} \subseteq \mathfrak{U}$, also net of the costs of exercising that option. Define the minimum replacement-and-stability budget by

$$\begin{aligned} \text{ReqBudget}(\mathfrak{U}) = & \min_{(\text{Tr}_j)_{j \in \mathfrak{U}} \in \mathbb{R}_+^{|\mathfrak{U}|}} \sum_{j \in \mathfrak{U}} \text{Tr}_j \\ \text{s.t.} \quad & \text{Tr}_j \geq \max\{\text{Pay}_j^*, \text{Bargain}_j\} \quad \forall j \in \mathfrak{U}, \\ & \sum_{j \in \mathfrak{H}} \text{Tr}_j \geq \text{Outside}(\mathfrak{H}) \quad \forall \mathfrak{H} \subseteq \mathfrak{U}. \end{aligned} \quad (13.18)$$

The constraints $\text{Tr}_j \geq \text{Pay}_j^*$ are replacement or sustainability floors; the outside-option constraints

are participation and stability constraints. Let $\text{PrivVal}_A(\mathfrak{U})$ be the AI platform's private incremental value from access to $\mathfrak{U}$, and let $\text{FeeCap}_A(\mathfrak{U}) \leq \text{PrivVal}_A(\mathfrak{U})$ be the maximum enforceable fee the clearinghouse can extract from the platform under the available access, contracting, or legal technology. Let $\text{MeasCost}(\mathfrak{U})$ denote administrative, reporting, audit, and non-sponsored measurement costs, including expected challenge rewards net of challenge fees. With outside financing $\text{PublicFunds} \geq 0$, the clearinghouse can implement replacement floors, source participation, coalitional stability, and measurement only if

$$\text{FeeCap}_A(\mathfrak{U}) + \text{PublicFunds} \geq \text{MeasCost}(\mathfrak{U}) + \text{ReqBudget}(\mathfrak{U}). \quad (13.19)$$

If the clearinghouse can freely choose transfers subject only to the constraints in (13.18), this condition is also sufficient. Hence the minimum outside financing is

$$\text{PublicFunds}^{\min} = [\text{MeasCost}(\mathfrak{U}) + \text{ReqBudget}(\mathfrak{U}) - \text{FeeCap}_A(\mathfrak{U})]_+. \quad (13.20)$$

In the best voluntary self-financing benchmark, $\text{FeeCap}_A(\mathfrak{U}) = \text{PrivVal}_A(\mathfrak{U})$. If access is not excludable and there is no legal obligation to pay, then $\text{FeeCap}_A(\mathfrak{U}) = 0$ for a purely voluntary access-fee clearinghouse; implementation must instead rely on a mandatory levy, compulsory licensing, or reporting-and-audit mandates. This subsection offers one implementable template under one set of technological assumptions, not a claim that it is the unique or universally optimal market design.

13.4 Keystone-topic compensation

Because K was derived above as a maximal no-transfer reproduction operator, the keystone formula should be read as a first-order rule for the maintained reproduction operator used by the mechanism. If the upper bound is loose because of fixed costs, entry frictions, credit constraints, pooling, or imperfect reinvestment, then the corresponding weights are diagnostic of fragility in the upper envelope rather than a complete sufficient statistic for feasible policy. In implementation, the mechanism should use an estimated feasible reproduction operator, or should interpret the K -based weights as a guide for where additional auditing and empirical measurement are most valuable.

The reproduction-number formulation gives a sharper market-design implication. Compensation should not be targeted only by current clicks. It should target topics whose marginal exposure most increases the reproduction number of costly human information. The marginal result below applies on the intensive margin: the topic must already be active, or its activation threshold must already have been funded, so that a small transfer can actually raise effective monetization. If fixed activation costs bind, the rescue problem is discrete rather than first order. A simple formulation is to choose a set of topics to activate and then allocate marginal funds

among the activated topics:

$$\max_{S \subseteq X, b_i \geq 0} \sum_{i \in S} b_i \frac{u_i(Kv)_i}{\kappa_i} \quad \text{s.t.} \quad \sum_{i \in S} F_i^{\text{act}} + \sum_{i \in S} b_i \leq B,$$

where F_i^{act} is the estimated activation payment needed before topic i can respond on the margin. Thus the Perron derivative is a local allocation rule, not a substitute for satisfying participation thresholds. Let $\mathcal{T}(\mathbf{m}) = D + \text{diag}(\mathbf{m})K$ and suppose $\mathcal{T}(\mathbf{m})$ is irreducible. Let $u \gg 0$ and $v \gg 0$ be Perron left and right eigenvectors, normalized by $u^\top v = 1$. Suppose a small transfer $b_i \geq 0$ to an active or threshold-funded topic i raises effective monetization by $b_i/\kappa_i + o(\|b\|)$, where $\kappa_i > 0$ is the cost of increasing m_i by one unit.

Theorem 5 (Keystone-topic compensation). *Assume the hypotheses of Proposition 4 for the reproduction operator maintained by the mechanism, and suppose transfers raise effective monetization according to the rule below. For a small budget B , the first-order increase in the open-web reproduction number from a transfer vector $b \geq 0$ with $\sum_i b_i \leq B$ is*

$$\sum_i b_i \frac{u_i(Kv)_i}{\kappa_i} + o(B). \quad (13.21)$$

Thus the first-order optimal rescue policy allocates the marginal budget only to topics maximizing

$$\frac{u_i(Kv)_i}{\kappa_i}. \quad (13.22)$$

The same Perron-Frobenius marginal values give a first-order formula for the minimum budget required to move a slightly subcritical no-transfer reproduction system to the replacement boundary. This is a local calculation: it identifies the cheapest way to raise the reproduction number to one, not a general proof that the system survives after that point.

Corollary 12 (Minimum budget to reach replacement at the margin). *Under the assumptions of Theorem 5, suppose*

$$\mathcal{R}^{\text{web}}(\mathbf{m}) = 1 - \Delta \quad (13.23)$$

with $\Delta > 0$ small, and suppose $\max_i u_i(Kv)_i/\kappa_i > 0$. To first order, the minimum budget needed to raise the reproduction number to one is

$$\frac{\Delta}{\max_i u_i(Kv)_i/\kappa_i} + o(\Delta). \quad (13.24)$$

A first-order optimal policy spends the marginal budget on topics attaining the maximum. If this maximum is zero, first-order rescue through the effective-monetization rates is impossible.

A (keystone) topic can deserve high compensation not because it has the most current clicks, but because additional monetization there most relaxes the web's reproduction constraint.

In implementation, the source-indexed provenance vector $s(y)$ can feed payments based on displaced visit value, marginal answer-quality contribution, topic diversity prices, and inverse-propensity-weighted exploration credits. Topic-level payments can be obtained by aggregating source-level provenance through $s_x^{topic}(y)$. The exploration credits are necessary because a pro-rata pool based only on observed AI use can reproduce the AI system’s own concentration and starve the tail.

I end this section by reiterating that the institutional point is as important as the formula. While I do not lay out a fully workable design, any future market design solutions to the problem that is the subject of this paper will have to address issues laid out above.

14 Boundary cases and scope

A few words of caution are called for here. The theorems are intentionally conditional. The paper identifies a missing-market force, not a law that every AI-mediated information system must collapse.

If the AI platform owns the relevant content producer or signs complete contingent contracts for content use, the missing-market wedge can disappear. In the model, this is represented by transfers that make the producer internalize displaced visit revenue and the relevant externality terms. The no-transfer collapse corollary then no longer applies.

Also, search degradation is not inevitable in systems with hard verified-source anchors. In Proposition 11, anchor exposure $\alpha_0(Y)$ is an asymptotic upper bound when contemporaneous genuine-source signals vanish. If a search engine reserves durable exposure for verified genuine sources, synthetic flooding need not eliminate genuine-source visibility through the ranking channel.

It is worth noting that the collapse of human production in a topic can be “efficient” if synthetic content genuinely supplies the same social information at lower cost. The inefficient-collapse theorems require positive nonappropriated value from costly human information. They do not claim that every reduction in human production is inefficient.

15 Empirical implications

The model yield some empirical predictions. AI-answer exposure should reduce publisher traffic most in query classes where answers are close substitutes for visits. Among exposed publishers, content like costly output, updates, expert labor, product testing, or maintenance should decline more than low-cost reposting or AI-generated fillers. Exit and non-entry should be concentrated in topics with high fixed costs and weak alternative monetization (e.g., local news, niche reviews, technical maintenance, or minority-language content).

The model makes predictions on the effects on diversity and measurement. Source entropy

should decline in exposed topics if AI retrieval concentrates attention. Topic support should decline when fixed costs bind. New backlinks, citations, and human-generated graph edges should fall in exposed categories. PageRank-like measures and other source-level ranking signals should become less predictive of current quality and more reflective of inherited authority, static text, or synthetic-volume effects. Licensing, pay-per-inference, pay-per-crawl, or answer-provenance payments should attenuate declines in output and freshness relative to uncompensated publishers.

If AI lowers website creation costs, raw URL counts and low cost synthetic pages may rise while human information investment, genuine source ranking exposure, PageRank mass, and source diversity fall. The model also predicts user migration. As conventional search quality deteriorates, more users should rely on AI answers, which should further lower non-AI visits and new human links.

I leave the empirical implementation for future work. We are still early in the game, and the relevant variation will likely come from a combination of AI-answer rollouts, changes in crawler policies, publisher licensing, and search-interface design. Promising designs include staggered rollout of AI answers across query classes or geographies, differences in exposure across devices and browsers, publisher-level variation in licensing, and changes in AI-crawler blocking policies. The outcomes to measure are traffic, article production, update frequency, author employment, source entropy, active topic coverage, backlinks, and content freshness.

16 Conclusion

Generative AI can make the web more convenient in the short run while weakening the economic and statistical loops that reproduce it in the long run. The reason is simple. The AI answer moves the attention event away from the publisher. If the revenue event moves far enough away from the publisher and is not replaced through transfers or other revenue sources, positive-cost commercial content can fall below its participation or reproduction threshold.

I show this through a strategic over-diversion theorem tied to a reproduction number. An AI platform that underinternalizes the continuation value of costly human information chooses weakly lower retained referrals than a social planner. If that private choice puts the no-transfer reproduction system below replacement (formally, if the reproduction number is below one), the maximal no-transfer reproduction system is subcritical and costly human information collapses. If a socially feasible policy satisfies the positive-stock viability condition and yields higher social planner value, the private equilibrium is inefficient. Avoiding the below-one region is not by itself a survival theorem; the social benchmark also requires an implementable policy that preserves a positive stock. The result is stronger than a conditional statement that content disappears when visits disappear. It identifies the strategic force that removes the monetizable event and the sufficient spectral condition under which the content system cannot reproduce itself.

The other results explain why the reproduction number is fragile. The effective-monetization

lemma shows that vanishing monetizable exposure kills positive cost content. The Blackwell theorem shows that AI diversion garbles source-level quality measurement. The reproduction-number theorem derives the reproduction system from referrals, subscribers, links, reputation, visit monetization, and reinvestment. The search theorem is conditional: modern ranking loses genuine-source exposure only to the extent that it depends on source-level web signals that vanish and lacks hard trust anchors. If AI also makes pages cheap, raw web volume can rise while the stock of costly human information falls. As conventional search becomes less useful, users can migrate further into AI, creating a feedback loop that further reduces non-AI visits and links.

The market design implication is not to ban AI answers or to recreate the old advertising web. The revenue and measurement events have to move to the new point of attention. A visitor replacement royalty restores the old publisher problem exactly. But that is only a benchmark. Provenance, diversity prices, exploration credits, and informative audits are needed to restore the broader ecosystem: quantity, quality, diversity, source-level signals, and the conventional-search discovery channel that prevents self-reinforcing migration into AI answers.

A Proofs

The appendix proves the stated results using only the assumptions named in each result. The collapse proofs use approximate optimality and the zero option rather than existence or uniqueness of exact maximizers. The signal proofs state pointwise total-variation convergence, adding uniform convergence only when the expected number of observations is uniformly bounded over states.

Proof of Lemma 1. Suppose not. Then there exist $\delta > 0$ and a subsequence, still denoted ℓ , such that $\|q_\ell\|_x \geq \delta$ for all ℓ . By Assumption 3,

$$C_x(q_\ell) \geq \underline{C}_x(\delta) > 0. \quad (\text{A.1})$$

By Assumption 2,

$$\pi_{x,\ell}(q_\ell) = m_\ell R_x(q_\ell) - C_x(q_\ell) \leq m_\ell \bar{R}_x - \underline{C}_x(\delta). \quad (\text{A.2})$$

Since $m_\ell \rightarrow 0$, for all sufficiently large ℓ the right-hand side is at most $-\underline{C}_x(\delta)/2$. The zero plan gives payoff zero, so the supremum payoff is nonnegative. Approximate optimality therefore requires $\pi_{x,\ell}(q_\ell) \geq -\varepsilon_\ell$, a contradiction because $\varepsilon_\ell \rightarrow 0$. Hence $\|q_\ell\|_x \rightarrow 0$.

If Assumption 4 holds, then for any active $q \neq 0_x$,

$$\pi_{x,\ell}(q) \leq m_\ell \bar{R}_x - F_x. \quad (\text{A.3})$$

For all large ℓ this is less than $-F_x/2$. The zero plan gives payoff zero, so no active plan can be ε_ℓ -optimal once $\varepsilon_\ell < F_x/2$. $\square$

Proof of Corollary 1. Apply Lemma 1 with $m_\ell = 1 - z_\ell$. $\square$

Proof of Corollary 2. Corollary 1 gives $\|q_\ell(x)\|_x \rightarrow 0$ for every $x \in X$. Since X is finite, summing over topics gives $Q_\ell \rightarrow 0$. $\square$

Proof of Proposition 1. For any active plan $q \neq 0_x$, Assumption 4 and bounded revenue imply

$$\pi_x(q; 1 - z_\ell) \leq (1 - z_\ell)\bar{R}_x - F_x. \quad (\text{A.4})$$

Since $z_\ell \rightarrow 1$, the right-hand side is less than $-F_x/2$ for all sufficiently large ℓ . The zero plan gives payoff zero, so $\Pi_x(1 - z_\ell) \geq 0$. For all sufficiently large ℓ , $\varepsilon_\ell < F_x/2$, so no active plan can satisfy $\pi_x(q_\ell; 1 - z_\ell) \geq \Pi_x(1 - z_\ell) - \varepsilon_\ell$. Hence $q_\ell = 0_x$ eventually. For the aggregate result, the active indicator converges to zero for every topic. Since X is finite, summing over topics gives $N_\ell \rightarrow 0$. $\square$

Proof of Corollary 3. If $q \neq 0_x$, then

$$\pi_x(q; 1 - z) \leq (1 - z)\bar{R}_x - F_x. \quad (\text{A.5})$$

If $\bar{R}_x = 0$, this upper bound is negative for every active plan. If $\bar{R}_x > 0$ and $1 - z < \underline{r} \leq F_x/\bar{R}_x$, the upper bound is again negative. The zero plan gives payoff zero. Hence no active plan can be an exact optimum. The approximate statement follows from the same inequality: if $(1 - z)\bar{R}_x - F_x < -\varepsilon$, then every active plan has payoff less than $-\varepsilon$ while the value is at least zero. $\square$

Proof of Corollary 4. Because $T_x(0_x, z) = 0$, the zero plan gives payoff zero and $\Pi_x^T(z_\ell) \geq 0$. Approximate optimality implies

$$(1 - z_\ell)R_x(q_\ell) + T_x(q_\ell, z_\ell) - C_x(q_\ell) \geq -\varepsilon_\ell. \quad (\text{A.6})$$

Using $R_x(q_\ell) \leq \bar{R}_x$ and $C_x(q_\ell) \geq \underline{C}_x(\delta)$ gives

$$T_x(q_\ell, z_\ell) \geq \underline{C}_x(\delta) - (1 - z_\ell)\bar{R}_x - \varepsilon_\ell. \quad (\text{A.7})$$

Taking the lower limit proves the claim. $\square$

Proof of Proposition 2. For the first claim, let $Q_x = \{0_x, q\}$, let $\|q\|_x = 1$, and set $R_x(q) = C_x(q) = 0$. Both 0_x and q are exact optima for every z , so there is an exact optimal sequence with nonzero content. For the second claim, add an exogenous non-visit payoff $B(q)$ with $B(q) > C_x(q)$ and payoff $(1 - z)R_x(q) + B(q) - C_x(q)$. Then production can survive even at $z = 1$. For the third claim, let $Q_x = [0, 1]$, $R_x(q) = q$, $C_x(q) = q^2$, and $\|q\|_x = q$. There is no activation cost because

$\inf_{q>0} q^2 = 0$. For $z < 1$, the unique maximizer is $q(z) = (1 - z)/2 > 0$, so the topic remains active for every $z < 1$, but $q(z) \rightarrow 0$. $\square$

Proof of Lemma 3. Upper semicontinuity follows from the stated upper-semicontinuity assumptions. Supermodularity follows because J^λ is a nonnegative linear combination of supermodular functions on the sublattice $\mathcal{R}_s$. To prove increasing differences, take $r' \geq r$ and $\lambda' \geq \lambda$. Then

$$\begin{aligned} J^{\lambda'}(r'; s) - J^{\lambda'}(r; s) - J^\lambda(r'; s) + J^\lambda(r; s) \\ = (\lambda' - \lambda) [\Delta V(\Gamma_s(r')) - \Delta V(\Gamma_s(r))] . \end{aligned}$$

Because Γ_s is increasing and ΔV is increasing, the bracket is nonnegative. This is increasing differences in (r, λ) . Therefore Assumption 6 holds with $\theta = \lambda$. $\square$

Proof of Proposition 5. By Assumption 6, J is upper semicontinuous and supermodular on a nonempty compact sublattice and has increasing differences in (r, θ) . Topkis's monotonicity theorem implies that the argmax correspondence is nonempty, is a sublattice, and is nondecreasing in θ in the strong set order. In particular, the smallest and largest optimal selections are nondecreasing in θ . Since $z = \mathbf{1} - r$, lower retained referral share is higher diversion. $\square$

Proof of Corollary 6. Fix $a > 0$. The separated maximum condition gives a positive gap

$$\Delta(a) = \Phi(0, s) - \sup\{\Phi(r, s) : r \in \mathcal{R}_s, \|r\|_\infty \geq a\} > 0. \quad (\text{A.8})$$

If an η_n -optimal sequence satisfied $\|r_n\|_\infty \geq a$ along a subsequence with $\eta_n \rightarrow 0$, it would lose at least $\Delta(a)$ relative to $r = 0$, a contradiction. Hence $r_n \rightarrow 0$, or $z_n \rightarrow \mathbf{1}$. The publisher conclusions follow from Corollary 1 and Proposition 1. $\square$

Proof of Proposition 6. By induction. If $h_t \geq h^*$, monotonicity of the nonnegative linear operator gives

$$\mathcal{T}(r, \xi)h_t \geq \mathcal{T}(r, \xi)h^* \geq h^*.$$

The implementability assumption gives $h_{t+1} \geq \mathcal{T}(r, \xi)h_t$, hence $h_{t+1} \geq h^*$. Since $h_0 \geq h^*$, the claim follows for every t . $\square$

Proof of Corollary 7. The proof is the same monotone induction in the augmented state. If $y_t \geq y^*$, nonnegativity of $\mathcal{M}(r \circ \xi)$ gives

$$\mathcal{M}(r \circ \xi)y_t \geq \mathcal{M}(r \circ \xi)y^* \geq y^*.$$

The implementability assumption gives $y_{t+1} \geq \mathcal{M}(r \circ \xi)y_t$, so $y_{t+1} \geq y^*$. Starting from $y_0 \geq y^*$, induction gives $y_t \geq y^*$ for all t . $\square$

Proof of Theorem 2. The inequality $r^P \leq r^S$ follows directly from Proposition 5 under the common monotone tie-breaking rule. Set $\mathbf{m}^P = r^P \circ \xi$ and $\mathcal{T}^P = D + \text{diag}(\mathbf{m}^P)K$. If $\mathcal{R}^{web}(r^P, \xi) < 1$, then $\rho(\mathcal{T}^P) < 1$. Hence any reduced no-transfer path satisfying $h_{t+1} \leq \mathcal{T}^P h_t$ converges to zero by Lemma 2. By Theorem 1, the corresponding augmented maximal no-transfer matrix $\mathcal{M}(\mathbf{m}^P)$ is also subcritical, so every content and attention-capital path dominated by the no-transfer financing system converges to zero from every initial state. At the social planner's choice, Proposition 6 or Corollary 7 supplies a feasible positive-stock path whenever its hypotheses hold. If that path yields social-planner value strictly above the private collapsing path, the private AI platform choice is not socially optimal. $\square$

Proof of Proposition 7. The zero plan is feasible and gives private payoff zero. Condition (7.15) implies that every plan with positive quality-adjusted content gives strictly negative private payoff. Therefore no such plan can be an exact private optimum, and every exact private optimum has zero quality-adjusted commercial content. In the social-planner problem, the zero plan gives value zero because $R_x(0_x) = C_x(0_x) = \mathcal{A}_x(0_x) = I_x(0_x) = 0$. Condition (7.16) gives a positive-content plan with strictly positive social-planner value. Hence zero content is not socially optimal. $\square$

Proof of Corollary 8. For any active plan $q \neq 0_x$, Assumptions 2 and 4 imply

$$(1 - z)R_x(q) - C_x(q) \leq (1 - z)\bar{R}_x - F_x. \quad (\text{A.9})$$

For every z sufficiently close to one, the right-hand side is strictly negative. Since the zero plan gives payoff zero, every exact no-transfer private optimum has zero quality-adjusted content. For the plan q^s , nonnegativity of R_x gives

$$(1 - z)R_x(q^s) + \mathcal{A}_x(q^s) + I_x(q^s) - C_x(q^s) \geq \mathcal{A}_x(q^s) + I_x(q^s) - C_x(q^s) > 0. \quad (\text{A.10})$$

Thus zero content is not socially optimal. $\square$

Proof of Corollary 9. By Proposition 1, $\mathbf{1}\{q_z(x) \neq 0_x\} \rightarrow 0$ for every $x \in X$ as $z \rightarrow 1$. Since X is finite and ω is bounded, summing over topics gives $D_\omega(z) \rightarrow 0$. $\square$

Proof of Proposition 8. If q is active, then Assumption 4 implies $C_x(q) \geq F_x$, while Assumption 2 implies $R_x(q) \leq \bar{R}_x$. Hence

$$\pi_x(q; 1 - z) \leq (1 - z)\bar{R}_x - F_x. \quad (\text{A.11})$$

An active exact optimum must have payoff at least zero because the zero plan is feasible and gives zero. Therefore $(1 - z)\bar{R}_x - F_x \geq 0$. If $\bar{R}_x > 0$, this is equivalent to $F_x/\bar{R}_x \leq 1 - z$. If $\bar{R}_x = 0 < F_x$, the inequality cannot hold. Thus all topics with $\tau_x > a$ are inactive whenever $z > 1 - a$. The weighted sum over H_a is then zero because the active indicator is zero on H_a . $\square$

Proof of Lemma 4. Let $p_\ell(\gamma) = T_\ell^\gamma(w)$ and $A(\gamma) = \sum_j w_j^\gamma$. Then

$$\log p_\ell(\gamma) = \gamma \log w_\ell - \log A(\gamma) \quad (\text{A.12})$$

and

$$\frac{dp_\ell}{d\gamma} = p_\ell(\log w_\ell - \mathbb{E}_\gamma \log w), \quad (\text{A.13})$$

where $\mathbb{E}_\gamma$ denotes expectation under $p(\gamma)$. Entropy can be written as

$$\text{Ent}(\gamma) = \log A(\gamma) - \gamma \mathbb{E}_\gamma \log w. \quad (\text{A.14})$$

Differentiating gives

$$\text{Ent}'(\gamma) = \mathbb{E}_\gamma \log w - \mathbb{E}_\gamma \log w - \gamma \frac{d}{d\gamma} \mathbb{E}_\gamma \log w. \quad (\text{A.15})$$

The derivative of $\mathbb{E}_\gamma \log w$ is $\text{Var}_{p(\gamma)}(\log w_\ell)$. Therefore

$$\text{Ent}'(\gamma) = -\gamma \text{Var}_{p(\gamma)}(\log w_\ell) \leq 0. \quad (\text{A.16})$$

The variance is zero if and only if all w_ℓ are equal. This proves the result. $\square$

Proof of Proposition 9. Shannon entropy is symmetric and strictly concave on the probability simplex, hence Schur-concave. Therefore if ν majorizes μ , then $\text{Ent}(\nu) \leq \text{Ent}(\mu)$. Strictness holds unless μ and ν differ only by a permutation. $\square$

Proof of Theorem 3. Let K_z be the thinning kernel that maps the old observation multiset to the retained multiset by retaining each old observation independently with probability $1 - z$. Then $\mathcal{E}_x^z = K_z \circ \mathcal{E}_x^0$. If $z_1 < z_2$, define an additional thinning kernel $K_{z_2|z_1}$ that retains each observation already retained under z_1 with probability

$$\frac{1 - z_2}{1 - z_1}. \quad (\text{A.17})$$

The resulting unconditional retention probability is $1 - z_2$. Hence $K_{z_2} = K_{z_2|z_1} \circ K_{z_1}$ and

$$\mathcal{E}_x^{z_2} = K_{z_2|z_1} \circ \mathcal{E}_x^{z_1}. \quad (\text{A.18})$$

The additional thinning kernel is independent of the state, so $\mathcal{E}_x^{z_1} \succeq_B \mathcal{E}_x^{z_2}$. Blackwell's theorem implies weakly lower value in every bounded decision problem. For convergence to the null experiment, let N be the number of old observations. The probability that at least one observation is retained under z is at most $(1 - z)\mathbb{E}_\vartheta[N]$ by Markov's inequality applied to the retained count. This converges to zero for every state because $\mathbb{E}_\vartheta[N] < \infty$. The total variation distance between the retained-multiset distribution and the point mass on the empty multiset is at most the probability of retaining at least one observation. Hence the retained multiset converges pointwise in

total variation to the empty multiset. If $\sup_{\vartheta \in \Omega_x} \mathbb{E}_{\vartheta}[N] < \infty$, the same bound is uniform in ϑ . $\square$

Proof of Proposition 10. Under ϑ and ϑ' , the distribution of all answer-level feedback is identical by assumption. Therefore every statistic of answer-level feedback has the same distribution under the two states. Since $\vartheta_i \neq \vartheta'_i$, source i 's quality differs across two observationally equivalent states. It is not identified from answer-level feedback alone. $\square$

Proof of Corollary 10. The argument is identical to the proof of Theorem 3, with current human-link observations replacing visit observations. The current-link experiment at higher diversion is obtained from the current-link experiment at lower diversion by additional state-independent deletion. Thus $\mathcal{L}_x^{z_1} \succeq_B \mathcal{L}_x^{z_2}$. Finite expected old link count implies pointwise convergence in total variation to the empty new-link experiment as $z \rightarrow 1$, by the same Markov bound. If the old link count has uniformly bounded expectation over states, the convergence is uniform over states. $\square$

Proof of Proposition 11. By the source-signal-dependence condition,

$$0 \leq \Xi_G(G_t, Y_t) \leq \alpha_0(Y_t) + L\|G_t\| \quad (\text{A.19})$$

for all sufficiently large t . Taking limits superior and using $\|G_t\| \rightarrow 0$ gives

$$\limsup_{t \rightarrow \infty} \Xi_G(G_t, Y_t) \leq \limsup_{t \rightarrow \infty} \alpha_0(Y_t). \quad (\text{A.20})$$

If $\alpha_0(Y_t) \rightarrow 0$, then $0 \leq \Xi_G(G_t, Y_t) \rightarrow 0$. If $\alpha_0(Y_t)$ is bounded away from zero, the displayed inequality only bounds exposure by the anchor term and does not imply collapse. This proves both statements. $\square$

Proof of Proposition 12. Suppose not. Then along a subsequence $\mathcal{H}_x(q_n) \geq \delta > 0$. By Assumption 7, $C_x^H(q_n) \geq \underline{C}_x^H(\delta) > 0$. Since $R_x(q_n) \leq \bar{R}_x$ and $C_x^S(q_n) \geq 0$,

$$\pi_x(q_n; z_n, \xi_n, \kappa_n) \leq m_n \bar{R}_x - \underline{C}_x^H(\delta). \quad (\text{A.21})$$

For all large n , this is at most $-\underline{C}_x^H(\delta)/2$. The zero plan is feasible and gives payoff zero, so the value is nonnegative. Since q_n is ε_n -optimal, its payoff is at least $-\varepsilon_n$, a contradiction for all sufficiently large n . $\square$

Proof of Corollary 11. For the synthetic sequence,

$$\pi_x(s_n; z_n, \xi_n, \kappa_n) = m_n R_x(s_n) - \kappa_n C_x^S(s_n) \geq -\kappa_n C_x^S(s_n) = o(1). \quad (\text{A.22})$$

Because costs are nonnegative and revenue is bounded, the value is at most $m_n \bar{R}_x = o(1)$ above the zero payoff. Hence the optimality gap of s_n is at most $m_n \bar{R}_x + \kappa_n C_x^S(s_n) = o(1)$. The human-information conclusion follows from Proposition 12. $\square$

Proof of Proposition 13. Suppose Y_n does not diverge. Then there is a finite $\bar{Y}$ and a subsequence, still denoted n , with $Y_n \leq \bar{Y}$. Since φ is strictly decreasing and positive at every finite argument, $\varphi(Y_n) \geq \varphi(\bar{Y}) > 0$ along that subsequence. But $\varphi(Y_n) = \kappa_n \rightarrow 0$, a contradiction. Hence $Y_n \rightarrow \infty$. The final statement combines this conclusion with Proposition 12 applied to the accompanying publisher sequence satisfying its hypotheses: synthetic volume diverges on the free-entry margin, while costly human information vanishes whenever the effective exposure-monetization multiplier converges to zero and the publisher sequence is approximately optimal. $\square$

Proof of Lemma 5. By the definition of v_n and the weight bounds,

$$v_n(G_n) = \frac{\sum_{i \in G_n} w_{ni}}{\sum_{i \in G_n} w_{ni} + \sum_{i \in \text{Syn}_n} w_{ni}} \leq \frac{\bar{w}_G |G_n|}{\bar{w}_G |G_n| + \underline{w}_S |\text{Syn}_n|}. \quad (\text{A.23})$$

If $|\text{Syn}_n|/|G_n| \rightarrow \infty$, the right-hand side converges to zero. $\square$

Proof of Proposition 14. Taking genuine-page mass in (11.8) gives

$$\pi_n(G_n) = (1 - \alpha)v_n(G_n) + \alpha \sum_{i \in G_n} \pi_n(i)P_n(i, G_n) + \alpha \sum_{i \in \text{Syn}_n} \pi_n(i)P_n(i, G_n) \quad (\text{A.24})$$

$$\leq (1 - \alpha)v_n(G_n) + \alpha\pi_n(G_n) + \alpha\delta_n. \quad (\text{A.25})$$

Rearranging yields

$$\pi_n(G_n) \leq v_n(G_n) + \frac{\alpha}{1 - \alpha}\delta_n. \quad (\text{A.26})$$

If $v_n(G_n) \rightarrow 0$ and $\delta_n \rightarrow 0$, the right-hand side converges to zero. Since each $\pi_n(j)$ for $j \in G_n$ is bounded above by $\pi_n(G_n)$, the final statement follows. $\square$

Proof of Proposition 15. Taking genuine-page mass in the PageRank equation gives

$$\pi_n(G_n) = (1 - \alpha)v_n(G_n) + \alpha \sum_i \pi_n(i)P_n(i, G_n). \quad (\text{A.27})$$

The second term is nonnegative. Hence $\pi_n(G_n) \geq (1 - \alpha)v_n(G_n) \geq (1 - \alpha)\underline{v}$. $\square$

Proof of Proposition 16. By Proposition 14, $v_n(G_n) \rightarrow 0$ and $\delta_n \rightarrow 0$ imply $\pi_n(G_n) \rightarrow 0$. Since $\xi_n \leq L_\xi \pi_n(G_n)$, the effective revenue multiplier

$$m_n = (1 - z_n)\xi_n \quad (\text{A.28})$$

converges to zero. Suppose, toward a contradiction, that an $o(1)$ -optimal sequence q_n satisfies $\|q_n\|_x \geq \delta > 0$ along a subsequence. By Assumption 3, $C_x(q_n) \geq \underline{C}_x(\delta) > 0$ on that subsequence, while bounded revenue gives

$$U_n(q_n) \leq m_n \bar{R}_x - \underline{C}_x(\delta). \quad (\text{A.29})$$

For all large n this is less than $-\underline{C}_x(\delta)/2$. The zero plan gives payoff zero, so an $o(1)$ -optimal plan must have payoff at least $-o(1)$, a contradiction. Hence $\|q_n\|_x \rightarrow 0$.

For the link statement, let M_n be the number of new human-link observations. If $\mathbb{E}_\vartheta[M_n] \leq L(1 - z_n)\xi_n$ for every fixed state ϑ and some finite L , then $\mathbb{E}_\vartheta[M_n] \rightarrow 0$. Markov's inequality gives $P_\vartheta(M_n \geq 1) \leq \mathbb{E}_\vartheta[M_n] \rightarrow 0$. The total variation distance between the link experiment and the null experiment is bounded by this probability, so the current-link experiment converges pointwise in total variation to the null experiment. $\square$

Proof of Proposition 17. We first prove the dynamic contraction. Since $\xi_0 \leq \bar{\xi}$, Assumption 8 gives

$$\xi_1 \leq \phi((1 - \zeta(\xi_0))\xi_0) \leq \lambda\xi_0 \leq \bar{\xi}. \quad (\text{A.30})$$

Inductively, if $\xi_t \leq \bar{\xi}$, then

$$\xi_{t+1} \leq \phi((1 - \zeta(\xi_t))\xi_t) \leq \lambda\xi_t. \quad (\text{A.31})$$

Thus $\xi_t \leq \lambda^t \xi_0 \rightarrow 0$. Since $z_t = \zeta(\xi_t)$ and ζ is continuous at zero, $z_t \rightarrow \zeta(0)$.

The effective revenue multiplier in the publisher payoff is

$$m_t = (1 - z_t)\xi_t. \quad (\text{A.32})$$

Because $0 \leq 1 - z_t \leq 1$ and $\xi_t \rightarrow 0$, we have $m_t \rightarrow 0$. Lemma 1 applies with effective multiplier m_t , so every $o(1)$ -optimal sequence satisfies $\|q_t\|_x \rightarrow 0$. Under Assumption 4, the activation part of the same lemma gives eventual inactivity.

For the link statement, let M_t be the number of new human-link observations. If $\mathbb{E}_\vartheta[M_t] \leq L(1 - z_t)\xi_t$ for a fixed state ϑ , then $\mathbb{E}_\vartheta[M_t] \rightarrow 0$. Markov's inequality yields $P_\vartheta(M_t \geq 1) \leq \mathbb{E}_\vartheta[M_t] \rightarrow 0$. The total variation distance between the current-link experiment and the null experiment is bounded by this probability. Hence convergence is pointwise in total variation under every state satisfying the stated finite bound. $\square$

Proof of Theorem 1. Fix a stationary exposure vector $\mathbf{m}$ and write $E = E(\mathbf{m})$. Since $\rho(P^A) < 1$, $I - P^A$ is invertible and $(I - P^A)^{-1} \geq 0$. The matrix $I - \mathcal{M}(\mathbf{m})$ is

$$I - \mathcal{M}(\mathbf{m}) = \begin{pmatrix} I - D - EQ & -EU \\ -B & I - P^A \end{pmatrix}. \quad (\text{A.33})$$

For a nonnegative matrix N , $\rho(N) < 1$ is equivalent to $I - N$ being a nonsingular M -matrix, or equivalently to $(I - N)^{-1} \geq 0$. Taking the Schur complement with respect to $I - P^A$, this condition is equivalent to

$$I - D - EQ - EU(I - P^A)^{-1}B \quad (\text{A.34})$$

being a nonsingular M -matrix. Since

$$D + EQ + EU(I - P^A)^{-1}B = D + E [Q + U(I - P^A)^{-1}B] = D + \text{diag}(\mathbf{m})K, \quad (\text{A.35})$$

this is equivalent to

$$\rho(D + \text{diag}(\mathbf{m})K) < 1. \quad (\text{A.36})$$

This proves the spectral equivalence.

For the convergence statement, let $y_t = (h_t, a_t)$. By nonnegativity of all matrices and the componentwise bound $\mathbf{m}_t \leq \bar{\mathbf{m}}$, for all $t \geq T$,

$$y_{t+1} \leq \mathcal{M}(\mathbf{m}_t)y_t \leq \mathcal{M}(\bar{\mathbf{m}})y_t. \quad (\text{A.37})$$

The first part gives $\rho(\mathcal{M}(\bar{\mathbf{m}})) < 1$. Hence $\mathcal{M}(\bar{\mathbf{m}})^s \rightarrow 0$. Iterating the preceding inequality from date T gives

$$y_{T+s} \leq \mathcal{M}(\bar{\mathbf{m}})^s y_T \rightarrow 0. \quad (\text{A.38})$$

Therefore $h_t \rightarrow 0$ and $a_t \rightarrow 0$ componentwise. $\square$

Proof of Proposition 3. Under the equality assumptions, $y_t = \mathcal{M}(\mathbf{m})^t y_0$. By Theorem 1, $\mathcal{R}^{web}(\mathbf{m}) < 1$ is equivalent to $\rho(\mathcal{M}(\mathbf{m})) < 1$, which is equivalent to $\mathcal{M}(\mathbf{m})^t \rightarrow 0$. Hence every initial state converges to zero.

If $\mathcal{R}^{web}(\mathbf{m}) > 1$, then $\rho(\mathcal{M}(\mathbf{m})) > 1$. Perron-Frobenius gives a nonnegative Perron eigenvector $w \neq 0$ with $\mathcal{M}(\mathbf{m})w = \rho(\mathcal{M}(\mathbf{m}))w$. Starting from $y_0 = w$ yields $y_t = \rho(\mathcal{M}(\mathbf{m}))^t w$, which does not converge to zero. If $\mathcal{M}(\mathbf{m})$ is irreducible, let $u \gg 0$ be a left Perron eigenvector. For every $y_0 \gg 0$,

$$u^\top y_t = u^\top \mathcal{M}(\mathbf{m})^t y_0 = \rho(\mathcal{M}(\mathbf{m}))^t u^\top y_0 \rightarrow \infty. \quad (\text{A.39})$$

Thus no strictly positive initial state can converge to zero. This argument does not require primitivity or convergence of normalized paths. This proves the benchmark threshold statement. $\square$

Proof of Lemma 2. Suppose first that $\rho(A) < 1$. Then $I - A^\top$ is invertible and

$$p = (I - A^\top)^{-1} \mathbf{1} = \sum_{k=0}^{\infty} (A^\top)^k \mathbf{1} \quad (\text{A.40})$$

is well defined and satisfies $p \gg 0$. Moreover $A^\top p = p - \mathbf{1}$. Let

$$\beta = 1 - \min_i \frac{1}{p_i}. \quad (\text{A.41})$$

Since $p_i \geq 1$ for every i , $0 \leq \beta < 1$, and $A^\top p = p - \mathbf{1} \leq \beta p$. This is equivalent to $p^\top A \leq \beta p^\top$.

Conversely, suppose $p \gg 0$ and $A^\top p \leq \beta p$ with $\beta < 1$. Let $D_p = \text{diag}(p)$ and define

$$\tilde{A} = D_p^{-1} A^\top D_p.$$

Then $\tilde{A} \geq 0$ and its i th row sum is

$$\sum_j \tilde{A}_{ij} = \frac{(A^\top p)_i}{p_i} \leq \beta.$$

Hence $\rho(A) = \rho(A^\top) = \rho(\tilde{A}) \leq \|\tilde{A}\|_\infty \leq \beta < 1$.

Finally, if $h_{t+1} \leq Ah_t$, then

$$p^\top h_{t+1} \leq p^\top Ah_t \leq \beta p^\top h_t.$$

Iterating gives

$$p^\top h_t \leq \beta^t p^\top h_0 \rightarrow 0.$$

□

Proof of Corollary 5. If $m_i = (1 - z)\xi$ for every i and $D = (1 - \delta)I$, then

$$D + \text{diag}(\mathbf{m})K = (1 - \delta)I + (1 - z)\xi K. \quad (\text{A.42})$$

The eigenvalues of $(1 - \delta)I + (1 - z)\xi K$ are $(1 - \delta) + (1 - z)\xi \lambda_j(K)$. For a nonnegative matrix K , $\rho(K)$ is an eigenvalue and every other eigenvalue has modulus at most $\rho(K)$. Since $1 - \delta \geq 0$ and $(1 - z)\xi \geq 0$, the spectral radius is

$$\rho\{(1 - \delta)I + (1 - z)\xi K\} = 1 - \delta + (1 - z)\xi \rho(K). \quad (\text{A.43})$$

Subcriticality, $\mathcal{R}^{web} < 1$, is therefore equivalent to $(1 - z)\xi \rho(K) < \delta$. When $\rho(K) > 0$, this can be written as $(1 - z)\xi < \delta/\rho(K)$. When $\rho(K) = 0$, the condition holds for every z and ξ as long as $\delta > 0$. □

Proof of Proposition 4. For an irreducible nonnegative matrix, the Perron root is simple. Standard eigenvalue perturbation gives

$$d\rho(A) = u^\top (dA)v \quad (\text{A.44})$$

under the normalization $u^\top v = 1$. Since

$$\mathcal{T}(\mathbf{m}) = D + \text{diag}(\mathbf{m})K, \quad (\text{A.45})$$

the derivative of A with respect to m_i is the matrix whose i th row is the i th row of K and whose

other rows are zero. Hence

$$\frac{\partial \mathcal{R}^{web}(\mathbf{m})}{\partial m_i} = u_i(Kv)_i. \quad (\text{A.46})$$

The derivatives with respect to z_i and ξ_i follow from $m_i = (1 - z_i)\xi_i$ and the chain rule. $\square$

Proof of Remark 1. Let $\pi(P)$ and $\pi(Q)$ denote PageRank vectors with the same damping parameter α and teleportation distribution v . Then

$$\pi(P) - \pi(Q) = \alpha[\pi(P)P - \pi(Q)Q] = \alpha\pi(P)(P - Q) + \alpha[\pi(P) - \pi(Q)]Q. \quad (\text{A.47})$$

Taking ℓ_1 norms and using that multiplication by a row-stochastic matrix is nonexpansive in total variation,

$$\|\pi(P) - \pi(Q)\|_1 \leq \alpha \max_i \|P(i, \cdot) - Q(i, \cdot)\|_1 + \alpha \|\pi(P) - \pi(Q)\|_1. \quad (\text{A.48})$$

Rearranging gives the bound. $\square$

Proof of Proposition 20. Substituting $T_x(q, d_x) = d_x R_x(q)$ into the compensated payoff gives

$$(1 - d_x)R_x(q) + d_x R_x(q) - C_x(q) = R_x(q) - C_x(q). \quad (\text{A.49})$$

This is exactly the old-web payoff. Therefore exact maximizers, approximate maximizers, and entry decisions coincide with those in the old-web benchmark. $\square$

Proof of Proposition 21. Substituting the stated transfer into the compensated publisher payoff gives

$$(1 - d_x)R_x(q) + \mathcal{A}_x(q) + I_x(q) + \Lambda_x(q) + \text{Expl}_x(q) - C_x(q) = \mathcal{W}_x(q; d_x). \quad (\text{A.50})$$

Thus the publisher and social planner maximize the same function over the same feasible set. Exact maximizers coincide. The same identity implies that ε -maximizers coincide for every $\varepsilon \geq 0$. If the objective adds externality terms to the old-web private payoff $R_x(q) - C_x(q)$, adding the extra term $d_x R_x(q)$ gives the second statement. $\square$

Proof of Theorem 5. For irreducible $\mathcal{T}(\mathbf{m})$, the Perron root is simple. With $u^\top v = 1$, first-order perturbation theory gives

$$d\rho(A) = u^\top (dA)v. \quad (\text{A.51})$$

If transfer b_i raises m_i by $b_i/\kappa_i + o(\|b\|)$, then

$$dA = \text{diag}(dm)K, \quad dm_i = \frac{b_i}{\kappa_i} + o(\|b\|). \quad (\text{A.52})$$

Hence

$$d\mathcal{R}^{web} = \sum_i \frac{b_i}{\kappa_i} u_i(Kv)_i + o(\|b\|). \quad (\text{A.53})$$

For fixed small budget B , the first-order problem is the linear program

$$\max_{b \geq 0, \sum_i b_i \leq B} \sum_i b_i \frac{u_i(Kv)_i}{\kappa_i}. \quad (\text{A.54})$$

A linear program over the simplex allocates all marginal budget to indices attaining the largest coefficient, with arbitrary splitting among ties. This proves the result. $\square$

Proof of Corollary 12. Theorem 5 implies that a budget B can raise the reproduction number by at most

$$B \max_i \frac{u_i(Kv)_i}{\kappa_i} + o(B)$$

to first order, and this bound is attained by spending on any maximizer. The maintained positivity condition ensures that this first-order gain can be strictly positive. Equating the first-order gain to Δ gives the displayed expression. $\square$

Proof of Proposition 18. Assumption 8 and $\xi_0 \leq \bar{\xi}$ imply by Proposition 17 that $\xi_t \leq \lambda^t \xi_0 \rightarrow 0$ and $z_t \rightarrow \zeta(0)$. Since $0 \leq 1 - z_t \leq 1$, $\mathbf{m}_t = ((1 - z_t)\xi_t)\mathbf{1} \rightarrow 0$. The map $\mathbf{m} \mapsto \rho(D + \text{diag}(\mathbf{m})K)$ is continuous. Because $D = \text{diag}(1 - \delta_i)$ with $\delta_i \in (0, 1]$, $\rho(D) < 1$. Hence there exists $\bar{\mathbf{m}} \gg 0$ small enough that $\mathcal{R}^{web}(\bar{\mathbf{m}}) < 1$. Since $\mathbf{m}_t \rightarrow 0$, there is T such that $\mathbf{m}_t \leq \bar{\mathbf{m}}$ for all $t \geq T$. Theorem 1 then implies $h_t \rightarrow 0$ and $a_t \rightarrow 0$. If $\|G_t\| \leq L_G \|y_t\|$, then $G_t \rightarrow 0$ as well. $\square$

Proof of Proposition 22. Fix j and suppress the subscript. Suppose first that

$$\text{Pay}^* - \widehat{\text{Pay}} > \frac{\text{TestCost}}{\text{RewardRate}}.$$

Under the maintained challenger-information assumption, an upward challenge has strictly positive expected payoff. Therefore a provisional vector with no profitable sponsored challenge cannot satisfy the displayed inequality. Similarly, if

$$\widehat{\text{Pay}} - \text{Pay}^* > \frac{\text{TestCost}}{\text{RewardRate}},$$

then the corresponding downward challenge has strictly positive expected payoff, so this inequality is also ruled out. Combining the two cases gives

$$|\text{Pay}^* - \widehat{\text{Pay}}| \leq \frac{\text{TestCost}}{\text{RewardRate}}.$$

If a challenge is sponsored, the mechanism updates the settlement value to Settle . The certified-experiment assumption (13.12) immediately gives

$$\Pr\{|\widehat{\text{Pay}} - \text{Pay}^*| \leq \text{Tol}\} = \Pr\{|\text{Settle} - \text{Pay}^*| \leq \text{Tol}\} \geq 1 - \text{ErrProb}.$$

For the final uniform statement, apply the preceding argument to every $j \in \mathfrak{U}$. For units not

challenged, (13.15) holds deterministically under the no-profitable-challenge condition. For units challenged, (13.16) fails with probability at most ErrProb_j . Since $\mathfrak{U}$ is finite, the union bound implies that all challenged-unit settlement bounds hold simultaneously with probability at least $1 - \sum_{j \in \mathfrak{U}} \text{ErrProb}_j$. Combining challenged and unchallenged units gives (13.17). $\square$

Proof of the clearinghouse financing claim. Any feasible implementation must pay $\text{MeasCost}(\mathfrak{U})$ and must choose transfers satisfying the constraints in (13.18). By definition of $\text{ReqBudget}(\mathfrak{U})$, every such transfer vector has total payment at least $\text{ReqBudget}(\mathfrak{U})$. Thus total required resources are at least

$$\text{MeasCost}(\mathfrak{U}) + \text{ReqBudget}(\mathfrak{U}).$$

Available resources are at most $\text{FeeCap}_A(\mathfrak{U}) + \text{PublicFunds}$, so feasibility requires (13.19). Conversely, suppose (13.19) holds. An optimizer in (13.18) exists whenever all outside values are finite: the feasible set is closed and nonempty after taking sufficiently large transfers, and the linear objective is coercive on the nonnegative orthant. The clearinghouse combines an enforceable platform fee no larger than $\text{FeeCap}_A(\mathfrak{U})$ with outside financing to cover measurement cost and the minimizing transfers. By construction, all replacement floors and outside-option constraints are satisfied. Solving the feasibility inequality for the smallest nonnegative outside financing gives (13.20). $\square$

Proof of Proposition 19. Because Y has the same distribution under G and M , any measurable transfer $T(Y)$ with finite expectation has the same expected value under both types:

$$\mathbb{E}[T(Y) \mid G] = \mathbb{E}[T(Y) \mid M]. \quad (\text{A.55})$$

If type G participates, then $\mathbb{E}[T(Y) \mid G] \geq c_G$, and therefore $\mathbb{E}[T(Y) \mid M] \geq c_G$. Type M 's expected rent is at least $c_G - c_M > 0$. Conversely, if type M receives nonpositive rent, then $\mathbb{E}[T(Y) \mid M] \leq c_M$. Equality of expected transfers implies $\mathbb{E}[T(Y) \mid G] \leq c_M < c_G$, so type G does not participate. This proves the screening impossibility. $\square$

Proof of Theorem 4. A pass-contingent payment T induces participation by type G if and only if

$$p_G T \geq c_G. \quad (\text{A.56})$$

It gives type M nonpositive rent if and only if

$$p_M T \leq c_M. \quad (\text{A.57})$$

If $p_M = 0$, the mimic constraint is automatic for every finite T , while $p_G > 0$ permits any $T \geq c_G/p_G$; this is exactly the stated infinite-likelihood-ratio case. If $p_M > 0$, such a T exists if

and only if

$$\frac{c_G}{p_G} \leq \frac{c_M}{p_M}, \tag{A.58}$$

which is equivalent to $p_G/p_M \geq c_G/c_M$. Necessity and sufficiency follow. □

References

- Aggarwal, Pranjal, Vishvak Murahari, Tanmay Rajpurohit, Ashwin Kalyan, Karthik Narasimhan, and Ameet Deshpande. 2023. “GEO: Generative Engine Optimization.” arXiv:2311.09735.
- Anderson, Simon P., and Stephen Coate. 2005. “Market Provision of Broadcasting: A Welfare Analysis.” *Review of Economic Studies* 72(4): 947–972.
- Athey, Susan, Markus Mobius, and Jeno Pal. 2021. “The Impact of Aggregators on Internet News Consumption.” NBER Working Paper No. 28746.
- Barabasi, Albert-Laszlo, and Reka Albert. 1999. “Emergence of Scaling in Random Networks.” *Science* 286(5439): 509–512.
- Baumann, Fabian, Erol Akcay, and Joshua B. Plotkin. 2026. “Generative Artificial Intelligence Reduces Social Welfare through Model Collapse.” arXiv:2604.21853.
- Berman, Abraham, and Robert J. Plemmons. 1994. *Nonnegative Matrices in the Mathematical Sciences*. Philadelphia: SIAM.
- Blackwell, David. 1953. “Equivalent Comparisons of Experiments.” *Annals of Mathematical Statistics* 24(2): 265–272.
- Brin, Sergey, and Lawrence Page. 1998. “The Anatomy of a Large-Scale Hypertextual Web Search Engine.” *Computer Networks and ISDN Systems* 30(1–7): 107–117.
- Chaney, Allison J. B., Brandon M. Stewart, and Barbara E. Engelhardt. 2018. “How Algorithmic Confounding in Recommendation Systems Increases Homogeneity and Decreases Utility.” *Proceedings of the 12th ACM Conference on Recommender Systems*.
- Chiou, Lesley, and Catherine Tucker. 2017. “Content Aggregation by Platforms: The Case of the News Media.” *Journal of Economics and Management Strategy* 26(4): 782–805.
- Cloudflare. 2025. “Introducing pay per crawl: Enabling content owners to charge AI crawlers for access.” Cloudflare Blog.
- Dohmatob, Elvis, Yunzhen Feng, Pu Yang, Francois Charton, and Julia Kempe. 2024. “A Tale of Tails: Model Collapse as a Change of Scaling Laws.” arXiv:2402.07043.
- Fader, Peter S., Bruce G. S. Hardie, and Ka Lok Lee. 2005. “RFM and CLV: Using Iso-Value Curves for Customer Base Analysis.” *Journal of Marketing Research* 42(4): 415–430.
- Gentzkow, Matthew, Jesse M. Shapiro, and Daniel F. Stone. 2015. “Media Bias in the Marketplace: Theory.” In *Handbook of Media Economics*, Vol. 1B, edited by Simon P. Anderson, Joel Waldfogel, and David Stromberg, 623–645. Elsevier.

- George, Lisa M. 2001. “What’s Fit to Print: The Effect of Ownership Concentration on Product Variety in Daily Newspaper Markets.” arXiv:cs/0108014.
- Gerstgrasser, Matthias, Rylan Schaeffer, Apratim Dey, Rafael Rafailov, Dhruv Pai, Henry Sleight, John Hughes, Tomasz Korbak, Rajashree Agrawal, Andrey Gromov, Daniel A. Roberts, Diyi Yang, David L. Donoho, and Sanmi Koyejo. 2024. “Is Model Collapse Inevitable? Breaking the Curse of Recursion by Accumulating Real and Synthetic Data.” arXiv:2404.01413.
- Google Search Central. 2025. “A Guide to Google Search Ranking Systems.” Google for Developers. <https://developers.google.com/search/docs/appearance/ranking-systems-guide>. Accessed May 15, 2026.
- Google Search. 2026. “How Search Systems Work.” <https://www.google.com/search/howsearchworks/how-search-works/ranking-results/>. Accessed May 15, 2026.
- The Guardian. 2026. “Publishers Fear AI Search Summaries and Chatbots Mean ‘End of Traffic Era’.” January 12.
- Gupta, Sunil, Donald R. Lehmann, and Jennifer Ames Stuart. 2004. “Valuing Customers.” *Journal of Marketing Research* 41(1): 7–18.
- Gyongyi, Zoltan, Hector Garcia-Molina, and Jan Pedersen. 2004. “Combating Web Spam with TrustRank.” *Proceedings of the 30th International Conference on Very Large Data Bases*: 576–587.
- Joachims, Thorsten, Adith Swaminathan, and Tobias Schnabel. 2017. “Unbiased Learning-to-Rank with Biased Feedback.” *Proceedings of the Tenth ACM International Conference on Web Search and Data Mining*.
- Kemeny, John G., and J. Laurie Snell. 1976. *Finite Markov Chains*. New York: Springer-Verlag.
- Kleinberg, Jon M. 1999. “Authoritative Sources in a Hyperlinked Environment.” *Journal of the ACM* 46(5): 604–632.
- Langville, Amy N., and Carl D. Meyer. 2006. *Google’s PageRank and Beyond: The Science of Search Engine Rankings*. Princeton: Princeton University Press.
- Laub, Rene, Klaus M. Miller, and Bernd Skiera. 2024. “The Economic Value of User Tracking for Publishers.” arXiv:2303.10906.
- Milgrom, Paul, and Chris Shannon. 1994. “Monotone Comparative Statics.” *Econometrica* 62(1): 157–180.
- Pew Research Center. 2025. “Google Users Are Less Likely to Click on Links When an AI Summary Appears in the Results.” July 22.

- Reinartz, Werner J., and V. Kumar. 2003. “The Impact of Customer Relationship Characteristics on Profitable Lifetime Duration.” *Journal of Marketing* 67(1): 77–99.
- Rochet, Jean-Charles, and Jean Tirole. 2003. “Platform Competition in Two-Sided Markets.” *Journal of the European Economic Association* 1(4): 990–1029.
- Seddik, Mohamed El Amine, Suei-Wen Chen, Soufiane Hayou, Pierre Youssef, and Merouane Debbah. 2024. “How Bad is Training on Synthetic Data? A Statistical Analysis of Language Model Collapse.” arXiv:2404.05090.
- Shumailov, Ilia, Zakhar Shumaylov, Yiren Zhao, Yarin Gal, Nicolas Papernot, and Ross Anderson. 2023. “The Curse of Recursion: Training on Generated Data Makes Models Forget.” arXiv:2305.17493.
- Shumailov, Ilia, Zakhar Shumaylov, Yiren Zhao, Nicolas Papernot, Ross Anderson, and Yarin Gal. 2024. “AI Models Collapse When Trained on Recursively Generated Data.” *Nature* 631: 755–759.
- Stiglitz, Joseph E., and Maxim Ventura-Bolet. 2025. “The Impact of AI and Digital Platforms on the Information Ecosystem.” NBER Working Paper No. 34318.
- Topkis, Donald M. 1998. *Supermodularity and Complementarity*. Princeton: Princeton University Press.
- Wu, Yihang, Jiajun Tang, Jinfei Liu, Haifeng Xu, and Fan Yao. 2026. “Do AI Overviews Benefit Search Engines? An Ecosystem Perspective.” arXiv:2601.22493.
- Yan, Shun Yao, Klaus M. Miller, and Bernd Skiera. 2020. “How Does the Adoption of Ad Blockers Affect News Consumption?” arXiv:2005.06840.
- Yu, Hongyeon, Dongchan Kim, and Young-Bum Kim. 2026. “Retrieval Collapses When AI Pollutes the Web.” arXiv:2602.16136.
- Zhang, Yukun, and Tianyang Zhang. 2026. “The Impact of Generative AI on Content Platforms: A Two-Sided Market Analysis with Multi-Dimensional Quality Heterogeneity.” arXiv:2410.13101.
- Zhang, Yukun, and Tianyang Zhang. 2026. “Generative AI as a Non-Convex Supply Shock: Market Bifurcation and Welfare Analysis.” arXiv:2601.12488.
- Zhao, Hangcheng, and Ron Berman. 2026. “Strategic Response of News Publishers to Generative AI.” arXiv:2512.24968.