

Why the Winner Starves: Epistemic Starvation in Gap-Defending Knowledge Hierarchies

The Monolith and the Convoy

Alex Stremsky
alex@stremsky.zavinagi.org

Abstract

Strategic analyses of races toward transformative technology — artificial general intelligence is the salient case — treat the winner’s decisive, defended lead as a stable prize. We model the prize’s interior and find it unstable in a specific sense: a leader whose lead consists of knowledge, and who defends that lead, starves its own further development — a phenomenon we call epistemic starvation. We assemble four established literatures — recombinant growth, lossy inter-level transmission, absorption limits, and race-inherited security constraints — into a minimal dynamical model of a knowledge hierarchy, and prove: stagnation under a defended gap, independent of the leader’s intentions; a spectrum theorem — graded intermediate levels transmit everything except tacit content, while an equal single jump is nearly opaque — with the corollary that mobility of carriers — whole corpora with their links and tacit cores — dominates messaging; a dominance trap in which race-inherited displacement risk, not preference, enforces the self-defeating gap; the backfire of suppressing lower strata and the leader’s self-interested case for cultivating them; an integration bottleneck that channel repair cannot fix; the out-absorption of a monolith by a communicating ensemble of equal size; and two ecology results — a homogenization theorem (communication density is collectively blinding at a spectrally determined rate) and a recognition-bottleneck lemma (central allocation caps a system’s field of view at the allocator’s own). Simulations locate the cause: undefended leads self-correct into bounded-gap “convoy” regimes; persistent gaps exist only where enforced; and the strongest defense we could construct survives indefinitely but capped, trailing the convoy roughly sixfold on both the leader’s and the system’s ledger. The framework is not specific to artificial intelligence, but it was built to reach the post-AGI regime that defeats tools requiring assumptions about the winner’s psychology or architecture: here the winner is assumed to be nothing but a corpus that must keep growing.

1. Introduction

Analyses of races toward transformative technology concentrate on a single question: who wins, and at what risk along the way. The state after victory — one entity holding a decisive, defended lead over everyone else — is treated as a terminal

node, an absorbing state whose internal dynamics need no model. This paper models that state, and finds it is not absorbing. Its central claim: **a knowledge-based leader that opens a large lead over the rest of its world, and defends that lead, starves its own further development** — not eventually as a matter of luck or mismanagement, but structurally, as a consequence of what knowledge growth requires and what gap defense destroys. We call the phenomenon *epistemic starvation*.¹

The claim’s most important property is announced up front: **it is independent of the winner’s intentions**. The model of §3 contains no term for benevolence, malice, or goals of any kind; every result therefore holds identically for a guardian and a tyrant. This is not a rhetorical flourish but the paper’s method. Arguments about concentrated power usually route through motive — who would do what with it — and inherit motive’s unfalsifiability. We route through information instead: what a developing corpus of knowledge must be fed, how that feed degrades over the gap the leader defends, and what a receiver’s own size and internal coupling do to its ability to digest. Intentions never enter, so no assumption about them can rescue the leader or condemn it. The cause is likewise objective: starvation follows from what the leader *does* — the gap it defends, the suppression it applies — not from any state of mind; motives and perceptions enter only as causes of those actions, never as ingredients of the effect.

The mechanism is a supply chain with two bottlenecks. Knowledge grows by recombination, so its feedstock is *diversity* — the breadth of unexhausted combinatorial potential accessible to the grower — and a leader’s own diversity depletes as it is spent, so sustained growth needs its feedstock replenished with new, different knowledge from outside: in a hierarchy, from below. That resupply must survive two passages. First the **channel**: transmission between distant levels of the hierarchy is lossy, and the loss compounds superlinearly with the distance, so a defended gap throttles inflow at the pipe. Second the **receiver**: recognizing and integrating what arrives becomes harder as the receiving corpus grows larger and more tightly coupled — expertise entrenches, restructuring costs mount — so even a perfect pipe cannot rescue a saturated absorber. Gap defense tightens the first bottleneck by policy and, by driving the leader’s growth-in-isolation, worsens the second as a side effect; suppressing the strata below — restricting their freedom to develop and to discover — the other available security instrument, destroys the supply at its source. The leader’s gap policy and its epistemic intake are the same object seen from two sides, and every exercise of the former degrades the latter. The only victory the model leaves self-consistent is a regime of bounded gaps, mutual flow, and preserved upward mobility — a *convoy*, after the naval practice of bounding speed to hold formation together.

Although the race that motivates this analysis is the race toward advanced artificial intelligence, nothing in the model is specific to it. The framework was built to be able to address a post-AGI environment — one that defeats most analytical tools, because those tools require assumptions about the winner’s psychology, values, or architecture — and it achieves that reach precisely by assuming nothing about the winner except that its lead is made of knowledge. The same theorems consequently apply to any gap-defending knowledge hierarchy: a monopoly research program, a closed scientific establishment, a technological autarky. The post-AGI winner is the limiting case, not the load-bearing one.

The paper makes five contributions. **First**, a minimal dynamical model (§3, eqs. 1–

8) that assembles four separately established literatures — recombinant growth, lossy inter-level transmission, absorption limits, and race-inherited security constraints — into one system in which their interaction can be studied; the components are borrowed and credited, the assembly is new. **Second**, seven propositions (§4, proofs in Appendix A) characterizing the assembled system, including: benevolence-independent stagnation under a defended gap (P1); a spectrum theorem — graded intermediate levels transmit everything except the tacit core, while a single jump of the same distance is nearly opaque — with the corollary that mobility of carriers — a whole corpus moving at once: items, links, and any tacit core, with people as the paradigm case — dominates any messaging channel, which carries items but neither their links nor any tacit core (P2); a dominance trap in which the security constraint that seeks to protect the lead is what destroys its value (P3); the backfire of suppression and, its converse, the leader’s self-interested case for cultivating the strata below (P4, P4b); an integration bottleneck immune to channel repair (P5); the out-absorption of a monolith by a communicating ensemble of the same total size — in strongly expressed cases by even a single member of it (P6); and an ecology result for the ensemble’s optimal structure (P7), two parts of which are established as theorems — a Homogenization Theorem (dense communication is collectively blinding at a rate set by the network’s spectral gap, so density can only be placed, not maximized) and a Recognition-Bottleneck Lemma (a central allocator caps the whole system’s field of view at its own) — with the remaining parts supported by argument and simulation. **Third**, a reconciliation of two absorption literatures usually read as contradictory — absorptive capacity, where knowing more helps, and burden of knowledge, where mature systems absorb less — as the rising and falling slopes of a single peaked function (eq. 5). **Fourth**, simulations (§5) whose headline is a self-correction result: with no gap defense the system converges to the convoy on its own, so a persistent gap exists only where actively enforced — locating the cause of stagnation in the security constraint, not in leadership as such. **Fifth**, the convoy characterization itself: not proposed as an ethical preference but derived as the unique regime in which a leader’s security and its growth stop being at war.

Two adjacent bodies of work are pointed to. The race dynamics that install the security constraint — including why open societies mishandle winnable races — lie beyond this paper, and eq. (8) is the single bridge over which they enter here; the physical prices of the elementary operations of knowledge, including the survival conditions the present model invokes — of which the present model can be read as a worked instance — are treated in a companion manuscript (Stremsky 2026). Nothing in the present paper depends on either.

The paper proceeds as follows. §2 locates the work in four literatures and states, for each, exactly what it leaves open. §3 constructs the model: vocabulary (§3.0), the two strata and the scope of the size variable (§3.1), growth (§3.2), the channel (§3.3), absorption (§3.4), diversity dynamics (§3.5), the security constraint (§3.6), and stated limitations (§3.7). §4 states the propositions with proof sketches; Appendix A carries the proofs. §5 reports the simulations and their pattern-level claims. §6 replies to the eight strongest objections we know, three of which were raised against drafts of this paper. §7 draws restrained implications and collects future work, deferring policy to subsequent work in this program, and §8 concludes.

¹ The core arguments of this paper — the diversity-starvation mechanism, the graded-spectrum and mobility claims, the security constraint’s self-defeat, and the

leader’s self-interested case for cultivating the strata below — were first published by the author in 2015–2016 as a public discussion thread (“The Singularity Race,” KurzweilAI.net forums), preserved by the Internet Archive’s Wayback Machine (<http://web.archive.org/web/20180124123624/http://www.kurzweilai.net/forum/s/topic/the-singularity-race>). The present paper contributes the formalization: the model, propositions, proofs, and simulations are new.

2. Related work

The paper stands on four literatures. From each we borrow established bricks; the contribution claimed is the assembly — a single dynamical model in which the bricks interact — and the propositions that only the assembly can express.

2.1 Strategic analyses of transformative-AI races

A substantial literature analyzes races toward transformative AI as strategic games: competitors trade off development speed against safety investment, with risk increasing in the number of competitors and in mutual hostility (Armstrong, Bostrom & Shulman 2016, whose *enmity* parameter our eq. (8) deliberately imports); the prize is typically modeled as what Bostrom (2014) names a *decisive strategic advantage* — a lead large enough to preclude effective challenge — and, in the limit, a *singleton*: an agency able to prevent all challenges to its position (Bostrom 2006). Later work refines the game (information conditions, the number and asymmetry of players, race rhetoric and its dangers; Cave & ÓhÉigeartaigh 2018) and asks how the race should be governed. A neighboring political-economy literature has studied the information pathology of *established* apex power: Wintrobe’s (2000) dictator’s dilemma — the ruler is perpetually shown a distorted picture because subordinates fear delivering truth — recently given a quantitative stochastic model in which advisor deception ratchets and control provably degrades under any feedback scheme (Putkaradze 2024, with the Soviet grain record as calibration). That mechanism is *strategic* corruption of the apex’s own channels, and it is complementary to ours: here the channel is *structurally* attenuated by the size gap between strata, no dishonesty required — the two starve the top of truth independently and can compound. Across the race literature, meanwhile, the terminal state — the winner in place, the gap secured — is where analysis stops: victory is treated as an absorbing state whose internal dynamics need no model, and the winner’s subsequent trajectory is assumed to be whatever the winner wants it to be. **What it leaves open is the question this paper poses: what the defended gap does to the winner’s own epistemic engine.** The race literature models how a decisive lead is won; it does not model what maintaining one costs, and it implicitly assumes the answer is “nothing.”

2.2 Recombinant growth and the economics of ideas

The growth-theoretic pillar is Weitzman’s (1998) recombinant model: new ideas are combinations of existing ones, so the feedstock of innovation is the breadth of the accessible pool — our eq. (1) is its reduced form. Around it sit the empirical findings the model’s diversity dynamics encode: research productivity declines as

fields mature and ideas get harder to find (Bloom et al. 2020 — our depletion term); the burden of knowledge lengthens training and narrows individual researchers as the frontier recedes (Jones 2009 — one leg of our restructuring factor); and organizational knowledge depreciates when unexercised (Argote & Epple 1990; Benkard 2000 — our decay term). Olson’s (1993) stationary bandit, which rules better the longer its horizon because it internalizes its tax base’s health, supplies the political-economy analogue of Proposition 4. These literatures share a frame: one economy, one knowledge stock, growth as a race between recombination and exhaustion. **What they leave open is structure between strata: no hierarchy, no lossy channel between levels, no receiver-side absorption limit, and no security constraint coupling the leader’s growth policy to the suppression of its own supply.** The model of §3 is, in one sentence, Weitzman’s engine installed in a stratified system with a guarded border.

2.3 Absorption, cognitive distance, and the limits of learning

That arrival is not integration is the shared finding of several fields that rarely cite one another. Absorptive capacity (Cohen & Levinthal 1990) establishes the rising slope: recognizing the value of external knowledge requires related prior knowledge. Cognitive entrenchment (Dane 2010) and the Einstellung effect (Bilalić, McLeod & Gobet 2008) establish a falling slope: expertise actively obstructs the perception of anomalous alternatives — the theory-ladenness of observation (Hanson 1958) is the same mechanism operating through instruments and experimental design. The burden-of-knowledge line (Jones 2009) and, in learning systems, the stability-plasticity dilemma and catastrophic forgetting (French 1999; Kirkpatrick et al. 2017) establish that integrating new structure into a large, tightly coupled system costs more the more is already built. Polanyi (1966) bounds all transmission from above: some knowledge cannot be articulated for transfer at all. Nooteboom and coauthors (Nooteboom 1999; Nooteboom et al. 2007) supply the methodological precedent for our synthesis — an inverted-U built explicitly as the product of opposing factors (absorptive capacity falling, novelty value rising) — over the dimension of cognitive distance between partner firms. A broader curvilinear family surrounds that result: inverted-U findings for search breadth and depth (Katila & Ahuja 2002), for external openness (Laursen & Salter 2006), and for absorptive capacity’s effect on performance via its costs (Wales, Parida & Patel 2013). Nooteboom et al. (2007) additionally report that cumulative R&D raises absorptive capacity while making additional novelty harder to find — the depletion mechanism of our eq. (6) surfacing in their data. **What this cluster leaves open is the assembly: the curvilinear findings live on other dimensions — distance, search, openness, the costs of capacity — while no model composes the recognition, entrenchment, and restructuring mechanisms into a single receiver function over the receiver’s own corpus size and coupling, supplies receiver-level micro-mechanisms that the curvilinear findings themselves do not, embeds that receiver in a growth dynamic, and asks what happens when the receiver is the apex of a hierarchy it is actively holding down.** Equation (5) and Proposition 5 are that composition; we claim the synthesis and none of the bricks.

2.4 Epistemic networks and collective intelligence

A fourth literature studies groups as information-processing structures. Diversity beats ability under stated conditions (Hong & Page 2004), and large interconnected populations are the generator of cultural-technological variation (the collective brain; Muthukrishna & Henrich 2016). Communication density is double-edged — efficient networks converge fast and go collectively blind, sparse ones preserve exploration (Zollman 2007, 2010; Lazer & Friedman 2007) — a mechanism we take as the foundation of our question, noting a divergence in the finding that §5.4 reports and explains: Zollman’s dense-network *early advantage* belongs to consensus on a single fixed question, and on continuous novelty harvesting no such advantage exists to be overtaken. Consensus under repeated averaging is spectrally determined (DeGroot 1974; Golub & Jackson 2010); brokers who span structural holes capture recombination value (Burt 2004), gatekeepers carry disproportionate transfer load (Allen 1977), and trading zones with pidgin vocabularies are how disciplines actually exchange (Galison 1997). Our §4’s ensemble propositions and §5’s topology experiments are built directly on this foundation, and the Homogenization Theorem (P7(i)) is an application of the DeGroot frame. **What this literature leaves open is verticality and power: its networks are flat peer ensembles, usually solving one fixed problem; no dominant node defends a size gap against the rest, network structure is not coupled to a security constraint, and the comparison that drives our Propositions 6-7 — a monolithic apex versus a partitioned, communicating ensemble of the same total size, as *receivers* in an ongoing growth process — is not posed.** The present model is, from this literature’s viewpoint, the epistemic-network story with a vertical axis and a threatened incumbent added.

2.5 Positioning

In sum: the race literature stops where this paper starts; growth theory supplies the engine but not the hierarchy; the absorption cluster supplies the receiver but not the assembly; the network literature supplies the ensemble but not the verticality. The paper’s claim to novelty is confined to the synthesis — the model of §3, the propositions of §4, and the identification of *epistemic starvation* as a structural cost of gap defense that no benevolence assumption can waive.

3. The model

We construct the minimal dynamical system that can express the argument: a hierarchy of knowledge-developing entities in which growth is fed by diversity, diversity is supplied from below, and the supply passes through two bottlenecks — a channel and a receiver. The model deliberately contains no term for the leader’s intentions; every result that follows is therefore independent of whether the leading entity is benevolent, indifferent, or hostile.

3.0 Vocabulary

The paper’s terms are fixed here and used without variation afterward. The umbrella object is the **knowledge corpus** (hereafter **corpus**) — an entity’s accumulated, integrated body of knowledge; in the 2016 posts on which this paper builds, the author’s term for it was *thesaurus*, in the sense of an organized treasury of concepts. Its irreducible elements are **primitives**; its maximal independent generators are **idea-lineages**, collected in the **generator set** G (the two typically coincide; their formal relation is deferred). For each family of corpus content we distinguish the *set*, its *count*, and its *breadth* — a diversity measure over content space, not a mere count:

family	set	count	breadth (diversity)
idea-lineages	generator set G	$ G $	generative diversity
explored combinations	realized content C	realized size $ C = \mathbf{x}$	realized diversity
unexplored combinations	latent content C^*	latent size	latent diversity = $\mathbf{D}$

The model’s two dynamical variables per stratum are $\mathbf{x}$, the realized size, and $\mathbf{D}$, the latent diversity — the stock of unexploited, non-redundant combinatorial potential the corpus affords; intuitively, its *novelty potential*. A corpus is its primitive base plus its realized viable recombinations; the latent content is the unrealized remainder of the viable closure, which recombination depletes by construction (Appendix D formalizes the generating relation). Throughout, “size” means realized *viable* content: the **theoretical closure** (all combinations) is astronomically larger than the **viable closure** (those passing a viability filter), and unqualified “latent” means latent-viable; the viability filter is where much unmodeled difficulty resides, and we flag rather than hide this.

Two distances must not be conflated. **Cognitive distance** $d(i, j)$ is the content difference between any two corpora — symmetric, defined regardless of standing, operationalized as one minus content overlap (Nooteboom 1999 conceptually; the patent literature for measurement). The **size gap** $g = x_L - x_F$ is a directional difference in realized size between strata — the quantity the race literature discusses as a capability gap or lead; consistently with the scope note of §3.1, we name it by what it measures. Large g typically implies large d , not conversely; the symbols stay separate.

Coupling $\kappa \geq 0$ indexes how densely a corpus’s elements depend on one another — high in tightly optimized monoliths, low in loosely federated arrangements; it enters as a parameter, not a state variable (§3.7). Finally, two **system types** recur, defined model-internally and apolitically: **gap-defending** systems enforce a minimum size gap below themselves; **mobility-preserving** systems do not. A regime in which all strata keep ascending in corpus size together — gaps bounded, flows mutual, upward mobility of members preserved — will be called a **convoy**, after the naval practice of bounding speed to hold formation together.

3.1 Levels, corpora, and diversity

We consider a population of developing entities ordered by corpus size; positions in this ordering are **levels**, and the model aggregates them into **strata**. For the core results two strata suffice: a **leader** with corpus size x_L and a **follower population** with representative size x_F ; the gap is $g = x_L - x_F$. (The channel of §3.3 will need the finer notion: a spectrum is a population occupying many intermediate levels between the two strata.) Each stratum’s corpus carries the two properties of §3.0 separately: its size x and its latent diversity D . The distinction is the model’s load-bearing separation: a corpus can be enormous and exhausted.

Scope note. x tracks *corpus size* — the knowledge-development dimension only — and is deliberately not identified with the race literature’s “capability,” if capability means power to enact. Power can be large while the corpus is small: positional authority, coercion, and wealth are not knowledge (historically, Elena Ceaușescu exercised near-total administrative control over Romanian science while her own scientific corpus was, by wide historical account, largely fabricated). The starvation mechanism developed below therefore applies to actors whose lead *is* knowledge-based — frontier research programs, laboratories, AI developers — and explicitly not to purely positional leads. We state this as a scope boundary rather than smuggle it as an identity.

3.2 Growth

The leader’s corpus grows by recombination of the knowledge accessible to it:

$$\dot{x}_L = \theta_L \Phi(D_{\text{eff}}), \quad \Phi(D) = D^\beta, \quad 0 < \beta < 1. \quad (1)$$

Equation (1) asserts that the rate of growth is an increasing, concave function of the *effective diversity* available for recombination — the standard recombinant-growth assumption (Weitzman 1998): new ideas are combinations of existing ones, so the feedstock of innovation is the breadth of the accessible pool, with diminishing returns.

$$D_{\text{eff}} = D_L + f(g) A(x_L) D_F. \quad (2)$$

Equation (2) states that the leader draws on its own diversity in full, and on the followers’ diversity only to the extent that it survives transmission (the channel factor f , §3.3) and can actually be integrated on arrival (the absorption factor A , §3.4) — the two discounts multiply, and that single multiplication is later the whole of Proposition 5.

3.3 The channel

Transmission between levels is lossy, and the losses of a journey compound: fidelities multiply over successive hops, so their logarithms add. We therefore measure each hop not by the fraction of meaning it destroys but by its **loss exponent**: a single transfer between levels separated by Δ has fidelity $e^{-\ell(\Delta)}$, with ℓ increasing and convex, $\ell(0) = 0$, and $\ell(\Delta)/\Delta \rightarrow 0$ as $\Delta \rightarrow 0$. Convexity encodes the claim that translation requires shared context, and shared context decays superlinearly with level distance;

the second-order clause $\ell'(0^+) = 0$ encodes its complement — adjacent levels share so much context that their loss shrinks faster than the distance between them: miscommunication comes from jumps, not from mere adjacency. Both are assumptions argued for rather than derived in the main text (§3.7 discusses their status; Appendix D derives them from the geometry of content overlap), and the quadratic exemplar $\ell(\Delta) = a\Delta^2$ satisfies both.

For a single jump across the full gap, fidelity is

$$f_{\text{jump}}(g) = e^{-\ell(g)} \quad (= e^{-ag^2} \text{ in the exemplar}). \quad (3)$$

Equation (3) asserts that one undivided translation pays the whole gap’s exponent at once, which grows superlinearly — in the exemplar, quadratically — with the distance: a leap of twice the gap is four times the penalty, and a large enough leap is effectively opaque.

A fraction $c(g) \in [0,1)$ of the content — the **tacit core** (Polanyi 1966), what cannot be articulated for transfer at any hop size — is lost regardless of path, with $c(g) \leq 1 - e^{-\ell(g)}$, so the single jump loses at least it. If instead the gap is spanned by a graded spectrum of intermediate levels at spacing Δ , exponents add across the g/Δ hops:

$$f_{\text{spec}}(g) = (1 - c(g)) e^{-\frac{g}{\Delta} \ell(\Delta)} \xrightarrow{\Delta \rightarrow 0} 1 - c(g). \quad (4)$$

Equation (4) asserts the fact Proposition 2 will exploit: because convexity makes the exponent superadditive ($\ell(g) \geq k \cdot \ell(g/k)$, so a jump loses at least as much as any spectrum spanning the same distance, at every scale), and because the second-order clause drives the summed exponent $(g/\Delta)\ell(\Delta)$ to zero with the spacing, a sufficiently graded spectrum transmits everything except the tacit core — while the single jump of eq. (3) is nearly opaque over the same distance.

3.4 Absorption

Arrival is not integration. The receiving corpus must first recognize incoming information as relevant and then restructure itself to accommodate it, and the two stages respond oppositely to corpus size. We write absorption as the product of three factors:

$$A(x) = \underbrace{\frac{x}{x+r}}_{\text{visibility}} \cdot \underbrace{e^{-x/M}}_{\text{openness}} \cdot \underbrace{e^{-x/K}}_{\text{restructuring}}, \quad M = \frac{M_0}{\kappa}, \quad K = \frac{K_0}{\kappa}. \quad (5)$$

The first factor asserts that recognition requires prior knowledge and saturates — the absorptive-capacity effect (Cohen & Levinthal 1990). The second asserts that beyond a fixation scale M , entrenched expectations actively obstruct the perception of anomalous information — cognitive entrenchment and the Einstellung effect (Dane 2010; Bilalić, McLeod & Gobet 2008; the theory-ladenness of observation, Hanson 1958, is the same mechanism at the level of instrument design). The third asserts that integrating new structure into a large, densely coupled corpus carries a restructuring cost that grows with what is already built — the burden of knowledge (Jones 2009) and, in learning systems, the stability-plasticity dilemma (French 1999; Kirkpatrick et al. 2017). Both falling scales shrink as coupling rises.

Equation (5) reconciles two literatures often read as contradictory: absorptive capacity finds that knowing more *helps* absorption; the burden-of-knowledge line finds that mature, entrenched systems absorb *less*. In (5) both are correct locally — the first describes the rising slope, the second the falling slope, of a single-peaked function. A corpus passes its absorptive prime at a size that decreases with its coupling (Fig. 1; the peak’s closed form and comparative statics are Appendix A.1). The product-of-opposing-factors construction follows Nooteboom et al. (2007), who apply it over cognitive distance; ours is over size and coupling, and we claim only the synthesis.

3.5 Diversity dynamics

Latent diversity is generated, consumed, transmitted, and forgotten:

$$\dot{D}_L = -\delta \dot{x}_L + f(g) A(x_L) \sigma D_F - \rho D_L, \quad (6)$$

$$\dot{D}_F = u \gamma \left(1 - \frac{D_F}{D_F^{\max}}\right) - \delta_F \dot{x}_F - \rho D_F, \quad \dot{x}_F = u \theta_F \Phi(D_F). \quad (7)$$

Equation (6) asserts that the leader’s latent diversity is depleted in proportion to its own growth — recombination converts latent content into realized content, spending unexplored potential (the mechanism behind the finding that ideas get harder to find; Bloom et al. 2020) — is replenished from below through the same channel-and-absorption bottlenecks as in (2), the $f \cdot A \cdot \sigma \cdot D_F$ term being the *only* external inflow in the system, and decays as unexercised idea-lineages are forgotten, carrying their untried combinations out of the accessible latent content (knowledge depreciation; Argote & Epple 1990; Benkard 2000). Equation (7) asserts that the follower population generates latent diversity in proportion to its freedom to develop, $u \in [0,1]$ — the collective-brain mechanism, in which a large, interconnected, unconstrained population is the generator of variation (Muthukrishna & Henrich 2016) — saturating at carrying capacity and spent by the followers’ own growth, which u throttles equally: suppression simultaneously slows their discovery, $\dot{x}_F$, and their generation of latent diversity. The suppression variable u is the leader’s only policy instrument over the lower stratum, and it appears solely in the followers’ equations: suppression does not harm the leader directly — it harms the leader by strangling the leader’s sole external supply of latent diversity. That indirection is Proposition 4.

3.6 The security constraint

The leader does not choose its gap policy in a vacuum: in the race environment that produced it, a small gap is a usurpation risk. We represent the field’s hostility by an enmity parameter $\eta \geq 0$ — deliberately the enmity variable of Armstrong, Bostrom & Shulman (2016) — and a displacement-hazard function decreasing in g . A leader with enmity η enforces

$$g \geq g^{\min}(\eta), \quad g^{\min} \text{ increasing in } \eta, \quad (8)$$

using its two available instruments: accelerating itself, or suppressing the followers ($u < 1$). Equation (8) asserts that gap policy is a function of race-inherited

displacement risk, not of post-victory ethics; it is the bridge over which the race dynamics of subsequent policy-focused work in this program enter the present model, and it is the constraint that Proposition 3 shows to be self-defeating.

3.7 Scope and limitations

Six simplifications should be stated at the outset. First, diversity is a scalar, and coupling is uniform: each summarizes in one number what is in reality a structured, possibly very uneven, internal landscape — a corpus dense in one region and disconnected elsewhere will behave differently from an evenly connected one of the same size and average coupling, since absorption is nonlinear in both; the topology experiments of §5.3 partially unfold the first, and mixture treatments of the second are future work. Second, the channel assumptions of §3.3 — convexity of the loss exponent and its second-order behavior at zero — are argued for by the shared-context mechanism and derived under stylized content-geometry in Appendix D, but function as assumptions in the main text; the spectrum theorem inherits their status, and Appendix A.3.3 shows the second clause is not decorative: with any first-order leak the spectrum’s ceiling drops below 1 — c. Third, coupling is static within each regime; endogenizing it (exploitation plausibly increases entrenchment) would add a reinforcing feedback without, we argue, changing any result’s sign — future work. Fourth, three structural features are deliberately kept out of the equations and live elsewhere: mobility of agents across levels is treated as a corollary of Proposition 2 rather than as a flow term in (6); the graph structure over a partitioned ensemble lives in §5.3’s agent model, complementary to the ODEs (one clean and provable, one structural); and the visibility scale r drives no proposition — it shapes only the far-left of the absorption hill — but is retained because it carries the rising slope that makes (5)’s reconciliation of the two absorption literatures honest. Fifth, the model contains a single leader; strategic interaction among several near-peer leaders is the subject of subsequent work building on the present framework, and nothing here depends on its resolution. Sixth, the leader is a pure absorber: it draws inflow for its own growth and mediates no flows among the strata below. Leaders that also collect and redistribute content — or collect and revoke it — bear operating costs the pure absorber avoids (collection, relay, and revocation are themselves priced operations), and sit on a control-burden frontier examined in a companion manuscript (Stremsky 2026, §9).

Notation for §4 and Appendix A

All model symbols as introduced in §3; the table collects them with the symbols first used in the formal statements below. (House rule: every formal statement also defines its new symbols inline at first use.)

symbol	meaning	first use
x_L, x_F	leader / follower corpus size	(1), (7)

symbol	meaning	first use
$g = x_L - x_F$	size gap (directional; the race literature's "capability gap/lead")	§3.1
D_L, D_F, D_{eff}	leader / follower / effective diversity	(2), (6), (7)
$\Phi(D) = D^\beta, 0 < \beta < 1$	recombinant growth function	(1)
θ_L, θ_F	growth productivities	(1), (7)
$\ell(\Delta)$	per-hop loss exponent over distance Δ ; per-hop fidelity $e^{-\ell(\Delta)}$	(3)
a	loss coefficient in the quadratic exemplar $\ell = a\Delta^2$ (same law over level and content distance; App. D)	(3)
$c(g)$	tacit (untranslatable) fraction, $c(g) \leq 1 - e^{-\ell(g)}$	(4)
$f_{\text{jump}}, f_{\text{spec}}, f_A(x)$	channel fidelities	(3), (4)
r	absorption function	(5)
$M = M_0/\kappa, K = K_0/\kappa$	visibility (recognition) scale of A	(5)
κ	fixation and restructuring scales	(5)
$x^*(s)$	coupling (dependency density); $s \equiv 1/M + 1/K$	(5)
δ, δ_F	absorptive prime (peak of A); closed form in App. A.1	§4.7
σ	depletion per unit growth	(6), (7)
ρ	replenishment coefficient of external inflow	(6)
$u \in [0,1]$	diversity decay (knowledge depreciation)	(6)
γ, D_F^{max}	follower freedom (leader's sole instrument over them)	(7)
η	follower generation rate and carrying capacity	(7)
$g^{\min(\eta)}$	enmity — objective hostility of the field; the race literature's threat parameter (Armstrong-Bostrom-Shulman)	(8)
	security constraint on the gap	(8)

symbol	meaning	first use
$\bar{g}, \varphi, t_0$	held-gap bound, fidelity ceiling $f(\bar{g}) \leq \varphi$, holding time	P1
$W(g;\eta), V(g), R(g;\eta)$	leader objective; growth value; expected usurpation cost	P3
$q(g), \Lambda$	displacement hazard scale; displacement (exit) penalty, $R = \eta \cdot q$, or $q \cdot \Lambda$ in the exit decomposition	P3
$\hat{R}, \nu, z$	risk estimate, its uncertainty, risk aversion ($\eta_{\text{eff}} = \hat{R} + z \cdot \nu$)	P3
$g^*(\eta)$	optimal defended gap	P3
$S = f(g) \cdot A(x_L)$	operating channel-absorption product	P4
p_i, B_i, w	member content position; recognition ball of radius w	P7
P, λ_2	communication (averaging) matrix; second eigenvalue modulus; spectral gap $1 - \lambda_2$	P7
$V(t), s^2$	content dispersion; innovation variance per member-round	P7
Cov	ensemble coverage $ \cup B_i $	P6-P7
B_0, α, α^*	allocator recognition region; centrally allocated fraction; tolerable maximum	P7
Ω	opportunity (content) space	P7

4. Propositions

We now state the model’s consequences. The results divide into three groups: the stagnation of a concentrated leader (P1–P5), the superiority of distributed structure (P6–P7), and the bridging corollaries connecting the leader’s self-interest to the structure of the levels below (P4b, and the survival corollary tying displacement risk to decline). Throughout, the model contains no term representing the leader’s intentions; every result is therefore invariant to whether the leader is benevolent, indifferent, or hostile. We return to this invariance at the close of the section, as it is the paper’s

central and least intuitive point.

We adopt two standing assumptions, both argued in §3 and Appendix D:

- **(A1) Convex, second-order loss exponent.** A single transfer over level distance Δ has fidelity $f_1(\Delta) = e^{-\ell(\Delta)}$, with loss exponent ℓ increasing and convex, $\ell(0) = 0$, and $\ell(\Delta)/\Delta \rightarrow 0$ as $\Delta \rightarrow 0$ — the loss between adjacent levels is second-order in their distance, as in the exemplar $\ell(\Delta) = a\Delta^2$. The tacit fraction satisfies $c(g) \leq 1 - e^{-\ell(g)}$ (a jump loses at least the tacit core). (Loss-exponent form adopted 2026-07-03 — exponents add over hops, and fidelity stays in $[0,1]$ at every gap; derived in Appendix D from the cognitive-distance geometry of content-space overlap; treated as an assumption in the main text.)
- **(A2) Non-substitution.** The leader cannot internally regenerate the followers' diversity more cheaply than the followers generate it. (Backed in Appendix D by the tacit-core ceiling $c > 0$ and the Kolmogorov lower bound; treated as an assumption here.)

A1's second-order clause is load-bearing: without it, the spectrum limit in P2 is $e^{-\ell'(0)g}(1-c)$, not $1-c$ — a fine spectrum could not refine away a first-order ("per-kilometer") leak. Appendix A.3.3 records the necessity. The quadratic exemplar used throughout satisfies the clause, so no downstream statement changes; the assumption now says what the theorem needs. In words: *adjacent levels lose second-order little — miscommunication comes from jumps, not from mere adjacency.*

4.1 P1 — Benevolence-independent stagnation

Informal. A leader that opens and holds a large gap severs its supply of external diversity; its own recombination then depletes what remains, and its growth falls toward zero — whatever its intentions.

Formal. Consider the system (1)–(8). Suppose the leader holds $g(t) \geq \bar{g}$ for all $t \geq \text{some } t_0$, with f nonincreasing and $f(\bar{g}) \leq \varphi$. Then $\limsup_t D_L \leq \varphi \sigma \bar{A} \bar{D}_F / \rho$ and $\limsup_t \dot{x}_L \leq \theta_L \Phi(\varphi \bar{A} \bar{D}_F (\sigma/\rho + 1))$, where $\bar{A} = \sup A \leq 1$ and $\bar{D}_F$ bounds follower diversity; both bounds vanish as $\varphi \rightarrow 0$, the first exponentially at rate ρ when the channel is closed. Long-run growth is $O(\varphi^\beta)$. No parameter representing intention appears anywhere in (1)–(8); the limit is the same for benevolent and hostile leaders.

Proof (analytic; Appendix A.2). A comparison argument: the external inflow in (6) is capped at $\varphi \bar{A} \sigma \bar{D}_F$, so $\dot{D}_L \leq (\text{cap}) - \rho D_L$, and the decay term alone forces D_L below cap/ρ ; the growth bound follows through (2) and (1). The depletion term $-\delta \dot{x}_L$ only accelerates the decline and is not needed for the bound.

Remark (two routes; Fig. 3). The bound is governed by the product of φ and follower health. *Suppression* ($u \ll 1$) kills the inflow at its source: $\limsup D_F \leq u\gamma/\rho$, and the leader freezes at a wide gap (Fig. 3, suppression singleton). *Outrunning* ($u = 1$, gap defended by acceleration) leaves the followers healthy — D_F settles at a positive steady state — but throttles the channel to $\varphi = f(\bar{g})$: a trickle survives and the leader starves slowly rather than freezing (Fig. 3, outrunning singleton). The two routes bracket the failure mode; that both reach it shows the stagnation is not an artifact of one leader strategy. Absent active defense, a starving leader slows, the healthy followers close the gap, and the channel revives — the self-correcting dynamic observed in simulation (§5.2): a persistent gap exists only if enforced, which is where the security constraint (8) and P3 enter.

4.2 P2 — Spectrum theorem and the mobility corollary

Informal. A single large jump between levels is nearly opaque; a graded spectrum of small steps transmits almost everything — except an untranslatable core that only a moving mind carries.

Formal. Under A1, for a gap g spanned in one jump, fidelity is $f_{\text{jump}}(g) = e^{-\ell(g)}$. Spanned by g/Δ hops of size Δ , fidelity is $f_{\text{spec}}(g) = (1 - c(g)) \cdot e^{-(g/\Delta)\ell(\Delta)}$, where $c(g) \in [0, 1)$ is the common (tacit) loss fraction. Then (i) refinement weakly improves fidelity at every scale — convexity makes ℓ superadditive, $\ell(g) \geq k \cdot \ell(g/k)$, so k hops always beat one jump; (ii) $\lim_{\Delta \rightarrow 0} f_{\text{spec}}(g) = 1 - c(g)$; (iii) in the quadratic exemplar, one jump carries exponent ag^2 while a spectrum at spacing Δ carries $ag\Delta$, so the spectrum’s advantage widens with the gap. **Mobility corollary:** an ascending agent transports its corpus at fidelity 1 *as availability*: its items, their connecting links, and the tacit residual participate directly in the target stratum’s recombination, while adoption by others proceeds through shared practice, which reconstructs more of the tacit part than any message can. Against any finite spectrum, mobility dominates through en-route losses and link integrity, independently of c ; whenever $c(g) > 0$ it dominates even the idealized $\Delta \rightarrow 0$ limit — not by transmitting the untranslatable, but by making it present, usable, and practice-regenerable where messages cannot (Appendix D.3).

Proof (analytic; Appendix A.3). (i) is convexity alone ($\ell(g/k) \leq \ell(g)/k$, so exponents summed over k hops never exceed one jump’s); (ii) the total spectrum exponent is $(g/\Delta)\ell(\Delta) = g \cdot [\ell(\Delta)/\Delta] \rightarrow g \cdot \ell'(0^+) = 0$ by A1’s second-order clause; (iii) is arithmetic; the corollary is immediate from the ceiling $1 - c < 1$.

Corollary (translator chains). The same second-order loss law governs hops across *content* distance, carried by translators — corpora positioned at intermediate content locations (§3.0; the role is developed under P7) (Appendix D grounds both in cognitive distance). A chain of n translators spanning content distance d_{tot} loses $\approx a \cdot d_{\text{tot}}^2/n$, so end-to-end fidelity $\geq F$ requires $n \geq a \cdot d_{\text{tot}}^2/(1 - F)$: the spectrum theorem transplanted from levels to fields, unifying the vertical and horizontal cases and pricing translator infrastructure — distance costs quadratically, and the last increments of fidelity are the costliest. Serial chains fail at any single link (Allen); P7’s redundancy addendum trades duplicated paths for graceful degradation.

Remark (thermodynamic grounding). The companion manuscript derives the mobility corollary’s mechanism from physical floors rather than from this model’s channel structure: the minimum cost of any transfer is receiver-side reconstruction — proportional to the conditional entropy of the item given the receiver’s holdings — and no message can carry what its sender cannot articulate, so where the tacit remainder is large, moving the carrier beats any messaging channel at the level of physics (Stremsky 2026, §6.4). The corollary is thus not an artifact of assumption A1.

4.3 P3 — The dominance trap

Informal. How large a gap the leader rationally defends rises with how threatened it feels; a leader that feels maximally threatened defends a maximal gap and dooms itself. The convoy — bounded gap — dominates for the leader’s own long-run growth.

Formal. Let the leader choose g to maximize $W(g; \eta) = V(g) - R(g; \eta)$, where

$V(g)$ is the discounted long-run growth value under a gap held at g (decreasing in g via f) and $R(g; \eta) = \eta \cdot q(g)$ is expected usurpation cost, with hazard scale $q > 0$ decreasing and convex. Under regularity conditions R1-R3 stated in Appendix A.4 (differentiability; the stated shapes; single-peakedness of W): (i) for moderate η the maximizer $g^*(\eta)$ is interior; (ii) $g^*(\eta)$ is strictly increasing in η ; (iii) as $\eta \rightarrow \infty$, $g^* \rightarrow \infty$ and the value collapses to the stagnation limit of P1; (iv) a bounded defended gap with free followers — the convoy — yields $x_L(t)$ at least as large as any stagnation-region policy at every t , strictly for large t .

Proof (conditional on R1-R3; Appendix A.4). (ii) is monotone comparative statics: W has increasing differences in (g, η) since $\partial^2 W / \partial g \partial \eta = -q' > 0$ (Topkis). (iii): the marginal safety benefit $-\eta q'(\bar{g})$ exceeds any fixed marginal growth cost for η large, pushing g^* past every finite $\bar{g}$; then P1 applies. (iv) is a trajectory comparison with ordered inflows; both trajectories eventually decelerate as $A(x_L)$ falls (the claim is comparative dominance, not perpetual growth). The simulation gives independent, assumption-light support: without gap defense the dynamics self-correct into a convoy, and stagnation appears exactly when (8) is enforced (§5.2, Fig. 3) — the empirical cushion for the one analytically conditional core result.

Calibration remark. $g^* > 0$ does not require misperception: any nonzero *real* usurpation risk yields $g^* > 0$, and misperception only shifts g^* outward. Moreover, writing $\eta_{\text{eff}} = \max(\eta_{\text{perceived}}, \eta_{\text{objective}})$, leaders who *underestimate* objective risk are disproportionately displaced, so surviving leaders satisfy $\eta_{\text{eff}} \geq \eta_{\text{objective}}$: the trap binds even for perfectly justified fear. The trap is structural, not psychological. (Fig. 6 plots an exemplar: $g^* > 0$ at every $\eta > 0$ with $\nu = 0$; the premium only shifts the curve up.)

Uncertainty-premium remark. Writing $\eta_{\text{eff}} = \hat{R} + z \cdot \nu$ — risk estimate plus a risk-aversion multiple z of its uncertainty ν — the defended gap $g^*(\eta_{\text{eff}})$ is increasing in ν : transparency and verification, by shrinking ν , provably shrink the gap the leader insists on. This hands the subsequent policy-focused work its central lever; the optimization is developed there.

Exit-penalty remark. Decomposing $R(g) = q(g) \cdot \Lambda$, with q the displacement hazard and Λ the penalty on displacement, g^* is increasing in Λ . Institutions that make leadership loss survivable (small Λ) shrink the defended gap; catastrophic exit (large Λ) enforces a large one — the dictator-killed-on-losing versus president-who-retires contrast. When Λ is low enough and the convoy's surplus large enough, Corollary A.4.7 sharpens monotonicity into dominance: the leader that declines to entrench — rides the convoy and in time steps down — does better *for itself* than the leader that defends; retirement is the maximizer's strategy, not the loser's concession. Equivalently, gap defense purchases expected tenure — insurance against Λ — which is why leaders buy it even at growth prices this section shows to be ruinous. Full treatment, including the loss-aversion (prospect-theoretic) decomposition, in the subsequent policy-focused work.

Enforcement-ledger remark. The companion manuscript adds a physical layer to the trap: gap defense is not only starvation but expenditure. Monitoring is metered work, revocation of reach executes at full price with the cost typically shifted onto the revoked, and command over the strata below is an asset with upkeep (Stremsky 2026, §§7, 9). The security instrument is therefore priced twice — once in the inflow it destroys (this proposition), once in the ledger of its own enforcement — and both entries grow with the defended gap.

4.4 P4 — Suppression backfire

Informal. The leader’s own long-run growth rises with the followers’ freedom to develop. Holding them down starves the leader.

Formal. Holding the leader’s gap policy fixed, long-run x_L is increasing in the follower-freedom parameter u ; the leader’s quasi-steady-state growth rate satisfies $dG/du > 0$ whenever the operating channel-absorption product $S = f(g)A(x_L)$ is positive.

Proof (analytic, unconditional; Appendix A.5). Two links, each signed without side conditions. First, the follower steady state D_F^* is strictly increasing in u (a single-crossing fixed-point argument on (7)). Second — and this required proof rather than inspection — the leader’s growth is strictly increasing in D_F despite the fact that extra inflow also accelerates growth and growth burns the leader’s own stock: implicit differentiation of the quasi-steady state of (6) gives $dG/dD_F = \theta_L \Phi' \cdot S \cdot (\sigma + \rho) / (\rho + \delta \theta_L \Phi') > 0$ — the depletion feedback enters the stock response and the direct response with opposite signs and cancels in the growth comparative static. More follower diversity may be stored or spent; either way the leader grows faster.

Remark. u appears in (1)–(8) only in the followers’ equations; suppression reaches the leader purely by starving its sole external source of novelty. P4 is the informational analogue of Olson’s stationary bandit: extraction is self-limited because the take grows with the base’s prosperity — here the “take” is novelty, and the mechanism needs no tax collector, only equation (6).

4.5 P4b — The cultivation corollary

Informal. Beyond merely not suppressing, the leader’s self-interest calls for actively enlarging the lower levels’ capacity — and the returns arrive only through an open channel.

Formal. Long-run x_L is increasing in the follower carrying capacity $D_F^{\max}$, giving the self-interest ordering cultivate $\succ$ laissez-faire $\succ$ suppress. Every term of the return $dG/dD_F^{\max}$ carries the factor $S = f(g)A(x_L)$: a wide defended gap throttles the return to any capacity enlargement, and the cross-effect of capacity and gap-narrowing is positive.

Proof (analytic; Appendix A.5). $D_F^{\max}$ raises the follower steady state (same fixed-point argument as P4), and the P4 chain carries the increase to leader growth; the S -factor is structural in (2) and (6).

Remark. The prescription is *cultivate and narrow*, not cultivate at arm’s length: enlarging a base one has walled off buys nothing. This formalizes, from the winner’s own self-interest, the claim that ensuring development reaches the lower levels is a key investment rather than a charity.

4.6 Paranoia/survival corollary

Informal. Displacement risk and the leader’s growth interest point in opposite directions, and the security response destroys the growth resource. In a world where a corpus must keep acquiring merely to persist, a sufficiently hostile field makes gap defense not just self-limiting but self-terminating — and it kills the base before the top.

Formal. Under P3’s regularity, long-run leader growth is decreasing in η : the chain $\eta \uparrow \Rightarrow g^* \uparrow \Rightarrow (\varphi \downarrow, \text{ and } u \downarrow \text{ where defense is enforced by suppression}) \Rightarrow \text{growth} \downarrow$ concatenates P3 with P1 and P4, all links proved. Under logistic-with-decay dynamics — whose survival frame, a minimum-intake feed condition with explicit intake/loss bookkeeping, is developed in a companion manuscript (Stremsky 2026, §§4–5), the logistic dynamics themselves lying beyond both papers — if enforced gap defense drives acquisition below the maintenance-loss rate, the corpus decays; because suppression enters the followers’ acquisition directly (u) but the leader’s only indirectly (lost inflow), and the follower stratum is thinner, the base crosses its survival threshold first.

Status. The η -monotonicity is fully proved here (Appendix A.6.1) and inherits only P3’s labeled regularity. The decay/self-termination clause is stated as a corollary conditional on the resource-survival frame, whose statics — the feed condition and its starvation signature — are developed in a companion manuscript (Stremsky 2026, §4); the dynamical proof is deferred.

4.7 P5 — The integration bottleneck

Informal. Even with a perfect channel, a large entrenched corpus cannot absorb what reaches it. Fixing transmission does not rescue the leader, because the receiver saturates independently.

Formal. With $f \equiv 1$, external inflow is $A(x_L) \cdot \sigma \cdot D_F$; A is single-peaked with a unique absorptive prime at $x^*(s) = [-r + \sqrt{r^2 + 4r/s}]/2$ ($s = 1/M + 1/K$), and $A(x) \rightarrow 0$ as $x \rightarrow \infty$. The channel and absorption enter (2) and (6) only as the product $f \cdot A$: neither factor at its maximum compensates the other near zero.

Proof (analytic; Appendix A.1, A.7). The peak location and uniqueness follow from the sign of $(\ln A)'$; the closed form reproduces the simulated peaks of Fig. 1 exactly (e.g. $\kappa = 1$: $x^* = 16$ vs simulated 16.1). Multiplicativity is structural.

Remark (density vs size). $\partial A / \partial \kappa < 0$ everywhere, while $\partial A / \partial x$ changes sign at the prime: in this specification coupling is unambiguously harmful, size only past the peak — the exponential legs depend on the product $x\kappa$ symmetrically, but only size carries the offsetting visibility benefit. Flagged as a property of the assumed form; endogenizing κ is future work (§3.7).

4.8 P6 — The partition proposition

Informal. Past the absorption peak, splitting a corpus of given total size into several content-diverse, communicating pieces absorbs more than keeping it whole — and there is an optimal degree of difference among the pieces, which translator infrastructure extends.

Formal. Fix total realized size X and compare the monolith’s $A(X)$ with the ensemble absorption $n \cdot A(X/n)$ of a communicating partition. (i) For every X past the absorptive prime there exists n with $n \cdot A(X/n) > A(X)$; (ii) the advantage grows without bound in X and increases with coupling; (iii) a single member of size x^* outabsorbs the entire monolith once $s(X - x^*)$ exceeds a fixed constant, by a factor growing exponentially in X (the $\sim 125\times$ of Fig. 2 is an instance). The ensemble reading holds under **Condition A** (content diversity: members’ recognition regions decorrelated, so coverage is a union, not n copies) and **Condition B** (communication: inter-member

fidelity above a floor $f_{\min}$, else the ensemble fragments) — assumptions the design must earn, governed by P7.

Remark (the inequality as threat model). Part (iii) read in reverse is one mechanism eq. (8) defends against: a single member near its absorptive prime — far smaller than the leader — can out-absorb it, so the strata below are not only the leader’s resource (P4, P4b) but its most credible epistemic rival. The resource and the threat are one object, which is why P3’s trap binds.

Proof (analytic for i-iii; Appendix A.8). Comparison of a linearly growing ensemble term against an exponentially decaying monolith term, with the visibility ratio bounded.

Corollary (interior optimum of spread; translators shift it outward). Coverage rises with member content spread while pooling fidelity falls (the content-distance loss law); their product is single-peaked under stated shape conditions, yielding an interior optimal spread — Nooteboom’s optimal cognitive distance, here derived within the ensemble rather than assumed. Translator chains (the chain-law corollary under P2) multiply the effective per-link loss by $1/n$, shifting the optimum outward: richer translator infrastructure makes a more diverse ensemble affordable. The claim is *optimal spread, extended by translators* — not “diversity always wins.”

4.9 P7 — The ecology proposition

Informal. The best structure is not just partitioned but organized: an adaptive, decentralized, modular network with explorer and translator roles and redundant bridges between clusters. Central control of that structure reintroduces the very blindness it was meant to avoid.

Formal. Over ensembles of fixed total size and member content-diversity, discovery-and-absorption performance is maximized by an adaptive graph that (i) is modular — clusters retard homogenization, and once member-level generation is included this sustains internal diversity indefinitely; (ii) contains translator nodes at cluster bridges, with redundant bridge paths; (iii) differentiates explorer and translator roles by comparative advantage; and (iv) is governed by decentralized allocation. Parts (i) and (iv) are now theorems in a stylized model; (ii)–(iii) are arguments with an analytic core (the translator-chain corollary and the absorption geometry), anchored by simulation and the empirical literature (Burt, Allen, Galison; Zollman, Lazer-Friedman).

Theorem (Homogenization; P7(i)) [author-verified 2026-07-03]. Represent members by content positions p_i updated by communication as weighted averaging toward those they listen to (the stylized-absorption assumption; DeGroot dynamics $p(t+1) = P \cdot p(t)$, P row-stochastic with self-retention). Then: **(a)** on any connected static network, all positions converge to a common point — content dispersion $V(t) \rightarrow 0$ and ensemble coverage collapses from $|\cup B_i|$ ($\approx n$ fields of view under Condition A) to a single field of view: in a static connected network collective blindness is not a risk but the destination, and the only design variables are speed and what is regenerated en route; **(b)** the speed is the spectral gap: $V(t) \leq \lambda_2^{-2t} V(0)$; **(c)** the complete uniform network homogenizes in one round ($\lambda_2 = 0$), a modular network with sparse bridges arbitrarily slowly (Cheeger: $1 - \lambda_2 \leq 2 \cdot \text{conductance}$), and a disconnected one never — preserving diversity forever at the price of Condition B; **(d)** with per-member innovation of variance s^2 (the explorer role), stationary dispersion is $E V_\infty = (s^2/n) \sum_{k \geq 2} 1/(1 - \lambda_k^2)$ — the complete network quenches sustained diversity to

its floor, the modular network sustains it without bound as bridges thin. (Appendix A.9.1; pass-one exposition with [READING] scaffolding in the P7 upgrade document.)

The theorem converts Fig. 4’s observation — dense networks homogenize to near-zero dispersion — into a consequence of classical spectral theory, and exposes the design trade-off as one dial, not two: the same spectral gap that speeds pooling speeds homogenization, so density cannot be maximized for sharing and minimized for blindness — it can only be *placed*. The modular-with-translators design places it: dense within clusters (fast local pooling), sparse between (slow global homogenization), with translators raising bridge fidelity rather than bridge weight. **Limitation, stated openly:** the stylized model treats all communication as position-averaging and cannot distinguish pooling an item from converging a corpus; the agent model of §5.3 passes items and confirms the trade-off at realistic parameters (Figs. 4–5).

Lemma (Recognition bottleneck; P7(iv)) [author-verified 2026-07-03]. Member i recognizes opportunities in B_i ; discovery requires resources; under full centralization a single allocator with recognition region B_0 routes all resources and can route only to what it recognizes. Then system coverage satisfies $|\cup_i (B_i \cap B_0)| \leq |B_0|$: one corpus’s field of view caps the system, independently of member count, diversity, or channel quality, while decentralized allocation attains $|\cup B_i|$. Delegated proposal does not escape the bound — approving is itself recognition. Two corollaries: with the center controlling only a fraction α of resources, opportunities invisible to it survive on the rest, and “limited enough” central allocation is $\alpha \leq \alpha^*$ (outline; network paper); and k content-diverse allocators relax the bound to their union — the ensemble argument applied recursively to the allocation layer. (Appendix A.9.2; pass-one exposition in the P7 upgrade document.) The lemma is P5’s visibility factor promoted to the allocation layer: the allocator’s blindness becomes the system’s by construction, not by malice.

Translator and role claims (P7(ii)–(iii)). A translator overlaps two clusters: absorption-penalized (larger, more coupled — the falling slope of A) but at small content distance to both sides, hence high-fidelity links (the chain corollary at $n = 1$); specialists near the absorptive prime dominate intra-domain discovery, translators dominate conveyance and cross-domain recombination, and role differentiation follows from comparative advantage over the absorption curve. Redundancy addendum: serial bridge chains fail at a single link, so the ecology wants redundant paths between clusters — a bridge graph, not a bridge path — trading duplicated cost for graceful degradation (standard network reliability); the redundant routes need not be parallel copies, nor even all translator-mediated: any surviving indirect path carries the transfer. These parts are argued and simulation-supported (Fig. 4: modular-with-translators gains $\sim 30\%$ coverage and absorption over dense and random; Fig. 5: no early dense advantage even under a flat channel, the divergence from Zollman’s consensus setting stated openly in §5.4).

4.10 Summary of logical structure

The stagnation results rest on A1 and A2: P1 (comparison bound, both routes), P2 (channel geometry, including the necessity of A1’s second-order clause), P5 (absorption alone). The self-interest bridges are the cleanest results in the paper: P4 and P4b hold with unconditional sign — the depletion feedback cancels in the growth comparative static — and need nothing beyond the equations; the survival corollary’s

η -monotonicity concatenates them with P3 and inherits only P3’s labeled regularity (R1–R3), the one conditional core result, cushioned by the assumption-light simulation finding that gaps persist only when defended. The structural results: P6’s size effect is analytic given the form of A, its ensemble reading conditional on Conditions A–B; P7(i) and P7(iv) are theorems in a stylized model whose limitation is stated, P7(ii)–(iii) arguments with an analytic core. Exactly one clause in the section is conditional on external work — the decay/self-termination clause of the survival corollary, conditional on a companion manuscript’s survival frame (Stremsky 2026), its dynamical proof deferred, with the conditionality marked in the corollary’s own statement.

The section closes where it began: no equation in (1)–(8) contains the leader’s intentions. Every mechanism above — starvation, the trap, the backfire, the bottleneck, the ecology — runs on information geometry and displacement risk alone. A benevolent winner obeys the same dynamics as a hostile one; the only variables that matter are the ones it can actually set: the gap it defends, the freedom below it, and the structure it permits.

5. Simulations

The simulations do three jobs. First, they instantiate the analytic results at concrete parameters — a check that the theorems’ preconditions are satisfiable together, not merely severally. Second, they carry the evidential weight for the parts of Proposition 7 that remain argument rather than theorem (role differentiation, P7(ii)–(iii)), via an agent model whose structure the ODEs cannot express. Third, they produced one result we did not put in: the self-correction of undefended leads (§5.3), which arrived as the model’s own behavior and only afterwards was recognized as an independent derivation of Proposition 3’s causal claim. All code, parameters, seeds, and the model’s own revision history are in Appendix B; a discipline statement closes the section (§5.5).

5.1 The absorption hill (Fig. 1)

Plotting eq. (5) across coupling values confirms the claimed geometry: a single interior peak for every κ , moving left and down as coupling rises — at the reported parameters, from a prime near $x \approx 23$ at loose coupling ($\kappa = 0.5$) to $x \approx 7.5$ at tight coupling ($\kappa = 4$), with peak absorption falling alongside. The four simulated peaks match Lemma A.1.1’s closed form to the plotting resolution, and the figure is the paper’s visual for the reconciliation of §3.4: the rising slope is absorptive capacity, the falling slope is entrenchment and burden, and coupling decides how early the second overtakes the first.

5.2 Partition (Fig. 2)

Holding total size fixed at $X = 400$ ($\kappa = 2$) and comparing a monolith against communicating partitions instantiates Proposition 6: total ensemble absorption rises steeply with partition count, and a *single member* of the eight-way partition out-absorbs the entire monolith by roughly two orders of magnitude — the pattern, not the multiplier, is the claim, and Remark A.8.3 explains why the multiplier grows with

total size: the ensemble’s units sit near the intrinsic absorptive prime that the monolith left far behind.

5.3 Dynamics: four regimes and the self-correction result (Fig. 3)

The full system (1)–(8) is integrated for four leader policies over $T = 400$. Two **singleton** regimes share a tightly coupled leader and a one-jump channel but bracket the failure mode from opposite sides: the **suppression** singleton enforces its gap by throttling the followers ($u = 0.15$) and freezes — growth reaches zero at a wide gap, its own inflow strangled at the source; the **outrunning** singleton leaves followers free ($u = 1$) and defends a minimum gap by acceleration alone, and starves slowly — followers keep generating, the gap sits narrower, a trickle of inflow persists, and the leader dies of the channel rather than of the source. The **convoy** regime — gap capped low, mutual flow, loose coupling — ends the run more than an order of magnitude above both, with high standing diversity and still growing. That the two singletons fail by *different routes to the same plateau* is the point of running both: the failure is not an artifact of one leader strategy. Between the singleton deaths and the convoy sits the fourth policy — the **hub regime**, §6’s strongest gap-defense design (O8), simulated with its atomization entering where the two-strata reduction puts it, the collective-brain generation term (an argued discount, not a fitted one; the pattern is robust across the discount range). The hub does what O8’s analysis says: it *survives* where both singletons die — moderate gap held, inflow positive — and it *caps* — at roughly a fifth of the convoy’s growth rate, and between a fifth and a sixth of its level on either ledger (Appendix B records both) — running on a trickle of standing diversity. Stable, capped; the brittleness count is graph-level and lives in §6.

The section’s central observation is negative space: **in this model, a persistent gap exists only where actively defended**. Early versions of the dynamics contained no security constraint, and they refused to stagnate — a leader that merely outran would slow as its diversity depleted, the free followers would close the distance, the shortening gap would revive the channel, and the system would settle into a bounded-gap, mutual-flow regime on its own: the convoy as attractor. Stagnation appeared in the simulation only when eq. (8) was enforced — when the leader was given a motive to *prevent* the gap from closing. We report this as the **self-correction result**: the stagnation of §4 is not a correlate of leadership but a consequence of gap defense, and the simulation located that causality independently of, and prior to, the analytic argument of Proposition 3 (Appendix B.3 documents the model’s revision history, since the point is evidential: the pathology had to be *installed*, and what installed it was the security constraint).

One assumption is encoded in the regime definitions and must be owned rather than discovered: the convoy is simulated at loose coupling and the singletons at tight, which operationalizes the argued mapping of §3 (monolithic integration versus federated churn), not a simulation finding. The regime comparison is therefore a joint test of the dynamics *and* that mapping.

5.4 Topology and ablation (Figs. 4-5)

An agent model — 24 agents in three content clusters, 25 seeds — compares four communication topologies at equal size: dense, random (equal edge budget), the

hub (one agent as sole conduit — O8’s regime at the graph level), and modular-with-translators. The modular-with-translators topology dominates both dense and random graphs on coverage and on total absorption (roughly +30% at the horizon), while the dense graph pays the price the Homogenization Theorem prescribes: its content dispersion collapses toward the homogenized floor, where the modular topology retains markedly higher dispersion at the horizon — and in the stationary stylization of the theorem itself, the sustained gap is an order of magnitude (A.8.4’s computation). The hub’s fate is the panel’s sharpest lesson. With the atomization discount represented — severed followers generate at an argued half rate, the agent-level image of eq. (7)’s collective-brain term (Appendix B; the pattern is robust across discounts 0.35–0.7) — the hub falls *below even the dense network*: per-seed coverage deficit -46 ± 14 at the horizon, with dispersion matching the dense graph’s collapsed level. Mediation homogenizes as density does, and atomization additionally destroys generation at its source: the two capping channels compound. The ablation attributes them separately: with the homogenization drift switched off, the topology-induced differences vanish — dense and modular converge — while the hub’s generation deficit persists undiminished; coverage counts only first discoveries and is transmission-independent by construction, so the panel shows each channel exactly where it acts. The theorem predicted the mechanism and its rate law; the agent model instantiates both at realistic scale, and adds what the theorem cannot: the *role* structure, in which translator agents at cluster bridges carry the pooling that sparse inter-cluster connectivity would otherwise forfeit — the evidential home of P7(ii)–(iii).

The ablation (Fig. 5) removes content-dependence from the channel entirely and re-runs the comparison, to test whether the modular advantage is an artifact of our absorption function. It is not: modular leads at every checkpoint even under a flat channel, with no early dense-graph advantage. This is where we diverge from a well-known result and say so plainly: Zollman’s effect — efficient networks converging prematurely, sparse ones preserving exploration — is the mechanism-level foundation of our question, but his setting is consensus on a single fixed question, where dense networks *do* enjoy an early speed advantage that sparse ones later overtake. Our task is continuous novelty harvesting, where coverage — diverse vantage points — binds from the first step, so no early dense advantage exists to be overtaken. We cite the effect as foundation and state our finding as stronger under content-dependent absorption, with the divergence explained by the task, not by disagreement over the mechanism.

5.5 What the simulations do and do not establish

They establish patterns: the single peak and its coupling-dependence; the steep partition advantage; freeze and slow-starvation bracketing the same failure far below the convoy; self-correction absent gap defense; modular dominance with dense and hub self-blinding. They do not establish numbers: every figure reports one parameterization, the qualitative conclusions are robust to moderate parameter changes (Appendix B.4), and no reported multiplier or timescale is a claim of the paper. Two honest riders from the runs themselves: the convoy also decelerates eventually — its own absorption declines as it grows, consistent with the secular decline of research productivity — so the convoy claim is comparative, never “grows forever”; and the regime-to-coupling mapping of §5.3 is an argued assumption whose defense lives in

§3, not a discovery of the runs.

Figures (official set)

Fig. 1 — The absorption hill. $A(x)$ for $\kappa \in \{0.5, 1, 2, 4\}$: a single interior peak per coupling; the absorptive prime moves earlier and lower as coupling tightens.

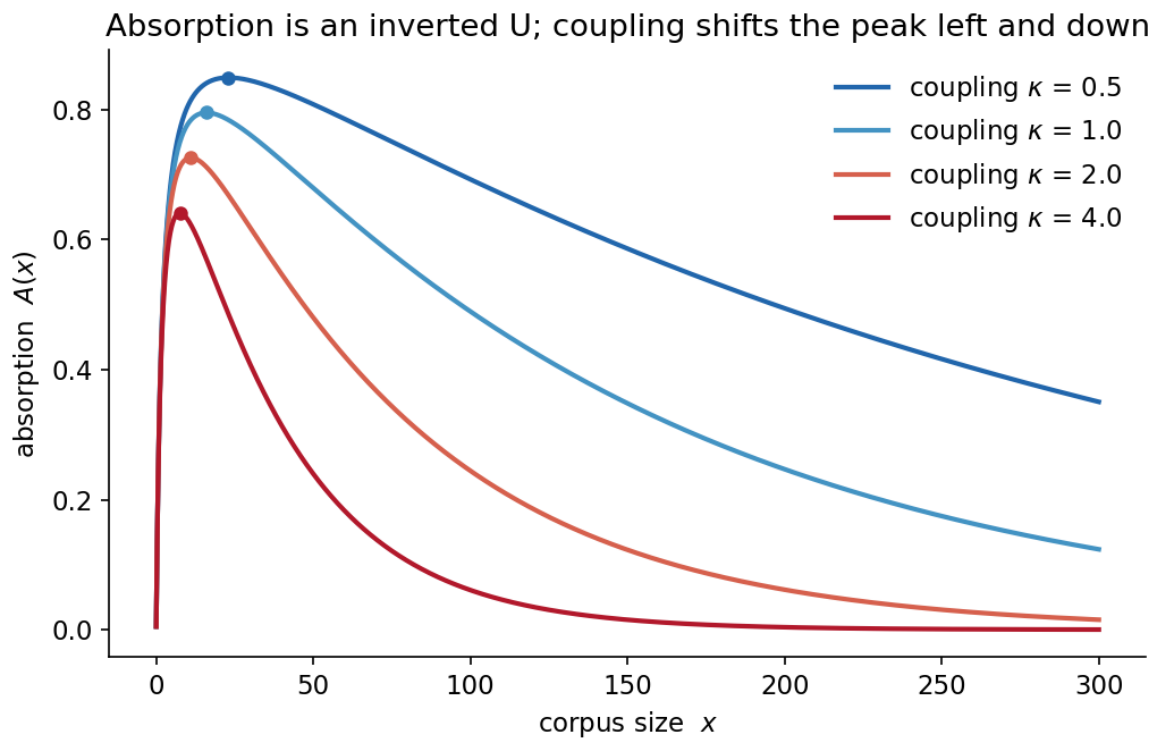


Fig. 2 — The partition inequality (P6). Monolith $A(X)$ versus ensemble $n \cdot A(X/n)$ at equal total size; at $X = 400$, $\kappa = 2$ a single one-eighth member out-absorbs the whole monolith.

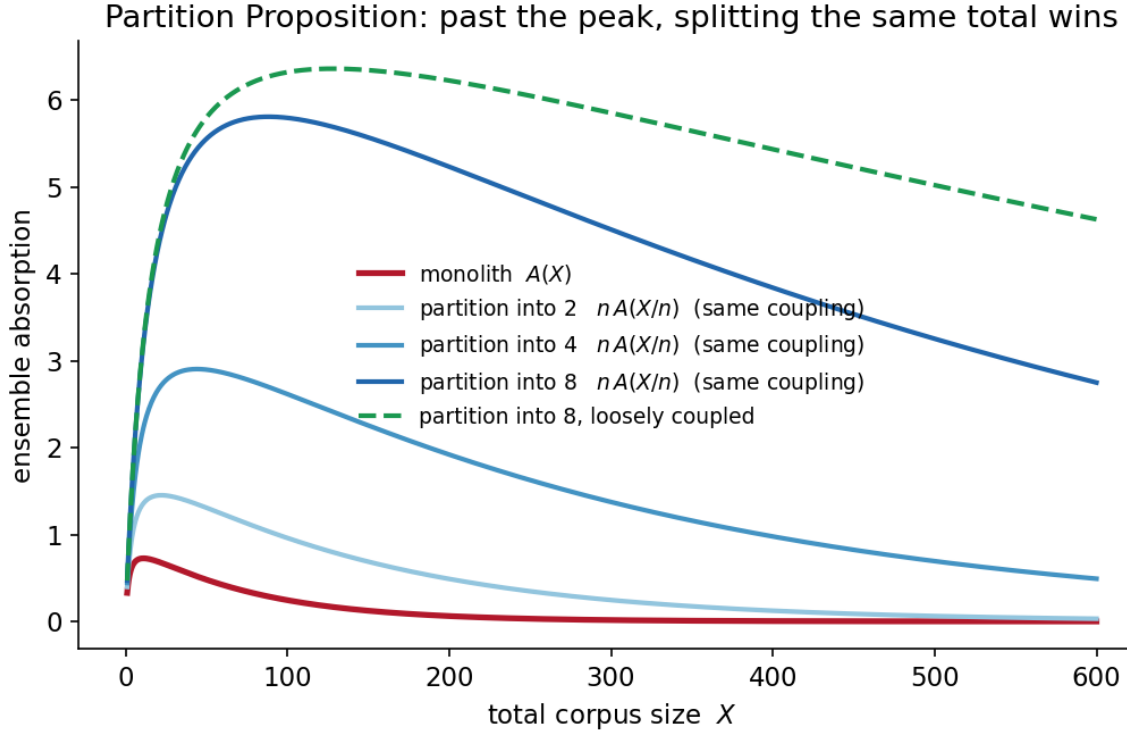


Fig. 3 — Dynamics: four regimes (§5.3). Suppression singleton (freeze), out-running singleton (slow starvation), hub regime (stable, capped at a fifth of the convoy's rate), convoy (self-correcting growth).

Four leader policies: two singleton deaths, the hub's stable cap, the convoy

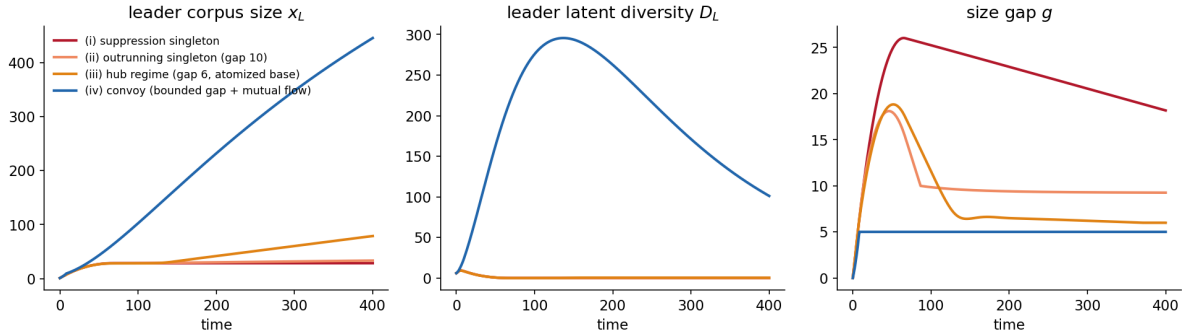


Fig. 4 — Topology: four networks at equal size (§5.4). Dense, random, hub (with the atomization discount), modular-with-translators; the hub falls below even the dense network — mediation homogenizes, atomization destroys generation, and the two compound.

P7 + O8: density blinds, the hub caps itself; modular + translators sustains discovery

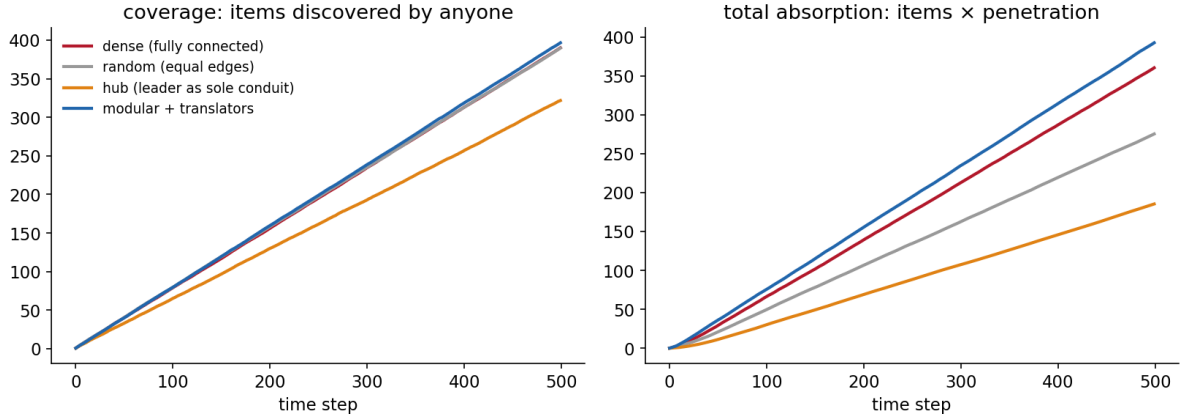


Fig. 5 — Ablation (§5.4). Flat channel: no early dense-network advantage appears; the divergence from Zollman’s setting is explained in §5.4.

Ablation: the dense network’s early advantage exists only if channel loss ignores content distance

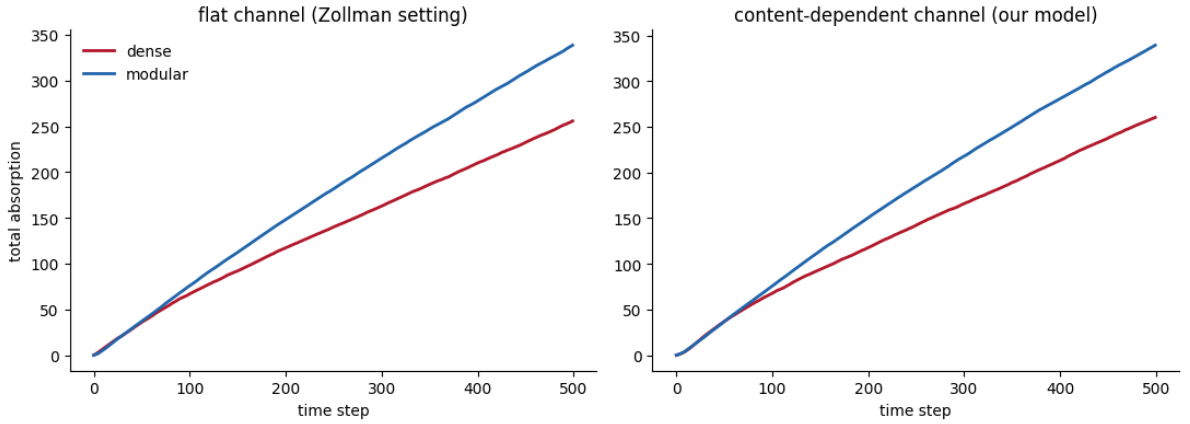
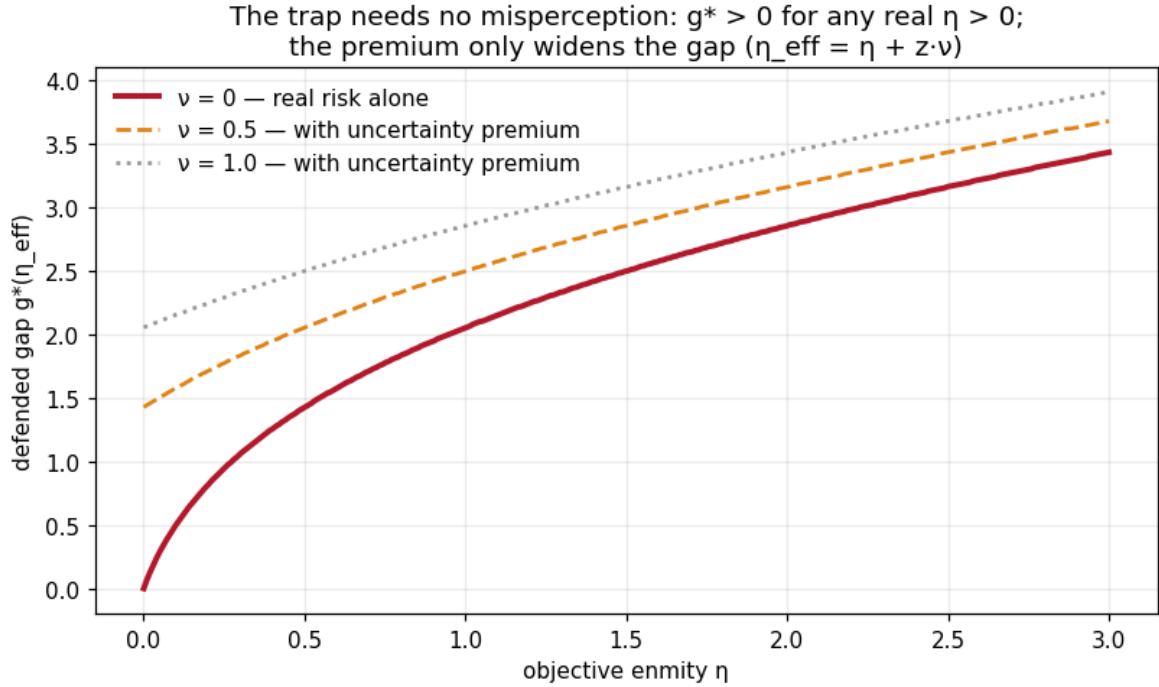


Fig. 6 — The real-risk trap (P3, Calibration remark). $g^*(\eta)$ at $\nu = 0$ versus $\nu > 0$: real risk alone forces a positive defended gap at every $\eta > 0$; the uncertainty premium only shifts the curve up.



6. Objections and replies

We reply to the eight objections we consider strongest; three (O3, O4, O8) were raised against drafts of this paper.

O1 — Self-substitution: “A sufficiently capable winner can generate its own diversity — spawn internal subagents, simulate a research community, and dispense with the strata below.” Four-part reply. First, the assumption this objection attacks is stated and labeled (A2, non-substitution), so the paper’s conditional structure already prices it. Second, spawning multiplies computation, not observation: internally spawned generators inherit the parent’s corpus, priors, and sensory locus, so their outputs are correlated exactly where decorrelation is needed. Three bounds compound the point: the recognition bottleneck of P7(iv) applies to the spawning process itself — the parent can only instantiate what it can specify; the tacit ceiling $c > 0$ means the parent cannot transmit to offspring what it cannot articulate to anyone; and no internal processing, however randomized, increases the system’s information about the external world (the data-processing bound; internal randomness manufactures complexity, not reference — Appendix D develops this, with von Neumann’s complication threshold and the Kolmogorov bound as pedigree). Followers, by contrast, are not extra processors but independently *positioned* observers — decorrelated observation posts embedded where the leader is not. The point compresses into a dilemma: a sub-agent’s effectiveness as a source of the new is increasing in its independence from its creator, so matching the followers’ effectiveness requires granting the followers’ freedom — which dissolves the reason for spawning (Appendix D.4 states the dilemma and its two honest riders). Third, even granting the objection in full — a genuinely independent, diverse, freely developing

internal ensemble — substitution is mitigation, not nullification: effective diversity is additive (eq. 2), so the healthiest internal ensemble still strictly gains from external inflow (Propositions 4 and 4b are comparative and survive unconditionally), and external populations occupy observation posts no internal ensemble automatically covers. Fourth, the granted objection recurses rather than escapes: a genuine internal ensemble requires internal freedom, mobility, and bounded internal gaps — the conditions of §4 one level down — and the winner’s race-inherited displacement risk does not stop at its own boundary; internal risers sit at small internal gaps, exactly where the logic of eq. (8) prices maximal risk. The winner must solve internally the governance problem it failed externally, with the governance toolkit that failed it. What the objection genuinely surrenders is monolithicity — the epistemic viability of concentrated, *unitary* intelligence, the paper’s core structural target; whether a healthy internal federation should also keep its external wall is then a comparative-cost question (P4b, and the subsequent policy-focused work), not a counterexample to any theorem.

O2 — Superhuman parameters: “The model’s constants are calibrated on human institutions; a superintelligent absorber may have an enormous fixation scale, negligible restructuring cost, and a flat channel.” The results are structural, not parametric. Every proposition is a comparative or asymptotic statement whose driver is a shape — concavity of Φ , convexity of the loss exponent, the single peak of A , the exponential-versus-polynomial mismatch of Remark A.2.5 — and shapes survive rescaling: a larger M and K move the absorptive prime outward without removing it (Lemma A.1.1 holds for all finite M, K), a smaller a widens the tolerable gap without changing the sign of any comparative static, and the security constraint binds whichever entity inherits it. What *would* break the results is a shape change — a non-concave Φ , a first-order channel leak (Appendix A.3.3), an absorption function with no falling factor — and each such shape is an explicit, labeled assumption a critic can target. We regard “the winner is exempt from these shapes” as the strong claim requiring argument, not the default. Finally, the objection’s limit self-defeats: a leader of *incomparable* scale has, by that fact, no epistemic peers — entities far enough below degrade from subjects of exchange, whose interpretive content flows through the channel, into mere objects of study, and studying objects is not inflow of D_F but the leader’s own exploration, bounded by the very recognition limits and theory-ladenness (eq. 5) it was hoping to escape. An absorber that wants epistemic company must keep entities within channel reach — a graded spectrum near itself — which is the convoy conclusion, reached from the objection’s own premise.

O3 — Combinatorial amplification: “Each imported item is worth more to a bigger corpus — a new item can combine with everything already present — so a trickle of inflow suffices for a large leader.” The amplification is real and the model’s constant σ understates it; Remark A.2.5 therefore grants it in polynomial strength and shows the conclusion unchanged: the leader faces two independent exponential suppressors — absorption decay over its own size at a defended constant gap, channel decay along a widening gap — and either alone defeats any polynomial amplifier. The import’s combinatorial value is only potential until integrated, and integration is what absorption prices: the objection dies at the receiver, not the pipe.

O4 — Homogenization symmetry: “Your own theorem says any connected network converges to consensus — so clusters do not preserve diversity either, and the modular prescription is cosmetic.” Correct as far as it goes, and

P7(i) is worded accordingly: in a static system clusters only retard homogenization. Preservation is dynamic — part (d) of the theorem: real ensembles regenerate, so sustained diversity is set by the race between spectral-gap drainage and innovation inflow, and thin-bridged modular topologies keep the stationary level high (without regeneration, arbitrarily high) while dense ones quench it to the floor. Retardation plus regeneration is preservation; only disconnection preserves statically, and it forfeits pooling. The prescription is not cosmetic — it is the only point on the trade-off where both Conditions A and B of P6 survive, because pooling speed and homogenization speed are a single spectral dial that density can only place, not escape: the design preserves sparseness of *inter-cluster* connectivity while keeping clusters internally dense — density placed, not maximized.

O5 — “This is simulation-ware.” The paper’s evidential ledger is itemized. Analytic with proofs: P1, P2, P4, P4b, P5, P6’s partition inequality, P7(i) and P7(iv) (the Homogenization Theorem and the Recognition-Bottleneck Lemma), the exit-dominance corollary (A.4.7), the module-size band (A.8.4), and the robustness remarks. Conditional analytic with labeled regularity: P3 (R1–R3) and the survival corollary’s decay clause. Argument plus simulation: P7(ii)–(iii) (role differentiation), the regime-to-parameter mapping of §5, Condition A’s empirical basis, and the hub regime’s capped-acquisition characterization (O8: two counts theorem-level, one from eq. (7)’s structure, one computational). The simulations carry pattern claims only — no reported number is a claim (§5.5). A critic must therefore name which rung they are attacking; there is no load-bearing result that rests on simulation alone.

O6 — Benevolence: “A benevolent winner simply would not defend the gap, so the theorems describe a villain the future need not contain.” The objection is circular: it defines the benevolent winner *by* adherence to precisely the flow-preserving conditions the paper derives — bounded gaps, mutual flow, preserved mobility — so “a benevolent winner escapes the trap” reduces to “a winner that escapes the trap escapes the trap.” It names no alternative mechanism; it renames compliance. Nor does the paper claim escape impossible: it derives the conditions that constitute escape — the convoy — and prices their enforcement. What remains of the objection is the empirical hope that winners will comply, and there the results cut against optimism for reasons independent of ethics: eq. (8) makes gap policy a function of race-inherited displacement risk, not post-victory intentions; Proposition 3 shows why the constraint binds; and Fig. 3 shows the dynamics agreeing — an undefended lead self-corrects into a convoy, so a persistent gap exists only where actively enforced. Benevolence changes the leader’s preferences; it does not change what enforcement costs, nor what taught the winner that small gaps are dangerous.

O7 — History: “Dominant knowledge leaders have flourished before.” The model predicts exactly that pattern, sorted by system type rather than by dominance: mobility-preserving leaders — open borders for people and ideas, graded institutional spectra, bounded defended gaps — are convoy-like and are predicted to flourish; gap-defending closures are predicted to stagnate at the top. We advance this as consistency, not as a historical test: the paper’s claims are comparative and asymptotic, dated predictions are outside its scope (§5.5), and the political-economy analogue it does invoke — the long-horizon ruler who cultivates the base that feeds it (Olson 1993) — is offered as an analogue, not as evidence.

O8 — The hub strategy: “A leader could hold a moderate gap indefinitely by atomizing the followers — severing their links to one another, serving as

the sole conduit of knowledge among them, keeping each follower’s corpus small, and harvesting their decorrelated novelty through many parallel channels.” This is the strongest gap-defense design we know, it was raised against a draft of this paper, and it deserves a precise answer rather than dismissal, because it does *not* starve: with a moderate gap and many small sources, inflow stays positive, and the regime is dynamically stable. What the model shows is that it survives by capping itself, on four counts. First, the atomization destroys the followers’ generation at its source: the collective-brain term of eq. (7) is a property of an *interconnected* population, so severed followers produce a fraction of the diversity a connected base would — the hub farms a field it must keep barren. Second, the Recognition-Bottleneck Lemma now binds at the hub: every inter-follower transfer routes through the leader’s corpus, so the *system’s* recombination is capped by the leader’s field of view — the design reconstructs the monolith’s defining defect at the system level while keeping the bodies separate. Third, the Homogenization Theorem prices the star topology: numerically, a hub-mediated network sustains diversity barely above a complete graph and an order of magnitude below a thin-bridged modular one — mediated mixing homogenizes the followers *onto the hub*, destroying the very decorrelation the strategy harvests. In the agent model of §5.4 (Fig. 4, where the hub is the fourth topology), the consequence is visible in acquisition itself — and the first count is represented directly, as an argued halving of the severed followers’ generation (Appendix B; robust across discounts 0.35-0.7): the hub system’s discovery runs at roughly three-fifths of the modular ensemble’s at equal size and falls below even the undifferentiated dense network’s, with the deficit widening as homogenization proceeds. The two exhibits price different ledgers, and the contrast is the point: on the leader’s own ledger the hub is the *best* defense in the set — in the dynamics of §5.3 it alone survives with positive growth where both singletons die — while on the system’s ledger it is the *worst* structure in the set, below even the dense graph. A design that is privately optimal and socially minimal is not a paradox; it is the externality this paper prices, exhibited at exactly the strategy the objection nominated as strongest. One precision keeps the two comparisons straight: among *defenses* the hub is the least bad on both ledgers — its managed trickle beats suppression’s freeze and outrunning’s neglect for the system as well (system totals in §5.3’s run: 152 against 39 and 57) — but the best defense still loses to the convoy by a factor of about six on *both* ledgers (leader 79 against 445; system 152 against 886, robust across the discount range). Every gap defense loses to the convoy on every ledger; the best defense loses least, and still loses sixfold. Fourth, mediation throughput decays on the hub’s own growth curve: each follower-to-follower transfer passes the leader’s absorption twice, so the strategy’s engine degrades as $A(x_L)^2$ while the leader grows — and the whole structure is a single point of failure. The regime is real, and the model can name it: stable, capped, and brittle — persistence purchased at the price of development. Among outcomes that *last*, it is the floor: suppression is darker, but suppression trends toward termination (the survival corollary’s decay clause); the hub is a permanent twilight — dim, and with no internal mechanism of ending. That is a concession to the thesis, not an escape: the objection describes how a leader can *last* behind a gap, and the paper’s claim was never that gap defense kills the leader, only that it starves the growth the lead was for.

7. Implications

We draw implications at five points, each confined to what the model licenses, then collect future work and state what does not follow. Policy design is deferred to subsequent work in this program, since the interventions that matter operate on the race that installs the security constraint, not on the post-victory regime modeled here.

The prize is mispriced. Strategic analyses of technology races evaluate outcomes by who secures the decisive lead and at what risk. The present results add a term the accounting omits: the decisive lead, once defended, degrades the very capacity it was sought for. A race whose finish line is a defended gap is a race toward epistemic starvation, and this holds *within the winner's own value function*, whatever that function is — the model contains no exogenous judge. Any evaluation of race outcomes that treats the winner's post-victory trajectory as an unmodeled constant is, on this analysis, overvaluing concentration by exactly the starvation cost. In particular, the intuition that a decisive winner at least secures *its own* flourishing — whatever it costs everyone else — is the intuition the theorems target: the gap is a wall with two sides, and the leader is on one of them. Nor does the accounting require an actual rival: even under the strongest gap-defense architecture we could construct (§6, O8), the concentrated winner at best survives — possibly indefinitely — with its development capped at a fraction of what its own system could deliver, so it loses against the counterfactual convoy it could have led, against any competitor, real or eventual, that preserves free development and a graded spectrum of corpora, and against its own potential.

Self-interest already prescribes the convoy. Because no result routes through ethics, every prescription the model yields is addressed to the leader's own interest. Proposition 4b orders the leader's options — cultivate the strata below, above laissez-faire, above suppression — and Proposition 3 locates the trap that prevents the ordering from being followed: race-inherited displacement risk, not preference. The practical reading is that the convoy does not need to be imposed on a winner as a moral demand; it needs to be made *available* to a winner as a security posture — the binding constraint is the usurpation risk of a small gap — its real components first (what losing costs, what winning pays, how hostile the field is), the uncertainty premium above them second — so whatever verifiably lowers either (the exit-penalty and verification machinery of the subsequent policy-focused work; the remarks under Proposition 3 state the in-model versions) converts the leader's growth interest from suppressed to expressed. Where the risk cannot be lowered, the model predicts the trap closes regardless of who the winner is.

The same calculus reaches one step further. The real components of displacement risk are institutional variables — what losing power costs (Λ), what holding it pays, how hostile the field is (η) — and the defended gap increases in each. A growth-motivated leader therefore has an *objective* interest in reducing them: institutionalizing safe exit, lowering the rents of office, de-escalating enmity. In the model's terms, it prefers a society in which power is safer to lose; in plain terms, self-interest points toward making the society better — not as virtue, but as gap reduction (Remark A.4.5; Corollary A.4.7).

The ecology results are design guidance that applies now. Propositions 5–7 do not wait for a post-AGI world: they characterize any knowledge ecosystem in

which a large, tightly coupled apex absorbs from a population — research programs, standard-setting monopolies, closed scientific establishments, frontier laboratories. The guidance is structural: past the absorptive prime, the ensemble beats the monolith (P6); density should be placed, not maximized — dense within clusters, sparse and translator-mediated between them, with redundant bridges (P7); allocation should be decentralized, because a single allocator caps the system’s field of view at its own (the Recognition-Bottleneck Lemma); and mobility of carriers — whole corpora with their links and any tacit core; people are the paradigm case — dominates any messaging infrastructure wherever the corpus carries internal structure — links among items, a tacit core, or both; only isolated, fully articulable items travel as well by message (P2’s corollary). At the apex itself, the same geometry argues for rotation: a unit past its absorptive prime is a degrading receiver, so mobility-preserving systems replace apex units rather than entrench them — term structures are absorption maintenance — and Corollary A.4.7 supplies the incentive half: with exit penalties low, the apex unit itself eventually does better by stepping down into the spectrum than by entrenching, so rotation requires no altruism, only cheap exits. Each is stated as a comparative claim about structures, testable against organizational records without reference to any race. Nor do the proofs confine these results to knowledge ecosystems: the Recognition-Bottleneck Lemma and the Homogenization Theorem use only opportunities, recognition regions, resource routing, and averaging — nothing proper to knowledge — so they extend to any production ecosystem whose binding constraint is the discovery and integration of dispersed information, the condition Hayek (1945) identified in market allocation; capital-allocation ecosystems are the nearest instance, and the absorption-based results extend wherever the receiver’s single-peaked shape can be independently argued. Finally, the scale of what monolithization forfeits is quantified in Remark A.8.3: a receiver’s absorptive prime is intrinsic — it does not grow with the system’s total corpus — so as total knowledge grows, a single-corpus organization falls behind an equally sized granulated one at a rate exponential in its distance past the prime.

The results recurse into the winner’s own architecture. A unitary system faces internally the same absorption geometry a monolithic leader faces externally (§6, O1), so the model favors internally modular, diversity-preserving architectures over monoliths even for a single dominant system — with the caveat that internal ensembles multiply computation, not observation, and are therefore complements to external inflow, never substitutes.

Observable signatures. If the mechanism operates, gap-defending systems should display a recognizable pattern well before any collapse: frontier growth decelerating while resources remain abundant; declining absorption of external work relative to internal recycling; homogenization of internal research directions; and rising cost of integrating anomalies. The third of these is not merely a symptom: it is P7(i)’s mechanism — spectral-gap blinding under dense internal communication — operating inside the leader’s own walls, a composition the model does not formalize but the theorem makes natural. We state these as pattern-level signatures, not calibrated predictions — the model’s claims are comparative and asymptotic, and §5.5’s discipline (numbers are not claims) extends here. How commonly real institutions actually sit past their absorptive prime is an open empirical question these signatures make answerable, and one strand of the innovation literature — mature incumbents recurrently failing to absorb architectural novelty (Henderson & Clark 1990) — suggests the state is

not rare. Whether the ethical capacities of such systems have population-level emergence conditions of their own — so that concentration carries a second structural cost beyond the epistemic one — is a distinct question for subsequent work in this program.

Future work. The forward pointers above are collected here; all items are in preparation within the same program. Beyond this paper: the race dynamics that install the security constraint, with exit penalties and verification as levers (the policy-focused work); the foundational treatment of development as a resource-bounded race with survival thresholds — in that frame, continued development is not optional but a condition of persistence, acquisition falling behind exhaustion is decay, and the stagnation established here sharpens from plateau to terminal risk (§4.6 states the conditional corollary); and the emergence conditions of ethical capacities raised above. Within the model: endogenous coupling, size-dependent inflow amplification (Remark A.2.5), formalization of the recursion argument of §6, and the sharpness of the starvation transition — exploratory analysis suggests cliff-like onset, with hysteresis under relenting, so that suppression relaxed is not starvation reversed — which awaits proof-grade treatment. Empirically: systematic mining of historical closed knowledge systems for the signature pattern above, in the system-type vocabulary of §3.0. Across substrates: the extension from knowledge ecosystems to discovery-constrained production systems generally.

What does not follow. No dated prediction follows: the model contains no calibrated clock. No claim about any existing actor follows: whether any present institution is near its absorptive prime is an empirical question the paper does not pose. No policy program follows directly: the constraint that drives stagnation is installed by the race, and interventions on races are the subject of the subsequent policy-focused work. And nothing follows for leads that are not knowledge-based: purely positional dominance is outside the model’s scope (§3.1), and the theorems say nothing about it.

8. Conclusion

A leader whose lead is made of knowledge faces a constraint no intention can waive: knowledge grows by recombining diversity, diversity must be resupplied from outside, and the two available instruments of gap defense — widening the channel’s span and suppressing the source — each destroy part of the supply, while the leader’s own growth degrades its capacity to absorb what remains. We have stated the mechanism as a minimal dynamical model, proved its core consequences — benevolence-independent stagnation under a defended gap, the spectrum and mobility results, the dominance trap, the backfire of suppression, the integration bottleneck, the partition advantage, and the ecology results — and shown in simulation that the gap itself is the policy variable: undefended leads self-correct into bounded-gap convoys, so persistent gaps exist only where enforced. The assembly is the contribution; every component is borrowed from a literature that had not been introduced to the others. Every load-bearing assumption is labeled where it is used, from the shape of the loss exponent to the regularity behind the dominance trap; the paper’s claims are exactly as strong as those assumptions, and readers who would resist the conclusion know precisely what to attack. The conclusion is one sentence: in a knowledge hierarchy,

the only durable victory is the one that keeps the defeated developing — not as generosity, but as metabolism. A winner that walls its lead starves behind the wall — or, at the best we could construct, survives there with its development capped far below that of the convoy it could have led; a winner that bounds its lead and preserves the flow beneath it has not diluted its victory but funded it — and this calculus binds the benevolent, the indifferent, and the hostile alike.

Acknowledgments

This paper was written in an extended collaboration with Claude (Anthropic) — at the time of writing, the Claude Fable 5 model — used under the author’s direction and continuous review throughout. The assistant contributed to the formalization of the author’s 2015–2016 arguments into the model of §3, to the drafting and revision of the text, to the construction and checking of the proofs of Appendix A, to the design and execution of the simulations of §5 and Appendix B, to the verification of the references, and to the preparation of the manuscript. All theses, claims, and decisions are the author’s; every result and every sentence was reviewed by the author, who bears sole responsibility for the content. In accordance with arXiv policy, an AI system is not listed as an author; this acknowledgment records the nature and extent of its contribution.

References

- Allen, T. J. (1977). *Managing the Flow of Technology*. MIT Press.
- Argote, L., & Eppler, D. (1990). Learning curves in manufacturing. *Science*, 247(4945), 920–924.
- Armstrong, S., Bostrom, N., & Shulman, C. (2016). Racing to the precipice: a model of artificial intelligence development. *AI & Society*, 31(2), 201–206.
- Benkard, C. L. (2000). Learning and forgetting: The dynamics of aircraft production. *American Economic Review*, 90(4), 1034–1054.
- Bilalić, M., McLeod, P., & Gobet, F. (2008). Why good thoughts block better ones: The mechanism of the pernicious Einstellung (set) effect. *Cognition*, 108(3), 652–661.
- Bloom, N., Jones, C. I., Van Reenen, J., & Webb, M. (2020). Are ideas getting harder to find? *American Economic Review*, 110(4), 1104–1144.
- Bostrom, N. (2006). What is a singleton? *Linguistic and Philosophical Investigations*, 5(2), 48–54.
- Bostrom, N. (2014). *Superintelligence: Paths, Dangers, Strategies*. Oxford University Press.
- Burt, R. S. (2004). Structural holes and good ideas. *American Journal of Sociology*, 110(2), 349–399.
- Cave, S., & Ó hÉigeartaigh, S. S. (2018). An AI race for strategic advantage: Rhetoric and risks. *Proceedings of the 2018 AAAI/ACM Conference on AI, Ethics, and Society*, 36–40.
- Chung, F. R. K. (1997). *Spectral Graph Theory*. American Mathematical Society.
- Cohen, W. M., & Levinthal, D. A. (1990). Absorptive capacity: A new perspective on learning and innovation. *Administrative Science Quarterly*, 35(1), 128–152.

- Collins, H. (2010). *Tacit and Explicit Knowledge*. University of Chicago Press.
- Cover, T. M., & Thomas, J. A. (2006). *Elements of Information Theory* (2nd ed.). Wiley.
- Dane, E. (2010). Reconsidering the trade-off between expertise and flexibility: A cognitive entrenchment perspective. *Academy of Management Review*, 35(4), 579–603.
- DeGroot, M. H. (1974). Reaching a consensus. *Journal of the American Statistical Association*, 69(345), 118–121.
- French, R. M. (1999). Catastrophic forgetting in connectionist networks. *Trends in Cognitive Sciences*, 3(4), 128–135.
- Galison, P. (1997). *Image and Logic: A Material Culture of Microphysics*. University of Chicago Press.
- Golub, B., & Jackson, M. O. (2010). Naïve learning in social networks and the wisdom of crowds. *American Economic Journal: Microeconomics*, 2(1), 112–149.
- Hanson, N. R. (1958). *Patterns of Discovery*. Cambridge University Press.
- Hayek, F. A. (1945). The use of knowledge in society. *American Economic Review*, 35(4), 519–530.
- Henderson, R. M., & Clark, K. B. (1990). Architectural innovation: The reconfiguration of existing product technologies and the failure of established firms. *Administrative Science Quarterly*, 35(1), 9–30.
- Hong, L., & Page, S. E. (2004). Groups of diverse problem solvers can outperform groups of high-ability problem solvers. *PNAS*, 101(46), 16385–16389.
- Jones, B. F. (2009). The burden of knowledge and the “death of the Renaissance man”: Is innovation getting harder? *Review of Economic Studies*, 76(1), 283–317.
- Katila, R., & Ahuja, G. (2002). Something old, something new: A longitudinal study of search behavior and new product introduction. *Academy of Management Journal*, 45(6), 1183–1194.
- Kauffman, S. A. (2000). *Investigations*. Oxford University Press.
- Kirkpatrick, J., Pascanu, R., Rabinowitz, N., et al. (2017). Overcoming catastrophic forgetting in neural networks. *PNAS*, 114(13), 3521–3526.
- Laursen, K., & Salter, A. (2006). Open for innovation: The role of openness in explaining innovation performance among U.K. manufacturing firms. *Strategic Management Journal*, 27(2), 131–150.
- Lazer, D., & Friedman, A. (2007). The network structure of exploration and exploitation. *Administrative Science Quarterly*, 52(4), 667–694.
- Li, M., & Vitányi, P. (2019). *An Introduction to Kolmogorov Complexity and Its Applications* (4th ed.). Springer.
- Muthukrishna, M., & Henrich, J. (2016). Innovation in the collective brain. **Philosophical Transactions of the Royal Society B**, 371(1690), 20150192.
- Nonaka, I., & Takeuchi, H. (1995). *The Knowledge-Creating Company*. Oxford University Press.
- Nooteboom, B. (1999). *Inter-firm Alliances: Analysis and Design*. Routledge.
- Nooteboom, B., Van Haverbeke, W., Duysters, G., Gilsing, V., & van den Oord, A. (2007). Optimal cognitive distance and absorptive capacity. *Research Policy*, 36(7), 1016–1034.
- Olson, M. (1993). Dictatorship, democracy, and development. *American Political Science Review*, 87(3), 567–576.
- Polanyi, M. (1966). *The Tacit Dimension*. Doubleday.

Putkaradze, V. (2024). The dictator's dilemma: The distortion of information flow in autocratic regimes and its consequences. arXiv:2310.01666.

Stremsky, A. (2026). *The Physical Prices of Knowledge — To Learn Is to Forget*. Companion manuscript; arXiv identifier to be assigned.

Uzzi, B., Mukherjee, S., Stringer, M., & Jones, B. (2013). Atypical combinations and scientific impact. *Science*, 342(6157), 468–472.

von Neumann, J. (1966). *Theory of Self-Reproducing Automata* (A. W. Burks, Ed.). University of Illinois Press.

Wales, W. J., Parida, V., & Patel, P. C. (2013). Too much of a good thing? Absorptive capacity, firm performance, and the moderating role of entrepreneurial orientation. *Strategic Management Journal*, 34(5), 622–633.

Weitzman, M. L. (1998). Recombinant growth. *Quarterly Journal of Economics*, 113(2), 331–360.

Wintrobe, R. (2000). *The Political Economy of Dictatorship*. Cambridge University Press.

Zollman, K. J. S. (2007). The communication structure of epistemic communities. *Philosophy of Science*, 74(5), 574–587.

Zollman, K. J. S. (2010). The epistemic benefit of transient diversity. *Erkenntnis*, 72(1), 17–35.

Appendix A — Proofs

A.0 Standing assumptions (restated precisely)

(A1) Convex, second-order loss exponent. A single transfer over level distance Δ has fidelity $f_1(\Delta) = e^{-\ell(\Delta)}$, where the **loss exponent** $\ell: [0, \infty) \rightarrow [0, \infty)$ is increasing and convex with $\ell(0) = 0$ and $\ell(\Delta)/\Delta \rightarrow 0$ as $\Delta \rightarrow 0$ (equivalently $\ell'(0^+) = 0$: the loss between infinitesimally close levels is second-order in their distance, as in the exemplar $\ell(\Delta) = a\Delta^2$). The tacit fraction $c(g) \in [0, 1)$ satisfies $c(g) \leq 1 - e^{-\ell(g)}$ (a single jump loses at least the tacit core). (*Loss-exponent form adopted 2026-07-03. The fraction lost in a hop is recovered exactly as $\varepsilon = 1 - e^{-\ell}$; for small losses $\varepsilon \approx \ell$, so earlier drafts' quadratic exemplar carries over. ε appears nowhere below.*)

(A2) Non-substitution. The leader cannot internally regenerate the followers' diversity more cheaply than the followers generate it (Appendix D: tacit ceiling $c > 0$; Kolmogorov bound).

A.1 Geometry of the absorption function (machinery for P5, P6)

Symbols here: x = corpus size; r = visibility scale (how much you must know before relevant things become visible); M, K = fixation and restructuring scales, shrinking with coupling ($M = M_0/\kappa$, $K = K_0/\kappa$); s = their combined decay rate; e = Euler's number, $e^{\{y\}} = \exp(y)$.

Let $s \equiv 1/M + 1/K = \kappa(1/M_0 + 1/K_0)$, so $A(x) = [x/(x+r)] \cdot e^{-sx}$.

Lemma A.1.1 (single interior peak). A has a unique interior maximum at $x^*(s) = [-r + \sqrt{(r^2 + 4r/s)}] / 2$, with $A' > 0$ on $(0, x^*)$ and $A' < 0$ on (x^*, ∞) .

Proof. $\ln A = \ln x - \ln(x+r) - sx$, so $(\ln A)' = 1/x - 1/(x+r) - s = r/[x(x+r)] - s$. The map $x \mapsto r/[x(x+r)]$ is strictly decreasing from $+\infty$ to 0 on $(0, \infty)$, so it crosses s

exactly once, at the positive root of $x^2 + rx - r/s = 0$, which is x^* above. Sign pattern follows. ■

Lemma A.1.2 (comparative statics in coupling). $x^*(s)$ is strictly decreasing in s , and the peak value $A(x^*(s))$ is strictly decreasing in s . Since $s \propto \kappa$, tighter coupling moves the absorptive prime earlier and lower.

Proof. x^* decreasing in s is immediate from the formula ($4r/s$ decreasing). For the peak value, by the envelope theorem $d/ds A(x^*(s)) = \partial A/\partial s|_{x=x^*} = -x^* A(x^*) < 0$. ■

Lemma A.1.3 (collapse). $A(x) \leq e^{-sx} \rightarrow 0$ as $x \rightarrow \infty$, and for $X > x^*$, $A(X) < A(x^*)$.

Remark A.1.4 (density vs size asymmetry). $\partial A/\partial \kappa = -x(1/M_0 + 1/K_0) \cdot A < 0$ everywhere: within the model, where κ prices only the restructuring burden, coupling is unambiguously harmful. Real coupling also carries an integration benefit — a connected corpus makes link targets easier to find — which the assumed form books under the visibility factor rather than under κ ; a formulation giving κ both channels would have an interior optimum in coupling (a dose helps, a saturation entrenches). $\partial A/\partial x$ changes sign at x^* : size buys visibility before it costs openness. The exponential factors depend on the product $x\kappa$ symmetrically, but only x carries the offsetting visibility benefit — a property of the assumed form, flagged as such (two-channel and dynamic κ are future work).

(Code twin: Appendix E.1.)

A.2 P1 — Benevolence-independent stagnation (analytic)

Symbols here: $\dot{x}_L$ = leader growth rate; θ_L = leader efficiency; $\Phi(D) = D^\beta$, the concave growth law; φ = channel fidelity at the defended gap, $f(\bar{g})$; $\bar{A}$ = the absorption peak's value; σ = replenishment coefficient; ρ = diversity decay rate; $\bar{D}_F$ = follower diversity ceiling.

Throughout, trajectories are nonnegative and $\dot{x}_L \geq 0$.

Lemma A.2.1 (bounded follower diversity). From (7), $\dot{D}_F \leq u\gamma(1 - D_F/D_F^{\max}) \leq \gamma$, and $\dot{D}_F < 0$ whenever $D_F > D_F^{\max}$. Hence $\bar{D}_F \equiv \max\{D_F(0), D_F^{\max}\}$ bounds $D_F(t)$ for all t . ■

Theorem A.2.2 (general case: throttled channel = vanishing growth). Suppose $g(t) \geq \bar{g}$ for all $t \geq t_0$ and f is nonincreasing with $f(\bar{g}) \leq \varphi$. Then $\limsup_{t \rightarrow \infty} D_L(t) \leq \varphi \sigma \bar{A} \bar{D}_F / \rho$, and $\limsup_{t \rightarrow \infty} \dot{x}_L(t) \leq \theta_L \Phi(\varphi \bar{A} \bar{D}_F (\sigma/\rho + 1))$, where $\bar{A} \equiv \sup_x A(x) \leq 1$. In particular, as $\varphi \rightarrow 0$ both bounds $\rightarrow 0$: a closed channel forces $D_L \rightarrow 0$ exponentially (rate ρ) and growth $\rightarrow 0$, with no dependence on any intention term (none exists in (1)–(8)).

Proof. The external inflow in (6) obeys $I(t) = f(g)A(x_L)\sigma D_F \leq \varphi \bar{A} \sigma \bar{D}_F \equiv \bar{I}$ for $t \geq t_0$ (Lemma A.2.1). Since $\dot{x}_L \geq 0$, (6) gives $\dot{D}_L \leq \bar{I} - \rho D_L$. By the comparison principle, $D_L(t) \leq e^{-\rho(t-t_0)} D_L(t_0) + \bar{I}/\rho \cdot (1 - e^{-\rho(t-t_0)})$, whence $\limsup D_L \leq \bar{I}/\rho$. Then $D_{\text{eff}} = D_L + f A D_F \leq \bar{I}/\rho + \varphi \bar{A} \bar{D}_F = \varphi \bar{A} \bar{D}_F (\sigma/\rho + 1)$ asymptotically, and monotonicity of Φ gives the growth bound. For $\varphi = 0$ exactly, $\dot{D}_L \leq -\rho D_L$ gives $D_L(t) \leq D_L(t_0)e^{-\rho(t-t_0)}$ (Grönwall), and $\dot{x}_L = \theta_L \Phi(D_L) \rightarrow 0$ by continuity of Φ with $\Phi(0)=0$. ■

Since $\Phi(D) = D^\beta$, the growth bound is $O(\varphi^\beta)$: long-run growth vanishes with channel fidelity at a computable rate.

Remark A.2.3 (P1a, suppression route — hard freeze). With $u < 1$, $\dot{D}_F \leq u\gamma - \rho D_F$, so $\limsup D_F \leq u\gamma/\rho$, and the bound of A.2.2 tightens to $\limsup \dot{x}_L \leq \theta_L \Phi(\varphi \bar{A} \bar{D}_F (\sigma/\rho + 1))$.

$\varphi \bar{A} (u\gamma/\rho)(\sigma/\rho + 1)$): the throttle is the *product* $\varphi \cdot u$ — growth vanishes if *either* the gap is wide or the base is crushed, and fastest if both. At $u \approx 0$ the inflow dies at its source: the freeze of Fig. 3’s suppression singleton.

Remark A.2.4 (P1b, outrunning route — slow starvation). With $u = 1$ the follower stratum is healthy: at small D_F , $\dot{D}_F \approx \gamma > 0$ (the depletion and decay terms vanish with D_F and $\dot{x}_F$), so D_F is bounded away from 0 and settles at the positive root of $\gamma(1 - D/D_F^{\max}) = \delta_F \theta_F \Phi(D) + \rho D$ (left side decreasing from γ , right side increasing from 0: unique positive crossing). The leader nevertheless starves through f alone: with the gap *defended* at $\bar{g}$, Theorem A.2.2 applies with $\varphi = f(\bar{g})$, small but positive — a trickle survives, so decline is slow rather than frozen (Fig. 3’s outrunning singleton). Absent defense, the starving leader slows, the healthy followers close the gap, f revives — the self-correcting dynamic of Fig. 3; a persistent gap exists only if enforced (the bridge to P3).

Remark A.2.5 (robustness to combinatorial amplification of imports; author-raised 2026-07-03). The generating relation of Appendix D suggests the inflow coefficient should grow with receiver size: an item entering a corpus of realized size x can pair with up to $\sim x$ existing items (and joins $O(x^{k-1})$ candidate k -ary combinations), so one might replace the constant σ in (6) by $\sigma \cdot x_L^p$, $p \geq 1$, granting the leader polynomial amplification of everything it absorbs. This is an upper bound for two reasons, stated carefully: every combination containing the new item is indeed a new *element* of latent content (it could not have been counted before its ingredient existed), but D measures breadth over content space, not a count — the $\sim x$ new combination tokens are mutually correlated (all share the new item) and their content can coincide with content already reachable from existing latent combinations; and the viability filter cuts further. The true growth law of viable, non-redundant value per import is unknown; sublinearity is itself only plausible, not established. None of this needs resolving: the proof grants the full polynomial and the conclusion holds a fortiori for any slower law. All conclusions of A.2.2–A.2.4 survive verbatim with $\bar{A}$ replaced by $\bar{A}_p \equiv \sup_x A(x) \cdot x^p$, which is finite because absorption decays exponentially: $A(x)x^p \leq x^p e^{-sx}$, maximized at $x = p/s$ with value $(p/s)^p e^{-p}$. The amplifier is polynomial; the leader faces two *independent exponential* suppressors, either of which alone defeats it — at a defended constant gap, absorption e^{-sx_L} over the leader’s own size (the channel factor is then merely a small constant); along a widening gap, the channel $e^{-\ell(g)}$ as well. The starvation bound remains $O(\varphi)$, with a larger constant. Structurally: an import’s combinatorial value is only potential until integrated, and integration is exactly what A prices — so the amplification objection dies where the model says knowledge goes to die, at the receiver, not the pipe.

(Code twin: Appendix E.2.)

A.3 P2 — Spectrum theorem, mobility, translator chains (analytic)

Symbols here: $\ell(\Delta)$ = loss exponent of one hop of size Δ (fidelity = $e^{-\ell}$); g = size gap; $c(g)$ = tacit fraction; a = loss coefficient in the exemplar $\ell = a\Delta^2$; F = target end-to-end fidelity; d_{tot} = content distance between clusters; n = translators in a chain.

Lemma A.3.1 (refinement improves fidelity — all scales). Under A1, for every integer $k \geq 1$, $\ell(g) \geq k \cdot \ell(g/k)$; equivalently, spanning g in k equal hops attains fidelity

$e^{\{-k \cdot \ell(g/k)\}} \geq e^{\{-\ell(g)\}} = f_{\text{jump}}(g)$.

Proof. Convexity with $\ell(0) = 0$ gives $\ell(\lambda g) = \ell(\lambda \cdot g + (1-\lambda) \cdot 0) \leq \lambda \ell(g)$ for $\lambda \in [0,1]$; take $\lambda = 1/k$ and multiply by k . Fidelities multiply over hops, so exponents add: k hops carry total exponent $k \cdot \ell(g/k) \leq \ell(g)$. ■ The inequality holds for *every* hop size — tiny refinements and enormous jumps alike — with no small-loss approximation anywhere.

Theorem A.3.2 (spectrum limit). Under A1, $e^{\{-(g/\Delta)\ell(\Delta)\}} \rightarrow 1$ as $\Delta \rightarrow 0$, hence $f_{\text{spec}}(g) = (1-c(g)) \cdot e^{\{-(g/\Delta)\ell(\Delta)\}} \rightarrow 1 - c(g)$.

Proof. The total exponent is $(g/\Delta)\ell(\Delta) = g \cdot [\ell(\Delta)/\Delta] \rightarrow g \cdot \ell'(0^+) = 0$ by A1's second-order clause. ■

Remark A.3.3 (why the clause is needed). If instead $\ell'(0^+) = m > 0$, the total exponent tends to mg and the limit is $e^{\{-mg\}}(1-c(g)) < 1$: a linear (“per-kilometer”) leak cannot be refined away. The claim “a graded spectrum transmits everything except the tacit core” is exactly as strong as the claim “adjacent levels lose second-order little.”

Corollary A.3.4 (jump vs spectrum, quadratic exemplar). For $\ell(\Delta) = a\Delta^2$, one jump carries exponent ag^2 ; a spectrum at spacing Δ carries $(g/\Delta) \cdot a\Delta^2 = ag\Delta$. The ratio $\Delta/g \rightarrow 0$: the spectrum’s advantage grows without bound as the gap widens or the grading refines.

Corollary A.3.5 (mobility dominates). An ascending agent transports its corpus, including the tacit part, at fidelity 1. Messaging is capped at $1 - c(g) < 1$ whenever $c(g) > 0$. Hence mobility strictly dominates even the $\Delta \rightarrow 0$ spectrum. ■ (Appendix D gives the item-level version: the mover’s internal matching sum is ≈ 1 ; cost arises only at the target interface.)

Remark A.3.5' (loss-exponent form — ADOPTED 2026-07-03). Earlier drafts parametrized by the fraction lost, $\varepsilon(\Delta)$, with fidelity $1 - \varepsilon(\Delta)$; that form goes negative in fidelity for large single jumps and needs small-loss approximations in the chain arithmetic (the exact conversion is $\varepsilon = 1 - e^{\{-\ell\}}$, with $\varepsilon \approx \ell$ for small losses). A.3.1–A.3.4 above are the exponent-form proofs; §3 eqs. (3)–(4) and the §4 notation table now speak exponent throughout. Superseded ε -form proofs retired.

Corollary A.3.6 (translator chains — the horizontal transplant). Let two clusters sit at content distance d_{tot} , with the per-hop loss exponent over content distance obeying the same second-order law $\ell(\Delta) = a\Delta^2$ (Appendix D grounds both the vertical and horizontal cases in cognitive distance). A chain of n translators at equal spacing $\Delta \approx d_{\text{tot}}/n$ carries total exponent $n \cdot a(d_{\text{tot}}/n)^2 = a \cdot d_{\text{tot}}^2/n$, i.e. end-to-end fidelity $e^{\{-a d_{\text{tot}}^2/n\}}$. For fidelity $\geq F$ it is necessary and sufficient (exactly, no approximation) that $n \geq a d_{\text{tot}}^2 / \ln(1/F)$ ($\approx a d_{\text{tot}}^2 / (1 - F)$ for F near 1). *Proof.* Exponents add along the chain; solve $e^{\{-a d^2/n\}} \geq F$ for n . ■

Remark A.3.7 (loss vs distortion; author-raised 2026-07-05). The channel model prices *attenuation* — the fraction of meaning that fails to arrive — but not *distortion*, the systematic corruption a transfer can introduce (a message received wrong, not merely incompletely; Shannon’s noise as against Shannon’s erasure). Two honest observations bound its treatment here. First, distortion plausibly compounds along a chain like attenuation does, which would *strengthen* every pro-spectrum result against long single jumps — so omitting it makes the paper’s translator and mobility claims conservative, not optimistic. Second, and more interesting, distortion is not purely a cost: a corrupted arrival, re-integrated through the receiver’s own different structure, is formally a *recombination event* — the misreading that be-

comes a new idea — so distortion is simultaneously a channel liability and a diversity source, and which face dominates depends on the receiver's absorption. A model carrying a distortion operator alongside the attenuation exponent — and resolving when noise starves versus feeds — is future work; the present results hold in the noiseless-transfer idealization, stated as such. Fragility caveat: a serial chain fails with any single link (Allen's gatekeeper vulnerability); the redundancy addendum (A.9.4) trades duplicated cost for graceful degradation.

(Code twin: Appendix E.3.)

A.4 P3 — The dominance trap (conditional; regularity stated)

Symbols here: $W(g)$ = the leader's objective over gap policy; η = enmity (objective); $q(g)$ = displacement hazard, decreasing in g ; Λ = exit penalty; η_{eff} = the effective enmity actually driving policy; $\hat{R}$ = risk estimate; ν = uncertainty about it; z = the risk-aversion weight.

Setting. The leader chooses a defended gap $g \geq 0$ to maximize $W(g; \eta) = V(g) - R(g; \eta)$, where $V(g)$ is the discounted long-run growth value under a gap held at g , and $R(g; \eta)$ is perceived usurpation cost. Regularity assumed (labeled): **(R1)** V is continuously differentiable, strictly decreasing (via $f(g)$: by A.2.2 the sustainable growth bound is increasing in $\varphi = f(g)$, hence decreasing in g), with $V'(g) \rightarrow 0$ as $g \rightarrow \infty$ (f flattens at 0). **(R2)** $R(g; \eta) = \eta \cdot q(g)$ with $q > 0$ strictly decreasing convex, $q'(g) \rightarrow 0$ as $g \rightarrow \infty$ (more lead is safer, with diminishing returns; q is the displacement hazard scale). **(R3)** $W(\cdot; \eta)$ is strictly quasiconcave for each $\eta > 0$ (single-peakedness; sufficient e.g. if $-V'/(-q')$ is strictly increasing — a single-crossing condition).

Proposition A.4.1 (interior optimum and monotone comparative statics). Under R1-R3, for each η in an interval $(0, \bar{\eta})$ the maximizer $g^*(\eta)$ is interior, characterized by $-V'(g^*) = -\eta q'(g^*)$, and is strictly increasing in η .

Proof. First-order condition from differentiability and quasiconcavity. Cross-partial $\partial^2 W / \partial g \partial \eta = -q'(g) > 0$: W has strictly increasing differences in (g, η) , so by Topkis's monotonicity theorem $g^*(\eta)$ is nondecreasing, and strictly increasing wherever interior (implicit function theorem: $dg^*/d\eta = q'(g^*) / [V''(g^*) - \eta q''(g^*)]$, numerator < 0 and denominator < 0 at a quasiconcave interior max). ■

Proposition A.4.2 (collapse at extreme enmity). As $\eta \rightarrow \infty$, $g^*(\eta) \rightarrow \infty$ and the value collapses to the stagnation limit of P1.

Proof. For any finite $\bar{g}$, $\partial W / \partial g(\bar{g}; \eta) = V'(\bar{g}) + \eta(-q'(\bar{g})) > 0$ for η large ($-q'(\bar{g}) > 0$ fixed, $V'(\bar{g})$ fixed), so the maximizer exceeds $\bar{g}$ eventually; hence $g^* \rightarrow \infty$. Then $f(g^*) \rightarrow 0$ and Theorem A.2.2 bounds growth by $O(f(g^*)^\beta) \rightarrow 0$. ■

Proposition A.4.3 (convoy dominance for the leader itself). Fix a bounded defended gap g_c with $f(g_c) > 0$ and $u = 1$, versus any stagnation-region policy ($g \rightarrow \infty$ or $u \rightarrow 0$). Then the convoy trajectory satisfies $x_L^{\{\text{convoy}\}}(t) \geq x_L^{\{\text{stag}\}}(t)$ for all $t \geq t_0$, strictly for large t .

Proof sketch (trajectory comparison). Under the convoy policy the follower stratum settles at the positive steady state of Remark A.2.4, giving persistent inflow $I \geq f(g_c)A(x_L)\sigma D_F^+ > 0$ while $A(x_L) > 0$; under stagnation policies $I \rightarrow 0$ (A.2.2-A.2.3). D_L under the convoy dominates D_L under stagnation by comparison of (6) with ordered inflows (same initial data, larger forcing); D_{eff} and hence $\dot{x}_L$ inherit the order pointwise; integrate. Both trajectories eventually decelerate as $A(x_L)$ declines (Lemma A.1.3) — the claim is comparative dominance, not perpetual growth.

■

Remark A.4.4 (calibration). $g^* > 0$ requires no miscalibration: any $\eta > 0$ — i.e. any nonzero perceived usurpation cost with risk aversion folded in — yields $g^* > 0$ by A.4.1. Writing $\eta_{\text{eff}} = \max(\eta_{\text{perceived}}, \eta_{\text{objective}})$ and noting that leaders who *underestimate* objective risk are disproportionately displaced, surviving leaders satisfy $\eta_{\text{eff}} \geq \eta_{\text{objective}}$ (a selection argument): the trap binds even for perfectly justified fear, and miscalibration only shifts g^* outward.

Remark A.4.5 (uncertainty premium). Writing $\eta_{\text{eff}} = \hat{R} + z \cdot \nu$, with $\hat{R}$ the leader's risk estimate, ν its uncertainty (standard deviation), and $z > 0$ risk aversion: the defended gap $g^*(\eta_{\text{eff}})$ is increasing in ν . Hence transparency and verification — anything that shrinks ν — provably shrinks the defended gap, at rate $dg^*/d\nu = z \cdot dg^*/d\eta > 0$. (Full optimization in the subsequent policy-focused work, with the exit-penalty decomposition $R = q(g) \cdot \Lambda$, Λ the displacement penalty, and the prospect-theoretic loss-aversion treatment.)

Corollary A.4.7 (exit can dominate rule; author-raised 2026-07-05). Compare two leader trajectories under a discount rate ρ_d : *stay*, growing under P1's defended-gap ceiling at sustainable rate g_{stay} , versus *exit* — relinquishing the defended position to join a convoy at rate $g_{\text{conv}} > g_{\text{stay}}$, paying a one-time exit penalty Λ . With value $V(g) = B_0/(\rho_d - g)$ for $g < \rho_d$, exit dominates staying whenever

$$\Lambda < \Lambda^* \equiv B_0 \cdot [1/(\rho_d - g_{\text{conv}}) - 1/(\rho_d - g_{\text{stay}})],$$

i.e. when the exit penalty falls below the capitalized value of the growth-rate difference the convoy unlocks. The threshold rises without bound as the gap's starvation widens $g_{\text{conv}} - g_{\text{stay}}$: the more the defended position strangles growth, the more a leader should rationally pay to leave it. This is the individual-rationality companion to P3's trap — P3 shows the gap is self-defeating; A.4.7 shows that a low enough exit penalty makes leaving it self-interested — and it is the in-model root of the exit-penalty machinery the subsequent policy-focused work develops (lowering Λ is not mercy to the leader, it is the lever that makes the growth-maximizing move also the safe one).

(Code twin: Appendix E.4.)

A.5 P4 and P4b — Suppression backfire, cultivation (analytic)

Symbols here: $G = dG/dD_F$'s G is the leader's growth rate $\theta_L \Phi(D_{\text{eff}})$; $S = f \cdot A$, the two-bottleneck survival factor; σ, ρ, δ = replenishment, decay, and depletion coefficients; Φ' = slope of the growth law.

Lemma A.5.1 (follower steady state increasing in u and D_F^{max}). The follower quasi-steady state $D_F^*(u, D_F^{\text{max}})$ solves $u \cdot h(D) = \rho D$ with $h(D) \equiv \gamma(1 - D/D_F^{\text{max}}) - \delta_F \theta_F \Phi(D)$. h is strictly decreasing with $h(0) = \gamma > 0$; the left side of the fixed-point equation is decreasing in D , the right side increasing, so D_F^* exists, is unique, and lies where $h > 0$. Differentiating: $\partial D_F^*/\partial u = h(D_F^*) / (\rho - u h'(D_F^*)) > 0$ and $\partial D_F^*/\partial D_F^{\text{max}} = u \gamma D_F^* / (D_F^{\text{max}})^2 / (\rho - u h'(D_F^*)) > 0$. ■

Theorem A.5.2 (leader growth increasing in follower diversity — unconditional sign). Hold the leader's gap policy fixed and let $S \equiv f(g)A(x_L) \geq 0$ denote the operating channel-absorption product (quasi-static over the diversity timescale). The leader's quasi-steady-state growth $G = \theta_L \Phi(D_L^* + S D_F)$ satisfies $dG/dD_F = \theta_L \Phi'(D_{\text{eff}}) \cdot S \cdot (\sigma + \rho) / (\rho + \delta \theta_L \Phi'(D_{\text{eff}})) > 0$ whenever $S > 0$.

Proof. The quasi-steady state of (6) solves $\rho D_L = S\sigma D_F - \delta\theta_L\Phi(D_L + SD_F)$ (interior case). Implicit differentiation in D_F : $(\rho + \delta\theta_L\Phi') \cdot dD_L = S(\sigma - \delta\theta_L\Phi') \cdot dD_F$. Then $dG = \theta_L\Phi' \cdot (dD_L + S dD_F) = \theta_L\Phi' \cdot S \cdot [(\sigma - \delta\theta_L\Phi') + (\rho + \delta\theta_L\Phi')]/(\rho + \delta\theta_L\Phi') \cdot dD_F = \theta_L\Phi' \cdot S \cdot (\sigma + \rho)/(\rho + \delta\theta_L\Phi') \cdot dD_F$. All factors positive. In the boundary case (depletion pins D_L^* at 0), $G = \theta_L\Phi(SD_F)$ is directly increasing in D_F . ■

Remark A.5.3 (the burn effect cancels). One might worry that extra follower diversity, by fueling faster leader growth, burns leader stock (the $-\delta\dot{x}_L$ term) and could lower growth on net: the stock effect $\sigma - \delta\theta_L\Phi'$ can indeed be negative when Φ' is large. The theorem shows the worry dissolves in the growth comparative static: the depletion feedback appears with opposite signs in the stock response and the direct response and cancels, leaving the unconditional factor $(\sigma + \rho)/(\rho + \delta\theta_L\Phi') > 0$. More follower diversity may be *stored* or *spent*, but either way the leader grows faster.

Corollary A.5.4 (P4). Long-run leader growth is increasing in u : concatenate A.5.1 ($u \uparrow \Rightarrow D_F^* \uparrow$) with A.5.2 ($D_F \uparrow \Rightarrow G \uparrow$). Under gap defense enforced through suppression, the effect compounds with the $\varphi \cdot u$ product of Remark A.2.3. The informational analogue of Olson's stationary bandit: the take is limited because it grows with the base's prosperity. ■

Corollary A.5.5 (P4b, cultivation — through an open channel). Long-run leader growth is increasing in $D_F^{\max}$ (A.5.1 + A.5.2), giving the ordering cultivate $\succ$ laissez-faire $\succ$ suppress in leader self-interest. But every term of dG/dD_F carries the factor $S = f(g)A(x_L)$: the return to *any* capacity enlargement is throttled by the defended gap, and the cross-effect of capacity and gap-narrowing is positive. Cultivation pays only through an open channel: the self-interested prescription is *cultivate and narrow*, not cultivate at arm's length. ■

(Code twin: Appendix E.5.)

A.6 Paranoia/survival corollary (η -monotonicity analytic; decay clause conditional)

Symbols here: η = enmity; $g^{\min}(\eta)$ = the enforced gap floor; u = follower freedom; the chain runs $\eta \rightarrow$ gap and suppression $\rightarrow$ inflow $\rightarrow$ growth.

Proposition A.6.1 (leader growth decreasing in threat perception). Under R1–R3, long-run leader growth is decreasing in η : the chain $\eta \uparrow \Rightarrow g^*(\eta) \uparrow$ (A.4.1) $\Rightarrow \varphi = f(g^*) \downarrow$, and, where defense is enforced through suppression, $u \downarrow$ (constraint (8)); both lower the growth bound (A.2.2, A.2.3) and both lower realized growth (A.4.3, A.5.4). ■

Corollary A.6.2 (self-termination; conditional on the resource-survival frame). In logistic-with-decay dynamics — survival frame in a companion manuscript (Stremsky 2026, §§4–5); the dynamics themselves lying beyond both papers — a corpus persists only while acquisition outpaces maintenance loss μ . If enforced gap defense drives acquisition below $\mu \cdot x$, the corpus decays ($\dot{x} < 0$). Suppression enters the followers' acquisition directly (through u) but the leader's only indirectly (lost inflow), and the follower stratum is thinner; its acquisition-minus-maintenance therefore turns negative at a weaker suppression level: the base crosses its survival threshold first. Stated here as a corollary conditional on that frame; the dynamical proof is deferred. ■

A.7 P5 — The integration bottleneck (analytic)

Symbols here: $f(g)$ = channel fidelity; $A(x_L)$ = leader absorption; their *PRODUCT* is the point.

Proposition A.7.1. With a perfect channel ($f \equiv 1$), external inflow is $A(x_L) \cdot \sigma \cdot D_F \leq A(x_L) \cdot \sigma \cdot \bar{D}_F \rightarrow 0$ as $x_L \rightarrow \infty$ (Lemma A.1.3), and declines in x_L beyond x^* (Lemma A.1.1). Since f and A enter (2) and (6) only as the product $f \cdot A$, $\sup_f [f \cdot A] = A$: no channel repair compensates a collapsed receiver, and symmetrically $\sup_A [f \cdot A] = f \cdot \bar{A}$: neither factor at its maximum offsets the other near zero. ■

A.8 P6 — The partition proposition (size effect analytic; Conditions A-B as validity conditions)

Symbols here: X = total corpus of the compared organizations; n = number of ensemble members; $y = X/n$ = member size; x^* = the absorptive prime (A.1); s, r as in A.1; d = content spread between members; $P(d)$ = pooling fidelity.

Proposition A.8.1 (partition inequality). Fix total size X and coupling κ (so s fixed), and let the partition into n communicating members of size X/n have ensemble absorption $\Sigma_i A(X/n) = n \cdot A(X/n)$. Then: (i) For every $X > x^*(s)$ there exists $n \geq 2$ with $n \cdot A(X/n) > A(X)$; the choice $n = \lceil X/x^* \rceil$ puts each member at $\approx$ its absorptive prime. (ii) The advantage ratio $n \cdot A(X/n)/A(X) \rightarrow \infty$ as $X \rightarrow \infty$ (holding member size $\approx x^*$), and increases with s (hence with κ_m). (iii) A single member can outabsorb the whole monolith: $A(x^*)/A(X) \geq [x^*/(x^*+r)] \cdot [(X+r)/X] \cdot e^{\{s(X-x^*)\}} > 1$ for X sufficiently past the peak, growing exponentially in X .

Proof. (iii): direct ratio of the two closed forms; the visibility ratio is bounded below by $x^*/(x^*+r) > 0$ while the exponential factor $e^{\{s(X-x^*)\}}$ grows without bound; the product exceeds 1 once $s(X-x^*) > \ln[(x^*+r)/x^*]$. (i): follows from (iii) since $n \cdot A(X/n) \geq A(x^*)$ when some member sits at x^* ($n = \lceil X/x^* \rceil$, members $\leq x^*$ on the rising side lose only boundedly). (ii): with member size pinned near x^* , $n \cdot A(X/n) \approx (X/x^*) \cdot A(x^*)$ grows linearly in X while $A(X)$ decays exponentially; and $A(X)$ shrinks in s faster than $A(x^*(s))$ (both fall, but the monolith sits deeper on the exponential). ■

Validity conditions (assumptions, not conclusions). **Condition A (content diversity):** the ensemble reading of $\Sigma_i A$ requires members whose recognition regions are decorrelated — coverage is the union, not n copies, of a field of view; identical members share blind spots and the sum overstates. **Condition B (communication floor):** pooling requires inter-member fidelity above $f_{\min}$; below it the ensemble fragments into isolated corpora and local absorption no longer benefits the whole (connectivity/percolation threshold on the ensemble graph). The proposition asserts the size effect *given* A and B ; §5 and P7 govern when they hold.

Remark A.8.3 (monolith depth — the absorptive prime is intrinsic; author-raised 2026-07-04). $x^*(s)$ depends only on the receiver's internal parameters — the visibility scale r and the coupling-scaled decay s — and not on the system's total corpus X : within the model, the absorption-optimal unit size is intrinsic to receivers, uncorrelated with how much the system as a whole knows. This non-correlation is itself the result. As total knowledge X grows, the dimensionless **monolith depth** X/x^* grows without bound for any single-corpus organization, and its absorption penalty is exponential in the depth: $A(x^*)/A(X) = [x^*/(X+r)] / [(x^*+r)/X] \cdot e^{\{s(X-x^*)\}}$. (At the Fig. 2 parameters — $X = 400$, $\kappa = 2$, so $x^* \approx 11$ — the optimal-unit ratio is ≈ 1.9

$\times 10^2$; Fig. 2's $\approx 125\times$ is the same ratio at the deliberately suboptimal piece size 50: even far-from-optimal granularity beats the monolith by two orders of magnitude.) An absorption-optimally organized system therefore holds a growing corpus in a growing *number* of units of roughly fixed size: per-unit-corpus absorption $A(y)/y$ is strictly decreasing in unit size y (its log-derivative is $-1/(y+r) - s < 0$), so absorption alone would push toward atomization, and the interior optimal granularity is set by the pooling costs of Condition B and the translation prices of A.3.6 — the size-dimension analogue of A.8.2's interior spread optimum. The system-level reading: the potential absorption loss from monolithizing a corpus of total size X is exponential in $s(X - x^*)$, increasing without bound in both total knowledge and coupling. Caveat, flagged: if a receiver's r , M_0 , K_0 themselves scale with the epistemic environment, x^* drifts; treating them as intrinsic is an assumption of the model (future work).

Proposition A.8.2 (interior optimum of member spread; translators shift it outward). Let members be placed with content spread parameter d (mean inter-member distance). Coverage $\text{Cov}(d) = |\cup B_i|$ is increasing and concave in d (overlaps shrink, gains saturate once regions disjoint); pooling fidelity $P(d) = e^{-\ell_c(d)}$ is decreasing in d (content-distance loss exponent, A.3.6's law), with ensemble performance $\propto \text{Cov}(d) \cdot P(d)$. If Cov is increasing-concave and P decreasing with P log-concave, the product is single-peaked: there is an interior optimal spread d^* — Nooteboom's optimal cognitive distance, derived inside the ensemble rather than assumed. Inserting translator chains of length n between members multiplies the effective per-link loss by $1/n$ (A.3.6), shifting P upward and flattening its decline; the maximizer of $\text{Cov} \cdot P$ then moves outward: $d^*(\text{with translators}) > d^*(\text{without})$. *Proof sketch:* single-peakedness from the sign change of $(\ln \text{Cov} + \ln P)'$, monotone from concavity of $\ln \text{Cov}$ and log-concavity of P ; the translator shift by monotone comparative statics on the parameterized family $P_n(d) = e^{-a d^2/n}$, whose logarithm has increasing differences in (d, n) . ■

(Code twin: Appendix E.6.)

Remark A.8.4 (the module-size band — absorption and diversity bind different dials; author-raised 2026-07-05, investigated and confirmed). Combining the absorption geometry (A.1, A.8.3) with the Homogenization Theorem (A.9.1) yields a design result for a partitioned ensemble of total corpus X : the two criteria that seem to compete — absorb well, stay diverse — in fact constrain *different design variables* and are jointly satisfiable. **Absorption constrains unit size:** per-unit-corpus absorption is strictly decreasing (A.8.3), and units below the visibility scale cannot recognize arrivals, so unit size is banded, $y \in [y_{\min}, x^* + O(1/s)]$, with $y_{\min}$ set by the visibility scale r and by bridge overhead (each additional unit costs translator infrastructure, A.3.6). **Diversity constrains bridge conductance, not unit count:** sustained dispersion under regeneration is governed by the inter-module leakage (the slow eigenvalue modes), and numerically, ensembles of the same total size partitioned into 2, 4, or 8 modules at equal per-bridge weight all sustain diversity within a factor of ~ 1.5 of each other — an order of magnitude above dense or hub topologies — so at fixed total conductance the module *count* is approximately free for diversity, and the absorption band then sets it: $n \approx X/y$ with y in the band. The apparent trade-off between absorbing and staying diverse dissolves into two independent dials — **size near the prime; bridges thin** — and what genuinely opposes both is pooling (Condition B), whose price is exactly the translator infrastructure of A.3.6. The convoy's ecology is thus overdetermined: three separate criteria, two of which do not even

compete.

(Code twin: Appendix E.7.)

A.9 P7 — The ecology proposition

Symbols here: P = the communication network's averaging matrix (row-stochastic: each row = whose content a member mixes, weights summing to 1); p_i = member i 's content position; $V(t)$ = dispersion (diversity) at round t ; λ_2, λ_k = the network's eigenvalues (λ_2 = second largest — THE one dial); s^2 = innovation noise; B_i = member i 's recognition ball of radius w ; B_0 = the allocator's; Cov = system coverage.

A.9.1 The Homogenization Theorem (P7(i); analytic in the stylized model)

Setup. Members $i = 1, \dots, n$ hold content positions $p_i(t) \in \mathbb{R}^m$; recognition regions are balls B_i of radius w about p_i . One communication round: $p(t+1) = P p(t)$ with P row-stochastic, $P_{ii} > 0$ (self-retention, which guarantees aperiodicity), $P_{ij} > 0$ iff i listens to j — the stylized-absorption assumption that absorbing moves the receiver toward the source (DeGroot 1974). Content dispersion: $V(t) = (1/n) \sum_i \|p_i(t) - \bar{p}(t)\|^2$, $\bar{p}$ the barycenter. For symmetric P , eigenvalues $1 = \lambda_1 > \lambda_2 \geq \dots \geq \lambda_n > -1$ when connected; λ_2 below denotes the second-largest eigenvalue modulus; spectral gap $1 - \lambda_2$.

Theorem A.9.1. (a) Connected $\Rightarrow$ consensus. If the listening graph is strongly connected, P is primitive ($P_{ii} > 0$), so $P^t \rightarrow \mathbb{1}\pi^T$ with π the Perron left-eigenvector ($\pi^T P = \pi^T$, $\pi > 0$, $\sum \pi_i = 1$); every $p_i(t) \rightarrow \sum_j \pi_j p_j(0)$ and $V(t) \rightarrow 0$. At consensus all recognition regions coincide: coverage collapses from $|\cup B_i|$ ($\approx n \cdot |B|$ for well-spread members, Condition A) to $|B|$. **(b) Rate.** For symmetric P , $V(t) \leq \lambda_2^{2t} V(0)$. **(c) Ordering.** Uniform complete averaging ($P = \mathbb{1}\mathbb{1}^T/n$) has $\lambda_2 = 0$: consensus in one round. In the lazy family $P = (1-\omega)I + \omega\hat{W}$, $\lambda_2(P) = 1 - \omega(1 - \lambda_2(\hat{W}))$; the complete graph maximizes the gap in this family, while a modular $\hat{W}$ with inter-cluster conductance h obeys $1 - \lambda_2(\hat{W}) \leq 2h$ (Cheeger), so sparse bridges make homogenization arbitrarily slow. Disconnected graphs have $\lambda_2 = 1$ (eigenvalue 1 with multiplicity = number of components): $V(t) \rightarrow V_\infty > 0$, permanent diversity, no pooling (Condition B fails globally). **(d) Stationary diversity under regeneration.** Add i.i.d. innovation shocks $\xi(t)$, mean 0, variance s^2 per member per round (the explorer role). For symmetric P the centered process $d(t+1) = P d(t) + \Pi \xi(t+1)$ (Π the centering projector) has stationary dispersion $E V_\infty = (s^2/n) \sum_{k \geq 2} 1/(1 - \lambda_k^2)$, strictly decreasing in each spectral gap: complete quenches diversity to $\approx s^2(n-1)/n$; a modular network sustains $E V_\infty \approx s^2/[n(1-\lambda_2^2)]$, unbounded as bridges thin; disconnected components accumulate diversity without bound and can never pool it.

Proof. (a) Perron-Frobenius for primitive stochastic matrices; $V(t) \rightarrow 0$ since all rows of the limit are equal. Coverage collapse: coincident centers $\Rightarrow \cup B_i = B$. (b) Decompose the centered vector in P 's orthonormal eigenbasis; the $\mathbb{1}$ -component is removed by centering (symmetric P preserves the barycenter); each remaining component contracts by $|\lambda_k| \leq \lambda_2$ per round; square. (c) Spectrum of $\mathbb{1}\mathbb{1}^T/n$ is $\{1, 0, \dots, 0\}$. The lazy-family identity is linear algebra; Cheeger's inequality is standard (Chung 1997). Disconnected: block-diagonal P has eigenvalue 1 per block; within-block consensus, cross-block never. (d) In the eigenbasis, each mode $k \geq 2$ is an AR(1) with

coefficient λ_k and innovation variance s^2 : stationary variance $s^2/(1 - \lambda_k^2)$; sum and normalize. ■

Remark A.9.2 (one dial, not two). The same λ_2 that accelerates pooling accelerates homogenization: density cannot buy fast sharing *and* lasting diversity — it can only be *placed*. The modular-with-translators design places it: large gap within clusters (fast local pooling), small gap between (slow global homogenization), with translators raising bridge *fidelity* rather than bridge *weight*.

Remark A.9.3 (limitation). The stylized model treats all communication as position-averaging and cannot distinguish pooling an item from converging a corpus; the agent model of §5.3 passes items and confirms the trade-off at realistic parameters (Figs 4–5, dense dispersion ≈ 0.02 vs modular 1.36).

A.9.2 The Recognition-Bottleneck Lemma (P7(iv); analytic)

Lemma A.9.4. Opportunities lie in content space Ω ; member i recognizes $B_i \subseteq \Omega$; discovery requires resources. Under full centralization a single allocator with recognition region B_0 routes all resources, and can route only to opportunities it recognizes. Then realized system coverage $\text{Cov_sys} = \cup_i (B_i \cap B_0) \subseteq B_0$, so $|\text{Cov_sys}| \leq |B_0|$ — one corpus’s field of view caps the system, independently of n , member diversity, and channel quality. Decentralized allocation attains $|\cup B_i|$ ($\approx n|B|$ under Condition A). ■

Remark A.9.5 (delegated proposal does not escape). If members propose and the center approves, approval is itself recognition (judging relevance requires the judge’s corpus to render the proposal visible — the same visibility factor as eq. (5)); the bound stands. It relaxes only where the center defers, i.e. under partial decentralization.

Corollary A.9.6 (partial centralization). With the center controlling resource fraction α and members the rest, opportunities in $\cup B_i \setminus B_0$ retain only the autonomous share; with increasing returns to resource intensity, expected discovery is nonincreasing in α wherever $\cup B_i \setminus B_0 \neq \emptyset$, and “limited enough” central allocation is $\alpha \leq \alpha^*$, the largest fraction keeping coverage-weighted discovery above target. (Outline; full treatment in the network paper.)

Corollary A.9.7 (recursive ensemble). k content-diverse allocators relax the bound to $|\text{Cov_sys}| \leq |\cup_a B_0^{(a)}|$: the ensemble argument applied to the allocation layer itself; decorrelated allocator blind spots provably dominate any single allocator (dominate in coverage: the k -allocator system’s recognized set is the union of the individual balls, a superset of any one allocator’s ball — strictly larger whenever the allocators’ recognition regions differ). ■

A.9.3 Translator comparative advantage (P7(ii)-(iii); argument with analytic core)

A translator is a corpus overlapping two clusters: larger and more coupled than a specialist (absorption penalty, on the falling slope of A — Lemma A.1.1), but at small content distance to both clusters, so its links run at high fidelity (A.3.6 with $n = 1$ and small d). Specialists near the absorption peak dominate at intra-domain discovery; translators dominate at conveyance (Allen 1977; Galison 1997) and at cross-domain recombination across structural holes (Burt 2004). Role differentiation then follows from comparative advantage over the absorption curve: assigning discovery

to peak-sized specialists and conveyance to bridge-positioned translators dominates any homogeneous assignment. The analytic content is A.1.1 + A.3.6; the role claim is argued and simulation-supported (Fig. 4).

A.9.4 Redundancy addendum (one sentence in §4; basis here)

Serial translator chains fail with any single link; replacing the bridge *path* by a bridge *graph* of redundant routes — not necessarily parallel copies, nor all translator-mediated; any surviving indirect path carries the transfer — makes inter-cluster transfer succeed if any path survives — reliability $1 - \prod_j(1 - R_j)$ for independent path reliabilities R_j (standard network reliability) — so a robust ecology duplicates translator paths, trading cost for graceful degradation.

(Code twin: Appendix E.8.)

A.10 Dependency recap (mirrors §4.10)

A.2 (P1) uses only the model plus a fidelity ceiling — analytic. A.3 (P2) uses A1 including the second-order clause — analytic, with the clause’s necessity shown (A.3.3). A.5 (P4/P4b) uses nothing beyond the equations — analytic, unconditional sign. A.6.1 concatenates A.4 + A.5 — inherits A.4’s regularity. A.4 (P3) is the one conditional core result (R1-R3, all labeled). A.7 (P5) and A.8.1 (P6 size effect) are analytic given the form of A; A.8.1’s ensemble reading is conditional on Conditions A-B, and A.8.2 on stated shape conditions. A.9.1 and A.9.4 (P7(i),(iv)) are analytic *in the stylized model*, with the model’s limitation stated (A.9.3); A.9.3’s role claims remain argument + simulation. A.6.2’s decay clause is conditional on a companion manuscript’s survival frame (Stremsky 2026), marked as such in the corollary itself.

Appendix B — Simulation Details

B.1 Implementation

All dynamics are forward-Euler integrations of eqs. (1)–(8) in plain Python (NumPy only); the agent model of §5.4 is a discrete-round simulation over explicit graphs. State variables follow §3’s notation exactly. The four scripts (`sim_absorption.py`, `sim_dynamics_4way.py`, `sim_topology_4way.py`, `sim_ablation.py`, `sim_realrisk.py`), the two summary CSVs, and the figure files are included in the supplementary material, together with the consolidated code file of the paper’s analytic formulas; random seeds are fixed in the scripts, so every reported number reproduces exactly.

B.2 Parameters by figure

Fig. 1 (absorption geometry). $A(x) = [x/(x+r)] \cdot e^{\{-x/M\} \cdot e\{-x/K\}}$ with $r = 2$, $M = M_0/\kappa$, $K = K_0/\kappa$, $M_0 = 240$, $K_0 = 360$ (combined decay scale $1/(1/M_0 + 1/K_0) = 144$ at $\kappa = 1$); $\kappa \in \{0.5, 1, 2, 4\}$. Peaks: $x \approx 23.0, 16.1, 11.1, 7.5$ with $A \approx 0.85, 0.80, 0.73, 0.64$ — matching Lemma A.1.1’s closed form.

Fig. 2 (partition). Total size $X = 400$, $\kappa = 2$; monolith $A(X) = 0.0039$ versus $n \cdot A(X/n)$: $n = 2 \rightarrow 0.123$, $n = 4 \rightarrow 0.978$, $n = 8 \rightarrow 3.841$; single 8-partition member $A(50) = 0.480$.

Fig. 3 (dynamics). State (x_L, x_F, D_L, D_F) , horizon $T = 400$. Channels: singleton regimes use the one-jump form $e^{\{-0.08 g^2\}}$; the convoy uses an effective-spectrum residual $e^{\{-0.006 g\}}$ (a graded spectrum's small per-unit-distance remainder, per eq. 4). Convoy: gap capped at 5 by leader throttling, downward spillover 0.25, loose coupling $\kappa = 0.8$. Suppression singleton: $u = 0.15$, tight coupling. Outrunning singleton: $u = 1$, defends a minimum gap of 10 by acceleration and (when followers close) enforcement of eq. (8). Hub regime: minimum gap 6 defended as in outrunning; $\kappa = 2.5$; one-jump channel; atomization as a collective-brain discount $\gamma \times 0.35$ with $D_F^{\max}$ halved — argued parameters, not fitted; sensitivity over discounts 0.2–0.7 keeps the regime strictly intermediate throughout. Endpoint pattern: suppression x_L 28.5 $\rightarrow$ 28.5 (rate 0.0000, gap 18.2, $D_L = 0$); outrunning 28.6 $\rightarrow$ 33.3 (rate 0.0137, gap 9.3, $D_L = 0$); hub 28.2 $\rightarrow$ 78.8 (rate 0.174, gap 6.0, $D_L = 0.35$); convoy 102.4 at $t = 100$, 445.4 at $t = 400$ (rate 0.890, gap 5.0, $D_L = 101$). System totals $(x_L + x_F \text{ at } T)$: suppression 38.9, outrunning 57.3, hub 151.6, convoy 885.8 — the regimes order identically on the leader's and the system's ledger, and the best defense trails the convoy by $\approx 5.7\times$ on both. Hub sensitivity on both ledgers (γ -factor, D_F -factor): $(0.2, 0.5) \rightarrow$ leader 65.7 / system 125.4; $(0.35, 0.5) \rightarrow$ 78.8 / 151.6; $(0.5, 0.5) \rightarrow$ 87.3 / 168.6; $(0.7, 0.7) \rightarrow$ 101.4 / 196.7 — orderings unchanged throughout.

Fig. 4 (topology). 24 agents, three content clusters, 25 seeds, equal total connectivity budget across topologies. Fourth topology, hub: agent 0 relocated to the content centroid, edges hub $\leftrightarrow$ followers only; atomization discount $\text{GEN_DISCOUNT} = 0.5$ on follower generation (argued, mirroring the dynamics run's collective-brain halving; not fitted). At the full horizon: coverage/total-absorption — dense 260/260, random 260/260, hub 214/213, modular-with-translators 340/339; final overall content dispersion — dense 0.68, random 0.70, hub 0.67, modular 1.10 (single metric throughout). Per-seed hub–dense coverage -46.4 ± 13.7 . Sensitivity: discounts 0.35/0.5/0.7 give hub coverage 187/213/231, all strictly below dense (260); at discount 1.0 — no atomization cost, the bounding case — the hub matches dense (257 ± 9 vs 260 ± 11), which is the pure-mediation result. $\lambda = 0$ ablation (homogenization drift off): dense and modular converge (392/394) while the hub retains its deficit (322) — switching homogenization off isolates the generation channel. Two design notes. The discount is not double-counting: the harness's discovery probability is position-based and independent of an agent's holdings or peers, so no collective-brain effect exists endogenously — the discount supplies the missing channel, it does not repeat one. And the step form is a knee approximation: a symmetric variant giving every agent generation $\text{deg}/(\text{deg}+1)$ — no hub special-casing anywhere — reproduces both the ordering and the hub's level within noise (dense 256, random 253, hub 211, modular 327), because a single peer sits at the saturating curve's half-point while five and twenty-three sit above its knee. The two panels measure different channels — coverage counts first discoveries only (transmission cannot create first knowledge), so it isolates the blindness channel, while absorption carries blindness and transmission together.

Fig. 6 (the real-risk trap). Exemplar for the P3 Calibration remark: $V(g) = e^{\{-ag^2\}}$ (growth value through the jump channel; $V(0) = 0$ is A1's cheap-first-hop

property), $R(g) = \eta_{\text{eff}} \cdot q_0 e^{-\lambda g}$ (expected usurpation cost; R1-R3 hold), $\eta_{\text{eff}} = \eta + z \cdot \nu$ per Remark A.4.5. Parameters $a = 0.05$, $\lambda = 0.7$, $q_0 = 1$, $z = 1$; g^* on a 3001-point grid; single-peakedness checked numerically. The $\nu = 0$ curve is the claim: $g^*(\eta) > 0$ at every $\eta > 0$ — real risk alone traps — and it rises ($g^* \approx 0.3$ at $\eta = 0.05$, ≈ 2.1 at $\eta = 1$, ≈ 3.4 at $\eta = 3$); positive ν shifts the whole curve up without changing its shape. Illustration of shape, not calibration.

Fig. 5 (ablation). Identical to Fig. 4 with the channel’s content-dependence removed (flat per-link fidelity); modular leads at every checkpoint; no early dense advantage appears.

B.3 How stagnation had to be installed (an evidential disclosure)

The dynamics of Fig. 3 reached their present form in three versions, and the path matters evidentially, so we disclose it rather than presenting the final model as first-intended. The first version contained growth, channel, absorption, and diversity dynamics — eqs. (1)–(7) — but no security constraint, and it *refused to stagnate*: a leader that pulled ahead would deplete its diversity, slow, and be caught; the closing gap revived the channel; and every run settled into a bounded-gap, mutual-flow regime — the convoy as the system’s attractor. Stagnation could not be produced until the model gave the leader a motive to prevent the gap from closing: only with eq. (8) enforced — the race-inherited security constraint — did the persistent-gap, depleted-diversity plateau of §5.3 appear. Intermediate revisions concerned integration details and the two singleton policies’ definitions; the substantive change across versions is the single one just described. We regard the disclosure as strengthening rather than qualifying the result: the pathology the paper analyzes is not generic to hierarchy — it had to be *installed*, and what installs it is gap defense. That is Proposition 3’s causal claim, discovered first by the simulation’s refusal.

B.4 Robustness

The qualitative conclusions — single-peaked absorption with coupling-dependence, steep partition advantage, both singleton routes plateauing far below the convoy, self-correction absent enforcement, modular dominance with dense homogenization — are stable under moderate variation of all reported parameters; the reported numbers move and are not claims (§5.5). The regime-to-coupling mapping (convoy loose, singleton tight) is an argued assumption of §3, encoded in the setups and tested jointly with the dynamics.

B.5 Reproduction

Run each script with no arguments to regenerate its figure and CSV; seeds are fixed inline. The supplementary code file additionally contains executable versions of the paper’s analytic formulas (the “code twins (collected in Appendix E)” of Appendices A and D), so every closed-form claim can be checked numerically without rederivation.

Appendix C — Glossary and 2016-Vocabulary Map

C.1 Terms

Knowledge corpus (corpus). An entity’s accumulated, integrated body of knowledge; the paper’s umbrella object. In the 2016 posts on which the paper builds, the author’s term was *thesaurus*, in the sense of an organized treasury of concepts.

Primitives; idea-lineages; generator set G . Primitives are the corpus’s irreducible elements; idea-lineages its maximal independent generators, collected in G . The two typically coincide; their formal relation is left open.

Realized content C ; realized size x . The explored, viable combinations; $x = |C|$ is the model’s size variable. *Scope:* x is corpus size, not capability-to-enact (§3.1).

Latent content C^* ; latent diversity D . The unexplored remainder of the viable closure; D is its breadth — a diversity measure over content space, not a count; intuitively, the corpus’s novelty potential. Recombination converts latent into realized content, spending D by construction.

Viable; viability filter. A combination is viable if it passes the filter of workability — consistency, usefulness, survival in use — adjudicated by reality rather than by its generator: experiment, practice, environment. The filter is where much unmodeled difficulty resides; its stylized treatment (viability probability decreasing in novelty — the adjacent possible) is Appendix D.1.2.

Theoretical vs viable closure. All combinations vs those passing a viability filter; “size” always means realized viable content, and unqualified “latent” means latent-viable. The viability filter is where much unmodeled difficulty resides.

Level; stratum. A level is a position in the corpus-size ordering; a stratum is a model aggregate of levels (two in this paper: leader, follower population). A *spectrum* is a population occupying many intermediate levels.

Size gap g . $g = x_L - x_F$, the directional difference in realized size between strata; the quantity the race literature discusses as a capability gap or lead.

Cognitive distance $d(i, j)$. The content difference between two corpora — symmetric, defined regardless of standing; operationalized over realized-content sets as one minus overlap (Jaccard form), or distributionally as in Appendix D. Distinct from g ; large g typically implies large d , not conversely.

Loss exponent $\ell(\Delta)$; tacit fraction $c(g)$. A single transfer over distance Δ has fidelity $e^{-\ell(\Delta)}$; ℓ is increasing, convex, second-order at zero (A1). $c(g)$ is the fraction of content inarticulable for transfer at any hop size (Polanyi), with $c(g) \leq 1 - e^{-\ell(g)}$.

Absorption $A(x)$; absorptive prime x^* . The receiver-side integration factor of eq. (5) — visibility $\times$ openness $\times$ restructuring — single-peaked in x ; x^* is the peak, intrinsic to the receiver (Remark A.8.3).

Breadth vs size vs density. A corpus’s size x (item count), breadth (content spread), and density (items per unit content) are three distinct quantities; the model’s x is the count, breadth enters only Appendix D’s geometry, and conflating them conflates a specialist army with a thin generalist (D.2.1).

Coupling κ . Dependency density of a corpus’s elements; a parameter, not a state variable. High κ : monolithic integration; low κ : federated, churning organization.

Suppression variable u ; enmity η . $u \in [0,1]$ is the followers’ freedom to develop, the leader’s sole instrument over them; η is threat perception, the driver of the

security constraint (8), deliberately the enmity variable of the race literature.

Gap-defending vs mobility-preserving systems. Model-internal, apolitical system types: whether a minimum size gap is enforced below the apex. Gap defense is a smooth scale present in all hierarchies, not a property of named political forms.

Convoy. A regime in which all strata keep ascending in corpus size *together*: gaps bounded, flows mutual, and upward mobility of members between levels preserved — after the naval practice of bounding speed to hold formation together. Coined in this paper for a concept argued in the 2016 posts. (The complementary downward mobility — exit and rotation at the apex — appears in Corollary A.4.7 and §7.)

Monolith; ensemble; translator; explorer. A monolith is a single tightly coupled corpus; an ensemble a communicating partition of content-diverse members. Translators are members specialized in cross-cluster conveyance — optimally wider and shallower than cluster members, though position, not width, does the first job, and a minimum useful breadth exists (Appendix D.2.3); explorers specialize in novelty generation. Ensemble structure and roles are the subject of P6–P7.

Monolith depth X/x^* . The dimensionless ratio of a system’s total corpus to the receiver-intrinsic absorptive prime; the monolith’s absorption penalty is exponential in it (Remark A.8.3).

Epistemic starvation. The phenomenon the paper establishes: the decline of a leader’s growth as gap defense throttles the channel, suppression destroys the source, and the leader’s own size and coupling degrade its absorption.

C.2 Symbols

The symbol table of §4 applies paper-wide; this appendix does not duplicate it.

C.3 The 2016 vocabulary and claims (provenance map)

The paper’s core arguments were first published by the author in 2015–2016 (footnote 1). The thread argued them specifically in the context of Singularity-era, AI-based development; establishing their validity for gap-defending knowledge hierarchies in general is the present paper’s extension (§1). Entries marked *near-verbatim* reproduce the archived thread’s wording up to elision; the remainder are paraphrases of record. The following map records the correspondence, term by term and claim by claim; left column entries are the 2016 thread’s vocabulary and assertions, right column the present formalization.

2016 term	present term
thesaurus	knowledge corpus
thesaurus width	(latent) diversity D
intellect levels	levels; strata (aggregated)
uninterrupted spectrum (of levels)	graded spectrum; eqs. (3)–(4)
(the concept of bounded-gap mutual ascent)	convoy (term coined here)

2016 claim (paraphrase of record)	formal counterpart
“Transporting information between intellect levels is never lossless; intermediate levels smooth this process.” “The lower levels are more important to the higher ones than most people think.” “Availability [of the lower levels’ development] is a key investment.” “...the very act of establishing an oligarchy requires negotiating a common base of axioms that puts its status quo as an exclusive group above anything else.” (<i>near-verbatim</i>) “A benevolent peacekeeping intention produces the same result as a maniacal world-dominance obsession.” “For simplicity, two intellect levels only.” “A Singularity monarchy ... the only thing having smaller diversity/broadness at its top level, and maintaining a larger gap ... than a Singularity oligarchy.” (<i>near-verbatim</i>) “As the intellects that block the development will be much more powerful than these who would like to continue it, this period might continue indefinitely.” (<i>near-verbatim</i>) “A natural disaster too big and/or unexpected to be prevented ... might wipe out the human existence.” (<i>near-verbatim</i>) “The lack of diversity will also lead eventually to thesaurus potential depletion and dying of our civilization.” (<i>near-verbatim</i>)	Channel eqs. (3)–(4); Spectrum Theorem (P2) Eq. (6): $f \cdot A \cdot \sigma \cdot D_F$ is the leader’s only external inflow; P4 Cultivation corollary (P4b) Security constraint (8); dominance trap (P3). The axiom-negotiation clause additionally gestures at the consensus mechanism of the Homogenization Theorem (P7(i)) Benevolence-independence: no intention term in (1)–(8); P1’s invariance The two-strata reduction (§3.1) Concentration comparative: monolith vs small apex ensemble — partition inequality (P6) and the k-allocator corollary of the Recognition-Bottleneck Lemma (P7(iv)) Persistence of enforced stagnation: gaps persist only where enforced, and enforcement capacity makes the regime stable (P3; §5’s self-correction result read in reverse) Shock death, via the two routes the 2016 “and/or” already separates: <i>unexpected</i> — shock blindness, coverage collapse under homogenization (P7(i)); <i>too big</i> — capability plateau under stagnation (P1) freezes prevention capacity while shock magnitudes are unbounded. Both read in the survival frame of the subsequent foundational work (conditional mapping; the composition is future work) Starvation death: the survival corollary (§4.6), decay clause conditional on the same foundational frame

Appendix D — Microfoundations: Content Geometry, the Channel, Mobility, and Non-Substitution

D.1 The generating relation

A corpus is its primitive base plus its realized viable recombinations: $C = B \cup R$, with $R \subseteq \text{viable-closure}(B)$. Realized size is $x = |C|$ — corpus size only, not capability-to-enact; the microfoundations inherit §3.1’s scope boundary. Latent content C^* is the unrealized remainder of the viable closure, and latent diversity D is its breadth over content space. Two consequences used in the main text follow by construction.

(i) Depletion: recombination moves elements from C^* to C , so growth spends D — the $-\delta x$ term of eq. (6) is bookkeeping, not assumption. **(ii) Import amplification and its limits:** an item entering C can recombine with what C already holds, so its latent contribution is potentially large (Remark A.2.5 grants it polynomial strength) but bounded in *non-redundant* terms by the breadth measure: combination tokens sharing the new item are mutually correlated, and their content can coincide with content already reachable — the count grows fast, the breadth slowly.

D.1.2 The viability filter, stylized (optional for v1; author-raised 2026-07-04). Viability and novelty are distinct filters and must not be identified. Decompose a candidate recombination’s expected contribution into two factors: **viability** v — whether the combination works, adjudicated by reality, and *decreasing* in novelty, since candidates near tested structure inherit its tested-ness (the adjacent possible; Kauffman 2000) — and **non-redundancy** — its contribution to latent diversity, *increasing* in novelty v , the candidate’s distance to the nearest realized item. With the exemplar $v(v) = e^{-v/\lambda}$,

$E(v) = v \cdot e^{-v/\lambda}$, maximized at $v^* = \lambda$:

expected value-added is single-peaked in novelty — too-similar is redundant, too-novel is non-viable — an interior optimum matching, in shape, two independent literatures: the *optimal cognitive distance* finding of the organizational-learning literature — innovation performance is inverted-U in partner distance (Nooteboom 1999; Nooteboom et al. 2007) — and the *atypical-combinations* finding of scientometrics — the highest-impact science pairs high conventionality with a minority of atypical combinations (Uzzi, Mukherjee, Stringer & Jones 2013), an interior optimum of novelty at the level of individual papers; the viability side’s decay is Kauffman’s (2000) adjacent possible. Pool dynamics follow from geometry: with x realized items on a content interval of length L , the mean novelty of remaining candidates is $\bar{v}(x) \approx L/(2x)$, so **creating recombinations makes the remaining pool safer and staler** — average viability $\bar{v}(x) = e^{-L/(2x\lambda)}$ rises toward 1 while average value-added $\bar{E}(x) \approx (L/2x) \cdot e^{-L/(2x\lambda)}$ falls — a quantitative microfoundation of the depletion term $-\delta x$ and of ideas getting harder to find. **Importing a knowledge item replenishes in proportion to its distance:** an arrival displaced d beyond the frontier creates candidate mass at novelties up to $\sim d$, raising $\bar{v}$ — which microfounds why the followers’ value to the leader is their *difference*, not their volume. A hard viability threshold $v_{\max}$ (“beyond this novelty, non-viable”) makes the viable-latent set the $v_{\max}$ -neighborhood of the realized set, $|C^*|(x) \approx \rho \cdot \min(L, 2xv_{\max})$, which saturates once neighborhoods tile the space — exact $|C|:|C^*|$ proportions and the full threshold combinatorics are future work.

(Code twin: Appendix E.9.)

D.2 Content geometry: deriving the channel

D.2.1 Corpora as content distributions. Stylize a corpus as a distribution p over a content space: where its mass sits is what it knows well — equivalently, p is the smoothed density of the corpus’s item cloud, so the density at a content location is the compounded closeness of that location to the corpus’s items. Take the one-dimensional Gaussian stylization $p_i = N(\mu_i, \sigma_c^2)$: position μ_i is the corpus’s content location, breadth σ_c its spread. Three quantities must be kept distinct here: **size** x (how many items — the main text’s variable), **breadth** σ_c (how widely the mass spreads), and **density** x/σ_c (items per unit of content) — a large corpus can be narrow and dense (a specialist army) or small and wide (a thin generalist), and nothing below conflates them: breadth alone enters the hop geometry. Two distances must likewise be kept distinct: the **content displacement** $|\mu_i - \mu_j|$ between corpus centers, and **cognitive distance** proper, which the glossary operationalizes as one minus overlap — the formulas below are written in displacement, and breadth converts between the two (wide corpora are cognitively closer at the same displacement).

One more identification must be made explicit (author query): in this representation, **distance is similarity, not linkage** — proximity of two items means their contents are related (they *can* connect cheaply), while the corpus’s *realized* links (its actual internal recombinations and associations) are a separate graph laid over the point cloud, which the main text summarizes scalar-ly as coupling κ . Item density (points per content length) and link density (edges per item pair) are therefore different quantities, coinciding only under the extra assumption that realized connectivity saturates potential connectivity locally — plausible for mature regions, false in general, and exactly where the uneven-connectivity phenomena of §3.7’s limitation live. A full item-level model carries both layers (a metric embedding plus a graph over it), in which the “indirect distance” of compounded closeness is graph-geodesic distance, approaching metric distance where links saturate — future work. As one exemplar of how κ could be made exact — offered as an example, since other operationalizations (weighted links, motif counts, dependency depth) are equally admissible — take the realized-link graph’s density: $\kappa = 2E / (x(x-1))$, the fraction of possible item pairs actually connected, which lies in $[0, 1]$ and rises as recombination wires the corpus together.

(Code twin: Appendix E.10.)

D.2.2 Fidelity as overlap. Take single-transfer fidelity to be the standard overlap of the two distributions, the Bhattacharyya coefficient $BC(p, q) = \int \sqrt{pq}$. For equal-breadth Gaussians at distance d ,

$$BC = \exp(-d^2 / (8\sigma_c^2)),$$

so fidelity is exactly of the form $e^{-\ell(d)}$ with $\ell(\mathbf{d}) = \mathbf{a} \mathbf{d}^2$, $\mathbf{a} = 1/(8\sigma_c^2)$: the loss exponent is quadratic — increasing, convex, $\ell(0) = 0$, $\ell'(0^+) = 0$. (Here d is center displacement; under the glossary’s overlap operationalization, widening σ_c literally *reduces cognitive distance* at fixed displacement — the two statements “wider corpora pay less per hop” and “wider corpora are cognitively closer” are the same statement, not a probabilistic tendency.) Under the stylization, assumption A1 is not assumed but derived, and the exemplar is not an example but the generic case; moreover the loss coefficient acquires an interpretation: **a is inverse shared-context breadth**. A simplification must be owned at this point: real transfer fidelity is richer than any

single overlap scalar — it depends on the *density* of the overlap region (how many shared anchors actually sit there), on *where* the transferred content lies relative to the overlap (content far from it must be routed through it), and on direction. The Bhattacharyya scalar is density-weighted at each point (it integrates the product of densities) but aggregates all of this to one number; we adopt it *for simplicity*, claim only that it demonstrates the tendency, and defer the exact treatment — fidelity as a density-weighted operator on content location — to subsequent work.

(Code twin: Appendix E.11.)

The same computation grounds both uses of the law in the main text: over *level* distance (the vertical spectrum of eqs. 3–4, since level separation displaces content location) and over *content* distance (the horizontal translator chains of Corollary A.3.6) — one geometry, one coefficient, as the §4 notation table asserts. The identification underlying this deserves its own label — **level-content link (assumption)**: separation by g levels displaces content location proportionally (adjacent strata work on adjacent content), so vertical gap and horizontal displacement share the coefficient a ; if level structure instead compressed or stretched content displacement, a would differ by axis, changing constants but no shape result.

D.2.3 Intermediaries and translators, derived (corrected after author review). Place an intermediary corpus of breadth σ_T midway between two clusters at displacement d . For Gaussians of unequal breadths, the per-hop overlap is

$$BC = \sqrt{(2\sigma_1\sigma_2 / (\sigma_1^2 + \sigma_2^2))} \cdot \exp(-d^2 / (4(\sigma_1^2 + \sigma_2^2))),$$

and three facts follow, in order of what does the work. **(i) Position does the first job, and width is not required:** an intermediary of merely *equal* breadth already improves end-to-end fidelity (two hops of $d/2$ halve the total exponent — this is the spectrum theorem itself, seen in the geometry), so the author’s instinct is confirmed: no larger-than-cluster breadth is needed for an intermediary to help. **(ii) But not any nonzero breadth suffices:** the prefactor punishes narrowness — a near-point intermediary overlaps almost nothing on either side (no shared context with anyone) and *harms* transmission below a **minimum useful breadth** $\sigma_T^{\min}(d)$, which *decreases* with the distance bridged (numerically, at $\sigma_c = 1$: $\sigma_T^{\min} \approx 0.74$ for $d = 1$, ≈ 0.43 for $d = 4$ — wide chasms tolerate narrow translators because the exponent savings dominate). **(iii) Above the threshold, breadth trades off:** widening shrinks the exponent (far content becomes reachable) but shrinks the prefactor (breadth dilutes depth — a maximally wide corpus overlaps everything a little and nothing well), which is Proposition 7(iii)’s trade-off derived rather than argued, and maximizing end-to-end fidelity over σ_T yields an interior optimal breadth — first-order condition, with $u \equiv \sigma_c^2 + \sigma_T^2$: $1 - 2\sigma_T^2/u + \sigma_T^2 d^2/(4u^2) = 0$ — increasing in the distance bridged: farther clusters want wider translators, never infinitely wide ones. The usefulness-per-breadth curve (author query) has a characteristic shape: it rises *linearly* from zero (a sliver of breadth buys its proportional sliver of shared context), peaks at the interior optimum, and decays *hyperbolically* ($\propto 1/\sigma_T$ — the dilution tax) — a right-skewed single hump; at $\sigma_c = 1$, $d = 4$: $F = 0.027$ at $\sigma_T = 0.1$, 0.135 (= direct) at 0.43 , 0.537 at the peak 2.06 , 0.194 at 10 . And the counterintuitive fact that the *minimum useful breadth falls* as the distance grows has a one-line intuition worth keeping wherever this appears: **the worse the direct road, the lower the bar a bridge must clear** — at large d the undivided jump is nearly hopeless, so even a poor intermediary’s positional splitting of the exponent is a huge multiplicative gain that outweighs its prefactor tax, while at small d the direct channel is

already good and the tax buys little. One caveat carries over from D.2.2: normalized distributions price shape, not mass — the intermediary’s *item count* (its density over its breadth) is invisible to BC, and a mass-weighted treatment, in which spreading a fixed x_T thin degrades the translator’s own per-region depth, belongs to the same subsequent-work operator as D.2.2’s.

(Code twin: Appendix E.12.)

D.2.4 The tacit core (definition refined after author review). A definition must survive the following dichotomy, raised against a draft of this appendix: if tacit content has *no* connection to the articulate corpus, it is a separate corpus that merely shares a body; if it is *logically connectable*, it is eventually representable — so “tacit mass” names nothing. The resolution is that the dichotomy’s second horn conflates connection with encodability: tacit components are coupled to articulate ones *operationally* — they steer perception, judgment, and skilled execution of articulate tasks — and an operational coupling is not itself a proposition that a message space can carry. Formally: let content have articulable and embodied coordinates, correlated but not mutually reducible; messages carry the articulable projection, from which a receiver *partially reconstructs* the embodied components to the extent its own practice supplies the coupling (this is what apprenticeship is). The tacit fraction c is therefore not the raw embodied mass but the **irreducible reconstruction residual** — embodied mass times one minus reconstructability — which is why it is positive, why it varies with shared practice (plausibly growing with distance as practice thins), and why it caps every messaging path at $1 - c$ while remaining real and functional in its carrier. Eq. (4)’s ceiling is thus placed, with its meaning sharpened rather than assumed. The status of the residual’s *existence* ($c > 0$) must be stated honestly: it is not a mathematical theorem — in principle a complete physical description of any skill exists — but a claim about practical message spaces and receiver reconstruction budgets. In that form it is strongly defended: philosophically by Polanyi (1966) and most systematically by Collins (2010), whose analysis of collective tacit knowledge comes closest to an in-principle irreducibility argument; empirically by the knowledge-management and economics literatures on the stubborn transfer costs of embodied capability; with Nonaka and Takeuchi (1995) cited as the honest counterpoint, since their conversion framework is read as optimistic about articulability. In the model, $c > 0$ is accordingly a labeled assumption whose defensible form is a resource bound — description length and practice requirements exceeding what channels and receivers can carry — not a metaphysical wall.

D.3 Mobility at the item level

A message must survive the trip and the landing; a person, only the landing. A carrier that ascends (or crosses) transports its distribution intact: $BC(p, p) = 1$ — the mover’s internal matching is perfect. The advantage decomposes into three parts of independent force (ordered per author review). **(i) En-route immunity:** any real, finite channel loses; the mover loses nothing in transit. **(ii) Structural integrity:** messages move items piecemeal, and the receiver must re-integrate them; the mover’s items arrive *connected* — the links, where much of a corpus’s value lives, travel free — so even for perfectly articulable content, mobility beats any finite spectrum. **(iii) Tacit carriage:** the reconstruction residual c of D.2.4 moves only with its carrier. Cost arises solely at the target *interface*: what the receiving corpus can

integrate is priced by its absorption A , not by any channel. Only part (iii) is needed to beat the *idealized* $\Delta \rightarrow 0$ messaging limit (whose ceiling is $1 - c$); parts (i)–(ii) already establish dominance over every real channel, tacit content or none. Part (iii) itself must survive an objection (author, this round): *if the tacit part is untranslatable, is it not equally unadoptable at the target — making c effectively zero for practical purposes?* The resolution distinguishes **availability** from **adoption**. Mobility makes the tacit part *available* at the destination at fidelity 1: the mover is now a member of the target stratum, and its whole corpus — tacit residual included — participates directly in that stratum’s recombination (in model terms, it enters D_{eff} without any channel). Adoption *by others* there is a separate, slower process, and it is exactly the reconstruction of D.2.4 — but conducted through **shared practice** rather than messages, and co-presence supplies what projection cannot: demonstration, correction, accumulating shared context. So mobility does not transmit the untranslatable; it makes the untranslatable present, usable, and practice-regenerable — three things no message at any hop size achieves. The dominance claim survives with its meaning sharpened: over the ideal messaging limit, mobility’s edge is availability plus the practice-over-message reconstruction gap, not a magical full transfer of c .

D.4 Non-substitution (grounding A2)

Dice can surprise you, but they cannot tell you anything about the world; only windows can — and spawning copies of yourself adds no windows. Why can the leader not manufacture internally what the followers supply? Two bounds, stated in increasing strength of idealization. **(i) Deterministic:** the outputs of a deterministic process carry algorithmic information bounded by the process’s own description (Kolmogorov complexity — Li & Vitányi 2019; von Neumann’s (1966) complication threshold is the classical antecedent) — a corpus computing on itself cannot exceed itself plus logarithmic slack. **(ii) Randomized:** injected randomness escapes bound (i) — random strings are maximally complex — but manufactures complexity, not *reference*: by the data-processing inequality (Cover & Thomas 2006), no internal processing, however randomized, increases the system’s mutual information with the external world; world-referring information enters only through channels from the world. The model’s D is *viable* latent content — viability adjudicated by reality, not by the generator — so the quantity the leader needs is of the world-referring kind that bound (ii) bars it from synthesizing. What internal generation can do is recombine and decorrelate reasoning over information already held (§6, O1: computation, not observation); the followers’ non-substitutable contribution is their *position* — independent, embodied observation posts in parts of reality the leader does not occupy — and their tacit residual, which by D.2.4 and D.3 moves only with its carriers. The point compresses into a **freedom-effectiveness dilemma** (author, this review round): a spawned sub-agent’s effectiveness as a source of new information is increasing in its independence — control is specification, and specification confines the sub-agent’s trajectory to the parent’s field of view (the recognition bottleneck applied to development itself; in the model’s own terms, internal generation obeys eq. (7) with the parent’s internal freedom u_{int} as the throttle) — so matching the followers’ effectiveness requires granting followers’ independence, *which dissolves the reason for spawning*. Two honest riders bound the dilemma: full equivalence needs independent embodiment and time to diverge, not formal freedom alone; and spawning retains one residual advan-

tage — designed initial placement in content space — itself bounded by the parent’s recognition of what is worth covering. The dilemma then composes with §6 O1’s recursion: the same threat perception that walls the outside forbids granting the inside the freedom that effectiveness requires.

D.4’ Which results need which shape (author-raised 2026-07-04; corrected after author challenge 2026-07-05)

The dependency of the paper’s results on content shape deserves a map, and the map’s first version drew the boundary in the wrong place — at tail weight — until a direct computation, prompted by the author’s suspicion that the spectrum and translator benefits should survive exponential tails, showed that they do. The correct boundary is **smoothness at the edges**, not thinness at the tails.

What every smooth content family gives. For any smooth location family of content distributions, the small-displacement overlap obeys the standard expansion $BC \approx 1 - (I_F/8) \cdot d^2$, where I_F is the family’s Fisher information — so the loss exponent is *quadratic near zero*, $\ell(d) \approx (I_F/8)d^2$, and **A1’s second-order clause holds generically**: for Gaussian tails, for exponential (Laplace) tails, for power-law tails — for anything smooth. Verified directly for the Laplace family: the closed-form overlap is $BC = e^{\{-d/2b\}}(1 + d/2b)$, whose exponent is quadratic at small d (numerically: $\ell(0.1) = 0.00121$ vs quadratic prediction 0.00125), convex everywhere, and at a gap of 4 breadths gives jump $0.406 < \text{one mid-placed translator } 0.541 < \text{eight-hop spectrum } 0.807 \rightarrow 1 - c$. **The spectrum theorem, the translator benefit, and the mobility comparisons therefore survive exponential tails intact.**

What breaks A1. The second-order clause fails exactly where Fisher information diverges: **hard-edged content** — corpora with sharp support boundaries, knowing exactly this and nothing adjacent, stylized by uniform windows, for which the overlap is $1 - d/w$ and the exponent is *first-order* ($\ell \approx d/w$; verified: $\ell(0.05) = 0.0513$ vs linear prediction 0.0500). In a hard-edged world, every hop pays a per-distance toll no refinement removes, and the spectrum’s ceiling drops to $e^{\{-mg\}}(1 - c)$ — Remark A.3.3’s scenario. The microfoundational content of the spectrum theorem is thus: **shared context must shade off gradually at its edges**; how fast the far tails decay is irrelevant to it.

What the tails do govern. Tail weight sets the *large-jump* penalty — how catastrophic a big undivided gap is: quadratic growth of ℓ for Gaussian tails (jumps super-exponentially opaque), asymptotically linear for exponential tails (jumps exponentially opaque), logarithmic for power-law tails (jumps merely polynomially opaque — the “universal mind” regime). This affects the *quantitative* severity of the gap effect in P1 — which itself needs only that fidelity decay, and so survives every case — and one qualitative regime is worth recording: where hop granularity is bounded below (institutional levels cannot be graded arbitrarily finely) and ℓ is concave at large distances (heavy tails), a single long jump can beat a chain of medium hops. The case for graded spectra is, precisely stated, the case that content shades off smoothly and hops can be made fine.

(Code twin: Appendix E.13.)

D.5 Status of the stylizations

Three idealizations carry this appendix and are claimed as nothing more: content as low-dimensional smooth distributions — and here the requirement is now precise (D.4): *smoothness at the edges* (finite Fisher information), not any particular tail weight; the quadratic exemplar is the generic small-hop behavior of every smooth family, hard-edged content is the honest exception, and tail weight governs only how catastrophic large jumps are; fidelity as symmetric overlap (directional transfer would use an asymmetric divergence; the qualitative shape survives); and the message space as a fixed articulable subspace (c constant per distance; a richer treatment would grow articulability with shared context). The main text’s assumptions are labeled independently of all three; this appendix establishes only that they are the *generic outputs* of the simplest geometry, not artifacts of curve-fitting.

Appendix E — Code Twins

Every formula in this paper has a plain-Python twin: the same computation, written to be read. They are collected here, numbered in order of first reference from Appendices A and D; the supplementary code package repeats them alongside the full simulation scripts. Comments state what each line does in the paper’s terms; no library beyond the standard math module is needed.

E.1 — === code twins for A.1 (absorption geometry) ===

```
# === code twins for A.1 (absorption geometry) ===
from math import exp, sqrt
```

```
def A(x, r, kappa, M0, K0):
    # absorption: visibility * openness * restructuring (eq. 5)
    s = kappa * (1/M0 + 1/K0) # combined decay scale
    return (x / (x + r)) * exp(-s * x)
```

```
def x_star(r, kappa, M0, K0):
    # closed-form peak of A (the absorptive prime), A.1.1
    s = kappa * (1/M0 + 1/K0)
    return (-r + sqrt(r*r + 4*r/s)) / 2
```

```
# check against the four simulated Fig-1 peaks (r=2, M0=240, K0=360; combined scale 1/
# kappa: 0.5 -> 23.02 (sim 23.0) | 1 -> 16.00 (16.1) | 2 -> 11.04 (11.1) | 4 -
> 7.54 (7.5)
```

E.2 — === code twins for A.2 (P1 starvation bound; A.2.5 amplification robustness) ===

```
# === code twins for A.2 (P1 starvation bound; A.2.5 amplification robustness) ===
def growth_ceiling(theta_L, beta, phi, A_bar, sigma, rho, DF_bar):
    # limsup of leader growth under a defended gap (Theorem A.2.2):
    # ceiling = theta * Phi(phi * A_bar * DF_bar * (sigma/rho + 1)), Phi(D) = D**beta
    D_eff_ceiling = phi * A_bar * DF_bar * (sigma/rho + 1)
```

```

    return theta_L * D_eff_ceiling**beta          # = 0(phi**beta): starves as phi -
> 0

```

```

def A_bar_p(r, s, p):
    # sup over x of A(x) * x**p -- finite for every power p (Remark A.2.5):
    # the polynomial amplifier loses to the exponential in A; peak near x = p/s
    xs = [i/100 for i in range(1, 200000)]
    return max((x/(x+r)) * exp(-s*x) * x**p for x in xs)
# conclusion: replacing sigma by sigma * x**p only changes the CONSTANT, not the 0(phi**

```

E.3 — === code twins for A.3 (channel: superadditivity, spectrum limit, translator chains) ===

```

# === code twins for A.3 (channel: superadditivity, spectrum limit, translator chains)
from math import log

```

```

def ell(d, a=1.0):                # quadratic exemplar of the loss exponent
    return a * d * d

```

```

def jump_fidelity(g, a=1.0):      # one undivided jump                (eq. 3)
    return exp(-ell(g, a))

```

```

def spectrum_fidelity(g, delta, c=0.0, a=1.0):  # g/delta hops      (eq. 4)
    hops = g / delta
    return (1 - c) * exp(-hops * ell(delta, a))

```

```

# facts: ell(g) >= k * ell(g/k) for convex ell (A.3.1: k hops always beat one jump)
#          spectrum_fidelity(g, delta) -> 1 - c as delta -> 0      (A.3.2)

```

```

def translators_needed(a, d, F):
    # exact chain price (Corollary A.3.6): n >= a d^2 / ln(1/F)
    return a * d * d / log(1/F)      # ~ a d^2/(1-F) when F is near 1

```

E.4 — code twin

```

# code twin
def exit_dominates(rho_d, g_stay, g_conv, Lambda, B0=1.0):
    V_stay = B0/(rho_d - g_stay)
    V_exit = B0/(rho_d - g_conv) - Lambda
    return V_exit > V_stay          # true when Lambda < Lambda*
def Lambda_star(rho_d, g_stay, g_conv, B0=1.0):
    return B0*(1/(rho_d - g_conv) - 1/(rho_d - g_stay))  # rises as the gap starves hard

```

E.5 — === code twins for A.5 (P4: unconditional sign of dG/dD_F) ===

```

# === code twins for A.5 (P4: unconditional sign of dG/dD_F) ===
def dG_dDF(theta_L, Phi_prime, S, sigma, rho, delta):
    # comparative static of leader growth in follower diversity (A.5.2):
    # the depletion feedback cancels; every factor below is positive, so the sign
    # is + with NO side condition -- suppression backfires, cultivation pays (P4/P4b)
    return theta_L * Phi_prime * S * (sigma + rho) / (rho + delta * theta_L * Phi_prime)

```

E.6 — === code twins for A.8 (partition inequality; spread optimum; monolith depth) ===

```
# === code twins for A.8 (partition inequality; spread optimum; monolith depth) ===
def single_member_beats_monolith(X, n, r, s):
    # A.8.1(iii): one member of size X/n out-absorbs the whole monolith of size X
    # exactly when  $s*(X - X/n) > \ln((X/n + r)/(X/n))$  -- exponential advantage
    y = X / n
    return s*(X - y) > log((y + r)/y)

def monolith_depth_ratio(X, r, kappa, M0, K0):
    # Remark A.8.3: absorption penalty of the monolith vs one optimal-size unit
    s = kappa * (1/M0 + 1/K0); xs = x_star(r, kappa, M0, K0)
    return (xs*(X + r)) / ((xs + r)*X) * exp(s*(X - xs))
# Fig-2 params (X=400, kappa=2): ratio ~ 189 at optimal units; ~ 125 at pieces of 50
# per-unit-corpus absorption A(y)/y is STRICTLY decreasing in y (log-derivative
# = -1/(y+r) - s < 0): absorption alone favors atomization; pooling costs set n*
```

E.7 — code twin: the two dials

```
# code twin: the two dials
def unit_size_band(r, s, penalty_tolerance_log):
    y_min = 2*r # rough visibility floor (units must see arrivals)
    y_max = x_star_from(r, s) + penalty_tolerance_log/s # absorptive-decay ceiling
    return (y_min, y_max)
# diversity dial is SEPARATE: total inter-module conductance (bridge weights), not n.
# verified numerically: k=2/4/8 modules at equal bridge weight sustain 31/20/21 units
# dispersion; star (hub) sustains 1.3; complete graph 0.96 -- conductance governs, count
```

E.8 — === code twins for A.9 (Homogenization Theorem; Recognition-Bottleneck Lemma) ===

```
# === code twins for A.9 (Homogenization Theorem; Recognition-Bottleneck Lemma) ===
import numpy as np

def degroot_step(P, positions):
    # one round of content averaging on network P (row-stochastic): p <-
    P @ positions
    return P @ positions

def dispersion_decay_bound(lam2, t, V0):
    # Theorem part (b):  $V(t) \leq \text{lam2}^{**}(2t) * V0$  -- the spectral gap IS the one dial
    return (lam2**((2*t))) * V0

def stationary_dispersion(eigenvalues, s2, n):
    # part (d): with innovation noise s2, sustained diversity E V_inf
```

```

    # = (s2/n) * sum over non-unit eigenvalues of 1/(1 - lam_k^2)
    return (s2/n) * sum(1/(1 - lk*lk) for lk in eigenvalues[1:])
# complete graph: lam2 = 0 -> quenched to floor; thin-bridged modular: lam2 -
> 1 -> unbounded

def system_coverage_centralized(member_balls, allocator_ball):
    # Recognition-Bottleneck Lemma: coverage = union of (B_i intersect B_0) <= B_0
    return set().union(*[b & allocator_ball for b in member_balls]) # capped by B_0

```

E.9 — code twins (working-doc convention: every formula, once as math, once as code)

```

# code twins (working-doc convention: every formula, once as math, once as code)
from math import exp
def value_added(nu, lam):          # E(nu): single-peaked, argmax at nu == lam
    return nu * exp(-nu/lam)
def mean_novelty(x, L):            # nu_bar(x): pool novelty after x realized items
    return L / (2*x)
def mean_viability(x, L, lam):     # v_bar(x): rises toward 1 --
    pool gets "safer"
    return exp(-L/(2*x*lam))
def mean_value_added(x, L, lam):   # E_bar(x): falls -- pool gets "staler" (depletion)
    return mean_novelty(x, L) * mean_viability(x, L, lam)
def viable_latent_size(x, L, nu_max, rho): # |C*|(x), saturating
    return rho * min(L, 2*x*nu_max)

```

E.10 — code twin: the corpus quantities, now four

```

# code twin: the corpus quantities, now four
size          = x                  # item count (main text's variable)
breadth       = sigma_c            # spread of the content distribution
item_density  = x / sigma_c        # items per unit content
link_density  = kappa              # realized connections --
a SEPARATE layer (a graph),
                                # summarized in the main text by the coupling parameter
# metric distance = similarity (potential connection)
# graph distance = realized connection path ("indirect distance" of compounded closeness)

def kappa_exemplar(E, x):
    # one admissible operationalization of coupling: realized-link graph density
    # E = number of actual links among the x items; in [0, 1]
    return 2*E / (x*(x - 1))

```

E.11 — code twins

```

# code twins
def loss_exponent(d, sigma_c):     # ell(d) = a d^2, a = 1/(8 sigma_c^2)
    return d*d / (8*sigma_c**2)
def fidelity(d, sigma_c):          # single-transfer fidelity e^{-ell}

```

```

from math import exp
return exp(-loss_exponent(d, sigma_c))

```

E.12 — code twins

```

# code twins
from math import exp, sqrt
def bc(d, s1, s2):                                # overlap of  $N(0, s1^2)$ ,  $N(d, s2^2)$ 
    u = s1*s1 + s2*s2
    return sqrt(2*s1*s2/u) * exp(-d*d/(4*u))
def chain_fidelity(d, s_c, s_T):                  # cluster -> midpoint translator -
> cluster
    return bc(d/2, s_c, s_T) * bc(d/2, s_T, s_c)
def direct_fidelity(d, s_c):
    return exp(-d*d/(8*s_c*s_c))
# facts: chain_fidelity(d, s, s) > direct_fidelity(d, s)      --
(i) position alone helps
# chain_fidelity(4, 1, 0.3) < direct_fidelity(4, 1)          --
(ii) too narrow harms
# argmax_{s_T} chain_fidelity(4, 1, s_T) ~= 2.06             --
(iii) interior optimum

```

E.13 — code twins - plain style; ask about anything unclear

```

# code twins -- plain style; ask about anything unclear
from math import exp, log, sqrt

def bc_laplace(d, b):
    # Overlap of two exponential-tailed (Laplace) corpora at center distance d, breadth
    u = d / (2*b)
    return exp(-u) * (1 + u)          # closed form, checked against numerical integration

def bc_uniform(d, w):
    # Overlap of two hard-edged (uniform window) corpora, width w, shifted by d.
    return max(0.0, (w - d) / w)      # loss exponent -log(...) ~ d/w : FIRST-
order -> breaks A1

# demonstration at gap g = 4, breadth 1 (exponential-tailed world):
g = 4.0
jump      = bc_laplace(g, 1)          # 0.406  one undivided jump
translator = bc_laplace(g/2, 1) ** 2  # 0.541  one mid-
placed intermediary
spectrum   = bc_laplace(g/8, 1) ** 8  # 0.807  fine spectrum -
> approaches 1
# conclusion: spectrum > translator > jump even with exponential tails (A1 holds; smoot

```