

The Coordination Cost of Light-Speed Delay in a Self-Replicating Probe Swarm Scales as v/c

Noah Hydén

Independent researcher

ORCID: 0009-0003-4523-0467

Abstract

A swarm of self-replicating probes fills a stellar field star by star. The standard exploration-timescale model grants every probe perfect, instantaneous knowledge of which stars are already settled, and names finite light-speed as future work, which we add and measure. First, a caution: a coarse timestep batches many launches into one stale snapshot, and the collisions appear to slow the fill by tens of percent – a discretization artifact that vanishes once each probe decides at its true arrival time, leaving a few percent at directed-energy speeds, nothing at slingshot speeds, and the slingshot speed-up intact. The dominant cost is redundant travel: the two regimes build essentially the same number of probes, but stale views send them to already-claimed stars. A short collision argument, confirmed in simulation, shows the delay multiplies wasted journeys by at most $1 + v/c$, so the coordination tax – extra wasted trips relative to a perfect-information swarm – scales linearly with the probe’s speed in units of light, v/c , with a measured coefficient below that derived ceiling of one (saturation removes waste the delay would otherwise add) that itself shrinks with system size: about $0.8 v/c$ on a few-thousand-star field – a sixth at directed-energy speeds – but only about $0.07 v/c$ on our quarter-million-star headline field, a fuel tax of about 1.3 journeys are redundant even with a perfect map – so this is a v/c surcharge on an already-saturated process. The fuel tax falls with system size, from about a fifth at a few hundred stars to one part in eighty at a quarter million over a 1024-fold range – a genuine bulk decline that survives both an interior-only and a periodic-box control – while a secondary fill-time cost does the opposite: negligible beside the fuel tax on small fields, it comes to exceed it at scale (about 101.3 trips). The tax grows with the branching factor, holds under a clumpy field, and is no hard floor: a probe that listens in flight recovers essentially all of the wasted-arrival tax, at a fill-time cost of its own. We denominate the tax in journeys, since whether the waste also costs energy is mission-dependent.

Keywords: self-replication; interstellar exploration; light-speed delay; multi-agent coordination; swarm dynamics

1 Introduction

A self-replicating probe does not explore a galaxy alone. It settles a star, builds copies of itself from local material, and sends them on to further stars, so a single launch grows into a front that sweeps outward and fills the reachable field. That a machine can build a copy of itself at all is von Neumann’s result (Neumann and ed. A. W. Burks, 1966), and a probe of this kind saturating the Galaxy in a few million years is the sharpest form of the Fermi question: if interstellar settlement is possible it should already be complete, so the observed silence is a paradox (Hart, 1975; Tipler, 1980). How long the filling takes is the exploration-timescale question, and it has a long modelling lineage. Early work framed each interstellar probe as an independent agent acting on priors fixed before launch (Freitas, 1980), and treated

settlement as a front: a density-dependent diffusion with a travelling-wave solution (Newman and Sagan, 1981), or a Monte Carlo migration filling the Galaxy in tens of millions of years (Jones, 1981). Later simulations made the field concrete, sending probes and subprobes through a defined stellar volume to measure fill times directly (Bjork, 2007; Cotta and Morales, 2009), and established the seeded Monte Carlo methodology that situates an exploration timescale within a distribution of realizations rather than a single number (Forgan, 2009).

The sharpest treatment for self-replicators is the slingshot line. Forgan et al. (2013) showed that a single probe stealing energy from the motion of the stars it passes, a gravitational slingshot, explores up to two orders of magnitude faster than a probe under its own power. Nicholson and Forgan (2013) extended that manoeuvre to self-replicating probes and confirmed the speed-up survives replication, with the surprising result that aiming for the nearest star can beat aiming for the largest boost. The probe speeds that make this regime interesting, from the powered cruise of $3 \times 10^{-5} c$ up toward the 0.1 to 0.2 c of beamed directed-energy concepts (Lubin, 2016), are exactly where the next assumption starts to hurt.

That assumption, which Nicholson and Forgan name as a limitation and leave for future work, is that every probe is granted perfect, instantaneous knowledge of which stars are already settled. No probe ever flies to a star another has just claimed, because it always knows. In a real galaxy no such oracle exists: news that a star has been settled can travel no faster than light, so a probe choosing its next target acts on a picture of the field that is, for distant stars, years or centuries out of date. This is the elementary consequence of a finite signal speed, that events have only a partial ordering and a remote event is unknowable until its signal arrives (Lamport, 1978).

This paper adds that light-speed limit to the self-replicating slingshot model and measures what it costs. The answer, built from a controlled paired-ensemble experiment, has two parts, and the first is a caution. The cost is easy to overstate: a simulation that advances in a fixed timestep coarser than the interval between probe launches makes many probes decide from the same stale snapshot at once, and they collide far more than they physically would. Read that way the delay appears to slow the fill of the field by tens of percent. That number is an artifact. When the dynamics are resolved so that each probe decides at its own arrival time, the tens-of-percent slowdown collapses to a few percent at most, and to nothing at slingshot speeds, while the slingshot speed-up survives contact with finite information intact, reproducing the source model to the same order of magnitude. The dominant cost shows up elsewhere. Every star still gets settled, and with a negligible change in how many probes are built, well under a percent; what the delay chiefly changes is how many wasted journeys the swarm flies to stars other probes have already claimed. That is a cost paid in fuel and effort, and – as the field grows large – in fill time, which a perfect-information model hides entirely. It is governed cleanly by a single dimensionless number, the probe’s speed in units of the speed of light. The delay multiplies the wasted journeys by at most $1 + \Lambda$, so the fuel tax – the extra wasted trips as a fraction of those a perfect-information swarm already flies – is negligible for a powered cruise, about a percent at slingshot speeds, and a sixth at the directed-energy speeds that beamed-propulsion concepts target on a few-thousand-star field; on our quarter-million-star headline field that same coefficient has fallen to about 0.07, a fuel tax of about 1.3% at directed-energy speed. That baseline is itself wasteful: even with a perfect map about nine in ten journeys are redundant, so the tax is a v/c surcharge on an already-saturated process. We derive from a short collision argument that this relative tax has Λ as its leading-order ceiling, and find in simulation that the measured coefficient sits below it – about 0.8 on a few-thousand-star field, 0.07 at a quarter million, declining with size as saturation removes waste the delay would otherwise add. The fill-time cost runs the other way, negligible beside the fuel tax on small fields but exceeding it at scale. We show the tax grows with the replication branching factor without saturating up to sixteen offspring and survives a clumpy (non-uniform) field, and – because a probe deciding only at its stops is pessimistic – that a probe which listens in flight recovers almost all of it,

bounding how much coordination can ever buy back. The cost is dominated by fuel and effort, with a secondary fill-time penalty that emerges only at the highest speeds; and it scales, cleanly, with how fast the swarm flies.

2 The light-delayed belief model

2.1 The knowledge gate

The change to the model is one gate on knowledge. A settled star is treated as an omnidirectional beacon that begins emitting in the year it is settled. A probe deciding its next target at star f in year Y believes a distant star i is settled only once that beacon’s light has had time to arrive:

$$\text{settled}(i) + \frac{\text{dist}(f, i)}{c} \leq Y. \quad (1)$$

Under the perfect-information assumption of the source model, (1) collapses to “settled at any non-negative year,” and the model reproduces our perfect-information baseline bit-for-bit; the light-speed term is what introduces stale views. When two probes both believe an unsettled star is the best target they both fly to it; the loser arrives to find it taken, its trip wasted, and must re-target from the even more stale picture it holds on arrival. This is availability chosen over consistency: rather than wait for a single agreed map that finite light-speed makes impossible, each probe acts now on its local view (Lamport, 1978). Stars carry velocity magnitudes that feed the slingshot boost but are held fixed in position, following the source model, so $\text{dist}(f, i)$ in (1) is static and the beacon geometry does not shift under the probe. Physical arrival always reads ground truth: what a probe finds on arrival is the star’s real status, whatever it believed when it aimed.

2.2 The three flight policies, and the speed axis

We carry three named policies, following the slingshot studies we extend (Forgan et al., 2013; Nicholson and Forgan, 2013). *Powered nearest-neighbour* flight moves under the probe’s own power to the nearest unsettled star. *Nearest-star slingshot* flight takes the same nearest target but gains speed from gravitational slingshots off the stars it passes. *Maximum-boost slingshot* flight instead selects, among the nearest believed-unsettled stars, the target that yields the largest velocity gain. The gravitational boost self-limits (its per-encounter gain falls off once the probe outruns the stellar escape speed), so both slingshot policies saturate at a mean speed of a few thousand km/s, about 10^{-2} of c . To reach the higher speeds that beamed directed-energy concepts target, 0.1 to 0.2 c (Lubin, 2016), a probe must be driven, not slung; we therefore sweep the powered cruise speed as a clean, direct control on v/c , from the source model’s $3 \times 10^{-5} c$ up to 0.2 c . Powered flight thus spans the whole speed axis, while the two slingshot policies sit near $v/c \approx 0.01$.

2.3 Event-driven resolution, and why it is necessary

The fold can advance either by a fixed timestep or by jumping to the next probe arrival. This is not a mere numerical convenience, and getting it wrong is the single largest source of error in this problem. A fixed timestep Δt processes together every probe that arrives within a window, so all of them decide against the same snapshot of the settled map. When Δt exceeds the interval between launches – which it does badly in the boosted regime, where a probe crosses a parsec in centuries against a source-model step of thousands of years – the window batches many independent decisions into one synchronized burst, and they collide far more often than they would if each decided at its true arrival time. The measured coordination cost then reflects the size of the timestep, not the physics. We therefore run the model event-driven throughout:

the fold jumps to the earliest outstanding arrival, advances the clock to it, and processes exactly the probes arriving then. This is the exact $\Delta t \rightarrow 0$ limit, dt-independent by construction, and it is what every number in this paper is computed on. Section 3 shows directly that the apparent fill-time tax is an artifact of a coarse step and collapses to zero as the step resolves.

2.4 The governing ratio

How much the gate matters is governed by one dimensionless number: the information delay across a hop divided by the travel time across it. For a hop of length d the delay is d/c and the travel time is d/v , so the ratio is

$$\Lambda = \frac{d/c}{d/v} = \frac{v}{c}, \quad (2)$$

the probe's speed in units of c , independent of the hop length. A slow probe ($\Lambda \ll 1$) outpaces its own ignorance only trivially, so stale views barely arise; a fast probe carries a large Λ and routinely decides before the relevant news can reach it. At the powered cruise speed $\Lambda \approx 3 \times 10^{-5}$ is negligible.

2.5 A collision model: why the tax should equal Λ

The same ratio predicts the size of the redundant-travel tax, not just its trend, from a short exposure argument. Fix a probe that has committed to a target star at distance d . Under perfect information it chose a truly unsettled star, so it is wasted only if some other probe claims that star during its flight – an exposure window equal to the travel time d/v . Under the light-speed gate the probe carries an *additional* blindness: the target may already have been claimed up to a light-time d/c before launch, with the beacon not yet arrived, and the probe would have launched anyway. Its total exposure to an unseen claim is therefore $d/v + d/c$. If claims of a given candidate arrive at some roughly steady local rate, the probability of a wasted trip is proportional to the exposure window, so the ratio of wasted trips with and without the gate is

$$\frac{p_{\text{lag}}}{p_{\text{perfect}}} \approx \frac{d/v + d/c}{d/v} = 1 + \frac{v}{c} = 1 + \Lambda, \quad (3)$$

and the coordination tax – the extra wasted trips as a fraction of the perfect-information baseline – is just $p_{\text{lag}}/p_{\text{perfect}} - 1 \approx \Lambda$. Two features of (3) are worth stating because the simulation independently reproduces both: the hop length d cancels, so the tax cannot depend on how far apart the stars are, and the local claim rate cancels, so it cannot depend on the stellar density. This is a leading-order estimate – it assumes claims of a candidate are independent and locally uniform in time, and it ignores both the geometry of which stars are in contention and the saturation of already-doomed hops – so it predicts the coefficient of proportionality is one, as a ceiling that those corrections can only lower. Section 3 shows the measured tax follows Λ in *form* but with a coefficient below that ceiling: a through-origin slope of 0.069 on the $N = 262,144$ headline field, resolvably under one, the endpoint of a decline with system size as saturation strengthens (from $a \approx 0.97$ at a few hundred stars down to this value at a quarter-million). The simulation confirms the d - and density-independence of the form directly.

2.6 Experimental design and parameters

The field is a box of stars at a uniform density of one star per cubic parsec, the convention of the source model. Each settled star launches offspring (the replication branching factor): branching is what makes two probes ever contend for a star, so it is the precondition for any coordination effect at all and the parameter most likely to change the answer, so we take two offspring as the default and sweep it explicitly (Section 3.4). Three properties keep the measurement honest. First, knowledge is evaluated only at the decision star at the moment of decision; news a probe's

path happens to cross in mid-flight is ignored, which deliberately undercounts what a probe could know and so makes the measured waste a conservative upper bound – we measure the opposite, optimistic bound (a probe that listens in flight) separately in Section 3.5. Second, a probe that loses a race re-targets, but only up to a fixed number of times before it is retired; this cap is applied identically under both coordination regimes, so the paired comparison is fair. On the few-hundred-star fields of earlier work the tax plateaued above about $\text{cap} = 8$, which is why we adopt $\text{cap} = 8$ as the default – but that plateau does not survive to larger fields (Section 3.8): at $\Lambda = 0.2$ and $N = 262,144$ the fuel tax is 0.27, 0.27, 1.31, 4.79, 8.97% for caps of 2, 4, 8, 16, 32, still climbing with no sign of levelling. The default cap of 8 is therefore a documented *lower bound*, not a converged value, and every fuel-tax figure it produces is an underestimate whose size grows with the field (negligible at small N , substantial at scale). Third, the whole simulation is a pure, seeded, event-driven fold: fixing the seed fixes the star field, every arrival, and every target choice, so a run with the light-speed gate and a run without it can be compared on the *identical* galaxy. That paired design attributes a difference to the delay alone rather than to the luck of the star layout.

Two dependencies are worth stating because they are easy to over-read. Changing the value of the uniform density rescales every inter-star distance by a common factor, multiplying travel time and light-travel time on every hop equally, so Λ and the relative penalties are unchanged: this is a scale-free property of a uniform field, exactly what (3) predicts when the hop length cancels, so confirming it holds to the printed digit from 0.14 to 5 stars per cubic parsec (RECONS and Gaia, 2018) is a consistency check on the derivation, not an independent robustness result. The case that would genuinely test density sensitivity is a clumpy, non-uniform field, which we examine directly in Section 3.7; what we do not probe is a full galactic disk with a large-scale radial density gradient. And because the field always fills under both regimes and each settlement launches a fixed number of offspring, the two regimes build essentially the same number of probes – they differ only in a handful of terminal launches as targets run out, well under a percent – so the resource the delay actually moves is the number of journeys flown, which is why the coordination cost is a travel cost.

The ensemble size varies by measurement, sized to each one’s event-mode cost; Table 1 lists them. Every run is event-driven and paired – each seed fills the identical field under both coordination regimes – and the seed counts are prefixes of one fixed 64-seed set, so the sweeps share seeds and differ only in the stated star count.

2.7 Baseline validation

By construction the “instant” regime is the $c \rightarrow \infty$ limit of (1), and it reproduces the plain perfect-information fold bit-for-bit. Against Nicholson and Forgan’s published model (Nicholson and Forgan, 2013) the event-driven fold reproduces both headline findings: gravitational slingshots explore about two orders of magnitude faster than powered flight, and their surprising result that aiming for the nearest star beats aiming for the largest boost on total exploration time holds, while the maximum-boost policy reaches the higher peak speed. The speed-up agrees to the same order of magnitude but not to the digit, as one should expect: we measure a factor of about 166 on a 400-star field, against their factor of about 100 (a reduction from $\sim 10^8$ to $\sim 10^6$ years) on a 200,000-star field, so with a field some 500 times smaller the exact multiple is not expected to match. A coarse fixed step had quantized the boosted hops and undercut even this to about twenty-fold; resolving the step recovers the full factor, which is the same correction that removes the spurious fill-time tax.

Table 1: The paired ensemble used by each measurement. All are event-driven and paired (each seed filled under both coordination regimes on the identical field); seed counts are prefixes of one fixed 512-seed set, sized to each measurement’s cost. Star counts are powers of two computed by the `run_fill_flat` fast path (Issue #38); the headline fixed- N sweeps now share the top of the ladder, $N = 262,144$, so every claim is read at the same maximal field (the clumpy-field sweep is the exception, kept on its historical $N = 500$ ensemble, its scale companion having no power-of-two rerun). The lower block carries each fixed- N sweep one power of two further, to $N = 524,288$ (the clumpy-field companion at $N = 200,000$ and the re-target-cap cross-check at $N = 262,144$ excepted), using the finite-size arm’s precision-target seed count, so every claim rests on a scale companion beyond its headline result – discussed together in Section 3.8.

Measurement	Seeds	Stars	Reported in
Speed sweep (headline)	128	262,144	Table 2, Fig. 3
Branching sweep	32	262,144	Fig. 5
Re-target-cap sweep	32	262,144	Section 3
In-flight floor	32	262,144	Fig. 6
Concurrency	32	262,144	Fig. 2
Finite-size sweep	32 down to 8	256–262,144	Fig. 7
Finite-size, interior edge test	32 down to 8	256–262,144	Section 3.6
Finite-size, periodic control	32 down to 8	256–262,144	Section 3.6
Clumpy-field sweep	48	500	Fig. 8
Timestep artifact	32	262,144	Fig. 1
<i>Scale companions, precision-target seed count (Section 3.8)</i>			
Speed sweep, at scale	8	524,288	Section 3.8
Branching, at scale	6	524,288	Section 3.8
Re-target-cap, at scale	8	262,144	Section 3.8
In-flight floor, at scale	8	524,288	Section 3.8
Concurrency, at scale	8	524,288	Section 3.8
Clumpy-field, at scale	8	200,000	Section 3.8

3 Results: the coordination tax

We measured the cost as a paired ensemble: 128 seeded galaxies of 262,144 stars, each filled twice on the *identical* field, once with perfect information and once with the light-speed gate of (1), all at the resolved event-driven timestep. Reporting the distribution across the ensemble follows the seeded Monte Carlo practice of the field (Forgan, 2009); we give the median with a seeded percentile-bootstrap 95% confidence interval. A stale view can only make a probe believe a settled star is still open, never the reverse, so the delay adds wasted journeys *in expectation*; but the sign is statistical, not structural, because retarget cascades make the paired difference non-monotone from seed to seed (at $\Lambda = 0.01$ a sizable minority of galaxies, 41 of 128, come out slightly negative, since at scale the $\Lambda = 0.01$ tax is only about 0.07%). We therefore report the median and its interval, not a sign test, as the finding – though a two-sided sign test, which the pipeline computes, agrees where the effect is large (87 of 128 seeds positive at $\Lambda = 0.01$, a bare majority for a near-zero effect, rising to all 128 at $\Lambda = 0.2$). What is informative is the magnitude and its scaling, which the intervals and sweeps below report. (The direction is itself regime-specific: the in-flight-relay variant of Section 3.5 can *lower* waste below the perfect-information baseline, by aborting doomed hops that perfect knowledge at the stops would not.) For the one genuinely null claim – no fill-time tax at low speed – we do not wave the test away but give an equivalence bound instead (below).

3.1 No fill-time tax at slingshot speed; the coarse-timestep one is an artifact

Light-speed lag does not measurably slow the fill of the field. Fig. 1 shows why the literature intuition says otherwise. Run with the source model’s fixed timestep the fill-100% penalty for nearest-star slingshot flight is a median of 26.1%, positive in all 32 seeds – a large, robust-looking tax. It is an artifact of the step. As the timestep is refined the penalty holds near 25% and then falls steeply – 25.4%, 22.4%, 18.6%, 12.2% at 2000, 1000, 500, 250 years – and at the event-driven limit it is a median of 0.2%, two orders of magnitude below the coarse-step value and at the edge of resolution. The coarse step had batched many launches into one window so they decided from the same stale snapshot and collided; resolving the step lets each probe decide at its own arrival time, and the collisions that cost time disappear. The same refinement restores the slingshot speed-up to its full value: against powered flight the event-driven fold reproduces Nicholson and Forgan’s roughly two-orders-of-magnitude figure to the same order (about $166\times$ on a 400-star field, against their factor of about 100 on a much larger field; Section 2), where the coarse step had reported only about $20\times$. The celebrated speed-up survives contact with finite information; the fill-time tax does not exist.

3.2 At scale the fuel cost shrinks below the fill-time cost

Fuel and time are hit unequally because the fill rate is carried by a large parallel population, and a wasted trip removes only one member of it. Fig. 2 tracks the number of probes in flight as the field fills: exponential branching keeps a median of about 91,000 probes aloft at the peak of a 262,144-star fill, and the light-delayed swarm carries essentially the same population as the perfect-information one (peaks of about 84,600 against 91,400), because both build the same probes. The front advances at the aggregate rate of tens of thousands of probes settling stars, and losing a fraction of them to wasted trips lowers that rate only in proportion, damped by a redundancy that only deepens as the field, and the population, grow. That is why the *fuel* tax falls with system size (Section 3.6) to about 1.3% here, an order of magnitude below its small-field value. The *fill-time* tax does not follow it down. At $\Lambda = 0.2$ it holds near four percent through the bulk of the fill – a median of 4.3% at half-fill, 4.2% at 90%, and 4.3% at 99% settled – and then spikes to 10.5% at the very last star, because completing the field waits on light-delayed news crossing the now quarter-million-star span, a tail lag the parallel population cannot absorb. So at this scale the small-field ordering is *reversed*: the fill-time tax – a few percent through the fill, 10.5% at completion – now exceeds the fuel tax of 1.3% rather than being a fraction of it. The large population dilutes the fuel waste but not the front lag; as the field grows the first shrinks while the second does not, so the cost, fuel-dominated on the few-thousand-star fields of earlier work, becomes fill-time-inclusive at scale. At low Λ , where few trips are wasted and the front barely lags, both taxes fall back toward zero.

3.3 The real cost is redundant travel, and it scales as Λ

What the delay changes is how far the swarm flies. Because the field always fills and each settlement launches a fixed number of offspring, both coordination regimes build essentially the same number of probes – differing by well under a percent, in a handful of terminal launches – so the resource the delay moves is journeys, and stale views add wasted ones. To read the numbers that follow it helps to fix the baseline first, because it is not small: even with a perfect, instantaneous map the swarm already wastes most of its journeys, since probes independently pick the same open star and all but the first arrive to find it taken. At the directed-energy configuration ($\Lambda = 0.2$, $N = 262,144$) that perfect-information waste is a median of about 90% of all arrivals. The coordination tax is the *extra* waste the delay adds on top of that already-wasteful baseline, expressed as a percent of it. Table 2 and Fig. 3 give this redundant-travel tax for powered flight swept across the speed axis: a median of 0.07% more wasted journeys than

perfect information at $\Lambda = 0.01$, rising to 1.35% [1.29, 1.41] at $\Lambda = 0.2$, with the confidence interval clear of zero at every speed and the per-seed spread widening with Λ (Fig. 4). This 128-seed, 262,144-star ensemble is our canonical configuration for the $\Lambda = 0.2$ tax, the largest field we run. The branching, re-target-cap, in-flight-floor, concurrency and timestep sweeps run on the *same* $N = 262,144$ field (Table 1), so the branching base case, for one, reproduces the headline (1.3% at two offspring), while the scale companions of Section 3.8 carry each claim one power of two further, to $N = 524,288$. About 1.3% more wasted trips is what the headline means – against a 90% baseline it lifts the wasted fraction only from about 90.0% to about 90.1% of all journeys, a v/c surcharge on an already-saturated process. (We count wasted arrivals; Fig. 6 instead measures redundant travel in parsecs, and the two track each other because the hops are near-uniform in length, so a count ratio and a distance ratio agree.)

These points sit close to the line the collision model predicts, but consistently *below* it, by a margin that grows with speed: the raw ratio of wasted trips $p_{\text{lag}}/p_{\text{perfect}}$ is 1.0007, 1.0025, 1.0041, 1.0074, 1.0135 at $\Lambda = 0.01, 0.03, 0.05, 0.1, 0.2$, against 1.010, 1.030, 1.050, 1.100, 1.200 predicted by (3) – so at this scale the tax runs far below the ceiling, about a fifteenth of it, the shortfall widening with speed. A through-origin fit of $\text{tax} = a \Lambda$ gives $a = 0.069$ on this $N = 262,144$ field, and the 128-seed intervals are tight enough to place it resolvably below the derived value of one. So the derivation’s $1 + \Lambda$ is an *upper bound* on the tax, not an equality: the linear-in- Λ *form* holds cleanly – the ratio tracks $1 + \Lambda$ with a fixed proportionality – but the coefficient sits below one because saturation removes waste the delay would otherwise add (a hop already likely to be scooped leaves the delay little room; Section 3.7), together with the non-stationary front (the extra pre-launch window d/c falls when the region was less contended than during the flight itself). Neither changes the law’s form, and both act downward, which is why the residual is a systematic shortfall rather than scatter. Crucially the coefficient is not a constant of the model but declines with field size, because a larger field saturates more: it falls from $a \approx 0.97$ on the few-hundred-star field of the clumpy-field sweep (Section 3.7) to 0.069 here at $N = 262,144$, and further to 0.040 on the $N = 524,288$ scale companion (Section 3.8). What the derivation fixes is the *form* and its ceiling of one; the measured coefficient is what saturation leaves of that ceiling at a given size, and the size dependence of that residual is exactly what the finite-size sweep of Section 3.6 maps. The same cancellation makes the *form* independent of hop length and density (both drop out of (3)), which Section 3.6, the density check, and the clumpy-field test (Section 3.7) confirm.

A fill-time tax rides alongside the fuel tax, and at this scale it is the *larger* of the two. At the slowest speed it is consistent with zero – a median of 1.5% with a bootstrap interval of $[-2.8, 3.1]$ spanning zero at $\Lambda = 0.01$, so we do not resolve a fill-time tax there; the completion-time tail is noisy enough at a quarter-million stars that this is an equivalence bound of a few percent rather than the razor-sharp null a smaller, higher-seed field would give. It resolves by $\Lambda = 0.05$ (3.8% [0.6, 5.8]), is unambiguous by $\Lambda = 0.1$, and reaches a median of 10.5% [8.9, 13.3] at $\Lambda = 0.2$ (Table 2) – roughly eight times the 1.35% fuel tax at the same speed. This is the reversal Section 3.2 explains: the parallel population dilutes the fuel waste as the field grows but not the front lag, so the cost, fuel-dominated on the few-thousand-star fields of earlier work, becomes fill-time-inclusive at galactic scale.

3.4 Contention grows with the branching factor, without saturating

The tax exists only because probes contend for stars, and contention is created by replication: each settlement launches offspring, and it is those offspring that can race. The branching factor is therefore the parameter most likely to change the answer, so we sweep it from two to sixteen offspring per settlement, at $N = 262,144$ (Fig. 5). At moderate speed ($\Lambda = 0.05$) the tax climbs gently, from about 0.4% at two offspring to 3.5% at sixteen. At directed-energy speed ($\Lambda = 0.2$) it grows more steeply and, importantly, does *not* saturate: a median of 1.3%, 3.7%, 4.8%, 7.0% and 8.8% at two, three, four, eight and sixteen offspring, with the intervals at eight and sixteen

Table 2: Coordination tax over the 128-seed paired ensemble at $N = 262,144$, event-driven timestep, powered flight, as a function of $\Lambda = v/c$: the redundant-travel (fuel) tax as a percent of the wasted journeys under perfect information, and the fill-time (completion) tax, each as a median with a bootstrap 95% confidence interval. At this scale the fuel tax has saturated to about a percent while the fill-time tax, set by the front lag across the field, is the larger cost.

$\Lambda = v/c$	Fuel tax (%)	Fill-time tax (%)
0.01	0.07 [0.04, 0.13]	1.5 [-2.8, 3.1]
0.03	0.25 [0.21, 0.29]	1.6 [-0.5, 4.4]
0.05	0.41 [0.37, 0.44]	3.8 [0.6, 5.8]
0.10	0.74 [0.70, 0.78]	4.5 [2.2, 7.6]
0.20	1.35 [1.29, 1.41]	10.5 [8.9, 13.3]

disjoint from those below. A sweep that stopped at four offspring, seeing the tax climb only to 4.8%, could have read the approach as levelling; carrying it to eight and sixteen shows it keeps rising, roughly linearly in the logarithm of the branching factor, because more offspring make more simultaneous races and the marginal collision rate does not level off over the range we can reach. Two consequences follow. First, the default of two offspring is a genuine lower bound, not a plateau: a factory that branches harder pays a larger coordination tax, about seven times the two-offspring value at sixteen offspring. Second, the growth is in the tax’s *coefficient*, not its scaling: at every branching factor the tax stays a modest fraction of Λ – the whole $\Lambda = 0.2$ column spans about 1 to 9%, all well below Λ – so the v/c law holds at fixed branching, and the branching factor sets how large its coefficient is.

3.5 The physical floor: how much survives relaying in flight

The tax reported so far is a conservative upper bound, because a probe here learns only at its decision stars and ignores news its path crosses in flight. The opposite extreme is a probe that listens continuously: the moment the beacon from a now-claimed target overtakes it – at the closed-form time set by the beacon closing on the probe at c while the probe closes on the target at v – it abandons the doomed hop and re-aims from wherever it then is. This “in-flight relay” mode is the optimistic bound on what any relay protocol can recover, and we measure it directly. Its effect is dramatic. Because a beacon travelling at c always overtakes a probe travelling at $v < c$ before the probe arrives, a listening probe essentially never completes a wasted arrival: the count of wasted arrivals falls from about 2.4 million under decision-site knowledge to zero at $\Lambda = 0.2$ (on this $N = 262,144$ field). Redundant travel falls too, and below even the perfect-information baseline – from about 4.9 million pc under the light-speed gate, and 5.2 million pc under perfect information, to about 410,000 pc – because listening in flight also lets a probe bail out of an “assignment collision” (two probes independently choosing the same open star) that perfect knowledge at the decision site does not prevent. The recovery is not free, and the price is larger than it looked at small fields: the mid-flight detours lengthen the fill by about 27% at $\Lambda = 0.2$. So the decision-site tax is an upper bound, not a floor: a swarm that relays in flight can drive completed wasted arrivals to zero and cut redundant travel below even the perfect-information baseline, but at a real cost in fill time of order a quarter at scale. This is an information ceiling rather than an engineering estimate – the beacon is idealized as isotropic, unattenuated, and detectable the instant its light arrives, so a real relay working against a finite link budget would recover less. What it establishes is that the decision-site tax is nowhere near a physical floor: most of it is protocol, not physics. What genuinely remains is the partial travel a probe flies toward a doomed star before the beacon reaches it, which grows with Λ as the probe covers more of the hop before the relatively slower news catches up.

3.6 Does the tax hold at scale?

A tax measured on a few-hundred-star field is only interesting if it survives to larger ones, so we sweep the system size at $\Lambda = 0.2$ over a range of about $1024\times$, from 256 to 262,144 stars, with 8 to 32 seeds per size (fewer at the largest fields, where the per-seed spread is tight; Fig. 7). This reach is what lets the question be answered rather than gestured at: it opens up once the nearest-not-yet-settled query – the one non-local step in the fold – is served by a dynamic spatial index over the unsettled set that skips already-settled space, which makes a single run near-linear in N rather than the quadratic cost that previously capped the sweep at a sixteen-fold range. The absolute count of wasted journeys grows with the field, as it must, from a median of about 190 extra trips at 256 stars to about 30,000 at 262,144. As a fraction of the perfect-information waste the tax does the opposite, and does so decisively: it holds near 19 to 21% up to about a thousand stars, then falls to 17.1% [14.7, 18.7] at 2,048 stars, 14.2% [13.2, 16.0] at 4,096, 11.1% [9.9, 12.7] at 8,192, 8.7% [7.0, 9.3] at 16,384, 5.7% [5.2, 6.7] at 32,768, 4.6% [3.2, 5.0] at 65,536, 2.1% [1.8, 3.0] at 131,072, and 1.3% [1.2, 1.7] at 262,144. A regression of the median tax on $\log_{10} N$ across all eleven sizes gives a slope of -7.1 percentage points per decade of N , with a seeded-bootstrap 95% interval of $[-7.8, -6.4]$ well clear of zero (the fit and its interval are computed in the measurement pipeline alongside the per- N medians, so they regenerate from source), and the confidence intervals at the ends of the range are separated by an order of magnitude. The decline is convex and monotone – gentle to a thousand stars, then steepening through the mid-range – so the single slope is a conservative local summary of a fall that carries the tax from roughly a fifth of the perfect-information waste to about one part in eighty.

Is that decline physical, or an artifact of the hard wall of a finite box? Edge stars have fewer neighbours and so draw less contention, and the edge fraction falls as $N^{-1/3}$, so in principle the fraction could shrink with N for a purely geometric reason. We test it two ways. First, tagging each contested arrival by its target’s distance to the nearest wall and recomputing the tax on interior stars only does shift the level up – edge stars carry a lower tax and dilute the all-stars value downward, so the interior tax sits at or above the all-stars value, clearly so once the field is large enough for the edge fraction to bite: 7.4% [5.9, 9.3] against 5.7% [5.2, 6.7] at 32,768 stars, and 2.8% [2.7, 3.4] against 1.3% [1.2, 1.7] at 262,144 (this masking test runs on the lighter interior-edge sub-ensemble of Table 1); at the smallest fields, where the edge fraction is largest but the tax is high everywhere, the two are within noise of each other. Crucially the interior tax declines just as steeply – a bulk-only slope of -5.6 $[-7.1, -3.6]$ percentage points per decade at the one-nearest-neighbour shell – so the shrinkage is not carried by the edge. The decisive control is a periodic (toroidal) box, which removes the boundary altogether: every star then sits in a full neighbourhood. Rerun that way the decline survives essentially unchanged in trend – a slope of -7.0 $[-7.7, -6.2]$ percentage points per decade, statistically indistinguishable from the hard-walled -6.9 $[-7.3, -6.3]$ – while the level sits a point or two higher at every size (20.1% against 18.8% at 256 stars, 1.5% against 1.3% at 262,144), exactly the small upward shift the edge dilution predicts. So the shrinkage is a genuine, bulk effect, not a boundary artifact: the wall moves the level, not the trend, and the effect holds over a $1024\times$ span with the boundary present, masked, and removed.

The direction is what matters for the one extrapolation a reader might be tempted to make: the fraction does not grow with scale – it genuinely shrinks, on a hard-walled box, an interior-only sub-field, and a periodic box alike – so nothing here supports carrying a fixed-percentage tax out to a 10^{11} -star galaxy. A $1024\times$ lever arm is far longer than we had, but it is still five orders of magnitude short of a real galaxy, so we report the scale behaviour over the range we can measure and make no quantitative claim beyond it.

3.7 Does the law survive a clumpy field?

The uniform field is the derivation’s friendliest case: the collision argument assumes claims of a candidate arrive at a locally uniform rate, which is what lets the hop length d cancel. One might reasonably worry that a clumpy, non-uniform field breaks this – there hop length and local claim rate correlate, since a long hop tends to cross a void between dense clumps, so the d -cancellation could fail and change the law’s *form* rather than merely rescale it. We test it directly. We place the stars with a Thomas cluster process – 25 cluster centres, each star scattered a Gaussian σ around its centre and reflected into the box, so the star count and the mean density are held exactly fixed and only the spatial arrangement changes – and sweep the scatter σ from the uniform limit to strong clustering, crossed with Λ . We report clumpiness by the measured Clark-Evans index R (observed mean nearest-neighbour distance over the Poisson expectation: $R = 1$ uniform, $R < 1$ clustered) (Clark and Evans, 1954), and fit the through-origin slope a of $\text{tax} = a \Lambda$ at each level (Fig. 8).

The law’s *form* survives – the tax stays linear in Λ throughout – and it is only the *coefficient* that moves, in the safe direction. At the uniform limit the slope is $a = 0.97$ [0.89, 1.19] on this $N = 500$ field – the largest coefficient we measure anywhere, and the small-field anchor of the size decline of Section 3.6, along which the coefficient falls to 0.73 [0.69, 0.78] at a power-of-two $N = 4,096$ rerun of this same clumpy-field generator and on to 0.069 at the $N = 262,144$ headline. (We keep the clumpiness narrative on the historical $N = 500$ ensemble because its scale companion has no power-of-two rerun.) At $N = 500$ the slope stays statistically consistent with one through moderate clustering ($a = 0.90, 0.86, 0.88$ at $R = 1.03, 0.97, 0.78$), and drops resolvably – the interval disjoint from the uniform one – only at extreme substructure ($a = 0.51$ [0.41, 0.70] at $R = 0.56$, more clumped than a dynamically relaxed stellar field would plausibly be, where the coefficient is roughly halved). Every deviation is *downward*: clumpiness makes the coordination tax smaller, never larger, so $\text{tax} = \Lambda$ is a conservative upper bound rather than a fragile equality. What this sweep does not probe is a full galactic disk with a large-scale radial density gradient, where hop length and claim rate would correlate over the disk’s scale rather than within a cluster; that is the specific open case, left for a disk-field slice.

The mechanism is saturation, and it explains both why the coefficient softens under clustering and why even the uniform fit sits a touch below one. Because v and c are global constants, the factor $1 + v/c$ pulls out of both the with-lag and the perfect-information exposure sums identically, so the correlation between d and the local claim rate cancels *exactly* in the linear regime – clumpiness cannot break the v/c form through that correlation. What it can do is push more hops into saturation: once a hop’s chance of being scooped already approaches one, the delay has little room to add waste, so the per-hop ratio $p_{\text{lag}}/p_{\text{perfect}}$ falls below $1 + \Lambda$. Stratifying that ratio by hop length confirms it – in both the uniform and the clumpy field the ratio declines from about 1.08 at short hops to about 1.00 at long ones, tracking the baseline per-hop waste probability as it climbs toward one, with the same shape in both. So clumpiness rescales the coefficient through saturation; it does not introduce a new hop-length dependence. This is also why even the $N = 500$ uniform fit gives 0.97 and not a clean 1.0: even a uniform field has long, already-doomed hops, and larger fields have more of them – which is the same saturation that drives the coefficient down with N .

3.8 The scale companions: what else survives at scale

The finite-size sweep (Section 3.6) settles one question at scale – the tax as a fraction of effort shrinks with N and does not extrapolate flat – but says nothing about whether each fixed- N subsidiary result also survives that reach. The dynamic spatial index that lets the finite-size sweep run to 262,144 stars also lets the other fixed- N sweeps of Table 1 be repeated there at their precision-target seed count. The picture is uniform: the qualitative shape of every result holds, and the redundant-travel tax rescales in step with the finite-size decline. The one nuance

is bookkeeping-side and does not touch a physical claim.

Linearity in Λ persists (Speed sweep beyond the headline). Pushing one power of two past the headline, to $N = 524,288$, the same Λ ladder returns fuel taxes of 0.07%, 0.19%, 0.26%, 0.50%, and 0.79% at $\Lambda = 0.01, 0.03, 0.05, 0.1, 0.2$; the $\Lambda = 0.01$ point is noise-consistent with zero and the rest track linearly, giving a through-origin coefficient of $a \approx 0.040$ (in the same fractional units as the $a = 0.97$ at $N = 500$), down from 0.069 at the headline. The decline in a tracks the fuel tax’s own finite-size shrinkage, so the linearity is preserved and only the slope rescales; the collision derivation remains the correct first-order description at scale, with the coefficient the quantity that carries the size dependence.

In-flight relay recovers everything (In-flight floor at scale). At the headline $N = 262,144$ the wasted-arrival count under the “inflight” regime is exactly zero at every $\Lambda \in \{0.05, 0.1, 0.2\}$ across all 32 seeds – every one of about 2.36 million wasted arrivals suffered under perfect information disappears when a probe redirects the instant an overtaking beacon reaches its current position. Some wasted *travel* remains (probes still burn distance before aborting; median wasted travel drops by about 79% to 95% versus perfect information across Λ), but no probe ever completes a journey to an already-claimed star. The $N = 524,288$ companion confirms it – still exactly zero completed wasted arrivals – so the light-speed number is an upper bound that a physical relay collapses to zero, and only more firmly as the field widens and beacons overtake in-flight probes ever more routinely.

More probes aloft, in step with the field (Concurrency at scale). At the headline $N = 262,144$ the median peak in-flight is about 91,000 under perfect information and about 84,600 under the light-speed gate, and the $N = 524,288$ companion rises to about 149,000 and 137,000; in both the two regimes track closely. This large, near-identical population is what dilutes the *fuel* waste – a wasted trip is one of tens to hundreds of thousands aloft, so losing it barely moves the fill rate, and the fuel tax saturates downward with N . What the population does *not* absorb is the completion-time lag at the tail (Section 3.2): that is why the fill-time tax, unlike the fuel tax, does not shrink with the field, and comes to exceed it at scale.

Clumpiness robustness holds (Clumpy-field at scale). At $N = 200,000$ with the same $n_{\text{clumps}} = 25$ Thomas layout (so the clumps become dense mega-clumps of about 8,000 stars each at the tightest scatter), the through-origin slope is flat across the clumpiness ladder: $a = 0.078$ at the uniform null, 0.076, 0.088, 0.072, 0.076 across the four Thomas levels, all with overlapping bootstrap intervals. This is the one scale companion without a power-of-two rerun, so it stays on the historical $N = 200,000$ field; its uniform-level slope sits within noise of the power-of-two speed-sweep companion’s 0.068 at $N = 262,144$. The exposure-cancellation argument for the v/c form therefore holds over a $400\times$ span in field size crossed with a $4.5\times$ span in Clark-Evans R : clumpiness rescales the coefficient at fixed N but does not break the form, at the smaller field or at scale.

The re-target cap is a lower-bound knob, not a settled default (Re-target-cap). A probe that loses a race re-targets, but only up to a fixed cap before it is retired; on the few-hundred-star fields of earlier work the tax plateaued above about $\text{cap} = 8$, which is why we adopt $\text{cap} = 8$ as the default. That plateau does not survive to larger fields, and this is the one place where moving every measurement onto a bigger field changes a methodological claim rather than just a number. At $N = 262,144$ the tax climbs monotonically across the whole ladder – 0.27%, 0.27%, 1.31%, 4.79%, 8.97% at caps 2, 4, 8, 16, 32 – with no sign of levelling; cap 8 captures only about 15% of the cap-32 value, a shortfall that has grown from the roughly 40% seen on the smaller $N = 32,768$ field, so the plateau is *more* absent at larger N . An independent eight-seed ensemble at the same field agrees, climbing from 0.27% at $\text{cap} = 2$ to 8.60% at $\text{cap} = 32$. So the default cap of 8 is a documented *lower bound*, not a converged value: its downward bias is negligible on the small fields where the plateau holds and grows with N , so the $N = 262,144$ headline, which uses $\text{cap} = 8$, carries this shortfall directly. Every fuel-tax figure in this paper that uses cap 8 should therefore be read as “at least this much,” with the shortfall small at

Table 3: Where real propulsion concepts sit on the Λ axis, and the fuel tax each incurs. The powered cruise is far too slow for lag to bite; gravitational slingshots self-limit near $\Lambda \approx 0.01$; only beamed directed-energy propulsion reaches the regime where the tax is large.

Propulsion regime	$\Lambda = v/c$	Fuel tax (%)
Powered cruise (Nicholson and Forgan, 2013)	3×10^{-5}	≈ 0
Gravitational slingshot (Nicholson and Forgan, 2013)	≈ 0.01	≈ 1
Directed-energy / light sail (Lubin, 2016)	0.1 to 0.2	9 to 16

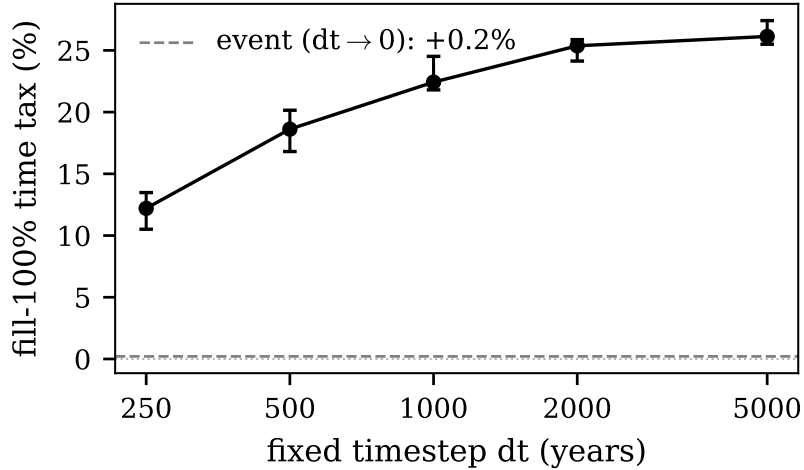


Figure 1: The fill-100% time penalty for nearest-star slingshot flight versus the fixed timestep (log axis), over the 32-seed ensemble at $N = 262,144$ (points: medians; bars: bootstrap 95% confidence intervals). The apparent penalty shrinks as the step is refined and collapses to a fraction of a percent (dashed line: the event-driven, $dt \rightarrow 0$ limit). The tens-of-percent “coordination tax” of a coarse-step model is a discretization artifact, not a physical slowdown.

small N and substantial at scale (the same headline field at $\text{cap} = 32$ reads about 9% rather than 1.3% at $\Lambda = 0.2$). Whether the cap-tax ever saturates at some field size is a natural next question.

3.9 Where the cost falls

Table 3 places the physical propulsion concepts on the speed axis. The source model’s powered cruise ($3 \times 10^{-5} c$) is far too slow for any race to matter, and pays nothing. The gravitational slingshot that motivated this study reaches only $\Lambda \approx 0.01$ before its boost self-limits, and pays under a tenth of a percent more wasted journeys – negligible. The fuel tax stays modest even for beamed directed-energy or light-sail propulsion at 0.1 to 0.2 c (Lubin, 2016), where at this scale it adds about three-quarters of a percent to 1.3% more wasted journeys than perfect information; the larger cost at directed-energy speed is the fill-time lag of Section 3.2. Every case still fills the field to completion: this is a property of the model, not a result, since with re-targeting and no loss term the reachable field must close, and delay supplies no term that would leave a permanent unsettled remainder of the kind an “Aurora” steady state requires (Carroll-Nellenback et al., 2019). The coordination failure shows up as wasted motion, not as an unfilled galaxy.

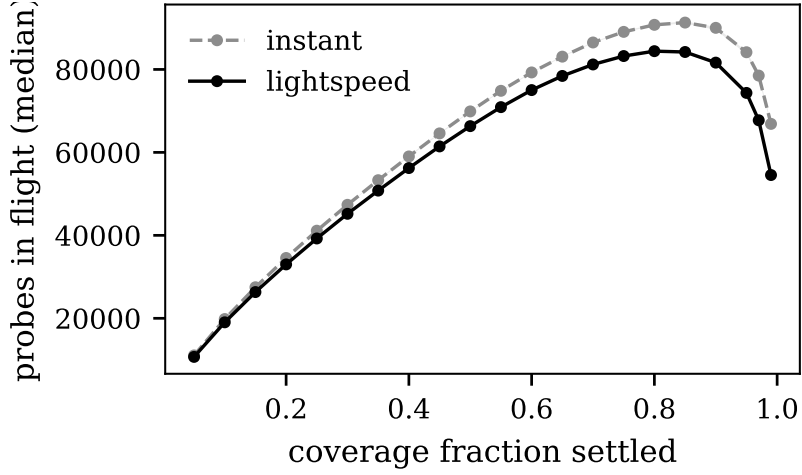


Figure 2: Probes in flight versus the fraction of the field settled, for perfect information and the light-speed gate (medians over 32 seeds with bootstrap 95% confidence bands, 262,144-star field, $\Lambda = 0.2$). The population stays near its peak of about 91,000 all the way into the final percent and the two regimes track each other closely; this near-identical population dilutes the fuel waste, while the fill-time cost is set instead by the front lag at the tail (Section 3.2).

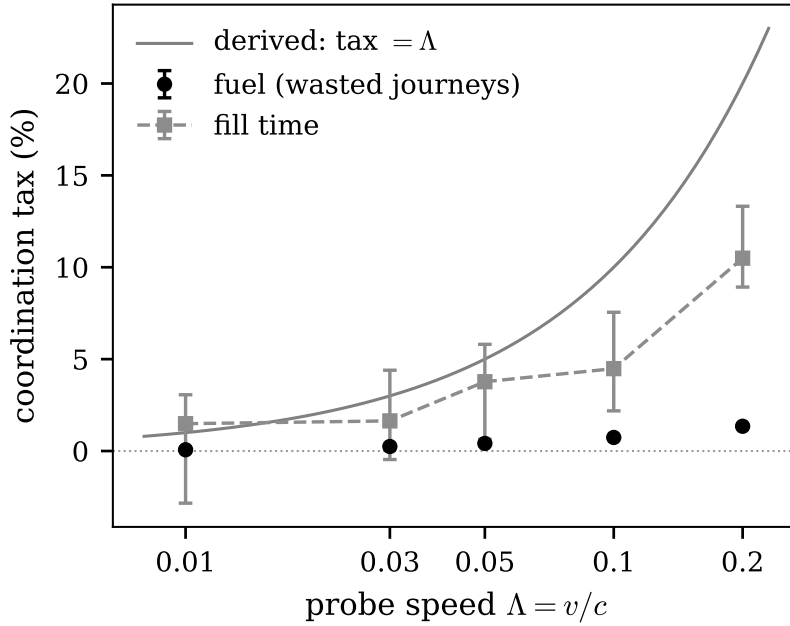


Figure 3: The coordination tax versus probe speed $\Lambda = v/c$ (log axis, so the linear law tax $= \Lambda$ appears as an upward curve), powered flight, over the 128-seed, $N = 262,144$ ensemble at the event-driven timestep (both as a percent over perfect information; points: medians; bars: bootstrap 95% confidence intervals). The redundant-travel (fuel) tax runs parallel to but far below the derived ceiling tax $= \Lambda$ (grey line, Eq. 3), a coefficient of 0.069 at this field; at this scale the fill-time tax exceeds the fuel tax, reaching about 10% at directed-energy speed.

4 Discussion

The result refines rather than overturns the slingshot picture, and it does so twice over. Nicholson and Forgan’s speed-up is real, large, and here reproduced to the same order of magnitude (Nicholson and Forgan, 2013); finite light-speed does not erode it, and does not measurably delay the fill of the field. What a perfect-information model hides is not lost time but lost

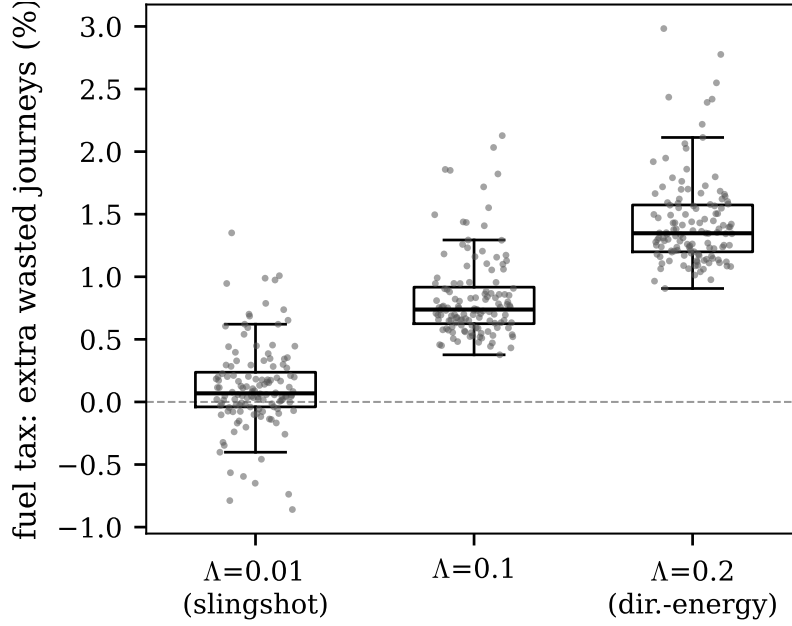


Figure 4: Per-seed distribution of the fuel tax (extra wasted journeys, as a percent of the perfect-information waste) at slingshot speed ($\Lambda = 0.01$), $\Lambda = 0.1$, and directed-energy speed ($\Lambda = 0.2$), over the 128-seed, $N = 262,144$ ensemble (box: interquartile range and median; points: individual seeds). The tax is small but consistently positive at slingshot speed and unambiguous at directed-energy speed.

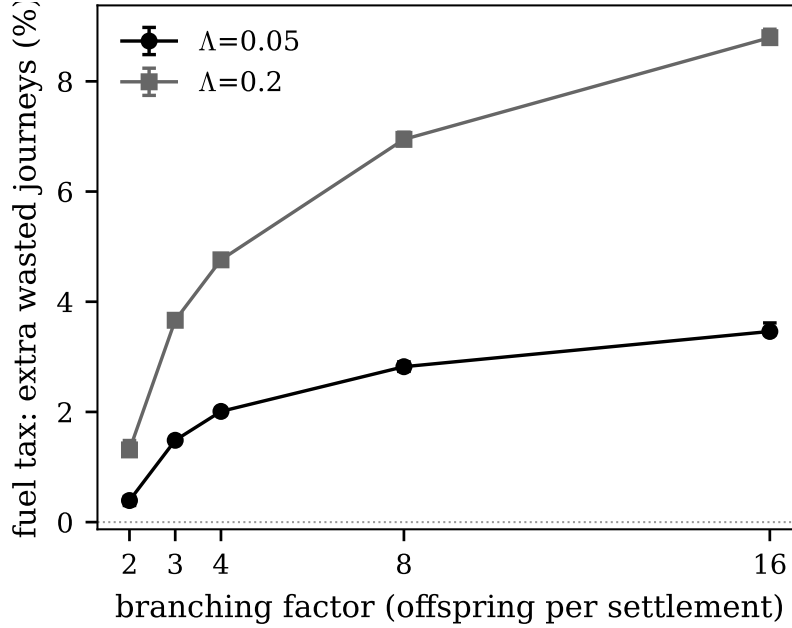


Figure 5: Fuel tax versus the replication branching factor (offspring per settlement) at moderate and directed-energy speed, over the 32-seed, $N = 262,144$ ensemble (points: medians; bars: bootstrap 95% confidence intervals). Contention grows with branching and does *not* saturate over the range tested (two to sixteen offspring): the tax keeps rising, roughly linearly in log branching, so the default of two offspring is a lower bound, not a plateau. The growth is in the coefficient; the tax stays a fraction of Λ at every branching factor.

motion: probes flying to stars that others have already claimed. That redundant travel is the coordination tax, it is governed by the probe's speed in units of c , and it matters only once a

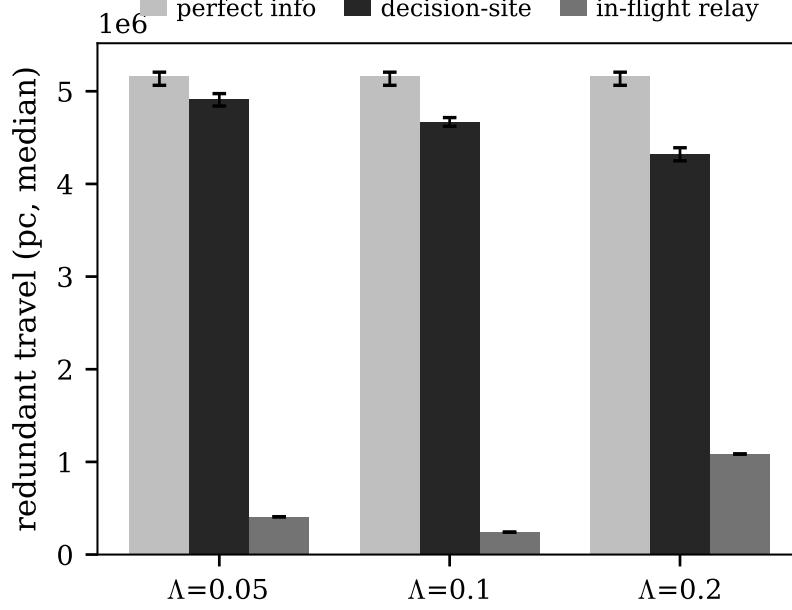


Figure 6: Redundant travel under perfect information, the decision-site light-speed gate, and in-flight relay, versus Λ (medians over 32 seeds with bootstrap 95% confidence intervals, 262,144-star field). Listening in flight drives completed wasted arrivals to zero and cuts redundant travel below even the perfect-information baseline, recovering the light-lag tax at a substantial cost in fill time (about 27% at $\Lambda = 0.2$ here).

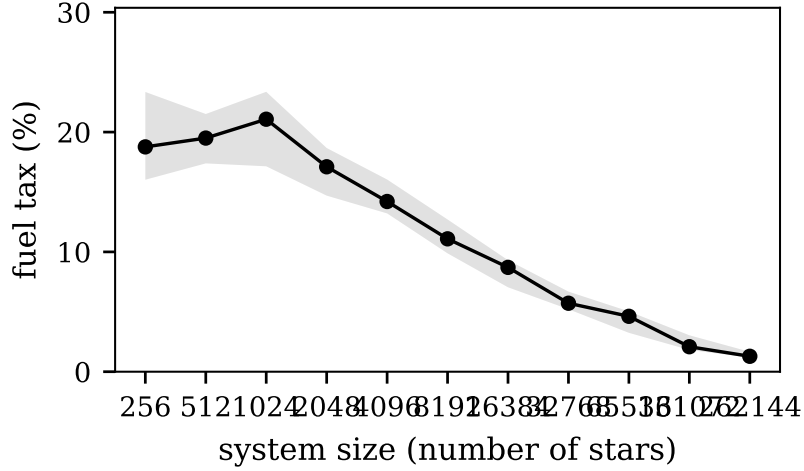


Figure 7: Fuel tax as a percent of the perfect-information waste versus system size (log axis), at $\Lambda = 0.2$ (point: median; shaded band: bootstrap 95% confidence interval; 8 to 32 seeds per point). Over a 1024-fold range (256 to 262,144 stars) the tax holds near 19 to 21% up to about a thousand stars and then declines steadily to about 1.3% at 262,144; the end intervals are separated by an order of magnitude, so it does not grow with scale but shrinks, which is why we make no fixed-percentage extrapolation to galactic size.

swarm flies fast enough for a decision to outrun its own news.

4.1 What the tax is denominated in

Because the field always fills and each settlement launches a fixed number of offspring, the two coordination regimes build essentially the same number of probes, differing by well under a

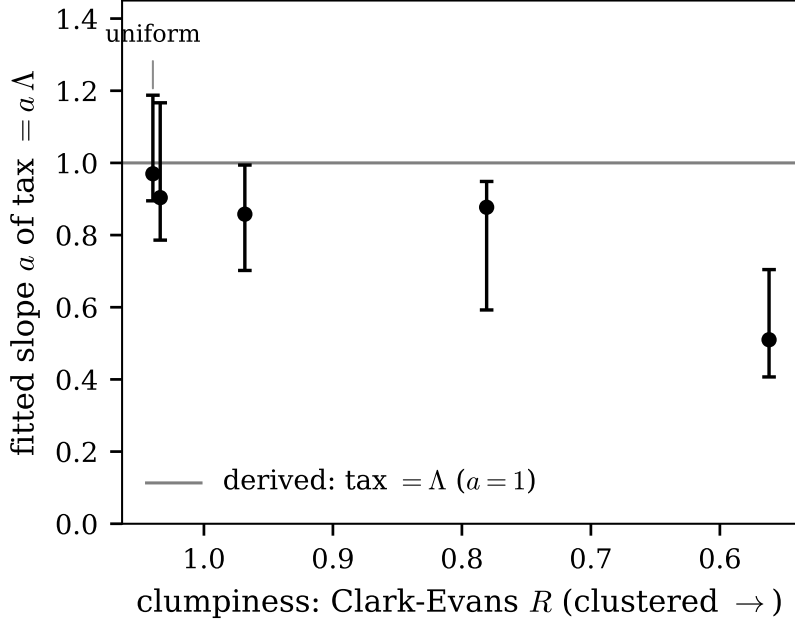


Figure 8: The fitted slope a of tax $= a\Lambda$ versus field clumpiness (Clark-Evans R ; $R = 1$ uniform, decreasing rightward as the field clusters), over the 48-seed ensemble at $N = 500$ (point: median; bars: bootstrap 95% confidence interval; grey line: the derived law $a = 1$). At this small field the uniform level sits near the derived $a = 1$; the slope stays consistent with one through moderate clumping and softens only at extreme substructure, always *below* the law, so tax $= \Lambda$ is a conservative upper bound under clumpiness. (The coefficient itself falls with field size; see Section 3.6.)

percent; the delay changes almost nothing about how much is manufactured, only how far it is flown. The tax is therefore a travel cost, and we deliberately denominate it in journeys rather than energy, because translating journeys into energy is mission-dependent in a way our model does not fix. A wasted journey’s kinetic energy is not straightforwardly a fuel cost: under the slingshot policies most of that energy was extracted gravitationally from stellar motion rather than spent as propellant, and a probe that re-targets on arrival keeps its speed and pays nothing further. The genuine energy penalty of a wasted trip is whatever a rendezvous at the doomed star would have cost – braking to rest and re-accelerating – which scales with v^2 for a fuel-braked craft but is near zero for a flyby or a gravitationally-braked one. We therefore do not headline an energy figure; the honest statement is that the journey tax becomes an energy tax only to the extent a mission must brake, and then it grows with v^2 . A caveat cuts the other way, too: even as a travel cost the tax is denominated in journeys, not in the swarm’s total resource budget, which is dominated by building the probes – a cost we hold fixed and identical across regimes and do not model. As a fraction of total mission resources the coordination tax is smaller than its share of journeys.

4.2 Limitations

The decision-site headline is a conservative upper bound: probes here learn only at their decision stars, ignoring news their paths cross in flight. We relax exactly that assumption in Section 3.5, where a probe that listens in flight recovers essentially all of the wasted-arrival tax, so the headline is an upper bound rather than a floor. A true probe-to-probe relay of the settled map, which we do not implement, would sit between the beacon-only in-flight bound and the decision-site one and is the natural next step. The floor beneath all of them – that no probe can know a distant star is taken until its light arrives – is set by physics, not protocol.

A further, inherited idealization concerns the stars themselves. They carry the velocity vectors that power each slingshot but are held fixed in position, exactly as in the two source models, both of which name this as their single most important simplification (Nicholson and Forgan, 2013; Forgan et al., 2013). It is not a small one: at the local circular speed a star sweeps about 225 pc in a million years, and even the peculiar velocity dispersion alone carries it some 30 to 40 pc (Kerr and Lynden-Bell, 1986; Nordstrom et al., 2004), far past the roughly one-parsec mean hop, so over a Myr-scale fill real proper motion would reshuffle which stars are near which. One might object that this touches the delayed-belief gate itself, not just the backdrop: the beacon leaves where the star *was*, so a moving field would need retarded positions in (1). But that specific correction is tiny, and for a clean reason: during a beacon’s light-crossing a star moves only the fraction $v_\star/c$ of the distance it is signalling across, where $v_\star$ is the *relative* velocity between emitter and observer – the bulk Galactic rotation is common to both and cancels, leaving the peculiar and shear velocity, ~ 40 km/s, so $v_\star/c \approx 10^{-4}$, a hundredth of a percent whatever the decision distance (and even the full circular speed, 220 km/s, gives 7×10^{-4} , under a tenth of a percent) – so the gate reads essentially the right geometry at the moment of decision. What the frozen field removes is the slow, Myr-scale drift of the layout, and that is exactly what the paired design controls: both coordination regimes run on the *identical* frozen field, so the difference still isolates the light-lag cost even where the absolute layout is idealized. This cancellation extends to the *dynamics*, not just the static geometry: what the tax multiplies is the baseline per-hop contention p_{perfect} , set by the local stellar density and the branching factor, both of which are mode-independent and both of which see the identical drift under either regime – so a moving field would perturb the paired difference only at second order, through any correlation between a star’s motion and its being contended, which the global constancy of v and c already suppresses in the same way it suppresses the density dependence (Section 3.7). A direct moving-field spot-check would confirm this and is the natural next test; the absolute fill times, meanwhile, do inherit the idealization. The slingshot is idealized in a second way we inherit from the source models: we track the scalar speed a probe gains from an encounter and let it then depart toward its chosen target, rather than resolving the outgoing *direction* that a real gravitational assist fixes by geometry. A directional model would couple boost to heading and so lower the achievable boosted speed, and thus Λ , in the slingshot regime; but because that regime already pays only about a percent, and the headline v/c law is driven by the directly-swept powered speed rather than by slingshot mechanics, the simplification does not touch the law.

One idealization we do *not* need to correct is relativity, and for a stronger reason than smallness. The tax counts journeys, and every quantity that enters it – the belief gate (1), the travel time d/v , the light-delay d/c , and the in-flight learn time of Section 3.5 – is a coordinate time in the one shared galaxy frame, the frame in which the star field is at rest and the beacons are defined. The governing ratio $\Lambda = (d/c)/(d/v) = v/c$ is a ratio of two coordinate times in that single frame; no Lorentz transformation to the probe frame enters, and $v < c$ throughout, so $\Lambda = v/c$ is *exact* for the coordination tax at any sub-luminal speed, not a small- β linearization. The Lorentz factor at $0.2c$ ($\gamma \approx 1.02$, a kinematic correction of about 2%) touches only the separate energy weighting, where it never enters a headline number, and not the journey count.

Three further bounds on scope. First, the field is uniform by default; changing its density only rescales the clock (confirmed), and a clumpy, non-uniform field leaves the law intact – the redundant-travel tax stays $\approx \Lambda$, softening if anything at extreme substructure (Section 3.7). What we do not probe is a full galactic disk with a large-scale radial density gradient. Second, the system-size sweep now spans about $1024\times$, from 256 to 262,144 stars: serving the nearest-not-yet-settled query with a dynamic spatial index over the unsettled set makes a single run near-linear in N , lifting the quadratic cost that once held the sweep to a sixteen-fold range. Over this longer arm the tax does not grow with scale – it holds near a fifth up to about a thousand stars and then falls steadily to about one part in eighty at 262,144, a genuine bulk

decline that survives both an interior-only masking and a periodic-box control and so is not a hard-wall boundary artifact (Section 3.6) – so we still make no quantitative extrapolation to galactic ($\sim 10^{11}$ -star) scale. A $1024\times$ lever arm is far longer than a sixteen-fold one but remains five orders of magnitude short of a real galaxy, and in any case a range over which the fraction only shrinks cannot support carrying a fixed percentage outward. Third, the dynamics have no loss term: settled sites never fail, so the field always closes. A settlement-death term of the kind that produces the “Aurora” steady state (Carroll-Nellenback et al., 2019) is what would be needed to turn coordination failure into a permanent unsettled remainder; delay alone does not.

4.3 Relation to delayed agreement

This is a coordination problem, and the distributed-systems literature names the regime it sits in – by analogy, not by derivation; we invoke it only to place that regime, and the result itself is our own. When communication delay grows past the decision cadence a system can no longer maintain a single agreed state and must trade consistency for availability, acting on local views (Gilbert and Lynch, 2002; Olfati-Saber et al., 2007). At parsec spacing the round-trip light-time is years, far longer than any plausible decision cadence, so a swarm that must keep acting has no option but availability – which is exactly what the gate (1) encodes, each probe an independent colony acting on priors. We borrow only this picture. The classical impossibility results assume fault models that are not ours (asynchronous crash faults for FLP (Fischer et al., 1985), partition for CAP (Gilbert and Lynch, 2002)) and none of them predicts the $\Lambda = v/c$ scaling, which is at bottom a geometric-probabilistic result about stale target selection, not a consensus theorem. The mechanism itself – agents on a stale shared map flying redundant trips to targets others have already claimed, and communication cutting that overlap – is familiar from multi-robot exploration, where coordinated frontier assignment exists precisely to suppress it (Burgard et al., 2005; Amigoni et al., 2017). What is specific here is not that stale views waste trips but that, when the delay is light-speed itself, the wasted fraction takes the clean form $\Lambda = v/c$, independent of hop length and density. The coordination tax is simply what that enforced independence costs when probes want the same stars.

4.4 An implication for design

The corollary is a threshold, not a preference. At slingshot speeds the fuel tax is well under a percent, so coordination machinery – relay, gossip, a shared map – returns less than its mass costs, and the coordination overhead is negligible against what speed buys. At the directed-energy speeds that make a probe fast enough to matter, the same physics that buys the speed also makes stale views expensive – a sixth more journeys wasted on a few-thousand-star field, and at galactic scale a saturated fuel tax of about a percent but a fill-time lag of order ten percent – of which, as Section 3.5 shows, a probe that merely listens in flight recovers the completed wasted arrivals, though at a fill-time cost of its own. The implication is that the value of coordination scales with how fast the swarm flies and how large the field it must fill, while the finite speed of light sets only the small residual – the travel already flown before the news arrives – that no protocol can ever recover.

5 Conclusion

Adding a finite speed of information to the self-replicating slingshot model turns a convenient assumption into a measurable cost, but not the cost a coarse model reports. Once the dynamics are resolved so that each probe decides at its true arrival time, the tens-of-percent fill-time penalty seen in a fixed-timestep model collapses – it was a discretization artifact – and the slingshot speed-up survives intact, reproducing Nicholson and Forgan to the same order of

magnitude. The genuine cost is redundant travel: every star is still settled, with a negligible change in the number of probes built, but stale views send probes to stars others have already claimed. That waste follows a simple law we both derive and measure – the delay multiplies the wasted journeys by at most $1 + v/c$, so the tax, the extra wasted trips as a fraction of those a perfect-information swarm already flies, scales with the probe’s speed in units of light, $\Lambda = v/c$, at a coefficient that saturation holds below one and that shrinks with system size – about 0.8 on a few-thousand-star field, where the tax runs from negligible at the powered cruise through about a percent at slingshot speeds to a sixth at directed-energy speeds, but only about 0.07 on our quarter-million-star headline field, where even the directed-energy fuel tax is about 1.3%. (The baseline is itself wasteful, about nine in ten journeys even with a perfect map, so this is a v/c surcharge on an already-saturated process.) The law holds under a clumpy field, where if anything it softens; it grows with the replication branching factor without saturating up to sixteen offspring; and whether it also costs energy is mission-dependent: only a rendezvous that must brake pays for a wasted trip’s kinetic energy, which for the fast policies was largely gravitationally sourced rather than fuelled. On the small fields of earlier work the cost is dominated by fuel and effort, with a smaller fill-time penalty alongside it; at galactic scale the two invert – the fuel tax saturates toward a percent while the fill-time lag, set by the front crossing an ever-larger field, grows to exceed it. It is not a hard floor: a probe that listens in flight, one of the two extensions the model invites, recovers essentially all of the wasted-arrival tax at a small cost in fill time, while a full probe-to-probe relay of the settled map remains open. Neither can remove the physical floor beneath the tax, and how much either can buy back is itself bounded by the finite speed of light.

Data Availability

The model is a deterministic, seeded, event-driven fold; every figure, table, and reported statistic regenerates from source by running the swarm module’s experiment and figure scripts, and every quantitative claim traces to a cited source in the project’s shared bibliography. The code and its citation metadata are archived with the concept DOI [10.5281/zenodo.21249296](https://doi.org/10.5281/zenodo.21249296).

Acknowledgment

This work used no external funding.

References

- F. Amigoni, J. Banfi, and N. Basilico. Multirobot Exploration of Communication-Restricted Environments: A Survey. *IEEE Intelligent Systems* 32(6):48-57, DOI [10.1109/MIS.2017.4531226](https://doi.org/10.1109/MIS.2017.4531226), 2017. URL <https://doi.org/10.1109/MIS.2017.4531226>.
- R. Bjork. Exploring the Galaxy using space probes (arXiv:astro-ph/0701238). *International Journal of Astrobiology* 6(2):89-93, 2007. URL <https://arxiv.org/abs/astro-ph/0701238>.
- W. Burgard, M. Moors, C. Stachniss, and F. E. Schneider. Coordinated Multi-Robot Exploration. *IEEE Transactions on Robotics* 21(3):376-386, DOI [10.1109/TRO.2005.844673](https://doi.org/10.1109/TRO.2005.844673), 2005. URL <https://doi.org/10.1109/TRO.2005.844673>.
- J. Carroll-Nellenback, A. Frank, J. Wright, and C. Scharf. The Fermi Paradox and the Aurora Effect (arXiv:1902.04450). *Astronomical Journal*, 2019. URL <https://arxiv.org/abs/1902.04450>.

- P. J. Clark and F. C. Evans. Distance to Nearest Neighbor as a Measure of Spatial Relationships in Populations. *Ecology* 35(4):445-453, DOI 10.2307/1931034, 1954. URL <https://doi.org/10.2307/1931034>.
- C. Cotta and A. Morales. A Computational Analysis of Galactic Exploration with Space Probes (arXiv:0907.0345). *Journal of the British Interplanetary Society* 62:82-88, 2009. URL <https://arxiv.org/abs/0907.0345>.
- M. J. Fischer, N. A. Lynch, and M. S. Paterson. Impossibility of Distributed Consensus with One Faulty Process. *Journal of the ACM* 32(2):374-382, DOI 10.1145/3149.214121, 1985. URL <https://groups.csail.mit.edu/tds/papers/Lynch/jacm85.pdf>.
- D. H. Forgan. A numerical testbed for hypotheses of extraterrestrial life and intelligence (arXiv:0810.2222). *International Journal of Astrobiology* 8(2):121-131, 2009. URL <https://arxiv.org/abs/0810.2222>.
- D. H. Forgan, S. Papadogiannakis, and T. Kitching. The Effect of Probe Dynamics on Galactic Exploration Timescales (arXiv:1212.2371). *Journal of the British Interplanetary Society* 66:171-178, 2013. URL <https://arxiv.org/abs/1212.2371>.
- R. A. Freitas, Jr. A Self-Reproducing Interstellar Probe. *Journal of the British Interplanetary Society* 33:251, 1980. Cited by reference; no online link.
- S. Gilbert and N. A. Lynch. Brewer’s Conjecture and the Feasibility of Consistent, Available, Partition-Tolerant Web Services. *ACM SIGACT News* 33(2):51-59, DOI 10.1145/564585.564601, 2002. URL <https://doi.org/10.1145/564585.564601>.
- M. H. Hart. An Explanation for the Absence of Extraterrestrials on Earth. *Quarterly Journal of the Royal Astronomical Society* 16:128-135, 1975. URL <https://ui.adsabs.harvard.edu/abs/1975QJRAS...16..128H/abstract>.
- E. M. Jones. Discrete calculations of interstellar migration and settlement. *Icarus* 46(3):328-336, 1981. URL <https://www.sciencedirect.com/science/article/abs/pii/0019103581901366>.
- F. J. Kerr and D. Lynden-Bell. Review of galactic constants. *Monthly Notices of the Royal Astronomical Society* 221:1023-1038, 1986. URL <https://ui.adsabs.harvard.edu/abs/1986MNRAS.221.1023K/abstract>.
- L. Lamport. Time, Clocks, and the Ordering of Events in a Distributed System. *Communications of the ACM* 21(7):558-565, DOI 10.1145/359545.359563, 1978. URL <https://doi.org/10.1145/359545.359563>.
- P. Lubin. A Roadmap to Interstellar Flight (arXiv:1604.01356). *Journal of the British Interplanetary Society* 69:40-72, 2016. URL <https://arxiv.org/abs/1604.01356>.
- J. von Neumann and ed. A. W. Burks. Theory of Self-Reproducing Automata. University of Illinois Press, Urbana, 1966. URL https://archive.org/details/theoryofselfrepr00vonn_0.
- W. I. Newman and C. Sagan. Galactic civilizations: Population dynamics and interstellar diffusion. *Icarus* 46(3):293-327, 1981. URL <https://www.sciencedirect.com/science/article/abs/pii/0019103581901354>.
- A. Nicholson and D. Forgan. Slingshot Dynamics for Self-Replicating Probes and the Effect on Exploration Timescales (arXiv:1307.1648). *Int. J. Astrobiology* 12, 337, 2013. URL <https://arxiv.org/abs/1307.1648>.

- B. Nordstrom, M. Mayor, J. Andersen, et al. The Geneva-Copenhagen survey of the solar neighbourhood. *Astronomy & Astrophysics* 418:989-1019, 2004. URL <https://arxiv.org/abs/astro-ph/0405198>.
- R. Olfati-Saber, J. A. Fax, and R. M. Murray. Consensus and Cooperation in Networked Multi-Agent Systems. *Proceedings of the IEEE* 95(1):215-233, DOI 10.1109/JPROC.2006.887293, 2007. URL <https://doi.org/10.1109/JPROC.2006.887293>.
- RECONS and Gaia. Nearby-star census (10-parsec sample and parallaxes). RECONS / ESA Gaia, 2018. URL <http://www.recons.org/census.posted.htm>.
- F. J. Tipler. Extraterrestrial Intelligent Beings Do Not Exist. *Quarterly Journal of the Royal Astronomical Society* 21:267-281, 1980. URL <https://ui.adsabs.harvard.edu/abs/1980QJRAS...21..267T/abstract>.