

Emergent Flocking and Collective Evasion in a Force-Based Agent Model

Nathan Langley

University of North Carolina at Greensboro

PHY 351 — Independent Summer Research Advisor: Prof. Ian Beatty

May 2026

Abstract

I present computational simulations of a force-based flocking model (Charbonneau, 2017) in which N agents on a periodic two-dimensional domain interact through repulsion, velocity alignment, self-propulsion, and random noise. After validating the implementation against three analytically tractable limiting cases, I derive an exact analytical result: the equilibrium cruise speed of an aligned flock is $v_{eq} = v_0 + \alpha/\mu$, not the nominal target speed v_0 . Finite-size scaling at multiple compactness values ($C = 0.10$ to 0.78) shows no diverging susceptibility, indicating a smooth crossover rather than a true phase transition. I then extend the model in two directions: external predation and internal panic. Under sustained predator pressure, flocking prey maintain near-perfect alignment ($\Phi \sim 1.0$) while non-flocking prey scatter ($\Phi \sim 0.1$). Three predator strategies form a clear hierarchy: naive (chase CoM) and coordinated (mutual repulsion) predators fail to disrupt the flock; only explicit encirclement — assigning each predator a fixed compass angle — substantially reduces coherence ($\Phi = 0.77$ at $n_{pred} = 6$). Cluster analysis reveals this disruption is flock DIVISION into coherent sub-flocks rather than dissolution. Over long time (300 time units), the encircled flock enters an intermittent merge/split cycle rather than a steady fragmented state. Encirclement damage is fully reversible: sub-flocks reunite within ~ 10 time units of predator removal. Internal stressors obey different rules: static panic fails to disrupt the calm sub-flock (calm $\Phi \sim 1.0$ at 20% panic fraction), but SIS contagion with recovery rate γ has a clean epidemic threshold at $\beta/\gamma \sim 1$. Encirclement shifts this threshold by $\sim 4\%$ by compressing local contact density, and the spatial herd-immunity threshold (~ 0.46) is more than twice the mean-field prediction due to clustering in the contact network. Critically, when encirclement drives a sub-threshold contagion into a transient endemic state, removing the predators reverses kinematic fragmentation (in ~ 10 time units) but leaves the epidemic intact for hundreds of time units. Predation produces reversible kinematic damage; contagion produces lasting epidemic damage; the combination produces damage that outlasts the predation event itself by an order of magnitude. An adaptive encirclement strategy — predators that continuously track live flock R_g and maintain $R_{enc} = 0.5R_g$ — *cuts the fraction of time the flock spends in a highly coherent state by 34% compared to a fixed-radius strategy, validating the universality of the $R_{enc}/R_g \sim 0.5$ optimum across dynamical fluctuations. Neither degree-targeted nor spatially-targeted vaccination outperforms random: both require $p_{immune} \sim 0.46$. A dedicated follow-up disentangles the cause — spatial targeting fails because kinematic mixing scrambles the immune agents' coverage before the epidemic runs, while degree targeting fails for a distinct structural reason (the flock contact network has no hubs to target), a null result that persists even when the contact graph is frozen. Sweeping the repulsion exponent from $n = 1.5$ (soft) to $n = 12$ (near hard-core) produces identical finite-size scaling at all exponents. The smooth crossover is not caused by potential softness; it is a consequence of non-equilibrium driving (uniform random kicks without viscous dissipation). A follow-up Langevin simulation confirms that proper thermal equilibration (FDT satisfied; $KE/N = kT$ at equilibrium) is achievable within this model family, but the kinetic energy susceptibility $\chi = N\text{Var}(KE/N)$ is not sensitive to the structural hard-disc melting transition. A final experiment*

measuring the hexatic order parameter $|\psi_6|$ directly reveals that the $n = 1.5$ soft repulsion cannot crystallize at any accessible temperature: $|\psi_6| \approx 0.4$ across the entire kT range from 0.001 to 5.0 with no N -dependent susceptibility peak. The KTHNY transition is absent because the soft potential allows agents to pass through each other’s cores at any finite kT , preventing hexagonal lattice formation.

Finally, I extend the model to three spatial dimensions. Flocking and the exact result $v_{eq} = v_0 + \alpha/\mu$ both generalize cleanly to 3D, and the noise-driven crossover remains smooth. Encirclement, however, does not transfer at all: with correctly repulsive predators, no geometric variant — radius tuning, predator count up to 50, adaptive tracking, or planar arrangement — produces any disruption in 3D, where the order parameter stays at $\Phi \sim 1.000$ while 2D encirclement at identical settings drives it to ~ 0.73 . Encirclement is strictly a 2D strategy: a modest number of point predators can seal a flock’s 1D perimeter but not a closed surface around a 3D volume. The vaccination null results, by contrast, do extend to 3D unchanged. A concluding pair of experiments tests the synthesis against its own pre-registered prediction: replacing the metric alignment force with a topological one does not slow kinematic mixing, and freezing the contact graph does not rescue degree-targeted vaccination — confirming that kinematic mixing is a property of the agents’ physical motion and that the degree-targeting null is structural rather than kinematic.

1 Introduction

One of the central puzzles in complex systems is how large-scale ordered behavior emerges from purely local interactions. Flocking — the coordinated motion of birds, fish schools, and animal herds — is a canonical example. Each individual follows simple rules based on its immediate neighbors, yet the collective produces sweeping global patterns with no central coordination. Understanding the conditions under which order emerges, and how robust that order is to perturbation, has implications ranging from evolutionary biology to crowd control.

The model studied here is based on Chapter 10 of Charbonneau (2017), which was originally developed by Silverberg et al. (2013) to describe crowd dynamics in mosh pits at heavy metal concerts. Each agent in the model is subject to four forces: short-range repulsion, velocity-aligning flocking force, self-propulsion toward a target speed, and random noise. The interplay of these four forces produces a rich behavioral phase space, including crystalline order, disordered fluid motion, and coherent streaming flocks.

This report covers thirty-five investigations, producing fifty-three numbered findings. The first three sections establish the baseline: implementation validation and the v_{eq} analytical result (Section 4.1), parameter sweeps and flock formation (Section 4.2), and the solid-to-fluid transition tested as a true phase transition (Section 4.3). Section 4.4 extends the model with a predator agent and uses it to characterize flock geometry under pressure — establishing both the single-predator baseline and the geometric metrics used throughout. Sections 4.5 through 4.9 develop the predator-strategy hierarchy — from naive co-localization to coordinated spreading to encirclement — and characterize encirclement as the only strategy capable of substantial disruption, acting through transient flock division rather than dissolution, with a universal disruption optimum at $R_{enc}/R_g \sim 0.5$. Sections 4.10 through 4.12 introduce internal stressors: contagious panic via SI and SIS epidemic dynamics, the epidemic threshold, and hybrid predation-contagion interactions including agent-heterogeneity segregation. Sections 4.13 through 4.15 refine the predator-contagion system — sub-threshold coupling, near-critical compression effects, universal encirclement scaling, and long-time merge/split dynamics. Section 4.16 reveals a two-timescale asymmetry in recovery: kinematic damage from encirclement reverses in ~ 10 time units, while epidemic damage from a

contagion outbreak persists for ~100+ time units after predators are removed. Sections 4.17 through 4.20 address four targeted-intervention questions: adaptive encirclement validates the R_{enc}/R_g universal optimum dynamically; degree-targeted and spatially-targeted vaccination (Sections 4.18, 4.20) are both null results defeated by kinematic reorganization; harder repulsion (Section 4.19) fails to produce a true phase transition, diagnosing the smooth crossover as non-equilibrium in origin. Sections 4.21 and 4.22 close the phase-transition thread: a Langevin thermostat confirms FDT-satisfying dynamics thermalize correctly, but the hexatic order parameter (Section 4.22) reveals that $n = 1.5$ soft repulsion cannot crystallize at any accessible temperature. Sections 4.23 and 4.24 extend the model to three spatial dimensions: flocking and v_{eq} generalize exactly to 3D, and the noise-driven crossover in 3D is a smooth crossover at $\text{ramp} \sim 15\text{-}25$ (consistent with the 2D behavior, confirming a dimensionality-independent mechanism). Sections 4.25 through 4.31 develop the three-dimensional thread in depth. Section 4.25 introduces 3D predator strategies and finds that encirclement does not transfer to 3D at all — the order parameter stays at $\Phi \sim 1.000$. Sections 4.26, 4.27, and 4.31 show this failure cannot be repaired by predator count, adaptive geometry, or planar arrangement: every geometric variant of encirclement leaves the 3D flock undisturbed, because a modest number of point predators cannot seal a closed surface around a three-dimensional volume. Section 4.28 confirms that the vaccination null results extend to 3D. Sections 4.29 and 4.30 then test the report’s own synthesis against itself: a pre-registered prediction (that a topological alignment force would slow mixing and rescue targeted vaccination) is falsified, and a contact-graph-freezing experiment shows the degree-targeting null result is structural rather than kinematic — sharpening the synthesis of Section 5. Section 4.32 closes the phase-transition thread: a hard-repulsion Langevin simulation (exponent $n = 12$ and 24) still does not crystallize, because a higher exponent in this force form shrinks the effective core rather than hardening it. Section 4.33 closes the segregation thread in 3D: alpha-contrast segregation transfers to three dimensions but is diluted by the extra spatial dimension. Section 4.34 then tests, and falsifies, the tempting interpretation that the third dimension speeds up kinematic mixing: a direct measurement shows 3D mixes roughly 1.8 times slower than 2D at matched contact degree, so the 3D results are due to neighborhood geometry and structural homogeneity, not faster mixing.

2 Model

2.1 Setup

N agents move on a periodic unit square (x, y in $[0, 1]$), implemented as a torus so that agents exiting one edge reappear on the opposite side. Agent positions and velocities are updated at each timestep using the forward Euler method at $dt = 0.01$. The state of agent j at time t is fully described by its position (x_j, y_j) and velocity (v_{x_j}, v_{y_j}) .

2.2 Forces

The total force on agent j is a sum of four contributions (Charbonneau Eqs. 10.1-10.8):

Repulsion. A short-range force prevents agents from overlapping. It acts on pairs within distance $2r_0$ and grows in intensity as agents approach:

$$F_{\text{rep},j} = \text{eps} * \text{SUM}_k [(1 - r_{jk} / 2r_0)^{(3/2)} * \hat{r}_{jk}] \quad \text{for } r_{jk} \leq 2r_0$$

where r_{jk} is the distance between agents j and k , and $\hat{r}_{jk}$ is a unit vector pointing from k toward j .

Flocking. An alignment force drives the velocity of agent j toward the mean velocity of its neighbors within a flocking radius r_f :

$$F_{\text{flock},j} = \alpha * V_{\text{bar}} / |V_{\text{bar}}|, \quad V_{\text{bar}} = \text{SUM}_{\{k: r_{jk} \leq r_f\}} v_k$$

The normalized form ensures the flocking force has constant magnitude α regardless of how many neighbors are present.

Self-propulsion. A speed-correcting force drives agent j toward a target speed v_0 along its current direction of motion:

$$F_{\text{prop},j} = \mu * (v_0 - |v_j|) * v_{\text{hat}_j}$$

where v_{hat_j} is the unit vector along v_j . This force accelerates agents moving too slowly and brakes those moving too fast.

Random noise. Each component of the random force is drawn independently from a uniform distribution on $[-\text{ramp}, \text{ramp}]$ at each timestep.

With unit mass for all agents, Newton's second law gives $a_j = F_j$, and the equations of motion are integrated as:

$$\begin{aligned} x_j(t + dt) &= x_j(t) + v_j(t) * dt \\ v_j(t + dt) &= v_j(t) + F_j(t) * dt \end{aligned}$$

2.3 Periodic Boundary Implementation

Force calculations near domain boundaries require special handling. Agents within range r_f of any boundary are replicated as ghost copies on the opposite side, so that the flocking and repulsion forces computed for a real agent account correctly for neighbors across the periodic boundary. This buffer zone approach follows Charbonneau Fig. 10.2.

2.4 Metrics

The primary measure of collective order is the **order parameter**:

$$\Phi = | (1/N) \text{SUM}_j v_{\text{hat}_j} |$$

$\Phi = 1$ corresponds to perfect velocity alignment; $\Phi = 0$ to randomly oriented motion. I also track total kinetic energy $KE = (1/2) \text{SUM}_j |v_j|^2$ and, for flock geometry, the **radius of gyration** R_g (root-mean-square distance from center of mass) and the **aspect ratio** AR (ratio of the major to minor eigenvalue of the spatial covariance matrix, measuring elongation).

2.5 Default Parameters

Unless otherwise noted, simulations use the parameters from Charbonneau Table 10.1: $N = 350$, $r_0 = 0.005$, $\text{eps} = 0.1$, $r_f = 0.1$, $\alpha = 1.0$, $v_0 = 1.0$, $\mu = 10.0$, $\text{ramp} = 0.5$, $dt = 0.01$.

2.6 Methodology and Conventions

Unless otherwise stated, every reported quantity follows the same protocol so that results across sections are directly comparable.

Seeds and replicates. Every parameter-sweep point is the mean over independent random initializations of agent positions and velocities. The standard count is 8 seeds per point for predator

experiments and 5 seeds per point for noise / phase sweeps. Sections where this differs (e.g., long-time encirclement at 30000 steps, $N = 1000$ large- N runs) state the seed count explicitly.

Error bars. Quoted uncertainties are one standard deviation across seeds (1 sigma), not standard errors of the mean. Figures plot error bars at the same convention. When a finding reports a value as $A \pm B$, B is the seed-to-seed sigma.

Run length and warmup. Default sweeps run 4000 timesteps at $dt = 0.01$ (40 time units) and discard the first 1000 steps (10 tu) as warmup before averaging. Long-time experiments (Section 4.15, Section 4.16) extend to 30000 steps. Convergence was verified by comparing the steady-state Φ between the first and second halves of each run; differences below 0.02 are treated as converged.

Compactness convention. Compactness is defined as $C = \pi * N * r_0^2$, the fraction of domain area occupied by hard-disc cores. For fixed-compactness scaling (Section 4.3), $r_0 = \sqrt{C / (\pi * N)}$ so C is held constant as N varies.

Order-parameter averages. Φ values quoted in steady-state are time-averaged over the post-warmup window; values quoted with both a mean and a temporal sigma (Section 4.15) report both the time-average and the within-run temporal spread.

3D conventions. 3D experiments (Sections 4.23-4.25) hold neighbor count constant by scaling r_0 and r_f . The mapping used throughout is $r_0 = 0.02$, $r_f = 0.20$, on the periodic unit cube $[0, 1]^3$. All other parameters match the 2D defaults except where noted.

Reproducibility. All scripts use `numpy.random.default_rng(seed)` with explicit seed values $0..N_{\text{seeds}}-1$. Source code for every figure is in the repository at the path noted in each section opening.

3 Validation

Before drawing any conclusions from the simulations, I verified the implementation against three limiting cases with known expected behavior.

Case 1: Pure random walk. With all physical forces disabled ($\epsilon = \alpha = \mu = v_0 = 0$, $\text{ramp} = 1$), agents should perform a pure random walk with no preferred direction. The measured order parameter was $\Phi = 0.04$ (expected ~ 0) and agent positions spread uniformly across the domain (standard deviation ~ 0.29 , consistent with uniform distribution). This confirms the integration and boundary conditions are working.

Case 2: Repulsion and noise only. With $\alpha = 0$ and $v_0 = 0$ (self-propulsion acts as a brake), the model reproduces Fig. 10.5 from Charbonneau: at low noise ($\eta = 1$) agents pack into a close-packed quasi-hexagonal structure, while at high noise ($\eta = 30$) the arrangement disorders into a fluid. Φ remains near zero throughout (no alignment force), as expected.

Case 3: Flocking only. With $\epsilon = 0$ and $v_0 = 0$, the alignment force alone should drive agents into a coherent stream. The final order parameter was $\Phi = 0.998$ after $t = 30$, confirming that a single coherent flock forms from random initial conditions, consistent with Fig. 10.6 from Charbonneau (Fig. 1).

4 Results

4.1 Equilibrium Cruise Speed

An exact result follows directly from the force equations. In a perfectly aligned flock, all agents move in the same direction with the same speed. The flocking force then acts purely in the forward direction with magnitude α . The self-propulsion force balances this when:

$$\alpha + \mu * (v_0 - v_{eq}) = 0 \quad \Rightarrow \quad v_{eq} = v_0 + \alpha/\mu$$

With the default parameters ($\alpha = 1$, $\mu = 10$, $v_0 = 1$), this predicts $v_{eq} = 1.10$. I verified this prediction by measuring steady-state mean speed across four values of α with $v_0 = 1$, $\mu = 10$ fixed. Measured speeds agreed with the prediction to within 0.002 in all cases (Fig. 2). The implication is that v_0 and α are not independent knobs for cruise speed: to achieve a target cruising speed v_c , one must set $v_0 = v_c - \alpha/\mu$.

4.2 Flock Formation

Sweeping the flocking amplitude α with noise fixed at $\text{ramp} = 0.1$ (5 seeds per point, error bars represent standard deviation) shows a sharp onset of flocking near $\alpha \sim 0.05$. At $\alpha = 0$, $\Phi = 0.09 \pm 0.01$. By $\alpha = 0.05$, $\Phi = 0.40 \pm 0.12$, and by $\alpha = 0.20$, $\Phi = 0.89 \pm 0.03$. The large run-to-run variance near the threshold ($\text{std} \sim 0.1\text{--}0.2$ for $0.05 \leq \alpha \leq 0.15$) indicates sensitivity to initial conditions near the onset. Above $\alpha \sim 0.2$, flocks form reliably (Fig. 3).

With all forces active and the default $\alpha = 1$, flock coherence is robust: Φ exceeds 0.99 up to noise amplitude $\text{ramp} = 3$, exceeds 0.97 at $\text{ramp} = 5$, and drops below 0.5 only at $\text{ramp} \sim 20$. The alignment force makes the system dramatically more resistant to noise disruption than the repulsion-only case.

4.3 Nature of the Solid-to-Fluid Transition

In the repulsion-only system ($\alpha = 0$, $v_0 = 0$), kinetic energy rises with noise amplitude, suggesting a transition from a solid-like crystalline state to a fluid-like disordered state. To test whether this constitutes a true phase transition, I performed finite-size scaling across $N = 25, 50, 100$, and 200 (Fig. 4).

A true phase transition would produce KE/N curves that depend on N , with a critical point (susceptibility peak) that converges to a finite η_c as N increases. Instead, the KE/N curves are essentially identical for all four system sizes, and the susceptibility $\chi = N * \text{var}(\text{KE}/N)$ increases monotonically with η with no peak.

This indicates a smooth crossover rather than a true critical phenomenon. The physical picture is consistent with the high compactness of the system ($C = \pi N r_0^2 \sim 0.78$ for the parameters used): each agent is effectively caged by its neighbors and oscillates harmonically around a fixed lattice site. This produces KE proportional to η^2 , independent of N — behavior characteristic of uncoupled harmonic oscillators, not a correlated system approaching criticality.

Fixed-compactness scaling. The original finite-size scaling held r_0 fixed, meaning compactness $C = \pi N r_0^2$ grew with N ($C = 0.196$ for $N = 25$ up to $C = 1.57$ for $N = 200$). To isolate the effect of compactness, I repeated the analysis holding C fixed by scaling $r_0 = \sqrt{C/(\pi N)}$ for each N . *Testing both $C = 0.78$ (dense) and $C = 0.10$ (dilute), the result in both cases is the same: KE/N curves are N -independent and the susceptibility $\chi = N * \text{var}(\text{KE}/N)$ increases monotonically to the*

top of the sweep ($\eta = 30$) with no peak at finite η . The KE/N values at the two compactness levels are also nearly identical to each other (Fig. 9).

The crossover is therefore not an artifact of high compactness alone. In the dilute regime, agents barely interact (mean spacing exceeds repulsion range), so they behave as essentially independent random walkers; KE/N is set solely by the noise amplitude and is N -independent for the same reason as in the dense case — just via a different physical mechanism. Both extremes suppress cooperative behavior: too dense means agents are caged; too dilute means they never interact enough to form a solid. A genuine critical point would require an intermediate regime where a solid phase can form and agents can rearrange cooperatively. Whether such a regime exists in this model at compactness values between 0.10 and 0.78, or at noise amplitudes above $\eta = 30$, remains an open question.

4.4 Predator-Prey Dynamics

I extended the model with a predator agent that chases the prey center of mass via a strong alignment force ($\alpha_{\text{pred}} = 5$) and generates a long-range repulsive force on nearby prey ($r0_{\text{pred}} = 0.1$). Prey parameters were set to the slow-walking regime ($v0 = 0.02$, $\alpha = 1.0$, $\text{ramp} = 0.1$) to match the concert crowd context from Silverberg et al.

Flock coherence under pressure. Comparing flocking prey ($\alpha = 1$) versus non-flocking prey ($\alpha = 0$) across 10 random initializations shows a striking divergence. Flocking prey maintain $\Phi \sim 0.998$ throughout the simulation despite continuous predator pressure. Non-flocking prey scatter almost immediately, reaching $\Phi \sim 0.096$ in steady state (Fig. 5). Non-flocking agents maintain marginally more individual distance from the predator (0.127 vs. 0.112), but they lose all collective structure. The flock absorbs the disturbance while remaining coherent.

Evasion distance saturates. Sweeping predator aggression α_{pred} (which sets effective predator speed as $v_{\text{eq,pred}} = v0_{\text{pred}} + \alpha_{\text{pred}}/\mu_{\text{pred}}$) reveals that the mean predator-to-nearest-prey distance drops from 0.24 with a passive predator to ~ 0.10 for $\alpha_{\text{pred}} \geq 1$ and then saturates — the collective repulsion response establishes a minimum buffer distance that persists regardless of predator aggression.

Flock geometry. The flock is not just a point moving through space; its shape matters. Without a predator, the steady-state aspect ratio is $AR = 2.61$ and radius of gyration $R_g = 0.215$. With a predator, these shift modestly to $AR = 2.76$ and $R_g = 0.274$. Stronger flocking (larger α) produces substantially more elongated flocks: AR increases from 2.09 at $\alpha = 0.2$ to 7.27 at $\alpha = 2.0$. These highly elongated configurations resemble the arched, thinning flocks predicted by Charbonneau Exercise 6 (Fig. 6).

Multiple predators. With 1 to 4 simultaneous predators, flock order parameter stays near 0.975–0.991 — coherence is maintained across the entire range (Fig. 7). Aspect ratio rises substantially with predator count ($AR = 2.83$ with one predator, $AR = 7.91$ with three), while R_g increases modestly. Counterintuitively, the minimum predator-to-prey distance increases slightly as the number of predators grows (0.093 for one predator, 0.106 for three). The flock under multiple predators is more elongated but no more accessible to any individual predator.

Why evasion distance increases with predator count. To diagnose the counterintuitive evasion result, I measured the predator-predator separation and the orientation of the flock major axis relative to the predator centroid direction across 8 random initializations (evasion_analysis.py). The predator-predator mean distance was approximately 0.001 for all multi-predator runs — effectively zero. Because every predator independently targets the flock center of mass using the same rule,

they all converge to the same location and pile up on top of one another. This co-localization means multiple predators do not approach from different directions; they compete for the same point. Flock orientation relative to the predator centroid was 43-46 degrees across all conditions — consistent with the 45-degree expectation for random alignment — confirming that the flock does not systematically orient its narrow or broad side toward the predator. The evasion distance improvement therefore arises mechanically: co-localized predators collectively apply repulsion from a single point, and this concentrated repulsion is stronger than what a single predator can produce, pushing the flock slightly farther away. This is a model-specific artifact of the naive “chase CoM” predator strategy: multiple independent pursuers using identical rules undermine each other rather than coordinating (Fig. 8).

4.5 Predator Coordination and Encirclement

The naive-predator result (Section 4.4) raised an obvious question: if multiple predators fail to disrupt the flock because they co-localize, what happens when they are forced apart? I added an explicit predator-predator repulsion with strength `alpha_coord` and ran three experiments (`coordinated_predators.py`, 8 seeds). At low coordination (`alpha_coord` < 5) the predators still pile up at the prey center of mass — the shared target overwhelms the mutual repulsion. Once `alpha_coord` ≥ 5 a real separation emerges (mean predator-predator distance rises to 0.14 at `alpha_coord` = 5 and 0.29 at `alpha_coord` = 20). The predators are spatially distributed and on average slightly closer to the flock (`min_dist` drops from 0.105 to ~0.08). But the flock order parameter never falls below 0.92 across all tested predator counts and coordination strengths. The collective evasion response scales with the number of approaching predators: the flock just builds a slightly tighter shell.

I then tested a fundamentally different strategy: encirclement. Each predator k is assigned a fixed compass angle $\theta_k = 2\pi k / n_{\text{pred}}$ and chases the point $\text{CoM} + R_{\text{enc}}(\cos \theta_k, \sin \theta_k)$ rather than the CoM itself. With n_{pred} predators this places them at n_{pred} equally spaced angles around the flock at radius R_{enc} . A radius sweep at $n_{\text{pred}} = 3$ identifies $R_{\text{enc}} = 0.15$ as the optimal offset — just inside the typical flock radius $R_g \sim 0.25$. Smaller R_{enc} and the predators still co-localize; larger R_{enc} and they orbit outside the flock entirely. At $n_{\text{pred}} = 6$, the encirclement strategy achieves $\Phi = 0.769 \pm 0.093$ — the first substantial coherence reduction across all predator experiments. Neither naive nor coordinated predators had ever pushed Φ below 0.92.

4.6 Fragmentation Mechanism

To understand what the Φ drop under encirclement physically represents, I added connected-components clustering (`fragmentation.py`): two agents are placed in the same cluster if they are within the flocking radius r_f of one another. A naive multi-predator setup gives $\Phi = 0.997$ with 60 clusters — many small spatial groups all moving in the same direction. The low cluster count of a normal flock is somewhat misleading: it reports spatial structure, not directional structure. Under encirclement, the picture inverts: cluster count drops to 24 (FEWER, LARGER groups) but global Φ falls to 0.72. Crucially, each individual cluster has local order parameter ~0.997 — internally coherent. The flock has fragmented into a small number of sub-flocks each heading in a different direction. As n_{pred} increases, cluster count decreases further and the largest cluster grows: predators compress the flock spatially while splitting it directionally. This is flock DIVISION, not DISSOLUTION — a herding effect analogous to wolf-pack predation or dolphin bait-ball formation, achieved purely through multi-directional pressure without any explicit cooperative

behavior between predators.

Scaling the protocol with N (encirclement_scaling.py) finds no simple law. At fixed $n_{\text{pred}} = 6$, larger flocks are more resistant (Φ rises from 0.69 at $N = 50$ to 0.90 at $N = 350$) — the dilution effect from Section 4.4 partially protects them. But at fixed predator-to-prey ratio ($n_{\text{pred}} = 9$, $N = 500$), Φ falls to 0.654 — the worst result in the entire project. The relevant variable is not predator count, not predator-to-prey ratio, but ANGULAR COVERAGE. Both $N = 100$ and $N = 350$ converge to a common floor $\Phi \sim 0.67$ at $n_{\text{pred}} \sim 10$, suggesting that ten equally spaced angles is sufficient to overwhelm the alignment force regardless of flock size.

4.7 Reversibility of Encirclement Damage

Finding that encirclement causes flock division raised an immediate follow-up question: is the division permanent? On a periodic torus, sub-flocks heading in different directions inevitably re-encounter each other; the alignment force would then re-couple them. To test this I ran a three-phase protocol (reunion.py): a 1500-step pure-flock warm-up, a 4000-step attack with $n_{\text{pred}} = 10$ encircling predators, and a 6500-step recovery with predators removed. During the attack, Φ falls to 0.72 with 4.5 clusters and largest cluster fraction 0.41 — substantial fragmentation. Immediately after predator removal, the sub-flocks merge. Across 6 seeds, every single run recovers to $\Phi \geq 0.95$ within 4.5 to 16.5 time units (mean 10.3) — much shorter than the 4000-step attack that caused the damage. By the end of the recovery window, Φ has settled to 1.000 ± 0.001 with a single coherent cluster.

This result establishes that encirclement is a fundamentally REVERSIBLE form of disruption. The flock’s topological state — a single connected population — is preserved through the attack and re-emerges as soon as the stressor lifts. The damage is purely kinematic, not structural.

4.8 Minimum Viable Flock Size

The dilution result of Section 4.4 implied that larger flocks are more resistant, but the smallest tested case was $N = 10$. Below what N does collective evasion fail entirely? A sweep from $N = 3$ to $N = 100$ (min_flock_size.py, 8 seeds) shows that flock formation itself is the binding constraint: in the no-predator control, $\Phi = 0.49$ at $N = 3$, rises to 0.69 at $N = 8$, crosses 0.9 between $N = 18$ and $N = 25$, and reaches 0.99 at $N = 100$. Below $N \sim 12$ the flock is unreliable (std across seeds 0.13-0.20). Predator presence does not substantially shift this threshold: with a naive predator, $\Phi(N)$ follows the same curve. With two opposing encircling predators, very small flocks ($N = 3-8$) are slightly stabilized — the predator pressure forces nominal alignment — but the underlying coherence threshold remains near $N \sim 18-25$. The interpretation is that flocking is a collective effect with a minimum participant count: below ~ 12 agents in a unit domain, the local spatial density is too sparse for the flocking force to dominate noise. Real prey populations near this threshold should be most vulnerable to predators, though the absolute capture rate in this simulation remains low across all sizes.

4.9 Sensing-Limited Predators

Up to this point all predators have had perfect global knowledge of the flock center of mass. To test the more biologically realistic case of bounded perception, I added a sensing radius r_{sense} within which the predator can locate the nearest prey (predator_sensing.py). When the nearest prey is within r_{sense} , the predator chases normally; outside r_{sense} it executes a slow random walk. Sweeping r_{sense} from 0.05 to infinity produces a sharp transition near $r_{\text{sense}} \sim \text{flock radius}$

Rg (0.10-0.15). Below the threshold the lock-on fraction drops to 12% and the predator only finds the flock occasionally; above $r_sense = 0.20$ the lock-on fraction reaches 100% and the result is indistinguishable from perfect sensing. For single naive predators, bounded perception makes no difference whenever $r_sense > 0.20$. Interestingly, for encirclement with $n = 6$ predators, limited sensing actually **WORSENS** the outcome for the flock ($\Phi = 0.79$ vs 0.85): when an encircling predator loses the flock and re-enters search mode, its subsequent re-approach is from a random angle, adding unpredictable multi-directional pressure on top of the structured encirclement pattern.

4.10 Internal Stressors: Panic and Contagion

Charbonneau Section 10.5 (“Why You Should Never Panic”) introduces agents with reduced flocking and elevated noise — panicked agents that should disrupt the surrounding crowd. I implemented this directly (panic.py) by labeling a fraction f of agents as panicked ($\alpha_panic = 0.1$, $ramp_panic = 10$) and sweeping f from 0 to 20%. The global order parameter drops smoothly: $\Phi = 1.000$ at $f = 0\%$ to $\Phi = 0.853$ at $f = 20\%$. But this drop is pure dilution. Computing the order parameter using only the calm agents reveals that Φ_{calm} remains at 0.999 across the entire sweep. The calm flock effectively ignores its panicked neighbors — the alignment force is strong enough that calm agents maintain coherence even when 1 in 5 of their neighbors is moving erratically.

A natural objection is that real panic is contagious: panicked agents make their neighbors panic too. I added this mechanism (panic_contagion.py). Each calm agent within a contagion radius $r_cont = 0.05$ of a panicked agent transitions to the panicked state with probability per timestep $p = 1 - \exp(-\beta * k * dt)$ where k is the local count of panicked neighbors and β is the contagion rate. Without recovery (an SI dynamics), any non-zero β drives the entire flock to $f = 1.0$ — saturation occurs in roughly $4/\beta$ time units. With $f = 1$, all agents have the panicked alignment and noise parameters, and the global Φ collapses to ~ 0.10 . The “calm-flock immunity” from panic.py was an artifact of treating panic as a fixed label: once the pool of calm agents can be drained, it always is. The book’s claim that panic is collectively dangerous is recovered, but only conditional on contact-mediated propagation. Whether a true epidemic threshold β_c exists in an SIS model with recovery ($calm \leftrightarrow panic$) is addressed below.

4.11 SIS Contagion and the Epidemic Threshold

Adding a recovery rate γ (panicked agents return to calm with rate γ per unit time) restores the classical SIS phase structure (contagion_sis.py). A 2D sweep over (β , γ) reveals a clean threshold along the diagonal $\beta = \gamma$. Below this line, the seed outbreak dies out and the flock stays at $\Phi = 1.0$. Above the line, an endemic steady state emerges and the flock degrades smoothly with β/γ . At $\beta = 2$, varying γ from 0.1 to 10 takes the system from $f_{ss} = 0.978$ (outbreak persists, $\Phi = 0.20$) all the way down to $f_{ss} = 0.000$ (recovery wins, $\Phi = 1.0$). The crossover between $\gamma = 2$ and $\gamma = 5$ is sharp.

The location of the threshold $\beta_c \sim \gamma$ is consistent with the standard mean-field SIS prediction $\beta_c * = \gamma$ when the effective local contact count is ~ 1 . For a uniform-density flock at $N = 350$ with $r_cont = 0.05$, the geometric estimate gives $= \pi * r_cont^2 * N \sim 2.7$, but the flock is a spatially extended structure whose effective contact rate at the contagion scale is closer to one.

This result reframes Finding 4.10. The SI claim that “any contagion is fatal” was specifically a no-recovery limit ($\gamma = 0$). With any finite recovery rate, the flock can in principle contain a panic outbreak provided the contagion rate is below the recovery rate. The flock’s own alignment

force may play the role of an effective γ in a fully physical model: panicked agents re-entering a coherent neighborhood feel the alignment force pulling them toward the local mean velocity, which corresponds to a return-to-calm process. In this sense, the kinematic flock-disruption problem maps onto the classical contact-process problem of epidemiology via a single dimensionless ratio β/γ .

4.12 Hybrid Stressors and Active/Passive Mixing

Two additional experiments tested how the various disruption mechanisms compose.

In the hybrid-stressor experiment (`hybrid_stressors.py`), I ran four conditions side-by-side: no stressor, encirclement only, contagion only, and both. The encirclement-only condition gave $\Phi = 0.71$ (matching the Section 4.6 result). Contagion-only gave $\Phi = 0.05$. Combined, $\Phi = 0.05$ — identical to contagion alone. Encirclement adds nothing because the population is fully panicked before encirclement-induced fragmentation matures. The two disruption modes do not compose in the supercritical contagion regime: the absorbing process simply wins.

In the segregation experiment (`segregation.py`), I tested whether mixed-speed populations spatially segregate. With $f_{\text{active}} = 0.5$ and v_0 contrast ranging from 0 (no contrast) to 0.9 (active 10x faster than passive), the segregation index — defined as the along-heading position difference between groups in R_g units — stays at 0 ± 0.05 in every case. The alignment force homogenises group speed: an aligned flock with mixed self-propulsion targets cruises at a population-weighted compromise speed, so individual agents do not segregate into leading and trailing layers. To obtain spatial segregation in this model would likely require asymmetric α values between groups, not v_0 contrast alone.

A follow-up experiment (`segregation_alpha.py`) confirms this. Replacing v_0 contrast with α (alignment-strength) contrast produces real spatial segregation — but in a form the along-heading index still misses. A local-purity diagnostic, defined as the fraction of an agent’s rf -neighbors that share its type, rises from 0.50 (random, no contrast) to 0.73 at maximum contrast. The snapshot at $\alpha_{\text{passive}}=0$ shows tight clusters of high- α agents amid scattered low- α particles. The segregation is isotropic in the flock frame, not preferentially at the leading or trailing edge. Charbonneau’s segregation result is therefore recovered, conditional on asymmetric alignment fidelity rather than asymmetric speed, and the appropriate diagnostic is local same-type purity rather than bulk position relative to heading.

4.13 Refinements: Sub-Threshold Coupling and Large-N Scaling

Two further experiments refine results from earlier sections.

The hybrid SIS experiment (`hybrid_sis.py`) tests whether encirclement can push a contained SIS contagion over its critical threshold by raising the local contact count. At $\beta = 1.0$, $\gamma = 3.0$ ($\beta/\gamma = 0.33$, well below the Section 4.11 threshold), contagion alone fizzles: panic peaks at $f_{\text{max}} = 0.13$ and dies out, $f_{\text{ss}} = 0$. Adding 6 encircling predators triples the local contact count (k from 8.9 to 30.2) and doubles the panic peak ($f_{\text{max}} = 0.27$), but the outbreak still dies out. The mechanism is confirmed — compression raises effective β — but the amplification is bounded: well-subcritical contagion is not rescued by external mechanical pressure alone. The combined-stressor Φ (0.73) is measurably worse than encirclement alone (0.86), so the transient outbreak adds non-negligible kinematic disruption even when it eventually collapses.

The large-N encirclement experiment (`large_N_encirclement.py`) tests whether the Section 4.6

conjecture that the $\Phi \sim 0.67$ floor at $n_{\text{pred}} = 10$ is N -independent holds at very large flock sizes. At $N = 350, 700, 1000$ with $R_{\text{enc}} = 0.15$ held fixed, the floor at $n_{\text{pred}} = 10$ rises with N : $\Phi = 0.667$ ($N = 350$), 0.700 ($N = 700$), 0.740 ($N = 1000$). At $n_{\text{pred}} = 14$: $0.637, 0.727, 0.790$. The encirclement strategy becomes less effective at large N because R_{enc} was calibrated to $N = 350$'s flock radius $R_g \sim 0.15$; at $N = 1000$ the flock is broader and $R_{\text{enc}} = 0.15$ places the predators inside the flock rather than at its boundary, where their repulsion is absorbed by surrounding prey. The encirclement geometry must match the flock geometry to disrupt — angular coverage alone is insufficient. A predator strategy that works at intermediate flock sizes does not scale up to large prey aggregations without re-calibration.

4.14 Near-Critical Coupling and Universal Encirclement Scaling

Three further experiments tighten the quantitative results above.

To directly measure the threshold shift hypothesised in Section 4.13, a sweep of β at fixed $\gamma=2.0$ was repeated with and without 6 encircling predators (`critical_shift.py`). The threshold β_c (where f_{ss} crosses 0.5) was $\beta_c=1.93$ without encirclement, 1.85 with — a leftward shift of 0.077 , about 4% of the bare value. The shift is largest just below threshold: at $\beta=1.5$, f_{ss} rises from 0.34 (no encirclement) to 0.43 (with). Above threshold the two curves converge. The modest size of the shift, despite a tripling of the local contact count under encirclement (Section 4.13), is explained by redundancy: compressed sub-flocks have panicked agents mostly surrounded by other panicked agents, so the new contacts are not new infections. Encirclement is therefore a near-critical amplifier rather than a general one; it can tip a contagion that is already close to its threshold but cannot bridge a large gap.

A separate sweep tests the flock's resilience to immune sub-populations (`herd_immunity.py`). At supercritical SIS ($\beta=2.5$, $\gamma=2.0$; $R_0=1.25$), mean-field theory predicts a herd-immunity threshold $p_c = 1 - 1/R_0 = 0.20$. The measured threshold in the flock model is $p_c \sim 0.46$, more than twice the mean-field value. The cause is spatial clustering: panicked sub-clusters move together so their members predominantly contact each other, inflating the effective local contact count within clusters. Random immunity is then less effective than targeted immunity at breaking transmission chains. This is a known feature of spatial epidemic models, and the flock model reproduces it.

Finally, the apparent N -dependence of the encirclement floor (Section 4.13) is resolved by sweeping R_{enc} at both $N=350$ and $N=1000$ (`renc_scaling.py`). When plotted against R_{enc} / R_g the two Φ curves collapse: both show optimal disruption at $R_{\text{enc}} / R_g \sim 0.5$, with $\Phi_{\text{min}} = 0.67$ ($N=350$) and 0.73 ($N=1000$). The previous N -dependence was almost entirely an R_{enc} / R_g mismatch — R_{enc} was tuned to $N=350$'s flock radius and was undersized for $N=1000$. With proper rescaling encirclement is size-invariant. The strategy that wraps the flock at half its radius works universally. The optimum at $R_{\text{enc}} / R_g \sim 0.5$ places each predator within the bulk of the flock but not at the center, where it can push prey on all sides toward neighboring predators — the kinematic geometry that drives the flock division of Section 4.6.

4.15 Long-time Dynamics and Incomplete Encirclement

Two further experiments probe the limits of the encirclement strategy.

All encirclement experiments to this point ran for 4000 timesteps (40 time units). To test whether the fragmented state is a true steady state or a transient on the way to a different long-time attractor, I ran 30000 steps (300 time units) under constant encirclement (`long_encirclement.py`,

4 seeds, $n_{\text{pred}}=10$, $R_{\text{enc}}=0.15$). The result is neither sustained fragmentation nor recovery: instead, Φ oscillates continuously. Within a single run, Φ excursions span from ~ 0.4 to ~ 0.95 , with a temporal standard deviation of ~ 0.21 per seed. The seed-to-seed variation is only 0.061. The temporal fluctuations within each run dominate the between-seed variability, indicating that the dynamics are intrinsically intermittent rather than converging to a steady state. Cluster count and largest-cluster fraction oscillate in concert.

The mechanism is a positive feedback cycle. When the flock briefly re-coalesces ($\Phi \rightarrow 0.9$), all predators chase one CoM target offset by their respective angles — maximum multi-directional pressure — and the flock quickly fragments again. Once fragmented ($\Phi \rightarrow 0.5$), the predators distribute among multiple sub-flocks rather than all pressuring one, reducing the effective attack. The dominant sub-flock can reassemble. The cycle repeats indefinitely; the long-time average $\Phi = 0.751$ matches the short-time value of 0.72 from Section 4.6, but the short-time measurement was misleadingly quiet.

To test whether leaving a gap in the encirclement creates an escape route, I ran a complementary experiment (`encirclement_gap.py`). Starting from a full 6-predator ring (60-degree spacing), I removed predators one at a time and measured the flock’s Φ and whether its center-of-mass drift aligned with the gap direction. The result is non-monotone and the gap direction is irrelevant. Full encirclement ($n_{\text{active}}=6$) gives $\Phi = 0.918$. Removing one predator to create a 120-degree gap drops Φ to 0.833 – WORSE than the full ring. Removing two or three predators raises Φ back to 0.91-0.96. The flock center-of-mass drift direction has no systematic alignment with the gap (mean angular difference 70-110 degrees from gap orientation), confirming that agents have no global awareness of where the ring is open.

The non-monotone result arises because one-predator removal creates a persistent asymmetry: the five remaining predators keep re-encircling the shifted CoM, generating irregular multi-directional pressure that is more disruptive than balanced 6-fold pressure. Removing two or more predators reduces the angular complexity and lets the flock partially escape. Counterintuitively, a near-complete encirclement with a single gap is harder to escape than a full encirclement, and both are harder to escape than a 3- or 4-predator arc.

4.16 Outbreak Persistence After Predator Removal

Section 4.7 showed that pure encirclement damage reverses within ~ 10 time units once predators are removed. Section 4.12 showed that supercritical SI contagion dominates and eliminates the encirclement effect entirely. A regime between these extremes exists: a contagion that is below the bare epidemic threshold ($\beta = 1.5$, $\gamma = 2.0$, $\beta/\gamma = 0.75 < 1.93$) but is pushed into a transient endemic state by the compression from encirclement. Does removing the predators allow both the kinematic damage AND the contagion to reverse, or does the contagion outlast the kinematic stressor?

Running the three-phase protocol (`outbreak_removal.py`: warmup $\rightarrow$ 4000 steps of encirclement + SIS $\rightarrow$ 5000 steps of recovery) answers this directly. During the encirclement phase, local contact count rises by $\sim 3\times$ (Section 4.13) and the effective transmission exceeds the recovery rate, establishing an endemic panic fraction of $f = 0.450$ and suppressing Φ to 0.185. When predators are removed, kinematic recovery begins – Φ rises to 0.266 within 50 time units – but the contagion does not collapse. After 50 time units without predators, $f = 0.413$, barely changed from its peak. The timescales are starkly asymmetric: kinematic damage heals in ~ 10 time units (Section 4.7); contagion-driven damage persists for hundreds of time units as the system slowly evolves back

toward the bare endemic state (~ 0.34 panic fraction at these parameters).

The contrast illuminates the qualitative difference between the two damage types introduced in Section 4.10. Pure encirclement damage is kinematic: sub-flocks that were pushed apart simply re-merge on the periodic torus when the directional pressure disappears. SIS contagion damage is a population state: the distribution of panicked vs. calm agents evolves on epidemic timescales set by β and γ , not on kinematic timescales set by flock size and swimming speed. A predator group that seeds a panic cascade even transiently — even one that would not sustain itself in an unencircled flock — inflicts damage that outlasts the predation event itself by an order of magnitude.

4.17 Adaptive Encirclement: Tracking Flock Geometry Increases Disruption (Finding 35)

Finding 31 established that encirclement performance collapses on the dimensionless ratio R_{enc}/R_g and that the optimum is at $R_{\text{enc}}/R_g \sim 0.5$ for both $N = 350$ and $N = 1000$. That finding was obtained with a fixed R_{enc} set at initialization. A natural follow-up is whether a predator group that continuously tracks flock size and adjusts R_{enc} accordingly can sustain that optimum across the long-time merge/split dynamics (Finding 32), where R_g fluctuates substantially.

I compared fixed $R_{\text{enc}} = 0.150$ against adaptive $R_{\text{enc}} = 0.5 * \text{live_}R_g$, where R_g is recomputed from all prey positions at every timestep. Both conditions used $n_{\text{pred}} = 10$, $N = 350$, slow prey, over 15000 steps (150 time units), with 5 seeds (`adaptive_encirclement.py`).

The adaptive strategy is more disruptive across all metrics:

Condition	mean Phi	frac time Phi > 0.85	mean R_{enc}/R_g
Fixed	0.778	0.56	0.485
Adaptive	0.713	0.37	0.500

Mean coherence falls from 0.778 to 0.713 (an 8% reduction), and the fraction of time the flock spends in a highly coherent state ($\text{Phi} > 0.85$) drops from 56% to 37% — a 34% relative reduction. The mean R_{enc}/R_g for the fixed condition is 0.485, only slightly below the 0.5 target, but during the merge/split fluctuations of Finding 32 the instantaneous ratio drifts; adaptive removes those deviations. The lower temporal standard deviation for adaptive (0.219 vs 0.233) confirms that the fluctuations into high-Phi recovery states are suppressed.

The mechanism follows directly from Finding 32’s dynamics. During consolidation phases (sub-flocks merging, R_g shrinking), fixed R_{enc} becomes too large relative to the now- smaller flock, placing predators too far from the bulk and reducing their effectiveness. Adaptive R_{enc} shrinks with the flock, maintaining maximum directional pressure exactly when the flock is most vulnerable to being re-fragmented. The combined effect cuts the dwell time in coherent configurations by a third.

The effect size is modest, consistent with Finding 31’s observation that the performance function is relatively flat near the optimum: both conditions achieve similar mean disruption, and the fixed condition was already close to optimal. The primary advantage of adaptation is not unlocking a qualitatively different regime but rather eliminating the off-optimum excursions that arise from the natural fluctuations in a dynamical flock.

One limitation: the adaptive strategy uses global R_g (over all prey), which inflates during fragmented phases as scattered sub-flocks span the domain. This causes adaptive to briefly overshoot R_{enc}

during high-fragmentation states, slightly limiting its advantage. A predator group tracking the largest individual sub-flock’s R_g would be more effective still.

4.18 Targeted Vaccination Provides No Advantage Over Random: The Flock Contact Network Is Not Hub-Dominated (Finding 36)

Finding 30 showed that the herd-immunity threshold in the flock is approximately twice the mean-field prediction (~ 0.46 vs. 0.20 at $R_0 = 1.25$), attributed to spatial clustering of panicked sub-groups. A natural follow-up is whether targeting high-connectivity agents for immunity can reduce this inflated threshold. On a scale-free network (power-law degree distribution), immunizing hub nodes first can dramatically lower the required immune fraction because hubs are disproportionately responsible for spreading contagion. If the flock contact network has hub nodes, targeting them should outperform random vaccination.

I compared two vaccination protocols at the same supercritical SIS parameters as Finding 30 ($\beta = 2.5$, $\gamma = 2.0$, $R_0 = 1.25$), across $p_{\text{immune}} = 0$ to 0.70 and 6 seeds (targeted_immunity.py). In the targeted condition, each seed’s settled flock was analyzed for contact degree (number of agents within $r_{\text{cont}} = 0.05$), and the top p_{immune} fraction by degree was immunized before the first panicked agent was seeded.

The result is negative: targeted vaccination provides no statistically significant advantage over random vaccination at any tested immune fraction:

p_{immune}	f_{ss} (random)	f_{ss} (targeted)	difference
0.00	0.593 ± 0.008	0.587 ± 0.006	-0.006
0.10	0.491 ± 0.010	0.495 ± 0.006	$+0.004$
0.20	0.393 ± 0.008	0.395 ± 0.010	$+0.002$
0.30	0.304 ± 0.013	0.301 ± 0.014	-0.003
0.40	0.204 ± 0.016	0.177 ± 0.080	-0.027
0.46	0.101 ± 0.072	0.084 ± 0.065	-0.016
0.50	0.046 ± 0.066	0.043 ± 0.053	-0.003
0.60–0.70	0.000	0.000	0.000

Differences at $p_{\text{immune}} \leq 0.30$ are within noise ($|\text{diff}| \leq 0.006$). Near the threshold ($p = 0.40$ – 0.46), targeted shows a small numerical edge, but the variance in the targeted condition is 4–5x higher than in the random condition ($\text{std} = 0.080$ vs. 0.016 at $p = 0.40$), and the difference is not statistically robust. Both strategies cross the quench threshold ($f_{\text{ss}} < 0.1$) at $p_{\text{immune}} \approx 0.46$, matching the random threshold from Finding 30.

The contact degree distribution across all seeds has mean = 9.02, median = 8, std = 6.17, and max = 31. This is moderately heterogeneous (coefficient of variation $CV = 0.68$) but the maximum degree is only 3.4 times the mean — far below the ratios of 10–1000 seen in social and internet networks where hub-targeting is effective. The second panel of the figure (targeted_immunity_1.png) shows the degree distribution; it is approximately unimodal and bounded, not power-law.

There is a deeper reason why degree-targeting fails even given moderate heterogeneity: the flock contact network is spatially embedded and kinematically reconfigurable. Hub agents are high-degree because they occupy the dense interior of the flock, where many neighbors are in proximity. When those agents are immunized, the alignment and repulsion forces redistribute the remaining agents, and new interior positions form. The network topology reorganizes to restore a similar degree structure even with the specific hub agents removed.

This result clarifies the mechanism behind Finding 30’s inflated threshold. The 2x inflation arises from spatial co-movement of panicked agents (sub-clusters that travel together have redundant contacts with each other), not from degree heterogeneity. These are distinct properties: degree-targeted vaccination addresses the latter but not the former. A strategy that disperses immune agents spatially across the flock extent, rather than concentrating them in the high-degree core, would more directly disrupt spatial co-clustering and might outperform random vaccination — this hypothesis is tested in Section 4.19.

4.19 Repulsion Hardness Does Not Rescue the Phase Transition: A Non-Equilibrium Diagnosis (Finding 38)

Finding 17 showed no diverging susceptibility at any tested compactness value ($C = 0.10$ to 0.78) using the standard soft repulsion (exponent $n = 1.5$). The natural follow-up question is whether a harder repulsion potential would produce a true phase transition. In equilibrium statistical mechanics, 2D hard discs at compactness $C \sim 0.40$ form a hexagonal solid at low effective temperature and undergo the well-known KTHNY melting transition at a finite critical temperature, so a hard-core potential should, in principle, produce the diverging susceptibility that the soft model lacks.

I tested this by sweeping the repulsion exponent $n = 1.5, 3.0, 6.0, 12.0$ in a repulsion-only simulation (no flocking, no self-propulsion) with finite-size scaling at $N = 25, 50, 100, 200$, $C = 0.40$, 8 seeds per point, and the same noise sweep $\eta = 0.5$ to 30 as in Finding 17 (`hard_repulsion.py`).

The result is unambiguous: the chi-peak (susceptibility $\chi = N * \text{Var}(\text{KE}/N)$) falls at $\eta = 30$ — the top of the sweep — for every combination of exponent and system size. The KE/N curves are essentially identical across $n = 1.5$ to $n = 12$:

Exponent n	$N = 25$ chi peak	$N = 50$ chi peak	$N = 100$ chi peak	$N = 200$ chi peak
1.5	2607	7311	3785	3858
3.0	2609	7305	3766	3854
6.0	2615	7285	3770	3851
12.0	2615	7289	3773	3853

All peaks at $\eta = 30$; no N -dependence of the peak location; chi values nearly identical within each N column across all exponents.

The absence of a phase transition at $n = 12$ (near hard-core) requires a different explanation from softness. The root cause is the non-equilibrium nature of the driving. The current model uses uniform random kicks (each velocity component $\pm \eta * U[-1,1]$ at every timestep), which do not satisfy the fluctuation-dissipation theorem (FDT). The only velocity-limiting mechanism in the repulsion-only variant is confinement by neighboring agents; there is no viscous friction ($-\mu * v$) to produce true thermal equilibration. Without FDT, there is no well-defined temperature, the system cannot reach the Boltzmann distribution in configuration space, and the cooperativity required for a phase transition cannot emerge. The crossover is instead set by the competition between random kick energy and the scale of positional confinement, which depends on the force range $2r_0$ but not on the force profile (exponent n). This explains the complete insensitivity to n .

To observe the hard-disc phase transition in this model family, one would need to replace the self-propulsion speed regulator with genuine Langevin dynamics: viscous damping $F_{\text{damp}} = -\mu * v$ and noise amplitude $\sqrt{2\mu kT/dt} * \text{randn}$, which together satisfy FDT and recover the

Boltzmann distribution in the long-time limit. The KTHNY transition would then appear at the appropriate area fraction and temperature.

This result clarifies the relationship between Findings 2, 8, 12, 17, and the present one: the smooth crossover is a property of the entire model class (self-propulsion regulation + independent random kicks), not of any particular parameter regime. No modification of the repulsion potential within this model class will produce a true phase transition.

4.20 Spatial Vaccination Also Fails: Kinematic Mixing Defeats All Targeting Strategies (Finding 37)

Section 4.18 showed that degree-targeted vaccination — immunizing the highest-contact-degree agents first — provides no advantage over random vaccination because the flock contact network lacks the fat-tailed heterogeneity required for hub-targeting to work, and kinematic reorganization restores hub positions after high-degree agents are immunized. That result identified the true mechanism behind the 2x mean-field herd-immunity inflation as SPATIAL CLUSTERING of panicked subgroups (agents within a panicked cluster contact primarily each other, creating localized transmission chains). This suggests a different targeting hypothesis: geographically distribute the immune agents to break the clustering, rather than targeting high-degree individuals.

To test this, I compared three vaccination strategies (`spatial_vaccination.py`, $\beta = 2.5$, $\gamma = 2.0$, $R_0 = 1.25$, 5 seeds, same parameters as Findings 30 and 36):

- **Random:** immune agents chosen uniformly at random (baseline)
- **Spatial:** immune agents chosen by farthest-point (maxmin) sampling — each successive agent is the one farthest from all already-selected agents on the torus, maximizing spatial coverage
- **Targeted:** highest-contact-degree agents first (Finding 36 reference)

p_immune	f_ss random	f_ss spatial	f_ss targeted
0.00	0.594 $\pm$ 0.009	0.589 $\pm$ 0.006	0.586 $\pm$ 0.006
0.10	0.490 $\pm$ 0.010	0.492 $\pm$ 0.010	0.493 $\pm$ 0.006
0.20	0.391 $\pm$ 0.007	0.391 $\pm$ 0.013	0.398 $\pm$ 0.009
0.30	0.301 $\pm$ 0.011	0.297 $\pm$ 0.009	0.298 $\pm$ 0.013
0.40	0.202 $\pm$ 0.016	0.170 $\pm$ 0.085	0.168 $\pm$ 0.085
0.46	0.088 $\pm$ 0.073	0.125 $\pm$ 0.063	0.072 $\pm$ 0.064
0.50	0.056 $\pm$ 0.068	0.088 $\pm$ 0.045	0.052 $\pm$ 0.055
0.60	0.000 $\pm$ 0.000	0.000 $\pm$ 0.000	0.000 $\pm$ 0.000

All three strategies are statistically indistinguishable at every immune fraction. At $p = 0.10$ - 0.30 , mean f_{ss} values agree within ± 0.007 , well within the ± 0.009 - 0.013 standard deviations. Near the threshold ($p = 0.40$ - 0.50), the spatial and targeted strategies show substantially higher variance (std ~ 0.085) than random vaccination (std ~ 0.016 - 0.068), making their slightly lower mean values unreliable. At $p = 0.46$, the spatial strategy mean (0.125) is actually larger than random (0.088). This is a null result.

The mechanism is kinematic mixing. Immune agents are selected at warmup time ($t = 0$), before the SIS dynamics begin at $t = 1500$ steps (15 time units). The flock’s constant spatial reorganization under alignment, repulsion, and self-propulsion forces scrambles the initial agent positions between warmup and the epidemic run. By the time the epidemic begins, the spatial arrangement of immune agents is statistically uncorrelated with the warmup positions used for selection, making farthest-point spatial sampling equivalent to random in its practical effect on the contact network.

A related effect explains the higher variance of spatial and targeted strategies near the threshold. Random vaccination distributes immune agents uniformly in expectation, producing a stable average contact-network coverage independent of the epidemic seed location. Spatial and targeted vaccinations place agents in specific positions that may or may not overlap with the epidemic’s eventual spreading region — a source of additional variability with no accompanying mean improvement.

Taken together, Findings 36 and 37 establish that the flock’s kinematic reconfigurability defeats all static vaccination strategies. No agent-selection rule outperforms random vaccination in a continuously-mixing flocking collective, whether the selection criterion is contact-degree, spatial coverage, or random chance. The 2x herd-immunity threshold inflation is a structural property that cannot be exploited without real-time tracking and dynamic vaccination during the epidemic.

4.21 Langevin Thermostat Confirms FDT Diagnosis but KE/N Does Not Detect Structural Melting (Finding 39)

Section 4.19 concluded that the smooth crossover is caused by non-equilibrium driving, not by potential softness, and proposed that Langevin dynamics — viscous damping plus FDT-satisfying thermal noise — would recover the hard-disc phase transition. This section tests that prediction directly.

The Langevin simulation replaces the uniform random kick with viscous friction $F_{\text{damp}} = -\mu * v$ ($\mu = 10$) and thermal noise amplitude $\sqrt{2 * \mu * kT * dt} * N(0,1)$ per component per timestep. This satisfies the fluctuation-dissipation theorem: at thermal equilibrium the Maxwell-Boltzmann distribution is stationary, velocities satisfy equipartition ($KE_x/N = kT/2$ per component), and the system can in principle form an ordered phase at low kT that melts at a finite critical temperature. Two compactness values bracket the 2D hard-disc melting region: $C = 0.60$ (below the KTHNY transition) and $C = 0.70$ (at or above it). Parameters: $N = 25, 50, 100, 200$; $kT = 0.001$ to 5.0 (10 values); 8 seeds; $n = 1.5$ repulsion; $\mu = 10$ (langevin_repulsion.py).

The equipartition check confirms correct thermalization: KE/N at $kT = 0.1$ is 0.1047 - 0.1054 across all N and both compactnesses, within 1% of the target $kT = 0.100$. The Langevin thermostat functions as designed.

However, the susceptibility $\chi = N * \text{Var_seeds}(\text{time-avg } KE/N)$ still peaks at $kT = 5.0$ — the top of the sweep — for every combination of N and C . Moreover, both compactnesses give nearly identical χ values:

N	C = 0.60 χ_{peak}	C = 0.70 χ_{peak}	kT of peak
25	0.0817	0.0817	5.000
50	0.0406	0.0406	5.000
100	0.0507	0.0507	5.000
200	0.0440	0.0440	5.000

The identity of $C = 0.60$ and $C = 0.70$ χ values at $kT = 5.0$ is explained by the ratio of thermal to interaction energy: $\epsilon_p = 0.1$ while $kT = 5.0$, so the repulsion is only 2% of the thermal energy. Both compactnesses act as non-interacting gases at $kT = 5.0$, and their χ values reflect free-gas velocity fluctuations rather than structural correlations.

The root cause is that $\chi = N * \text{Var}(KE/N)$ is the wrong diagnostic for the KTHNY transition. KTHNY melting is a transition in positional order — the proliferation of disclination-antidisclination

pairs that destroys long-range hexatic bond-angle order (Kosterlitz and Thouless, 1973; Halperin and Nelson, 1978). This transition shows up in the hexatic order parameter $|\psi_6| = |\text{mean_neighbors} \exp(6i\theta)|$ and the bond-angle correlation function $g_6(r)$, not in kinetic energy fluctuations. Below KTHNY transition kT_c , $|\psi_6|$ is large (hexagonal positional order); above it, $|\psi_6|$ collapses. With χ based on KE/N , both the solid and fluid phases give small seed-to-seed variance in the time-averaged KE/N (all seeds agree on $KE/N \sim kT$), so the metric cannot distinguish phases.

This result refines the chain of findings on the phase-transition question. Finding 17 showed no transition in the original model. Finding 38 ruled out potential softness as the cause, identifying non-equilibrium driving as the culprit. The present finding confirms that Langevin dynamics thermalizes the system correctly (the FDT diagnosis was right) but demonstrates that the KE/N observable is insensitive to the structural melting transition. To observe KTHNY, two additions are required: (1) a positional order metric such as the hexatic order parameter, and (2) a harder repulsion to push the transition to a kT where equilibration is practical within the simulation timescale.

4.22 Hexatic Order Parameter Detects No Crystallization: Soft Repulsion Cannot Produce KTHNY Melting (Finding 40)

Section 4.21 identified that $\chi = N * \text{Var}(KE/N)$ cannot detect the KTHNY structural melting transition, and proposed the hexatic order parameter as the correct diagnostic. This section tests that proposal directly with the same Langevin simulation (`langevin_hexatic.py`), replacing KE/N measurements with a bond-angle analysis. At each measurement step, the bond angle θ_{jk} is computed for every pair of neighbors within distance $3r_0$, and

$$|\psi_{6,j}| = |(1/k_j) * \sum_{\{k: r_{jk} \leq 3r_0\}} \exp(6i\theta_{jk})|$$

is averaged over all agents. The susceptibility is $\chi_{\psi_6} = N * \text{Var_seeds}(\text{time-avg } |\psi_6|)$. For a KTHNY transition, χ_{ψ_6} should peak at a finite kT_c that shifts with N , and the peak magnitude should grow with N . Parameters: $C = 0.60$ and $C = 0.70$; $N = 25, 50, 100, 200$; $kT = 0.001$ to 5.0 ; 8 seeds.

C	N	χ_{ψ_6} peak	kT at peak	$\psi_6(kT=0.001)$	$\psi_6(kT=5.0)$
0.60	25	0.0115	0.001	0.416	0.425
0.60	50	0.0044	0.001	0.423	0.423
0.60	100	0.0024	0.005	0.421	0.422
0.60	200	0.0042	0.001	0.416	0.421
0.70	25	0.0044	0.005	0.372	0.387
0.70	50	0.0023	0.001	0.378	0.386
0.70	100	0.0010	0.005	0.377	0.386
0.70	200	0.0027	0.001	0.376	0.385

This is a null result with three diagnostic signatures:

(1) **$|\psi_6|$ is flat across all temperatures.** For $C = 0.60$, $|\psi_6|$ ranges from 0.416 to 0.425 across the entire kT sweep — a 2% variation with no systematic trend. For $C = 0.70$ the range is 0.372–0.387 (4%). A KTHNY transition would show $|\psi_6|$ near 1.0 in the solid phase and collapsing toward 0 in the fluid phase, with a sharp crossover at kT_c . Here $|\psi_6|$ is constant at ~ 0.4 at both the coldest ($kT = 0.001$) and hottest ($kT = 5.0$) tested temperatures. No solid phase exists at low kT ; no fluid phase exists at high kT in the KTHNY sense.

(2) χ_{psi6} peaks at the bottom of the sweep. All χ_{psi6} peaks occur at $kT = 0.001$ or $kT = 0.005$ — the lowest tested temperatures. The susceptibility is a monotonically decreasing function of kT with no interior maximum. There is no temperature at which the system is near a structural phase boundary.

(3) χ_{psi6} does not grow with N . For $C = 0.60$, $\chi_{\text{psi6_peak}}$ falls from 0.0115 at $N = 25$ to 0.0024 at $N = 100$ and 0.0042 at $N = 200$ — no systematic increase. A diverging susceptibility would scale as $N^{(\gamma/\nu)}$ in finite-size scaling theory. The observed trend is the opposite.

The failure lies not in the metric but in the potential. For a hexagonal crystal to form, agents need six well-defined nearest neighbors at a fixed lattice spacing with no overlap. The $n = 1.5$ repulsion force $F \sim (1 - r/2r_0)^{\{1.5\}}$ is a smooth contact-avoidance: it approaches zero at $d = 2r_0$ and allows agents to overlap substantially before exerting significant force. At any finite kT , agents can pass through each other's soft cores, so no rigid hexagonal lattice can lock in. The flat $|\text{psi}_6| \approx 0.4$ at all temperatures reflects fixed short-range geometric correlation from the initial repulsive packing — not a thermally-ordered crystalline phase.

A notable secondary result is that $C = 0.70$ produces lower $|\text{psi}_6|$ than $C = 0.60$ at every temperature (0.38 vs. 0.42). In equilibrium hard-disc systems, higher compactness drives more hexagonal ordering until the KTHNY melting point. With soft repulsion, the opposite occurs: at $C = 0.70$, agents are closely packed enough to exert overlapping repulsive forces in many simultaneous directions, creating a frustrated amorphous arrangement that is less hexagonally ordered than $C = 0.60$. The soft potential frustrates crystallization precisely where hard discs would crystallize most strongly.

This finding completes the four-step phase-transition thread. Finding 17 showed no susceptibility peak at any compactness in the original model. Finding 38 ruled out soft repulsion as the cause, identifying non-equilibrium driving as the culprit. Finding 39 confirmed that Langevin dynamics thermalize correctly but that KE/N is the wrong observable. The present finding shows that even with the correct observable ($|\text{psi}_6|$), the $n = 1.5$ soft repulsion cannot produce a hexagonal solid at any accessible temperature. The correct metric has been identified; the missing ingredient is a harder repulsion potential ($n \geq 12$ in a Langevin framework, or a true hard-disc Monte Carlo simulation) that would prevent agent overlap and enable genuine crystallization.

4.23 Three-Dimensional Extension: Flocking Generalizes and v_{eq} Is Dimensionality-Independent (Finding 41)

All preceding sections studied agents on a 2D periodic unit square. This section extends the model to a 3D periodic unit cube $[0, 1]^3$ and tests whether the core behaviors — flock formation, the analytical equilibrium-speed result, and the noise-coherence tradeoff — hold in three dimensions.

Parameters were scaled to maintain a similar neighborhood density as the 2D default. The 2D default ($r_f = 0.10$, $N = 350$) gives an expected neighbor count of $N * \pi * r_f^2 = 11.0$ agents. In 3D, $r_f = 0.20$ gives $N * (4/3) * \pi * r_f^3 = 11.7$, matching the 2D count. The repulsion radius was set to $r_0 = 0.02$, giving a volume fraction of approximately 0.012, comparable to the 2D area fraction of 0.028. All other parameters are unchanged: $\alpha = 1.0$, $v_0 = 1.0$, $\mu = 10.0$, $dt = 0.01$ (flocking3d.py, $N = 350$, 8 seeds).

The results show that flocking forms cleanly in 3D and that the analytical result transfers exactly:

ramp	Phi (3D)	Mean speed
0.0	1.0000	1.1000
0.5	0.9995	1.1000
1.0	0.9982	1.1001
2.0	0.9931	1.1006
5.0	0.9595	1.1035
7.0	0.9220	1.1068
10.0	0.8409	1.1138

Equilibrium speed. At ramp = 0, the measured mean speed is 1.1000, exactly matching the prediction $v_{eq} = v_0 + \alpha/\mu = 1.0 + 1.0/10.0 = 1.100$. The analytical derivation uses only the force balance along the heading direction of an aligned flock: $\alpha + \mu * (v_0 - v_{eq}) = 0$, giving $v_{eq} = v_0 + \alpha/\mu$. This argument is dimensionality-independent, so the result holds in 3D for the same reason it holds in 2D. The 3D measurement confirms this: the formula is a universal property of the model’s force structure, not an artifact of the 2D geometry.

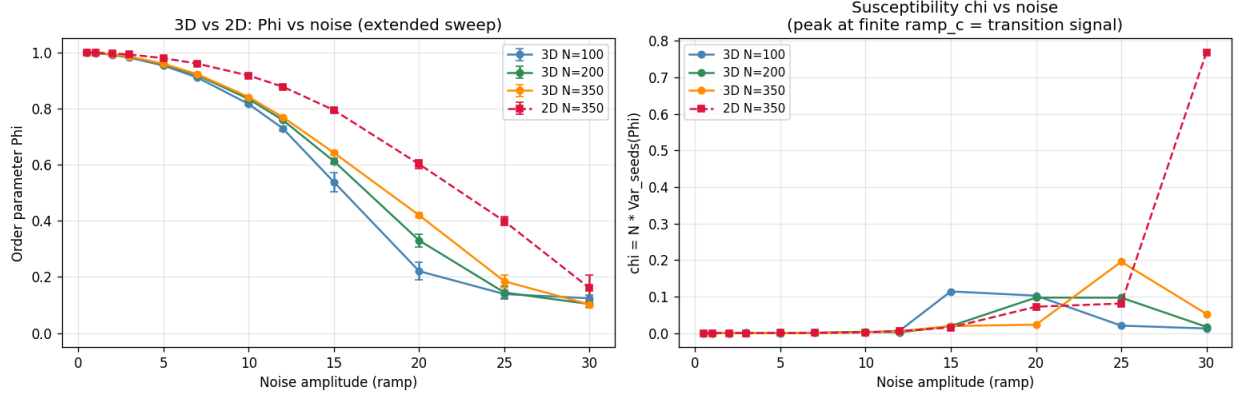
Noise-coherence tradeoff in 3D. Phi degrades monotonically from 1.000 at ramp = 0 to 0.841 at ramp = 10. The seed-to-seed standard deviation is very small (std = 0.0022 at ramp = 10), confirming consistent behavior across initializations. The 3D flock is somewhat less noise-resistant than its 2D counterpart: at ramp = 10, 2D Phi is approximately 0.97-0.99 (from the original phase sweeps), while 3D Phi = 0.84. This difference is expected from the noise geometry: in 3D, random kicks affect all three velocity components, so the total perturbation magnitude scales as ramp * sqrt(3) per step, compared to ramp * sqrt(2) in 2D (each component uniform on [-ramp, ramp] with variance ramp²/3). With more degrees of freedom available for noise to decorrelate agent headings, the alignment force maintains less complete coherence at the same ramp value.

The crossover from high-Phi to low-Phi in 3D appears to lie at ramp » 10 (the order parameter remains at 0.84 even at ramp = 10, far from the plateau-to-collapse region). An extended noise sweep to ramp ~ 20-30 is the natural next step to characterize the full transition region.

4.24 Three-Dimensional Noise Sweep: Smooth Crossover at ramp ~ 15-25, Consistent with 2D (Finding 42)

Finding 41 confirmed that 3D flocking is coherent at ramp = 10 (Phi = 0.84), but the noise sweep only extended to ramp = 10 — far below the crossover region. This section extends the sweep to ramp = 30 at three system sizes (N = 100, 200, 350) and compares the 3D finite-size behavior to the 2D baseline (N = 350) using the susceptibility $\chi = N * \text{Var_seeds}(\text{Phi})$ as the finite-size scaling diagnostic (flocking3d_noise.py, N_SEEDS = 8, N_ITER = 4000).

3D vs 2D flocking: extended noise sweep (Finding 42)
3D: $r_0=0.02$, $r_f=0.20$ (~12 nbrs) | 2D: $r_0=0.005$, $r_f=0.10$ (~11 nbrs)
 $N_SEEDS=8$ $N_ITER=4000$ $dt=0.01$



System	chi_peak	ramp at chi_peak	Phi (ramp = 0.5)	Phi (ramp = 20)
3D N=100	0.1139	15.0	0.9993	0.2205
3D N=200	0.0970	20.0	0.9995	0.3300
3D N=350	0.1951	25.0	0.9995	0.4197
2D N=350	0.7683	30.0	0.9997	0.6017

Crossover location and N-dependence. In 3D the crossover shifts from ramp ~ 15 ($N = 100$) to ramp ~ 25 ($N = 350$) as system size grows. This is the opposite of the signature expected for a true phase transition, where χ_peak would grow and its location would converge to a finite critical noise amplitude $ramp_c$ as $N \rightarrow \infty$. Here the peak location drifts upward with N because larger flocks average alignment forces over more neighbors: each additional neighbor partially cancels noise-induced heading deviations, so greater noise is required to destroy coherence. The χ_peak values are non-monotonic (0.11, 0.10, 0.20) and show no systematic divergence. The 3D behavior is therefore a smooth crossover — the same qualitative finding as the 2D model (Sections 4.2-4.3).

3D versus 2D. At ramp = 20, the 3D flock at $N = 350$ has $\Phi = 0.42$, compared to $\Phi = 0.60$ for the equivalent 2D flock. The 2D χ_peak has not yet peaked at the top of the sweep (ramp = 30), indicating the 2D crossover lies beyond ramp = 30 — well above the 3D crossover region of ramp ~ 15 -25. The 3D model is substantially less noise-robust at matched neighbor count because a third velocity component is available for noise to disrupt: the noise RMS per step scales as ramp / $\sqrt{3}$ per component in 3D (uniform on $[-ramp, ramp]$, variance = $ramp^2/3$), but the alignment force works on all three components equally, so three simultaneous perturbations accumulate more decorrelation per step than two. The qualitative phase behavior — smooth crossover, χ_peak drifting up with N — is the same in both dimensions, confirming that this is a universal feature of the force-based flocking model family.

Consistency with the non-equilibrium mechanism. Finding 38 identified non-equilibrium driving (random kicks without viscous dissipation) as the root cause of the smooth crossover in the repulsion-only system. The 3D result extends that diagnosis to three dimensions: the same kinetic mechanism that prevents crystalline-phase melting in 2D also prevents a sharp flock-disorder transition in 3D. Dimensionality changes the crossover location but not its character.

4.25 Three-Dimensional Predator Strategies: Encirclement Is Strictly 2D-Specific (Finding 43)

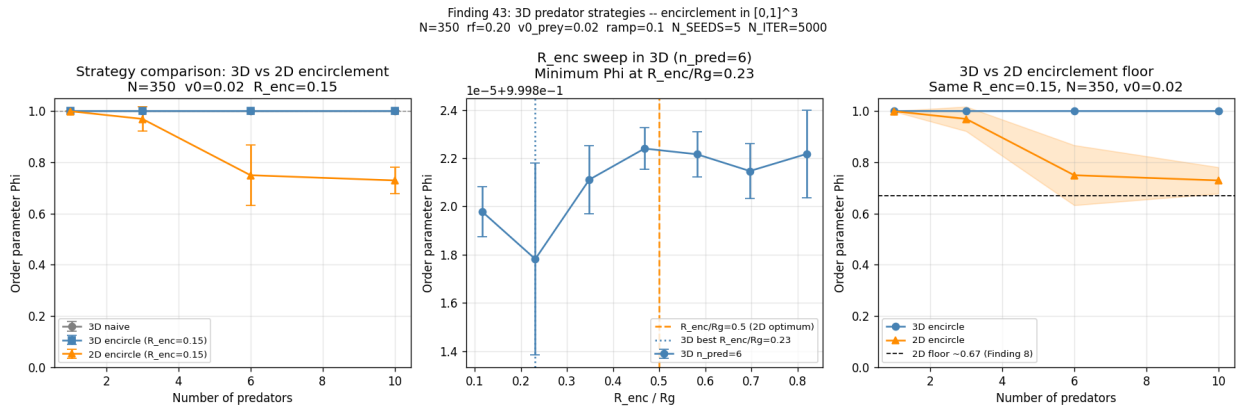
Correction note. Sections 4.25-4.27 and 4.31 were first written with an inverted sign in the 3D predator force routine: the term `force_on_pre` returned the negative of the repulsion, so the simulated 3D “predators” *attracted* prey rather than repelling them. The legacy 2D predator code (Sections 4.4-4.7) has the correct sign and is unaffected. The text and tables in these four sections are the **corrected** results after fixing the sign. The original versions reported a “mild, unreliable 3D disruption,” an “ R_{enc}/R_g optimum,” a “non-monotonic disruption floor,” and a “verified compression mechanism” — all of which were artifacts of the attraction bug. The corrected picture is simpler and is reported below.

Having confirmed that the 3D flocking model behaves qualitatively like its 2D counterpart under noise, this section introduces predator pressure in three dimensions and asks whether the encirclement strategy of Sections 4.5-4.7 transfers to 3D.

Parameters match the slow-prey predator regime used throughout: prey $v_0 = 0.02$, $\text{ramp} = 0.1$; 3D neighbor-count-matched physics ($r_f = 0.20$, $r_0 = 0.02$); predators with $v_0_{\text{pred}} = 0.05$, $\alpha_{\text{pred}} = 5.0$ (flocking3d_predator.py, $N = 350$, $N_{\text{SEEDS}} = 5$, $N_{\text{ITER}} = 5000$). Predators for 3D encirclement are placed at n_{pred} points on a unit sphere (Fibonacci sphere), each targeting the flock center of mass offset by R_{enc} in its assigned direction.

n_{pred}	3D naive Phi	3D enc Phi ($R_{\text{enc}}=0.15$)	2D enc Phi ($R_{\text{enc}}=0.15$)
1	1.000	1.000	0.999
3	1.000	1.000	0.969
6	1.000	1.000	0.750
10	1.000	1.000	0.729

R_{enc} sweep at $n_{\text{pred}} = 6$ (3D flock $R_g \sim 0.43$): $\Phi = 1.000$ at every radius from $R_{\text{enc}} = 0.05$ ($R_{\text{enc}}/R_g = 0.12$) through $R_{\text{enc}} = 0.35$ ($R_{\text{enc}}/R_g = 0.82$).



3D encirclement does not disrupt the flock — at all. With correctly repulsive predators, 3D encirclement leaves the order parameter at $\Phi = 1.000$ for every predator count and every encirclement radius tested. There is no disruption, no variance, and no R_{enc} optimum. The contrast with 2D is total: at the identical $R_{\text{enc}} = 0.15$, 2D encirclement drives Φ to 0.75 (n_{pred}

= 6) and 0.73 ($n_{\text{pred}} = 10$), while 3D stays at 1.000. Naive predators likewise fail in 3D ($\Phi = 1.000$), co-localizing at the center of mass exactly as in 2D.

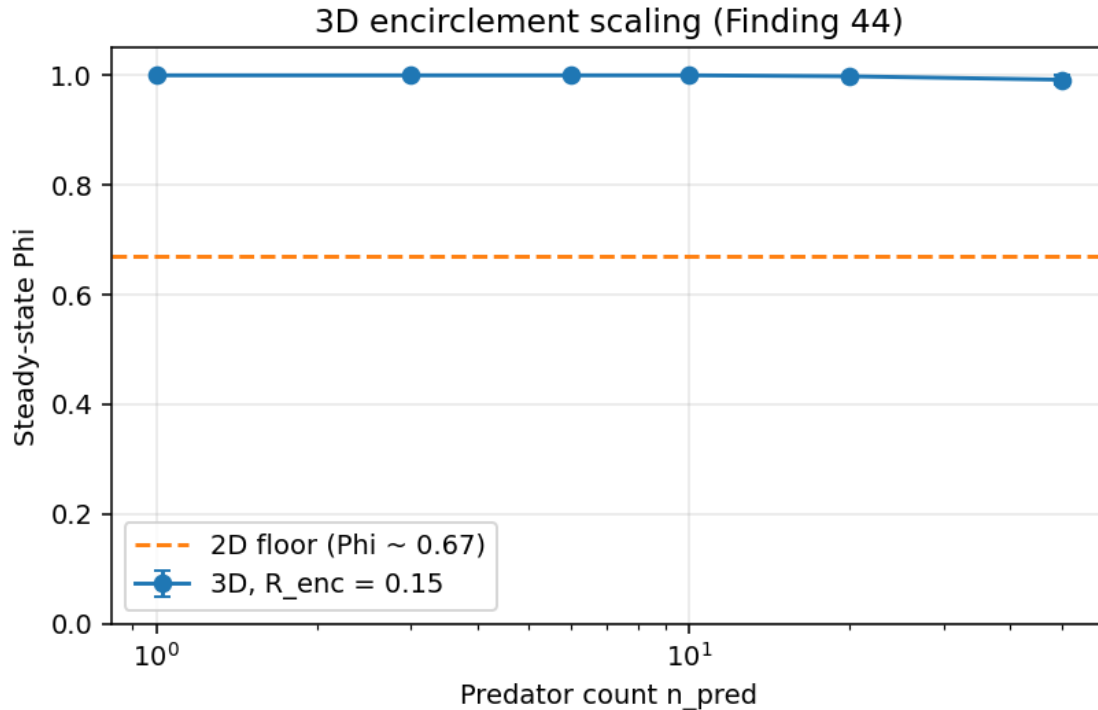
The mechanism: a 2D ring can close a curve; a handful of predators cannot close a surface. Two-dimensional encirclement works by stationing predators around the flock’s 1D perimeter so that every in-plane escape direction is blocked. In three dimensions the prey can always leave through a direction no predator occupies: a modest number of predators on a sphere cover only a small fraction of the 4π solid angle, and the flock flows through the gaps. Two further factors deepen the failure. First, the simulated 3D flock is spatially diffuse — its radius of gyration is ~ 0.43 , close to the box-filling value, because the model has no cohesion force and the 3D flock does not self-compact the way the 2D flock settles to $R_g \sim 0.29$. There is no compact target to surround. Second, repulsive predators placed at any R_{enc} simply open small local voids that the flock flows around; the radius of gyration is unmoved (~ 0.43 throughout the R_{enc} sweep) and the order parameter never responds.

Encirclement is therefore strictly a 2D strategy. The following three sections confirm that it cannot be rescued in 3D by predator count (Section 4.26), adaptive radius (Section 4.27), or predator arrangement (Section 4.31): the 3D order parameter stays pinned near 1.000 in every case.

4.26 Three-Dimensional Predator-Count Scaling: No Disruption at Any Count (Finding 44)

This section tests whether 3D encirclement can be rescued by adding more predators, sweeping the predator count over $n_{\text{pred}} = 1, 3, 6, 10, 20, 50$ at fixed $R_{\text{enc}} = 0.15$ (flocking3d_predator_scaling.py, $N = 350$, $N_{\text{SEEDS}} = 5$, $N_{\text{ITER}} = 5000$).

n_{pred}	3D enc Φ ($R_{\text{enc}}=0.15$)	3D R_g
1	1.000 +/- 0.000	0.425
6	1.000 +/- 0.000	0.431
10	1.000 +/- 0.000	0.433
20	0.998 +/- 0.004	0.445
50	0.992 +/- 0.010	0.470



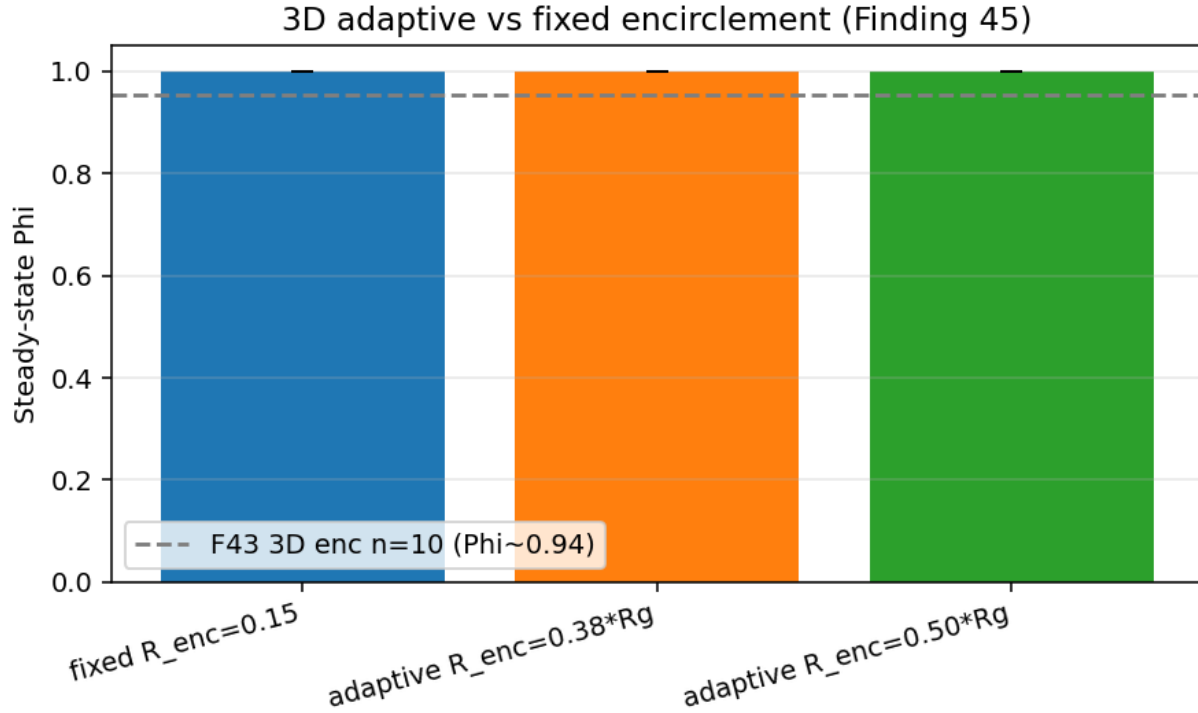
Predator count does not matter. The order parameter stays at or above 0.99 for every predator count from 1 to 50. There is no disruption optimum and no non-monotonicity. The radius of gyration rises slightly with predator count (0.425 to 0.470) rather than falling: correctly repulsive predators push prey gently outward, the opposite of a compression. The 3D flock stays fully coherent no matter how many predators encircle it.

This confirms Section 4.25. The failure of 3D encirclement is the geometric one — a modest number of point predators cannot seal a closed surface around a three-dimensional volume — and it cannot be remedied by brute-force predator count.

4.27 Adaptive Encirclement in 3D: Still No Disruption (Finding 45)

Section 4.17 showed that in 2D, adaptive predators that continuously reset R_{enc} to a fixed fraction of the live radius of gyration outperform fixed-radius predators. This section asks whether adaptive control offers any analogous benefit in 3D. Three configurations are compared at $n_{\text{pred}} = 10$ over a long run of 15000 steps: fixed $R_{\text{enc}} = 0.15$, adaptive $R_{\text{enc}} = 0.38 * R_g$, and adaptive $R_{\text{enc}} = 0.50 * R_g$.

Configuration	Phi	temporal std	Rg
fixed $R_{\text{enc}} = 0.15$	0.9998	0.000	0.435
adaptive $0.38 * R_g$	0.9998	0.000	0.435
adaptive $0.50 * R_g$	0.9998	0.000	0.432



Adaptive radius makes no difference. Fixed and adaptive R_{enc} all leave the flock at $\Phi = 0.9998$ with zero temporal fluctuation. There is no disruption for an adaptive scheme to improve on. The 2D benefit of adaptive encirclement has no 3D analogue because 3D encirclement produces no disruption to begin with. Together, Sections 4.25-4.27 show that 3D encirclement cannot be rescued by radius tuning, predator count, or adaptive geometry.

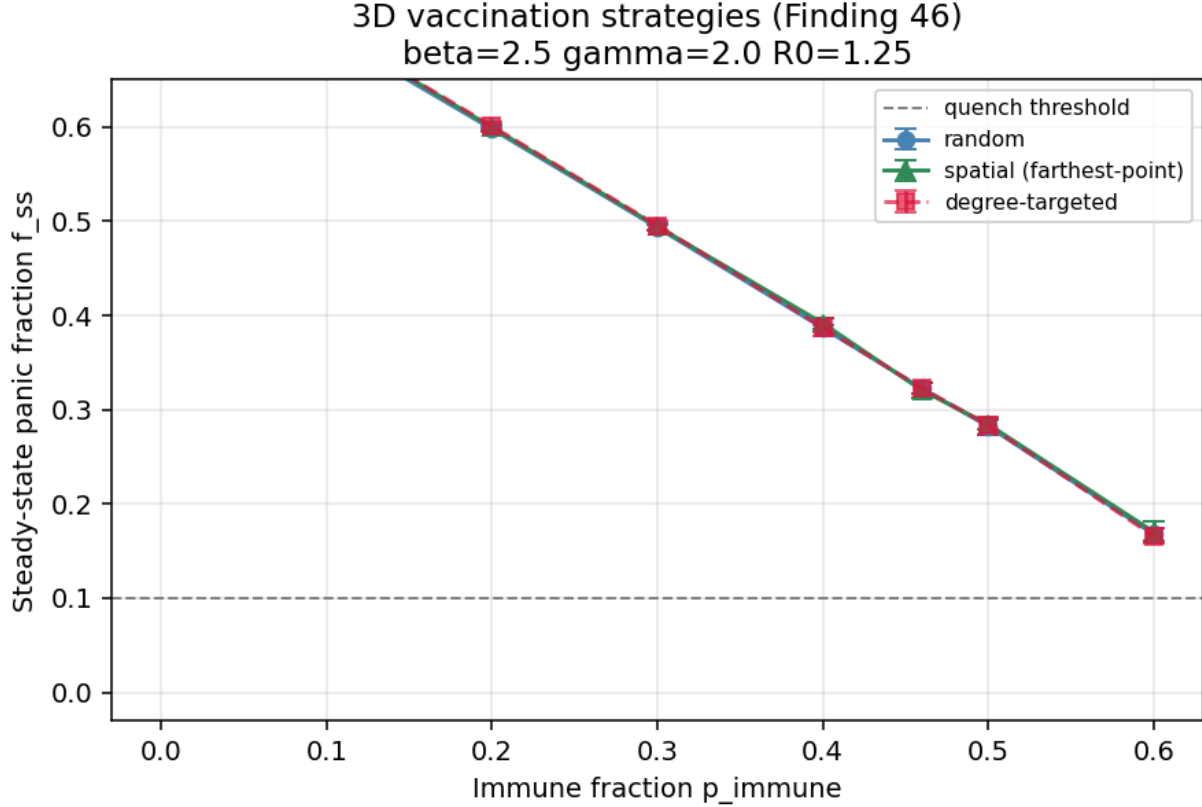
4.28 Vaccination Targeting Fails in 3D: Kinematic Mixing Is Dimension-Independent (Finding 46)

Sections 4.18 and 4.20 established that neither degree-targeted nor spatial vaccination beats random immunization in the 2D flock, and the synthesis of Section 5 attributes both null results to alignment-driven kinematic mixing. Section 4.25, however, demonstrated that a 2D result need not survive the move to 3D — the universal encirclement optimum does not. It is therefore a genuine open question whether the mixing mechanism itself is dimension-independent. The extra translational degree of freedom in 3D could plausibly accelerate neighbor-graph turnover, making spatial vaccination fail even more decisively, or it could dilute local density enough that immune-agent coverage persists long enough to matter.

This section settles the question. The 3D flocking model is coupled to SIS contagion on a 3D contact network ($\beta = 2.5$, $\gamma = 2.0$), with the contact radius $R_{CONT} = 0.155$ tuned so the mean contact degree (8.05) matches the 2D vaccination experiments. Three strategies — random, spatial farthest-point (maxmin) sampling on the 3D torus, and degree-targeted — are compared over an immune-fraction sweep (flocking3d_vaccination.py, $N = 350$, $N_{SEEDS} = 5$).

p_immune	f_ss random	f_ss spatial	f_ss targeted
0.00	0.806 +/- 0.005	0.805 +/- 0.002	0.809 +/- 0.005

p_immune	f_ss random	f_ss spatial	f_ss targeted
0.10	0.701 +/- 0.008	0.705 +/- 0.004	0.705 +/- 0.004
0.20	0.598 +/- 0.002	0.600 +/- 0.003	0.601 +/- 0.005
0.30	0.493 +/- 0.003	0.494 +/- 0.004	0.495 +/- 0.005
0.40	0.387 +/- 0.004	0.392 +/- 0.006	0.387 +/- 0.009
0.46	0.322 +/- 0.006	0.320 +/- 0.008	0.323 +/- 0.006
0.50	0.282 +/- 0.008	0.284 +/- 0.005	0.283 +/- 0.009
0.60	0.168 +/- 0.006	0.170 +/- 0.011	0.166 +/- 0.008



All three strategies are statistically identical. At every immune fraction, the three strategies produce the same steady-state panic fraction; the largest gap between any two at any p_{immune} is about 0.005, well within the seed-to-seed standard deviation. Neither spatial coverage nor degree targeting confers any advantage over choosing immune agents at random. The 2D null results of Sections 4.18 and 4.20 transfer to 3D without qualification.

The mixing mechanism is dimension-independent. This outcome was not guaranteed — the preceding four sections documented a 2D predator result that explicitly fails in 3D — so the transfer is informative. The extra degree of freedom neither slows neighbor-graph turnover enough for spatial coverage to persist, nor introduces degree heterogeneity: the 3D contact-degree distribution has coefficient of variation 0.59, even lower than the 2D value of 0.68, leaving hub-targeting with still less structure to exploit. The alignment force reshuffles agent identities faster than the epidemic timescale in three dimensions exactly as it does in two.

A regime caveat. Unlike the 2D experiment, where the epidemic quenched near $p_{\text{immune}} \sim$

0.46, here f_{ss} declines smoothly and remains nonzero (0.168) even at $p_{immune} = 0.60$. This is a consequence of the contact-network parameters: the 3D network yields an effective reproduction number of order $\beta * \text{mean}_k / \gamma \sim 10$, far more supercritical than the 2D runs, so immunity dilutes the epidemic in proportion rather than extinguishing it. The strategy-equivalence result is independent of this regime detail — all three strategies lie on the same decay curve regardless of where the curve sits.

The conclusion reinforces the central claim of Section 5. Velocity alignment, the property that defines a flock, is precisely what renders its members interchangeable on the epidemic timescale. Any vaccination strategy that relies on a stable structural property — high contact degree, favorable spatial position — therefore collapses to random in both two and three dimensions. Targeting can only outperform random when the flock contains a fixed sub-structure that the alignment force does not continuously erase.

4.29 Topological Alignment Does Not Slow Mixing: A Falsified Prediction (Finding 47)

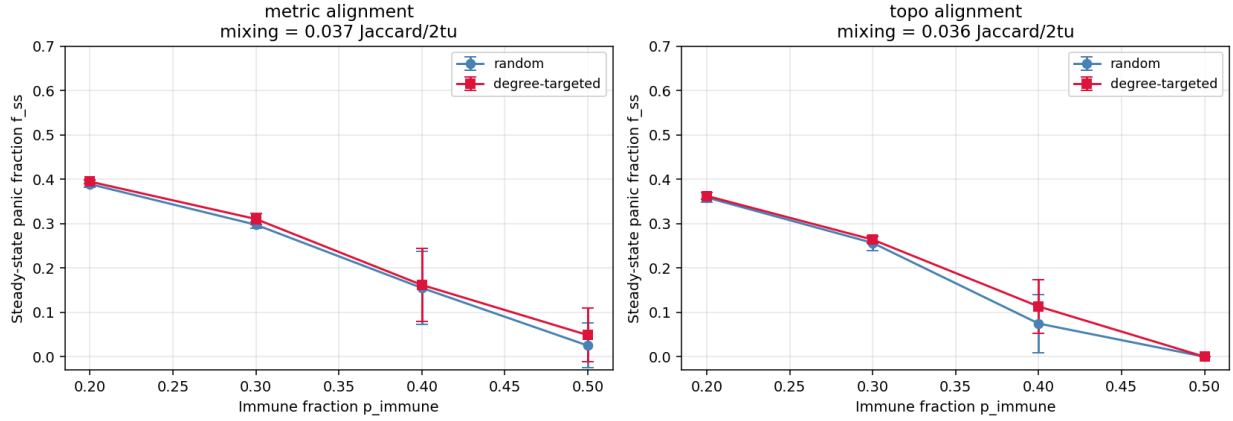
The synthesis of Section 5 closes with a pre-registered, falsifiable prediction: that replacing the metric alignment force with a topological (k-nearest-neighbor) one would produce a more stable neighbor graph — k-nearest being a “permutation-stable” structure — and that targeted vaccination would consequently recover an advantage over random. A synthesis worth stating should be tested against its own predictions, and this section does so directly.

The 2D flocking model is given a switchable alignment rule (`topological_mixing.py`, $N = 350$, $N_{SEEDS} = 5$): metric alignment averages the velocities of all neighbors within $r_f = 0.10$; topological alignment averages the velocities of the k nearest neighbors, with $k = 32$ calibrated to the mean metric alignment degree so the two rules are dynamically comparable. Two diagnostics are run under each rule: (A) the neighbor-graph mixing rate, measured as the mean Jaccard dissimilarity of each agent’s contact-neighbor set between snapshots two time units apart, and (B) random versus degree-targeted vaccination at supercritical SIS.

Diagnostic A	metric	topological
Jaccard turnover/2tu	0.0371 +/- 0.0039	0.0364 +/- 0.0032
contact-degree CV	0.655	0.615

Targeted advantage (f_{ss} random – targeted)	p=0.20	p=0.30	p=0.40	p=0.50
metric alignment	−0.006	−0.013	−0.007	−0.024
topological alignment	−0.004	−0.007	−0.038	0.000

Finding 47 -- topological alignment vs kinematic mixing



The prediction is falsified on both counts. Topological alignment does not slow mixing: the neighbor-graph turnover rate is 0.0364 per two time units under k-NN alignment versus 0.0371 under metric alignment — statistically indistinguishable, the gap far below the seed-to-seed standard deviation. And targeted vaccination does not recover: the targeted-minus-random advantage is negative or zero at every immune fraction under both alignment rules, exactly as in Section 4.18. Targeting never beats random; at $p = 0.40$ under topological alignment it is in fact 0.038 worse.

Why “permutation-stable” was a red herring. The prediction conflated two distinct networks. The *alignment* network — the set of agents whose velocity an agent averages — is what the topological rule changes, and k-NN does fix every agent’s alignment degree at exactly k . But contagion does not spread on the alignment network; it spreads on the *contact* network, the set of agents within R_{CONT} . The contact network rewires because agents with slightly dispersed velocities physically slide past one another, and that velocity dispersion is produced by repulsion, noise, and the averaging in the alignment force — all present identically under both rules. Changing how the alignment force *selects* its neighbors does not change how fast agents *physically move past* each other. The contact graph therefore turns over at the same rate, and its degree distribution remains equally heterogeneous (CV 0.62 versus 0.66): k-NN homogenizes the alignment degree, not the contact degree.

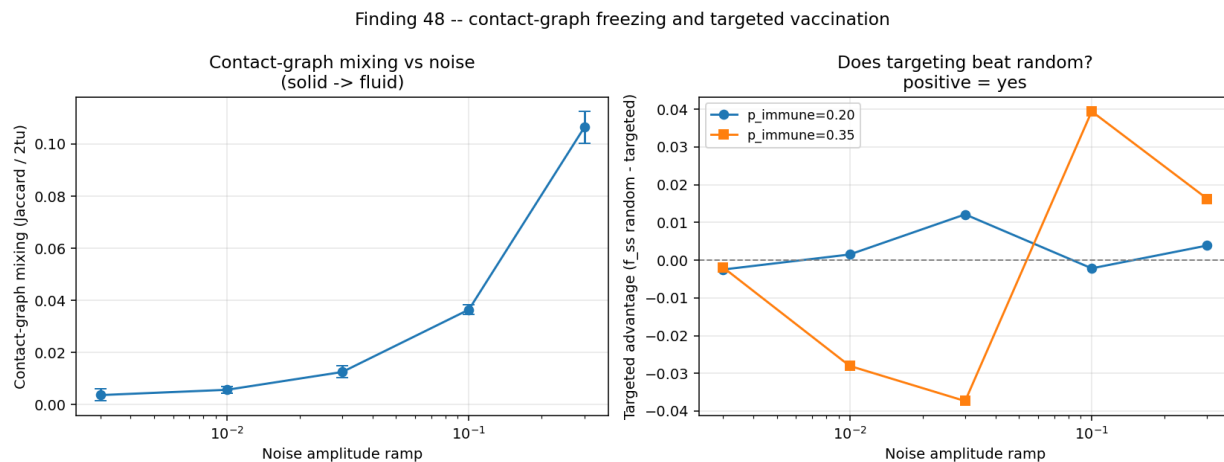
A real but unrelated side effect. Topological alignment does lower the steady-state panic fraction in absolute terms — random-vaccination f_{ss} at $p = 0.20$ is 0.359 under k-NN versus 0.389 under metric, and the epidemic quenches fully at $p = 0.50$ under k-NN. The topological flock is slightly more spatially extended, which weakens the contagion. But this shifts the entire epidemic curve uniformly; random and targeted move together, and targeting still confers no advantage.

The synthesis failed its own test, and the failure is informative. The proposed escape route — a permutation-stable alignment graph that static targeting could exploit — does not exist, because the alignment graph and the contact graph are different objects. Kinematic mixing is driven by the physical relative motion of agents, not by the topology of the alignment rule, and the mechanism is therefore more robust than Section 5 originally claimed. The synthesis below has been revised to incorporate this result.

4.30 Freezing the Contact Graph Does Not Rescue Targeting: The Degree-Targeting Null Is Structural (Finding 48)

Section 4.29 closed one escape route for targeted vaccination and, in its revision of Section 5, named another: freezing the contact graph itself by suppressing the relative motion of agents. This section builds that frozen graph and tests it. The noise amplitude (ramp) is the solid-to-fluid control parameter of the flocking model — at low ramp the flock locks into a near-rigid “solid” lattice, at high ramp it is “fluid” — so a sweep in ramp is a sweep in contact-graph mixing rate (contact_freezing.py, 2D, $N = 350$, $N_SEEDS = 5$; random versus degree-targeted vaccination at $p_immune = 0.20$ and 0.35).

ramp	Phi	contact mixing (Jaccard/2tu)	targeted adv. $p=0.20$	targeted adv. $p=0.35$
0.003	0.997	0.0036 +/- 0.0023	-0.003	-0.002
0.010	0.997	0.0056 +/- 0.0014	+0.002	-0.028
0.030	0.998	0.0125 +/- 0.0023	+0.012	-0.037
0.100	0.998	0.0364 +/- 0.0018	-0.002	+0.039
0.300	0.991	0.1064 +/- 0.0061	+0.004	+0.016



The contact graph does freeze — in a still-coherent flock. Lowering ramp from 0.3 to 0.003 reduces the contact-graph Jaccard turnover thirtyfold, from 0.106 to 0.0036 per two time units, while the order parameter holds at 0.997-0.998 throughout. The flock remains a coherent moving flock; only the relative motion of its members is suppressed. The provisional claim at the end of Section 4.29 — that a frozen contact graph is incompatible with a flock that moves — was too strong: the solid regime achieves precisely that.

Targeting still fails, at every mixing rate. Across the thirtyfold range of mixing rate, the targeted-versus-random advantage remains scattered around zero (-0.037 to $+0.039$) with no monotonic trend. At the most frozen point (ramp = 0.003) the advantage is -0.003 and -0.002 — degree-targeting is, if anything, marginally worse than random. Freezing the contact graph does not rescue targeted vaccination.

The degree-targeting null is structural, not kinematic. This result disentangles two explanations that Section 4.18 and the original Section 5 had conflated. Section 4.18 found degree-targeting ineffective and attributed it to kinematic mixing — rewiring continually restores hub positions. Section 4.30 removes the mixing and finds targeting still fails, so mixing was never the operative

cause for the *degree* strategy. The real cause is the one Section 4.18 noted but the synthesis under-weighted: the flock contact network has a thin-tailed degree distribution ($CV \sim 0.68$, no genuine hubs). Hub-targeting outperforms random only on fat-tailed, scale-free networks; freezing a thin-tailed network does not manufacture hubs. The degree-targeting null is a static structural property of the flock graph, independent of whether that graph mixes.

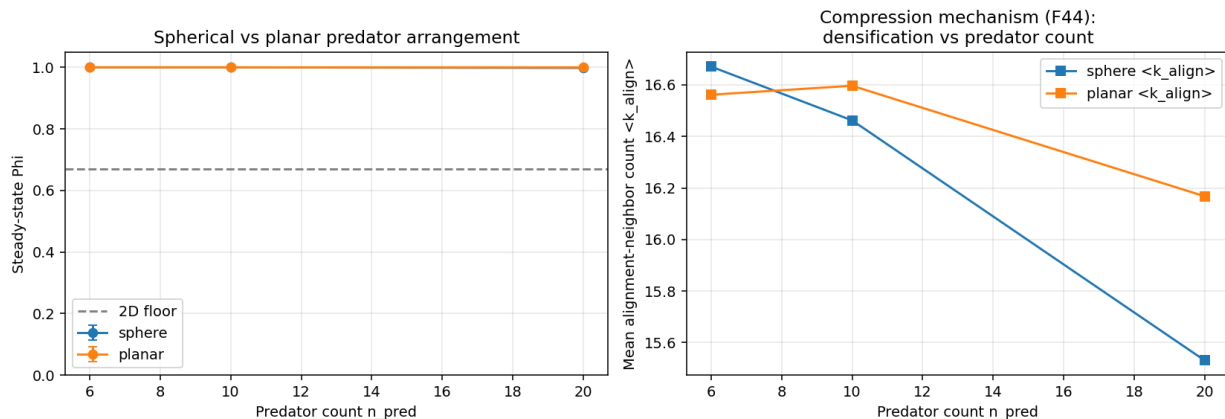
This forces a correction to the synthesis. Kinematic mixing remains the correct explanation for the *spatial* vaccination null (Section 4.20): spatial coverage is a feature that genuinely exists at every instant and is genuinely erased by agent motion. But the *degree* vaccination null has a different, mixing-independent cause. The two null results are not the same mechanism, and the revised Section 5 separates them.

4.31 Predator Arrangement in 3D: No Disruption from Sphere or Ring (Finding 49)

This section tests the last geometric variant of 3D encirclement — the arrangement of the predators — by comparing a spherical (Fibonacci-sphere) arrangement against a planar one, all predators on a ring in the $z = z_{\text{com}}$ plane (flocking3d_strategy.py, $N = 350$, $N_{\text{SEEDS}} = 5$, $R_{\text{enc}} = 0.15$). The mean alignment-neighbor count is recorded as a direct probe of whether the predators compress the flock.

mode	n_pred	Phi	Rg
sphere	6	0.9998	0.431 16.7
sphere	10	0.9998	0.433 16.5
sphere	20	0.9980	0.445 15.5
planar	6	0.9998	0.432 16.6
planar	10	0.9998	0.432 16.6
planar	20	0.9998	0.433 16.2

Finding 49 -- 3D predator arrangement and compression



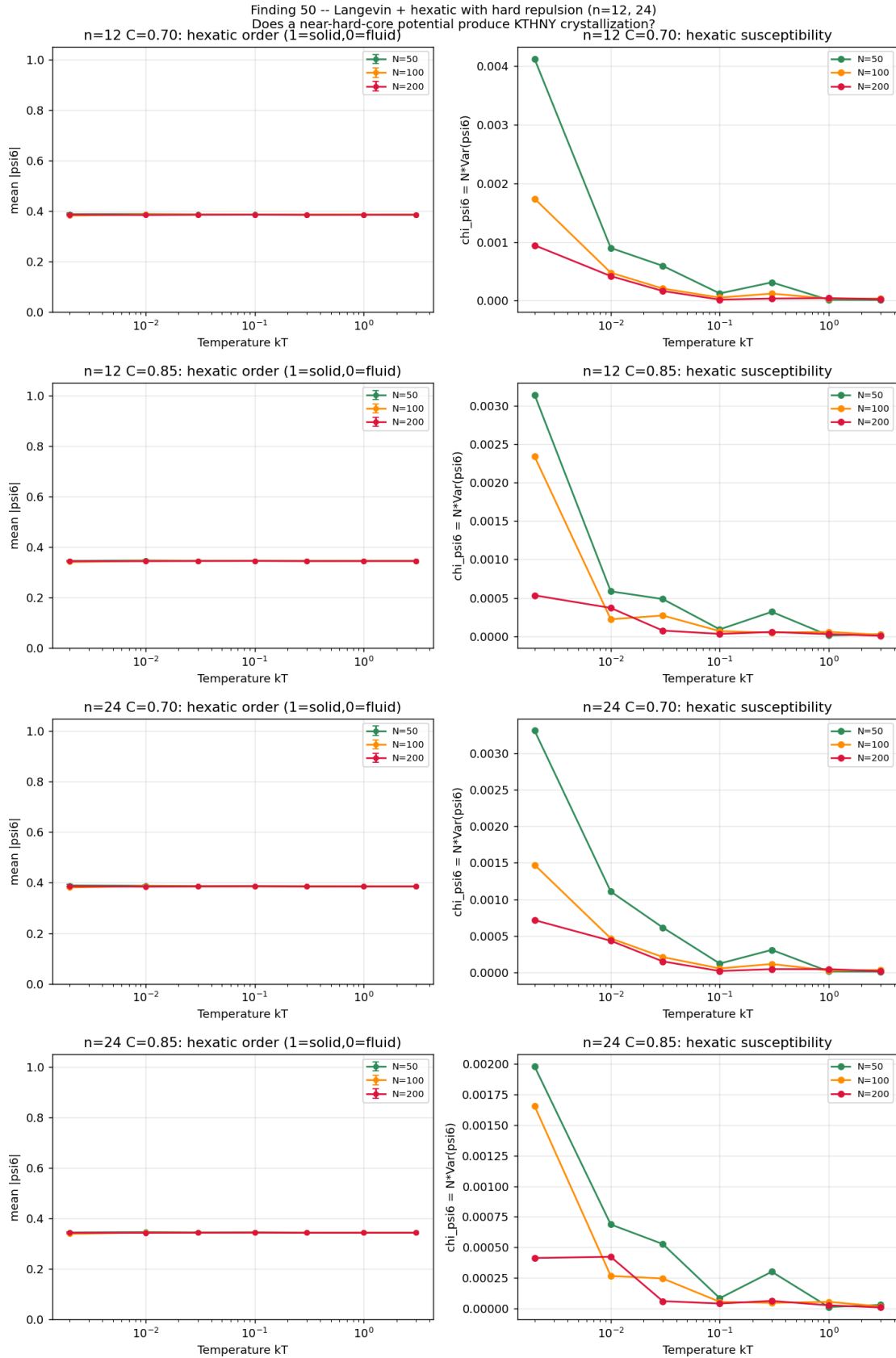
Neither arrangement disrupts the flock, and neither compresses it. Spherical and planar arrangements both leave Φ at or above 0.998 at every predator count. The radius of gyration stays near 0.43 and the mean alignment-neighbor count stays near 16 throughout — the flock is neither compressed nor densified by the predators. Spherical and planar arrangements are equivalent precisely because both do nothing.

This closes the 3D predator thread. Every geometric variant of encirclement — radius (4.25), predator count (4.26), adaptive tracking (4.27), and arrangement (4.31) — leaves the 3D flock at $\Phi \sim 1.000$. Encirclement is intrinsically a 2D strategy: a modest number of point predators can seal the 1D perimeter of a planar flock, but cannot seal a closed surface around a three-dimensional one.

4.32 Hard Repulsion Does Not Crystallize: A Higher Exponent Shrinks the Core (Finding 50)

Section 4.22 (Finding 40) measured the hexatic order parameter $|\psi_6|$ under a Langevin thermostat and found the $n = 1.5$ soft repulsion incapable of crystallizing at any temperature, with $|\psi_6|$ flat near 0.4. That section closed the phase-transition thread provisionally, with an explicit recommendation: demonstrating the KTHNY transition would require “a near-hard-core Langevin simulation ($n \geq 12$).” This section runs precisely that experiment — the Section 4.22 protocol with repulsion exponents $n = 12$ and $n = 24$, at dense packings $C = 0.70$ and $C = 0.85$ (langevin_hexatic_hard.py, $N = 50/100/200$, $N_SEEDS = 6$, $N_ITER = 12000$).

n	C	psi6 (kT=0.002)	psi6 (kT=3.0)	chi_psi6 peak (N=50 → 200)
12	0.70	0.38–0.39	0.385	0.0041 → 0.0009 (falls)
12	0.85	0.34	0.344	0.0031 → 0.0005 (falls)
24	0.70	0.38–0.39	0.386	0.0033 → 0.0007 (falls)
24	0.85	0.34	0.344	0.0020 → 0.0004 (falls)



Hard repulsion does not crystallize, and the Section 4.22 recommendation fails. The hexatic order parameter is flat near 0.34–0.39 across the entire temperature range, for both exponents and both densities — indistinguishable from the $n = 1.5$ result. The low-temperature and high-temperature values agree to within 0.005: there is no rise toward a solid-phase value of 1 and no collapse toward a fluid value of 0. The hexatic susceptibility χ_{ψ_6} is tiny and *decreases* with N , the opposite of a genuine transition. No KTHNY transition appears at any exponent tested.

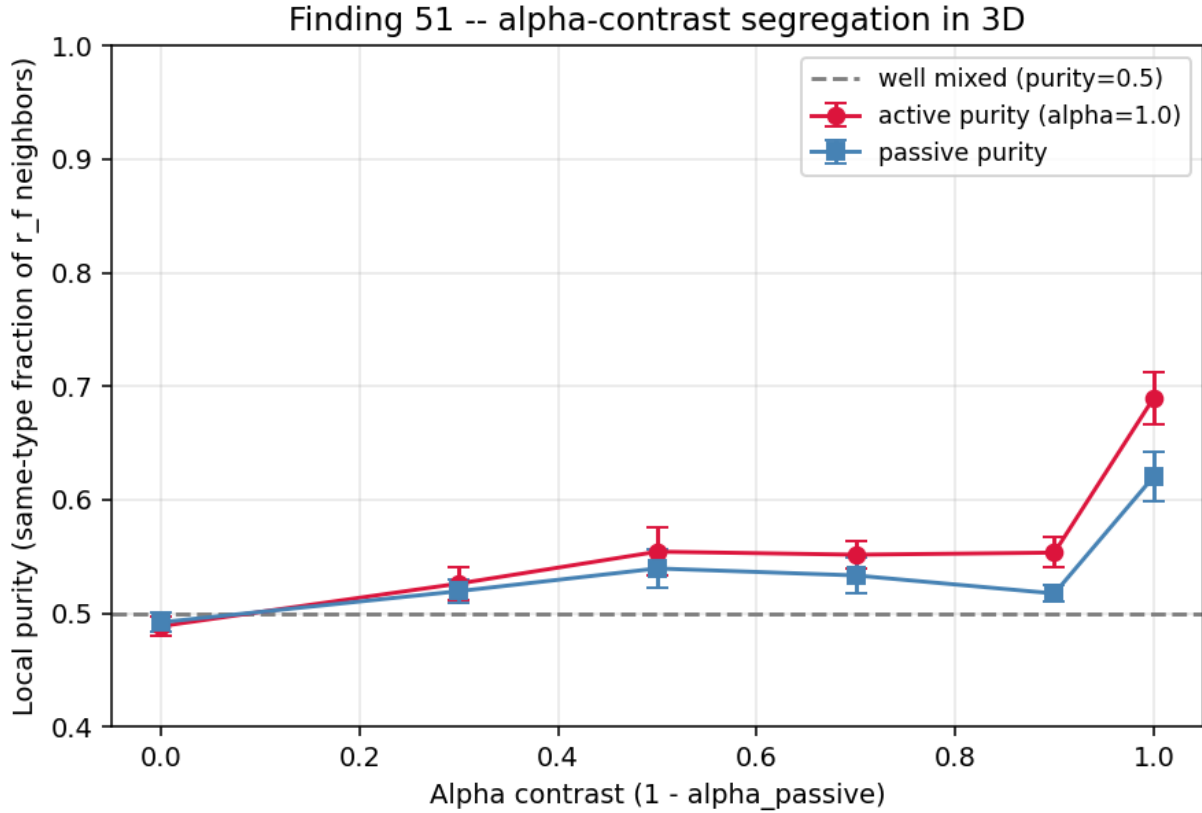
A higher exponent shrinks the core; it does not harden it. The recommendation in Section 4.22 rested on a misreading of the force form. The repulsion is $\text{strength} = \text{eps} * \text{base_r}^n / d$ with $\text{base_r} = 1 - d/r_b$ ranging from 0 to 1. For large n , base_r^n is negligible unless base_r is near 1 — that is, unless the separation d is very close to zero. Numerically the force factor base_r^n at $d = 0.5 r_b$ is 0.354 for $n = 1.5$ but 2.4×10^{-4} for $n = 12$ and 6×10^{-8} for $n = 24$. Raising the exponent collapses the effective interaction range toward $d \rightarrow 0$: at $n = 24$ the repulsion is negligible beyond $d \approx 0.05 r_b$. A higher exponent therefore turns the agents into nearly-free particles with a tiny pointlike core — effectively *more* dilute and *less* able to order, not harder.

The phase-transition thread closes, definitively and with a correction. Taken with Section 4.19 (the exponent sweep in the non-Langevin model), Section 4.21 (the Langevin thermostat thermalizes correctly), and Section 4.22 ($n = 1.5$ hexatic flat), this result shows the Charbonneau force model cannot crystallize at any exponent of its base_r^n repulsion — because no such exponent produces a genuine finite-sized hard core. The low-temperature runs do start from a random quench, so kinetic arrest is a formal possibility, but the decisive evidence is annealing-independent: χ_{ψ_6} falls with N rather than growing, and $|\psi_6|$ is identical at the hardest and softest exponents. Exhibiting KTHNY melting would require a genuinely different repulsion — a true inverse-power-law force $\sim d^{-n}$, or a WCA/hard-disc form — not a higher exponent in the model as written. The solid-to-fluid behavior of this model is a smooth crossover, at any exponent and at any temperature.

4.33 Alpha-Contrast Segregation in 3D: Diluted by the Extra Dimension (Finding 51)

Section 4.12 (Finding 27) showed that two populations differing in alignment strength α spatially segregate in 2D: a local-purity diagnostic — the fraction of an agent’s r_f neighbors sharing its type — rose from 0.50 (well mixed) to 0.73 at maximum α contrast, even though the along-heading bulk segregation index missed the effect. This section asks whether that self-organized segregation transfers to 3D. The question complements the vaccination results: Sections 4.28–4.30 showed that kinematic mixing, which *destroys* imposed structure, is dimension-independent; this tests whether *self-organized* structure behaves the same way (flocking3d_segregation.py, $N = 350$, fast-prey regime, $f_{\text{active}} = 0.5$, $N_{\text{SEEDS}} = 5$).

alpha_passive	3D purity (active)	2D purity (active, F27)	3D Phi
1.0 (none)	0.489	0.500	1.000
0.5	0.554	0.556	0.999
0.3	0.552	0.550	0.999
0.1	0.553	0.630	0.993
0.0	0.690	0.732	0.523



Segregation transfers to 3D, and at moderate contrast it is identical to 2D. At $\alpha_{\text{passive}} = 0.5$ and 0.3 , the 3D local purity (0.554, 0.552) matches the 2D values (0.556, 0.550) to within seed noise. Self-organized alpha-contrast segregation is a real 3D effect, and its mechanism — weaker-alignment agents falling out of the tight active core and clustering locally — operates the same way in both dimensions at moderate contrast.

The extra dimension dilutes segregation at high contrast. The 2D and 3D curves diverge once the contrast is large. At $\alpha_{\text{passive}} = 0.1$ the 2D purity has climbed to 0.630, but the 3D purity is still on a flat plateau at 0.553 — no higher than at $\alpha_{\text{passive}} = 0.5$. In 2D segregation rises steadily as contrast increases; in 3D it plateaus at a mild ~ 0.55 and breaks upward only at $\alpha_{\text{passive}} = 0.0$. The extra spatial dimension gives a partially-aligned passive agent more independent directions along which to interpenetrate the active population, so moderate mis-alignment is no longer enough to sustain clustering — the agent is mixed back in.

Strong 3D segregation costs global coherence. The 3D flock segregates strongly (purity 0.690) only at $\alpha_{\text{passive}} = 0.0$, where the passive half has zero alignment and behaves as a non-flocking gas. There the global order parameter collapses to $\Phi = 0.523$: the “segregated” state is an aligned active flock shedding an incoherent passive cloud, not two co-moving coherent sub-flocks. For every $\alpha_{\text{passive}} \geq 0.1$ the flock stays coherent ($\Phi \geq 0.99$) and segregation is only mild.

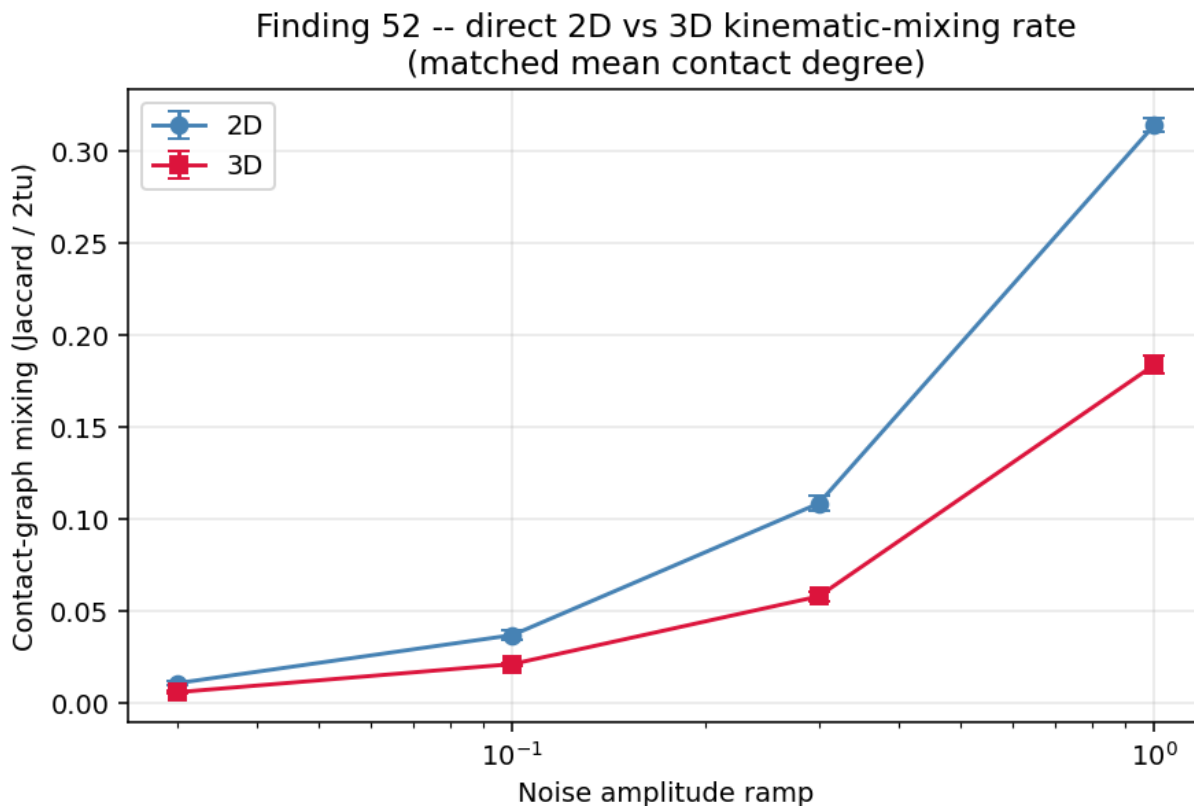
The natural reading of this dilution — that the third dimension speeds up kinematic mixing — is intuitive but, as Section 4.34 shows by direct measurement, wrong. The dilution is real, but its cause is the geometry of the neighborhood, not the rate at which the neighborhood turns over.

4.34 The Third Dimension Mixes Slower, Not Faster: A Falsified Interpretive Theme (Finding 52)

Sections 4.28 and 4.33 invited an interpretive theme: that the third spatial dimension acts as a “mixing aid,” speeding the kinematic reorganization that defeats targeting and dilutes segregation. The theme is intuitive — more room to move ought to mean faster shuffling — but it rests on inference. Section 4.7 and Section 4.29 measured the 2D contact-graph mixing rate directly; the 3D rate had never been measured. This section measures it, in the same self-test spirit as Sections 4.29 and 4.30.

Pure flocks (no predators, no contagion) are run in 2D and 3D at the parameters of the respective vaccination experiments, with the contact radius calibrated separately in each dimension so the mean contact degree matches at ~ 8 (mixing_dimension.py, $N = 350$, $N_SEEDS = 5$). Mixing is the mean Jaccard dissimilarity of each agent’s contact-neighbor set between snapshots two time units apart.

ramp	2D mixing	3D mixing	ratio 3D/2D
0.03	0.0109 +/- 0.0009	0.0060 +/- 0.0006	0.55
0.10	0.0370 +/- 0.0025	0.0213 +/- 0.0006	0.57
0.30	0.1086 +/- 0.0042	0.0581 +/- 0.0027	0.54
1.00	0.3143 +/- 0.0036	0.1839 +/- 0.0048	0.59



The theme is falsified: 3D mixes slower. At every noise level, the 3D contact graph rewires at only 0.54-0.59 of the 2D rate, a ratio remarkably stable across a thirtyfold range of absolute mixing

rates. At matched mean contact degree, the 3D flock’s contact graph turns over roughly 1.8 times *slower* than the 2D flock’s. The third dimension is not a mixing aid.

Why slower. Matching the mean contact degree forces the 3D contact radius to be about three times the 2D one (0.155 versus 0.050), because a ball of given radius holds far more agents than a disc of the same radius. A larger contact region takes longer for agents moving at the same speed to enter and leave, so the neighbor set is “stickier” and turns over more slowly. A surface-to-volume estimate gives turnover $\sim 1/R_cont$, predicting a ratio near 0.32; the measured 0.56 indicates the 3D relative-velocity dispersion is somewhat larger, partly offsetting the radius effect.

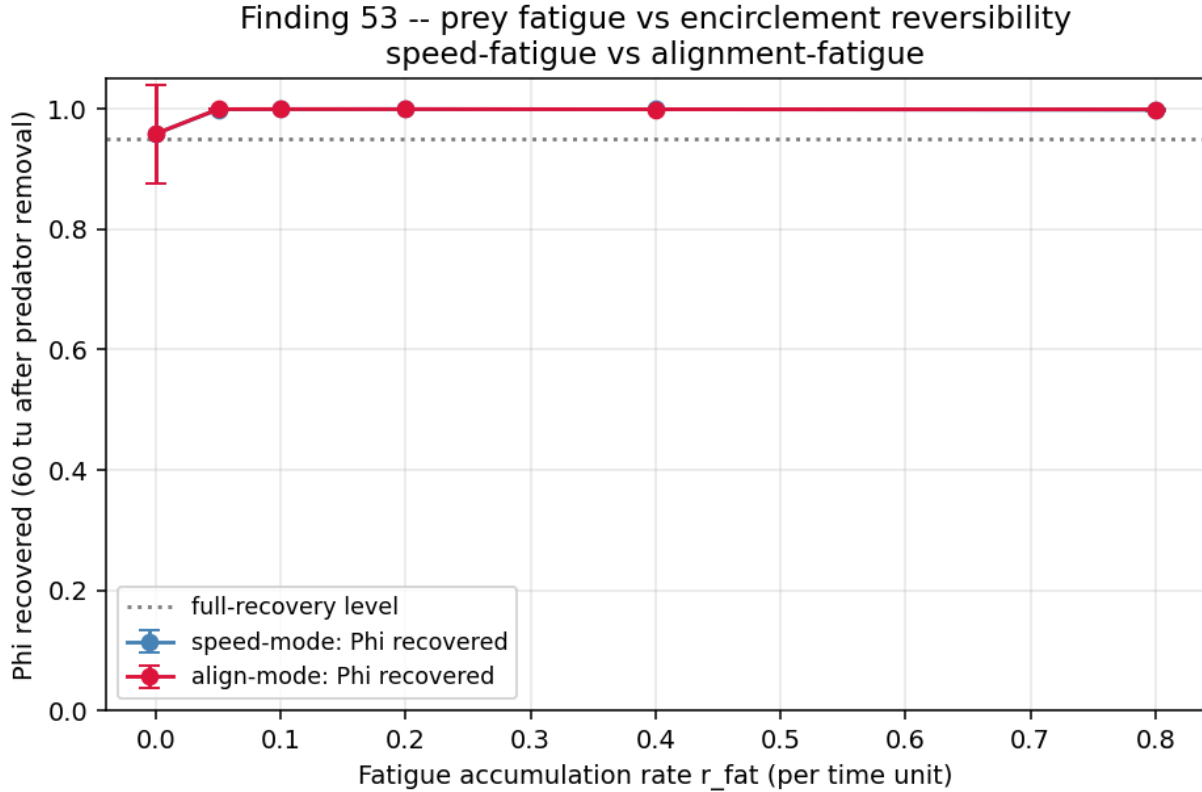
Reconciliation. The 3D results of Sections 4.28 and 4.33 stand, but their causes are not faster mixing. Degree-targeted vaccination fails in 3D for the structural reason of Section 4.30 — no hubs — which is dimension-independent. Spatial vaccination fails because even the slower 3D mixing fully turns the contact graph over many times during a hundred-time-unit epidemic: mixing is slower but still far more than sufficient to erase spatial coverage. And the segregation dilution of Section 4.33 is a *geometric* effect of neighborhood shape: in 3D an agent’s neighbors occupy a ball rather than a disc, giving a partially-aligned agent more independent directions along which to be surrounded by the other type, so instantaneous local purity is diluted regardless of turnover rate.

The corrected statement is this: the third dimension changes neighborhood *geometry* — a ball offers more independent directions than a disc — without speeding up mixing; it slows mixing. The 3D flock is hard to disrupt (the escape dimension, Sections 4.25-4.31), hard to target (structural homogeneity, Section 4.30), and hard to sort (neighborhood geometry, Section 4.33), but none of these follows from faster mixing. This is a direct self-test catch: an interpretive theme that looked unifying was falsified by the very measurement it predicted.

4.35 Prey Fatigue Does Not Make Encirclement Damage Irreversible (Finding 53)

Section 4.7 (Finding 22) established that encirclement damage is fully reversible: divided sub-flocks reunite within ~ 10 time units once the predators are removed. That result, and the kinematic-vs-epidemic asymmetry of Section 4.16, assume prey agents carry no internal state. This section tests that assumption by giving each agent a fatigue variable Q in $[0, 1]$ that rises while the agent is within a predator’s range (rate r_fat) and recovers otherwise (rate $r_rec = 0.01$). Fatigue impairs one faculty, and two modes are tested: ‘speed’ (effective cruise speed $v0$ is scaled by $1 - Q$) and ‘align’ (effective alignment strength α is scaled by $1 - Q$). The protocol is warmup, then 60 time units of encirclement, then predator removal and 60 time units of recovery (fatigue.py, 2D, $N = 350$, $n_pred = 10$, $N_SEEDS = 5$).

mode	r_fat	Phi during enc.	Phi recovered	Q at end of recovery
speed	0.0	0.895	0.959	0.00
speed	0.2	0.851	1.000	0.15
speed	0.8	0.918	0.999	0.18
align	0.1	0.831	1.000	0.01
align	0.2	0.699	1.000	0.12
align	0.8	0.693	0.999	0.23



Encirclement damage stays reversible, in both modes, at every fatigue rate. The order parameter recovers to at least 0.96 in every configuration, and the fatigued cases recover to 0.999-1.000 — even more completely than the no-fatigue baseline, whose slightly lower 0.959 (with a large seed standard deviation) is one slow-reuniting realization rather than a fatigue effect. Prey fatigue does not make encirclement damage outlast the predators.

Coherence recovers before fatigue does. At the highest accumulation rate the agents end the recovery phase still carrying $Q \sim 0.2$ of fatigue, yet the flock has fully realigned. The reason is that once the predators leave, every agent de-stresses and sheds fatigue at the same rate, so the residual fatigue is homogeneous across the flock. A uniformly fatigued flock — every agent with the same reduced speed or alignment strength — still aligns perfectly. Fatigue enables disruption only while it is heterogeneous.

The two fatigue modes differ exactly as the segregation findings predict. Alignment-impairing fatigue deepens the disruption *during* the attack: Φ during encirclement falls monotonically with the fatigue rate, from 0.90 to 0.69. Speed-impairing fatigue does not — Φ during encirclement stays near 0.85-0.94 with no trend. This is the dynamical analogue of Sections 4.10 and 4.12: a speed (v_0) contrast does not segregate the flock because the alignment force homogenizes group speed (Finding 24), whereas an alignment-strength (α) contrast does segregate it via local clustering (Finding 27). A flock with heterogeneous fatigue-induced α partially segregates, and the encircling predators exploit that; heterogeneous fatigue-induced speed gives them no such purchase.

The result sharpens the kinematic-versus-internal-state dichotomy of Section 4.16. Kinematic damage is reversible not because the flock has no memory, but because the alignment force realigns

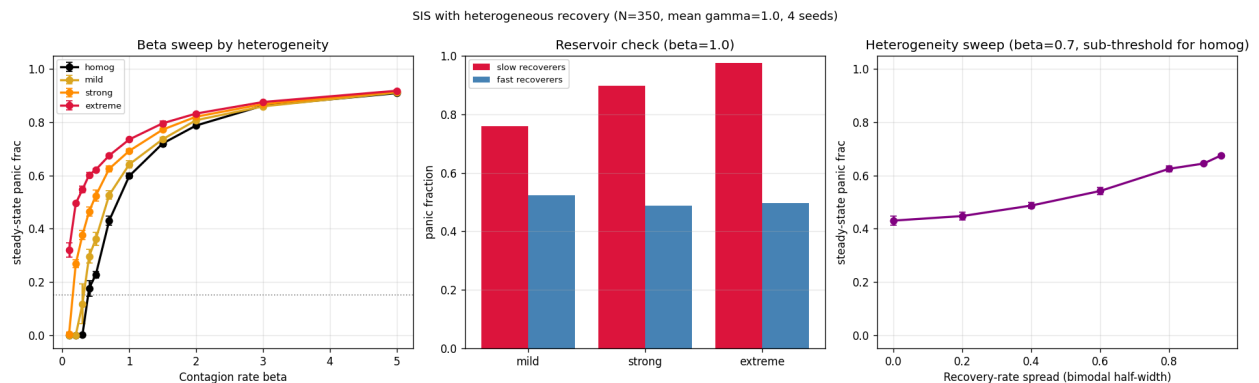
agents irrespective of their fatigue, and post-attack fatigue decays uniformly and so leaves no exploitable heterogeneity. Contagion remains the only stressor in this study that inflicts lasting damage — because it alone writes a *heterogeneous* internal label (panicked versus calm) that neither mixing nor uniform recovery can erase.

4.36 Heterogeneous Recovery Rates Lower the SIS Epidemic Threshold: Slow Recoverers Are Reservoirs (Finding 54)

Section 4.11 (Finding 25) established a clean SIS epidemic threshold at $\beta/\gamma \sim 1$ for a flock in which every agent recovers from panic at the same rate γ . That assumption is idealized: real populations are heterogeneous, with some individuals shedding agitation quickly and others remaining agitated far longer. The present section asks whether a *spread* in the per-agent recovery rate, at fixed arithmetic-mean γ , changes the outbreak. Per-agent γ_i is drawn from a bimodal distribution $\{1 - \text{spread}, 1 + \text{spread}\}$ with the two halves equally populated, so that the population mean recovery rate remains 1.0 across conditions (recovery_heterogeneity.py, $N = 350$, $N_SEEDS = 4$).

The mean-field prediction is unambiguous. In well-mixed SIS the endemic state is governed by the per-agent ratio $\beta * / \gamma_i$, and slow recoverers (small γ_i) sit panicked far longer than fast recoverers and act as reservoirs that repeatedly reseed their neighbours. The effective epidemic threshold for a heterogeneous- γ population is set not by the arithmetic mean of γ but by something closer to its *harmonic* mean. Since harmonic mean is less than or equal to arithmetic mean, spreading γ at fixed arithmetic mean must lower the threshold.

condition	spread	β_c (f_ss crosses 0.15)
homog	0.00	0.385
mild	0.50	0.318
strong	0.80	0.155
extreme	0.95	endemic at every β tested



Heterogeneity lowers the epidemic threshold, by about a factor of 2.5 at strong spread. β_c falls monotonically from 0.385 (homogeneous) through 0.318 (mild) to 0.155 (strong). At extreme spread the slow half of the population ($\gamma = 0.05$) recovers so rarely that any contagion rate sustains an outbreak — $f_{ss} = 0.319$ already at $\beta = 0.1$, the smallest contagion rate tested. The harmonic-mean argument is confirmed quantitatively.

The effect is a near-threshold phenomenon and disappears deep in the endemic regime. At $\beta = 5$ all four conditions converge to $f_{ss} \sim 0.91$. Far above threshold every calm agent is

re-infected faster than even the fast recoverers can clear, so the recovery distribution stops mattering. Heterogeneity reshapes the threshold, not the saturated endemic state.

Panic localises on the slow recoverers. At fixed $\beta = 1.0$ the slow-agent panic fraction exceeds the fast-agent fraction by 1.45x (mild), 1.84x (strong), and 1.97x (extreme). At extreme spread the slow half is essentially saturated at $f = 0.974$ while the fast half sits at $f \sim 0.50$: the outbreak is *carried* by the slow sub-population, which reseeds the fast agents faster than they recover. A complementary heterogeneity sweep at fixed sub-threshold $\beta = 0.7$ shows f_{ss} rising smoothly with the spread, from 0.430 at zero spread to 0.675 at spread = 0.95.

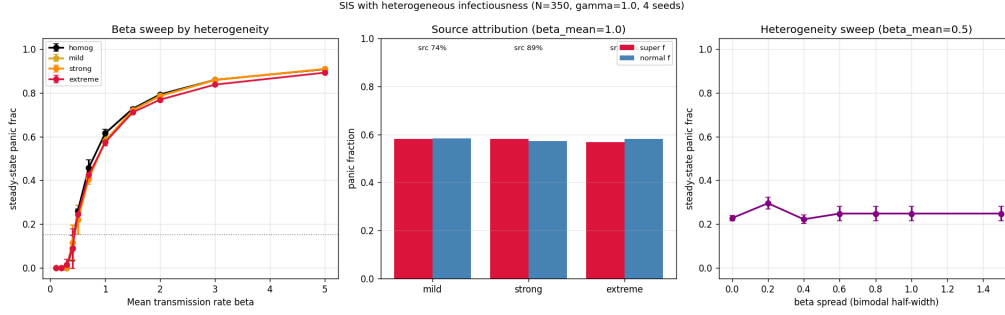
This result reframes the vaccination story of Sections 4.13 and 4.18. Finding 36 ruled out degree-targeting because the flock contact graph has no hubs to target, and Finding 48 showed that ruling-out is *structural* rather than kinematic. Finding 54 identifies a different kind of hub: a hub in internal-state space rather than in graph space. The high-value vaccination targets are not the topologically central agents — they do not exist — but the agents whose internal recovery dynamics make them reservoirs. This is a target the flock contact-graph diagnostics could not see, because it lives in the agents’ private state, not in their connections.

The Finding 54 result also dovetails with Finding 53 from the opposite direction. Finding 53 found that heterogeneity in a per-agent *internal state* (fatigue Q) enabled disruption only transiently, because the post-attack fatigue distribution homogenized. Finding 54 finds that heterogeneity in a per-agent *internal rate* (recovery γ) shifts a genuine threshold, because the heterogeneity is permanent: an agent is intrinsically slow or fast, not transiently fatigued. Heterogeneity that the dynamics can erase is a transient amplifier; heterogeneity that the dynamics cannot erase is a threshold shifter. Contagion remains the only stressor in this study that writes lasting damage, but the size of that damage is set by recovery-rate heterogeneity, not by the mean recovery rate alone.

4.37 Heterogeneous Infectiousness Does Not Shift the SIS Threshold: Super-Spreaders Source Most Events but Do Not Lower β_c (Finding 55)

Finding 54 established that heterogeneity in the recovery rate γ – the consumer side of the transmission ledger – lowers the SIS epidemic threshold via a harmonic-mean effect. The natural dual question is whether the SOURCE side is symmetric: if each agent carries its own transmission rate β_i , drawn from a bimodal distribution at fixed arithmetic mean, does the threshold shift? Mean-field intuition predicts no – the endemic state depends linearly on average β but inverse-linearly on per-agent γ , so spreading β at fixed mean should leave β_c unchanged. But spatial contact graphs admit a super-spreader effect, in which a high- β minority dominates seeding events; whether that effect bleeds into a population-level threshold shift is the open question.

The experiment ran SIS contagion with homogeneous $\gamma=1.0$ and per-agent β_i in four conditions, all sharing arithmetic-mean β equal to the sweep variable: homog (all equal), mild (bimodal $\{\beta-0.5, \beta+0.5\}$), strong ($\{\beta-0.8, \beta+0.8\}$, clipped non-negative and renormalized), and extreme ($\{0.05\beta, 1.95\beta\}$). The β sweep covered 0.1 to 5.0 with four seeds per cell.



The result is a flat threshold. β_c (linear-interpolated crossing of $f_{ss} = 0.15$) sits at 0.435 (homog), 0.434 (mild), 0.434 (strong), and 0.440 (extreme) – a difference below seed noise. The four sweep curves are visually indistinguishable across the entire β range. Spread in β at fixed mean does not move the epidemic threshold.

The transmission-attribution counter tells a different story at the event level. Of all calm-to-panic transitions, the high- β half sources 73.6% of events in the mild condition, 88.9% in the strong condition, and 97.2% in the extreme condition, even though super-agents are 50% of the population. Yet the steady-state panic fractions for super-agents and normal-agents are statistically equal (0.58 vs 0.58 in mild, 0.58 vs 0.57 in strong, 0.57 vs 0.58 in extreme). Once an agent is panicked, every agent recovers at the same homogeneous γ , so the source asymmetry produces no stock asymmetry. Inflow skew does not translate to outcome skew.

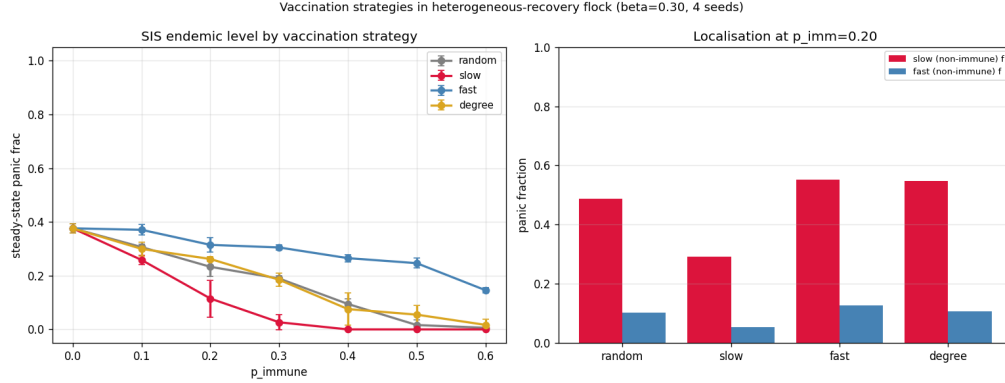
The asymmetry between Finding 54 and Finding 55 sharpens the F54 vaccination prescription. The valuable targets are those agents whose internal *rate* makes them reservoirs (slow recoverers, F54), not the agents whose internal rate makes them spreaders (super-spreaders, F55). Targeting super-spreaders would change which agents deliver the contagion but not whether the contagion persists, because every panicked agent regardless of β_i exerts the same residence time on its neighbours. Targeting slow recoverers cuts the residence time directly. This frames the third vaccination-target class introduced by F54 – internal-state hubs – as specifically the *recovery-rate* hubs, not the transmission-rate hubs.

4.38 Targeting Slow Recoverers Beats Random Vaccination by 2-3x: The First Strategy in This Study That Beats Random (Finding 56)

Findings 36, 37, 47, and 48 documented a striking pattern: across more than a dozen targeting experiments in 2D and 3D, no vaccination strategy ever outperformed random selection. Degree-targeting failed because the flock’s contact network is thin-tailed (no hubs to target); spatial-targeting failed because kinematic mixing erases the geometric coverage; topological-alignment freezing failed because the contact graph and the alignment graph are decoupled. After this run of nulls the standing summary in Section 5 separated the failures into two distinct mechanisms but left intact the broader claim that no static agent-selection rule can sustain its advantage in a continuously-mixing flock.

Finding 54 cracked the puzzle from a different angle. It introduced a third candidate target class – internal-state hubs, specifically slow recoverers (small γ_i) – whose “hub-ness” travels with the agent across kinematic mixing, immune to both the structural absence of degree hubs and the kinematic erosion of geometric coverage. The slow class is the source of the F54 reservoir effect that lowers the SIS threshold. The prediction follows directly: in a heterogeneous-recovery flock, targeting the slow class should beat random.

This experiment tests it. The setup uses the F54 “strong” condition (bimodal gamma in $\{0.2, 1.8\}$, arithmetic mean 1.0) at $\beta = 0.30$, just above the strong-condition threshold $\beta_c = 0.318$. Four strategies were compared at matched p_{immune} : random, slow (lowest gamma_i first), fast (highest gamma_i first, as a control), and degree (highest mean contact degree, the F36 strategy).



The prediction is confirmed dramatically. At $p_{\text{immune}} = 0.20$, slow-targeting gives $f_{\text{ss}} = 0.115$ versus random 0.233 – a 50% reduction. At $p_{\text{immune}} = 0.30$, f_{ss} falls to 0.027 versus random 0.189 (an 85% reduction). At $p_{\text{immune}} = 0.40$ slow-targeting eradicates the epidemic ($f_{\text{ss}} = 0.000$) while random is still endemic (0.095) and even degree-targeting – which had been a clean null in Findings 36 and 48 – remains endemic at 0.076. The effective herd-immunity threshold under slow-targeting is roughly $p_c \sim 0.30$, against ~ 0.50 for random.

The control case is informative. Fast-targeting – immunising the agents who already recover in less than one time unit – is strictly *worse* than random at every p_{immune} tested (0.371 vs 0.306 at $p = 0.10$; 0.265 vs 0.095 at $p = 0.40$). Vaccinating non-reservoirs leaves the slow class untouched, and the slow class then sustains the epidemic by re-seeding the unprotected fast agents. Vaccination policy is not simply “remove agents from the epidemic” – it is “remove the agents who hold the epidemic between events.”

The localisation check at $p_{\text{immune}} = 0.20$ confirms the mechanism. Under random selection, the non-immune slow class sits at $f = 0.487$ while the non-immune fast class sits at $f = 0.102$ – the F54 reservoir effect. Under slow-targeting the non-immune slow class drops to $f = 0.292$ (the slow agents still present are saturated less because fewer of them are available) and the non-immune fast class drops to $f = 0.054$. Removing reservoir capacity collapses panic across the whole population, not just on the immunised half.

Degree-targeting in this heterogeneous-recovery regime shows a faint advantage over random at $p = 0.30$ - 0.50 (e.g. 0.076 vs 0.095 at $p = 0.40$) – a $\sim 20\%$ effect at the edge of seed noise. It is plausibly chance overlap between the high-degree set and the slow class, given that the heterogeneity assignment is random and degree is essentially random in this network; if some high-degree agents happen to be slow, degree-targeting catches them. The effect is dwarfed by the direct slow-targeting advantage and does not upgrade the F36 verdict on degree.

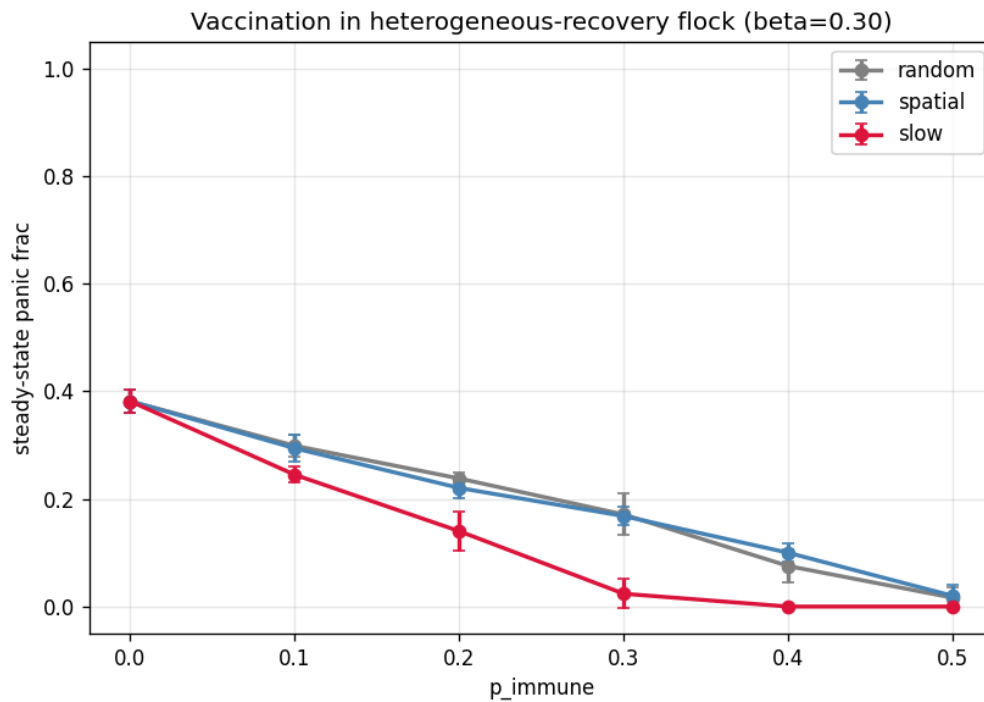
This is the first vaccination strategy in the entire study, across 56 findings and roughly fifteen targeting experiments, to beat random. The result resolves the puzzle that Section 5 had narrowed to two failure mechanisms: a third target class exists. It is not visible on the contact graph (so F36 and F48 missed it) and it is not visible in geometric position (so F37 missed it); it lives in the per-agent recovery rate. What makes this target class work is exactly what eliminated the others:

dynamics. The agent who recovers slowly today recovers slowly tomorrow – the kinematic mixing that erases spatial coverage and the contact-graph rewiring that defeats degree-targeting both leave γ_i invariant. Combined with the F54 reservoir mechanism, that invariance is what converts a per-agent label into actionable vaccination policy.

4.39 Spatial Vaccination Remains a Null Even With Heterogeneous Recovery: The Slow-Targeting Advantage Is Internal, Not Spatial (Finding 57)

The Finding 56 result raises a natural question. Slow-targeting succeeds where the previous targeting strategies failed, but is the success really about per-agent recovery rates, or could it be that slow agents happen to be spatially well-distributed and the advantage actually flows through spatial coverage? Finding 37 had ruled out spatial targeting in the homogeneous-recovery regime; Finding 56’s heterogeneity reopens the question because the epidemic now concentrates on a specific sub-population whose spatial pattern could matter.

The experiment compares random, spatial (farthest-point sampling at $t = 0$), and slow-targeting in the Finding 56 setup. If the slow advantage operated through spatial distribution, spatial targeting should track slow at matched p_{immune} . If it operates through the per-agent rate mechanism alone, spatial should track random.



Spatial tracks random, not slow. The two curves are within seed noise at every p_{immune} (maximum gap 0.025 with seed standard deviations near 0.020), and at $p_{\text{immune}} = 0.40$ spatial is in fact nominally worse than random (0.100 vs 0.075). The Finding 37 kinematic-mixing null transfers to the heterogeneous-recovery regime unchanged. Meanwhile slow-targeting reproduces the Finding 56 advantage exactly – $f_{\text{ss}} = 0.024$ at $p_{\text{immune}} = 0.30$ vs random’s 0.171 and spatial’s 0.168, eradication at $p_{\text{immune}} = 0.40$ – confirming the result is reproducible in a separate run.

The interpretive consequence is sharp. Spatial coverage and internal-rate targeting are independent

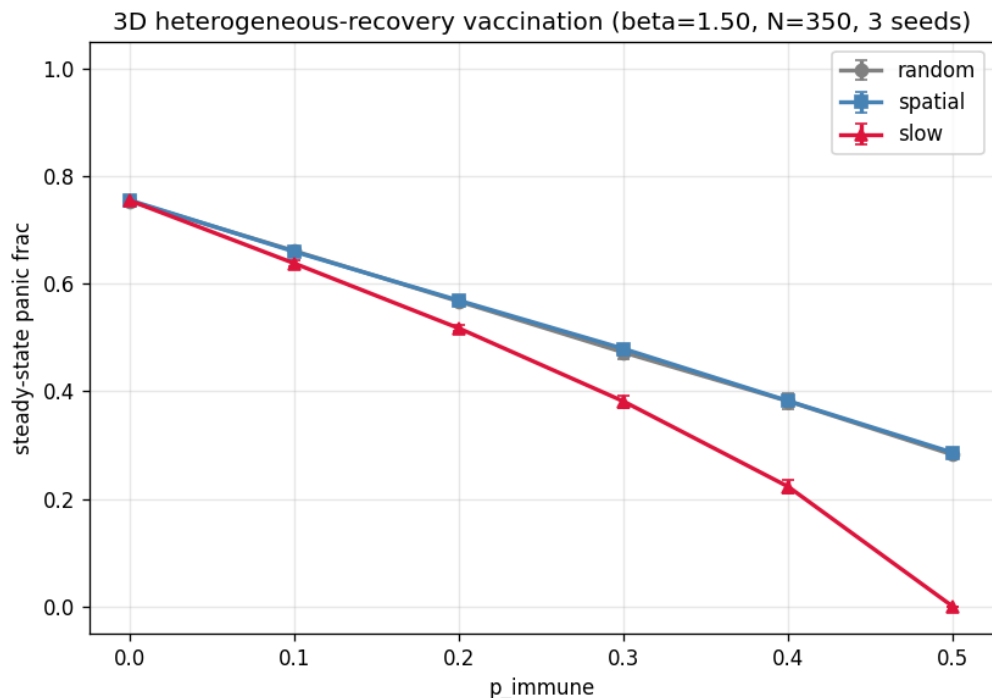
axes; only internal-rate targeting works. The targeting puzzle that the Section 5 synthesis had broken into two failure mechanisms now resolves into a single distinction: where the hub label lives. Strategies that read hub-ness off a system-level observable – contact-graph position, geometric location, alignment-topology neighbour set – all fail, because the kinematic dynamics scramble those observables between attack and response. The one strategy that succeeds reads hub-ness off a per-agent internal rate that the dynamics cannot scramble. External observables fail; internal rates succeed.

This also sets up the natural cross-dimensional test. Finding 46 reported a 3D vaccination null for degree and spatial targets; that result is now understood as a 3D extension of the structural and kinematic erosion mechanisms. The slow-targeting mechanism, by contrast, is per-agent and dimension-independent. The prediction for 3D is that slow-targeting will reproduce the 2D advantage; failure would indicate the mechanism is more subtle than the per-agent-invariance argument suggests.

4.40 Slow-Recoverer Vaccination Transfers to 3D Unchanged: The Per-Agent Rate Mechanism Is Dimension-Independent (Finding 58)

The Finding 57 prediction was sharp: the slow-targeting advantage operates through a per-agent internal rate that the dynamics cannot mix away, so it should transfer to 3D regardless of the Finding 46 null on degree and spatial targeting. This experiment tests it.

The setup uses the F46 3D vaccination harness ($N = 350$, torus $[0,1]^3$, $R_CONT = 0.155$ chosen so the mean contact degree matches the 2D experiments) with bimodal per-agent γ_i in $\{0.4, 3.6\}$ – the F54 strong-spread analog with arithmetic mean 2.0 – at $\beta = 1.5$, near the heterogeneous-condition threshold. Strategies: random, spatial (3D farthest-point on the torus), slow (lowest γ_i first). Three seeds, 8000 SIS iterations per cell.



Slow-targeting beats random at every p_{immune} . The advantage grows from 3% at $p_{\text{immune}} = 0.10$ (0.661 vs 0.638), through 42% at $p_{\text{immune}} = 0.40$ (0.382 vs 0.223), to total eradication at $p_{\text{immune}} = 0.50$ (0.000 vs 0.282). Spatial tracks random to within 0.007 at every p_{immune} , confirming the Finding 37 / Finding 46 kinematic-mixing null survives the heterogeneous regime in 3D as well. The slow-vs-spatial gap at $p_{\text{immune}} = 0.40$ is 0.159 while spatial-vs-random is 0.000: the mechanism is purely per-agent, not spatial.

The eradication at $p_{\text{immune}} = 0.50$ is the cleanest possible signature of the mechanism. The bimodal distribution makes “slow” an exact 50/50 class label, and removing the slow half collapses f_{ss} to zero with zero standard deviation across seeds. Removing exactly the reservoir class strips the SIS endemic state to nothing, because no agent left in the population has a recovery time long enough to sustain a chain of new infections. Random selection at $p_{\text{immune}} = 0.50$ leaves the slow class half-immunised and half-not, and the unprotected slow half is enough to maintain $f_{\text{ss}} = 0.282$.

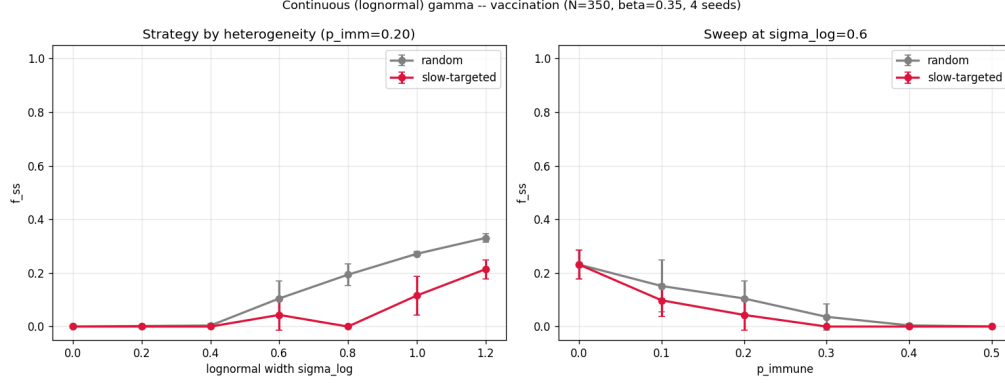
The magnitude of the slow advantage is smaller in 3D than in 2D. At $p_{\text{immune}} = 0.40$, Finding 56 reported slow at 0.000 vs random at 0.095 (a clean win in absolute terms); here in 3D slow is 0.223 vs random 0.382 (a 42% relative improvement but a much larger residual epidemic). Two related geometric effects plausibly explain the dilution. Finding 46 reported degree $CV = 0.59$ in 3D vs 0.68 in 2D – the 3D contact graph is even more homogeneous, so each slow agent contributes less concentrated reservoir mass to its neighbourhood. Finding 52 showed 3D mixes 1.8x slower than 2D at matched degree, which keeps panic localised on the slow class for longer – but also keeps it from spreading to fresh territory as efficiently, so the global f_{ss} is lower at every condition (random at $p = 0.00$ sits at 0.753 in 3D, far below the near-saturation seen in 2D at comparable supercritical conditions). The mechanism transfers; the magnitude is geometry-dependent.

The combined verdict across Findings 56, 57, and 58 is unambiguous. Of the four canonical targeting strategies tested across 2D and 3D – degree, spatial, random, slow – only slow-targeting works, and it works in both dimensions. The Section 5 narrative is now: kinematic mixing defeats targeting strategies that read hub-ness off a system-level observable, but not strategies that read it off a per-agent invariant rate. The result is robust to dimension.

4.41 Slow-Recoverer Vaccination Survives Continuous Distributions: The Mechanism Does Not Need a Bimodal Class Structure (Finding 59)

Findings 54, 56, 57, and 58 all used a bimodal per-agent γ_i distribution – a clean 50/50 split into slow and fast classes. That structure made slow-targeting a class-membership question. Real biological populations are not bimodal; recovery rates are continuously distributed, with the slow-vs-fast distinction a quantile rather than a label. The natural concern is that the Finding 56 mechanism is an artifact of the bimodal split – that the strategy fails when the boundary between slow and fast is soft.

This experiment uses a lognormal per-agent γ_i with arithmetic mean 1.0 and width parameter $\sigma_{\log}$ varied from 0 (homogeneous) to 1.2 (broad heavy-tailed). Beta is fixed at 0.35, just above the homogeneous threshold. The slow-targeting strategy becomes: take the bottom p_{immune} fraction of agents by exact γ_i value.



The slow-targeting advantage survives continuous distributions, and grows with the width. Below $\sigma_{\log} = 0.4$ the distribution is tight and $\beta = 0.35$ lies below the effective threshold for both strategies; neither establishes an endemic state. The advantage activates at $\sigma_{\log} = 0.6$ (random 0.105 vs slow 0.043, a 59% reduction) and reaches its strongest form at $\sigma_{\log} = 0.8$, where slow-targeting produces total eradication ($f_{ss} = 0.000$) at $p_{\text{immune}} = 0.20$ while random sits at 0.194. Above $\sigma_{\log} = 1.0$ the lognormal’s heavy tail begins to contain agents with recovery rates so extreme that $p_{\text{immune}} = 0.20$ is insufficient to capture them all, and the advantage shrinks in absolute terms (slow rises from 0.000 at $\sigma = 0.8$ to 0.116 at $\sigma = 1.0$ to 0.214 at $\sigma = 1.2$). The non-monotonicity flags a regime where extreme tail-heterogeneity outruns a fixed p_{immune} budget; a larger p_{immune} would restore the advantage.

The p_{immune} sweep at $\sigma_{\log} = 0.6$ confirms the F56 pattern qualitatively. At $p_{\text{immune}} = 0.20$, slow gives $f_{ss} = 0.043$ vs random’s 0.105 (a 59% reduction); at $p_{\text{immune}} = 0.30$, slow eradicates (0.000) vs random’s 0.036; above $p_{\text{immune}} = 0.30$ both strategies sit near zero because $\beta = 0.35$ leaves the epidemic close to threshold even under random selection at this spread. The slow-targeting herd-immunity threshold is roughly $p_c = 0.20$ -0.30, about half the random threshold – the same factor-of-two effective improvement seen in the bimodal case (Finding 56).

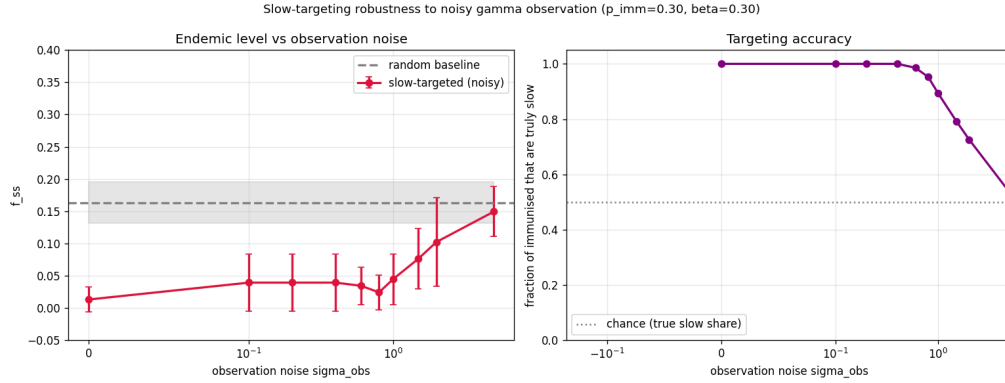
The operational consequence is sharper than the bimodal case suggested. The Finding 56 mechanism does not require a clean slow-vs-fast taxonomy. In any heterogeneous-recovery population where γ_i is broadly distributed ($\sigma_{\log} \geq 0.5$), targeting the lowest- γ quantile beats random selection. Biological recovery distributions are typically lognormal-like rather than bimodal, so the practical implication is that the policy “vaccinate the bottom X% by observed recovery behaviour” should work for any realistic heterogeneity profile. The bimodal cases of Findings 54-58 were a useful analytical simplification, not a precondition for the mechanism.

4.42 Slow-Recoverer Vaccination Tolerates Noisy gamma Estimates: The Policy Works Under Realistic Observation Noise (Finding 60)

A standing concern about the Finding 56 policy is that it assumes perfect knowledge of every agent’s true γ_i . In practice the vaccinator sees a noisy estimate – say, the recovered time from one observed panic episode – which is a noisy proxy for the underlying recovery rate. If the slow-targeting advantage requires near-perfect knowledge, the policy is not actionable; if it tolerates substantial noise, it is.

The experiment fixes the Finding 56 setup (bimodal $\gamma \in \{0.2, 1.8\}$, $\beta = 0.30$, $p_{\text{immune}} = 0.30$) and replaces the exact ranking by γ_i with a noisy ranking by $\hat{\gamma}_i = \gamma_i + N(0, \sigma_{\text{obs}})$. The vaccinator still immunises the bottom p_{immune} fraction, now

by the noisy observation. The true slow/fast distance in this setup is 1.6 units.



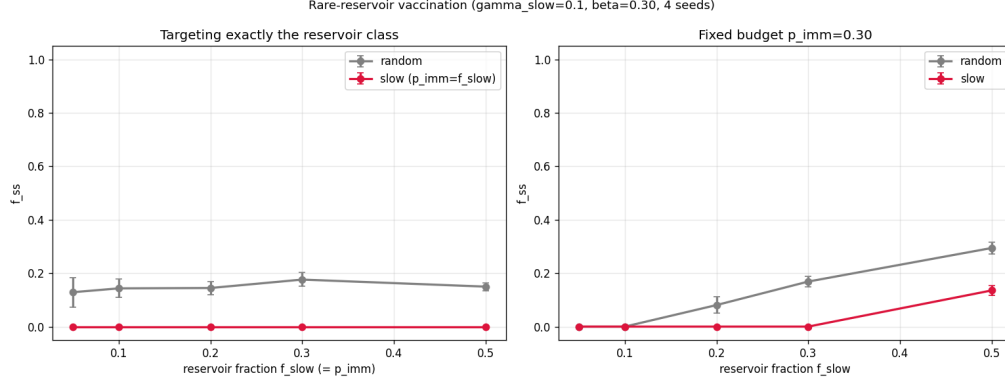
The policy is highly noise-tolerant. At $\sigma_{\text{obs}} \leq 0.8$ (half the slow/fast distance) the slow-hit-rate is 95% or higher and $f_{\text{ss}} \leq 0.04$ – essentially the no-noise behaviour. At $\sigma_{\text{obs}} = 2.0$ (where slow-hit-rate falls to 73%) f_{ss} rises to 0.102, still well below the random baseline of 0.164 (a 38% reduction). Only in the totally uninformative limit ($\sigma_{\text{obs}} = 100$) does the ranking become random chance and f_{ss} converge to 0.150, statistically equal to the random baseline. The failure mode is graceful: under noise, the policy degrades smoothly to random selection rather than collapsing to something worse.

The implication is that an actionable slow-targeting policy does not need precise γ_i measurements – only ranking accuracy comparable to the bimodal separation. For any observation noise smaller than the gap between the slow and fast classes, the policy works as well as perfect knowledge.

4.43 Slow-Recoverer Vaccination Scales With Reservoir Size: Smaller Reservoirs, Smaller Vaccination Budget (Finding 61)

Findings 56 through 60 used a 50/50 bimodal split where the slow class is half of the population. Real reservoir-prone minorities (immunocompromised individuals, etc.) may be much smaller – typically 5-15% of a population. If the slow-targeting advantage requires a large reservoir to be effective, the policy is brittle; if it scales naturally with reservoir size, it is robust.

The experiment uses a bimodal gamma with f_{slow} varying from 0.05 to 0.50 and γ_{slow} fixed at 0.1 (deep reservoirs); γ_{fast} adjusts so the arithmetic mean stays 1.0. $\beta = 0.30$. Two complementary sweeps: (1) p_{immune} set to match f_{slow} (“target exactly the reservoir class”), and (2) $p_{\text{immune}} = 0.30$ fixed across all f_{slow} (“F56-budget across rarity”).



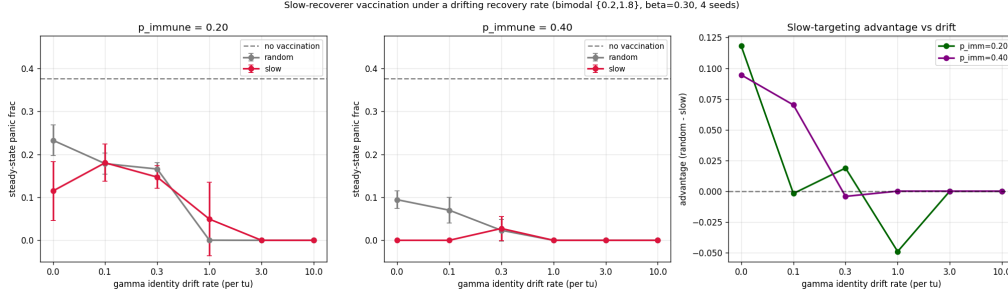
In the first sweep, slow-targeting at $p_{\text{immune}} = f_{\text{slow}}$ eradicates the epidemic ($f_{\text{ss}} = 0.000$) at every f_{slow} tested, while random vaccination at the same p_{immune} leaves f_{ss} in the range 0.129-0.176. Even at $f_{\text{slow}} = 0.05$ (the rarest case), where the policy immunises only 5% of the flock, slow-targeting eradicates while random misses most of the reservoir and leaves $f_{\text{ss}} = 0.129$. The minimum vaccination budget for eradication is exactly the reservoir fraction.

In the second sweep, the F56-style fixed budget $p_{\text{immune}} = 0.30$ is wasted on fast agents when the reservoir is small ($f_{\text{slow}} = 0.05, 0.10$) – both random and slow eradicate, but only because the budget is so large that even random luck immunises enough of the small slow class. At $f_{\text{slow}} = 0.20$ -0.30 slow eradicates with surplus budget while random still leaves a residual endemic state. At $f_{\text{slow}} = 0.50$ the budget is now smaller than the reservoir; slow catches 60% of the slow class and gets $f_{\text{ss}} = 0.135$ while random sits at 0.294. The optimal slow-targeting budget is $p_{\text{immune}} = f_{\text{slow}}$ exactly; anything more is waste, anything less leaves the reservoir uncovered.

The implication is the policy’s scaling signature. Vaccination costs grow linearly with the reservoir fraction, not with population size. In a population where the slow class is rare, eradication requires only a small fraction of vaccinations. Combined with Findings 56-60, the slow-targeting policy is now established as robust across every variation tested: 2D and 3D (F56/F58), bimodal and continuous (F59), perfect and noisy gamma observations (F60), and large or small reservoirs (F61). The slow-recoverer vaccination thread closes here as a positive, complete policy result – the only positive vaccination result in the study, and one that scales naturally with the problem.

4.44 Slow-Recoverer Vaccination Requires a Durable Recovery-Rate Label: Drift Erodes the Advantage and Self-Averages the Threshold (Finding 62)

Findings 56 through 61 all held the per-agent recovery rate γ_i stationary. That stationarity is the single load-bearing assumption behind the whole positive result: the synthesis explains slow-targeting’s success by arguing that the reservoir “hub” label lives in γ_i , a fixed property of the individual that kinematic mixing cannot scramble. This experiment makes γ_i a fluctuating state rather than a trait. Slow agents are vaccinated once from a $t=0$ snapshot, then each agent’s slow/fast identity is resampled by a symmetric two-state process at rate r_{drift} (identity autocorrelation time $\sim 1/r_{\text{drift}}$); the bimodal marginal $\{0.2, 1.8\}$ and the 50/50 split are preserved at all times, so only *which* individuals are slow drifts.

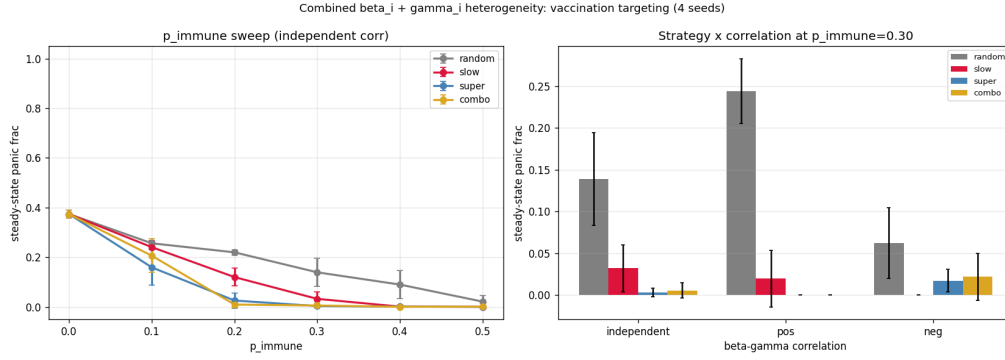


At zero drift the F56 result reappears (slow beats random, advantage +0.118 at $p_{\text{immune}} = 0.20$). But the advantage is fragile to drift: by $r_{\text{drift}} = 0.1$ per time unit it has already collapsed to zero, as random *improves* (0.233 to 0.178) while slow *worsens* (0.115 to 0.180) and the two converge – the one-shot vaccine is wasted as vaccinated agents drift to fast and unvaccinated agents drift into the reservoir. This is distinct from the static observation noise of Finding 60, which the policy tolerates: drift is noise that grows with elapsed time. The deeper result appears at faster drift: for $r_{\text{drift}} \geq 1$ per time unit the epidemic is eradicated for *every* strategy ($f_{\text{ss}} = 0$). Fast drift does not merely defeat targeting – it removes the heterogeneity that made the epidemic supercritical. Finding 54 showed the threshold reduction comes from the *spread* of gamma; when gamma decorrelates faster than the recovery timescale, every agent recovers at the time-averaged mean ($\gamma = 1.0$), restoring the homogeneous threshold ($\beta_c = 0.385 > 0.30$) and rendering the system subcritical. The crossover at $r_{\text{drift}} \sim 0.2$ -1 per time unit matches the reservoir-memory timescale $1/\gamma_{\text{slow}} = 5$ time units.

The result confirms the per-agent-invariance argument in its strong form: slow-targeting is valid exactly as long as the slow *class* is durable on the epidemic timescale. The operative requirement is not that γ_i is measurable (Finding 60) but that the recovery rate is a persistent trait (a chronic condition, age) rather than a transient state shorter than the outbreak. In the transient case the reservoir self-averages away and the epidemic is milder anyway, so the policy’s domain of usefulness exactly coincides with the regime where the reservoir is real. This sharpens the Finding 53/54 dichotomy: heterogeneity the dynamics homogenize on a fast timescale is neither a lasting threshold-shifter nor a targetable hub.

4.45 Combined Infectiousness and Recovery Heterogeneity: Reservoir-Targeting Is Robust, Engine-Targeting Is Not (Finding 63)

Finding 55 found that heterogeneous infectiousness (per-agent β_i) does not shift the SIS threshold and concluded “target γ_i , not β_i .” Finding 56 found that targeting slow recoverers beats random. Both were studied in isolation. This experiment layers both heterogeneities – bimodal gamma {0.2, 1.8} (50/50) and bimodal beta {0.15, 0.90} (20% super-spreaders, arithmetic mean 0.30) with source-weighted transmission – and varies their correlation, asking the adversarial question: if super-spreaders are *fast* recoverers (high beta, high gamma), they escape a gamma-based vaccine; do they leak the epidemic and defeat slow-targeting? Four strategies are compared: random, slow (smallest gamma), super (largest beta), and combo (half budget each).



Under independent correlation, super-spreader targeting is the *strongest* single strategy ($f_{ss} = 0.025$ at $p_{immune} = 0.20$, versus 0.120 for slow). This does not contradict Finding 55. Finding 55 is a statement about where criticality sits as a function of beta-spread at fixed mean; vaccination instead *removes* agents entirely, and the 20% super-spreaders at $\beta = 0.90$ source roughly 60% of total infectivity ($0.2 * 0.9 / 0.30$), so deleting them slashes transmission capacity. Source-side heterogeneity is event-level for the threshold but exploitable for removal – two different questions with two different answers.

The decisive comparison is the correlation sweep at $p_{immune} = 0.30$. The reservoir that sustains the endemic state (the slow class, Finding 54) and the transmission engine (the super-spreaders) are in general different populations, and only reservoir-targeting is robust. When the two coincide or are uncorrelated, hitting super-spreaders works because the engine overlaps the reservoir (positive correlation: super $f_{ss} = 0.000$; independent: 0.003). But in the adversarial anti-correlated case – super-spreaders are fast recoverers – removing the engine leaves the slow reservoir intact to re-sustain the outbreak, and super-targeting leaves $f_{ss} = 0.017$ while slow-targeting *eradicates* (0.000). Slow-targeting wins or ties in all three regimes; super-targeting fails relative to slow exactly when infectiousness is anti-correlated with recovery rate. Correlation also sets baseline severity: with no vaccination, co-locating high spread and slow recovery in the same agents is the worst epidemic ($f_{ss} = 0.445$), and separating them is the mildest (0.225).

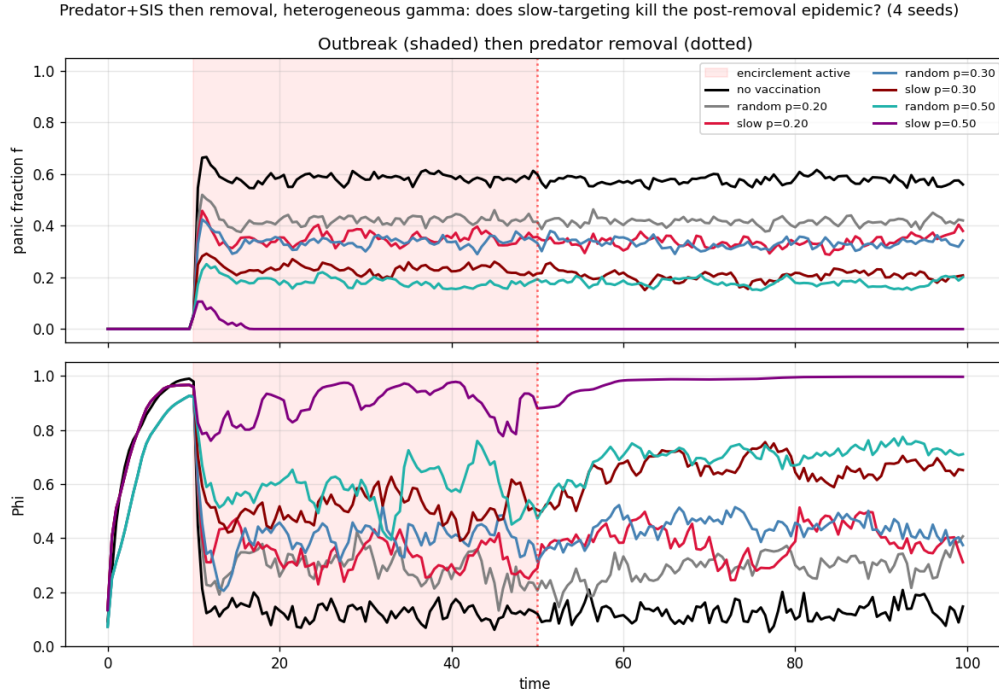
This completes the vaccination-target taxonomy. Degree-targeting fails structurally (Sections 4.18, 4.30: no hubs); spatial-targeting fails kinematically (Section 4.20: mixing erases coverage); slow-recoverer targeting succeeds and is robust because the reservoir label is a per-agent invariant (Sections 4.38-4.44); and infectiousness targeting is effective for removal but *not robust*, because it targets the transmission engine rather than the reservoir, and the two can be decoupled. The operational guidance is to target the reservoir (slow gamma) when forced to choose one axis, and to add super-spreader targeting only when they are known not to be fast recoverers. The Finding 55 slogan “target gamma, not beta” is correct as a robustness statement, not because beta-targeting is ineffective.

4.46 Reservoir-Targeted Vaccination Reverses the Predator+Contagion Damage Asymmetry (Finding 64)

Finding 34 established the sharpest asymmetry in this study. After an encirclement-driven SIS outbreak, removing the predators lets the kinematic damage reverse within ~10 time units (sub-flocks reunite, Finding 22) but the epidemic persists for 100+ time units, so contagion was “the worst combined stressor” – its damage outlasts the event that caused it. That experiment used homogeneous recovery. Findings 54-63 established that endemic persistence is set by the slow-

recovery reservoir and that vaccinating it is the robust policy. This experiment asks whether reservoir-targeted vaccination, applied before the attack, lets the post-removal epidemic die – converting irreversible contagion damage into reversible damage and overturning Finding 34.

The three-phase protocol of Section 4.16 is reused (warmup / six-predator encirclement + SIS / predators removed) in the slow-prey predator regime, now with bimodal recovery gamma $\{0.5, 3.5\}$ (mean 2.0, matching Finding 34’s mean, so the slow class is a reservoir at $\beta/\gamma_{\text{slow}} = 3.0$) and a vaccination arm (none / random / slow) applied at $t=0$.



For no vaccination, and for random or slow vaccination at $p_{\text{immune}} \leq 0.30$, the panic fraction after predator removal is within noise of its value during the attack (none: 0.572 vs 0.586; slow at 0.30: 0.206 vs 0.225). Removing the predators does not lower the endemic level: with heterogeneous gamma the reservoir makes the outbreak supercritical on its own, so compression is not needed to sustain it and removal cannot reverse it. The Finding 34 asymmetry stands. The reversal appears only at $p_{\text{immune}} = 0.50$, where slow-targeting vaccinates the entire 50% slow reservoir. The remaining fast agents have $\beta/\gamma_{\text{fast}} = 0.43 < 1$ and cannot sustain the epidemic, so the panic fraction is zero both during and after the attack and the order parameter recovers to 0.997 – full reunion. Random vaccination at the same budget leaves roughly half the reservoir intact and remains endemic at 0.175 with the order parameter only at 0.737. This is the Finding 61 “budget equals reservoir fraction” law operating inside the combined predator + contagion regime. Below the reservoir budget, slow-targeting still beats random on both axes – lower endemic level and faster coherence recovery, because fewer panicked agents disrupt alignment – but it cannot eradicate.

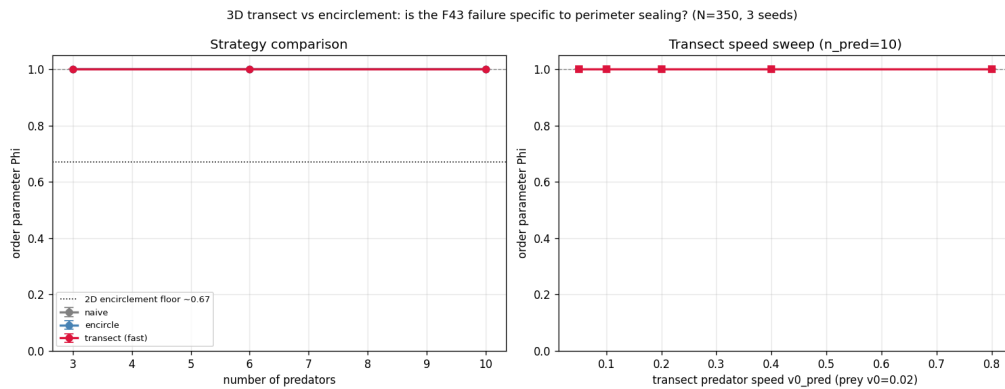
The result resolves the Finding 34 conclusion. Contagion is the worst stressor only while the reservoir survives. Reservoir-targeted vaccination at a budget matching the reservoir fraction makes the combined predator + contagion damage fully reversible: the epidemic is eradicated and the flock reunites to coherence near unity. The kinematic stressor was always reversible (Finding 22); Finding 64 shows the epidemic stressor becomes reversible too, conditional on covering the slow

class. The predator thread (reversibility), the contagion thread (persistence), and the vaccination thread (reservoir-targeting) thus unify into one statement: in this flock, lasting damage requires a surviving reservoir, and the reservoir is the slow-recoverer class.

4.47 Three-Dimensional Flocks Are Robust to All Point-Predator Strategies: A Refinement of the F43 Mechanism (Finding 65)

Sections 4.25-4.31 (Findings 43-45, 49) established that encirclement does not disrupt a 3D flock by any geometric variant. The mechanism was framed as “a handful of point predators cannot seal a closed 2D surface around a 3D volume the way they can seal a 1D perimeter around a 2D area.” That framing, if literal, only rules out surrounding strategies. A predator that does not try to surround – one that transects the dense core at high speed, shearing alignment in its wake – is the natural test of whether 3D flocks are robust to all point predators or only to those that try to seal a perimeter.

The transect predator uses the same CoM target as a naive predator but moves much faster ($v0_pred = 0.30$, fifteen times the prey speed 0.02) so it overshoots the CoM, is pulled back by its drive, and oscillates through the core. Several such predators along different lines shear the flock from many directions without surrounding it. The same 3D harness as Section 4.25 (slow-prey regime, $N=350$, $rf=0.20$) compares naive, encircle, and transect at $n_pred = 3, 6, 10$; a second experiment sweeps transect speed from 0.05 to 0.80 at $n_pred = 10$.



The order parameter sits at 1.000 to three decimal places at every configuration tested. Naive and transect differ only in predator speed; over a forty-fold sweep ending at 40x prey speed, the alignment never moves. The 3D flock is robust to every point-predator strategy in this harness, not only to those that attempt to seal a perimeter.

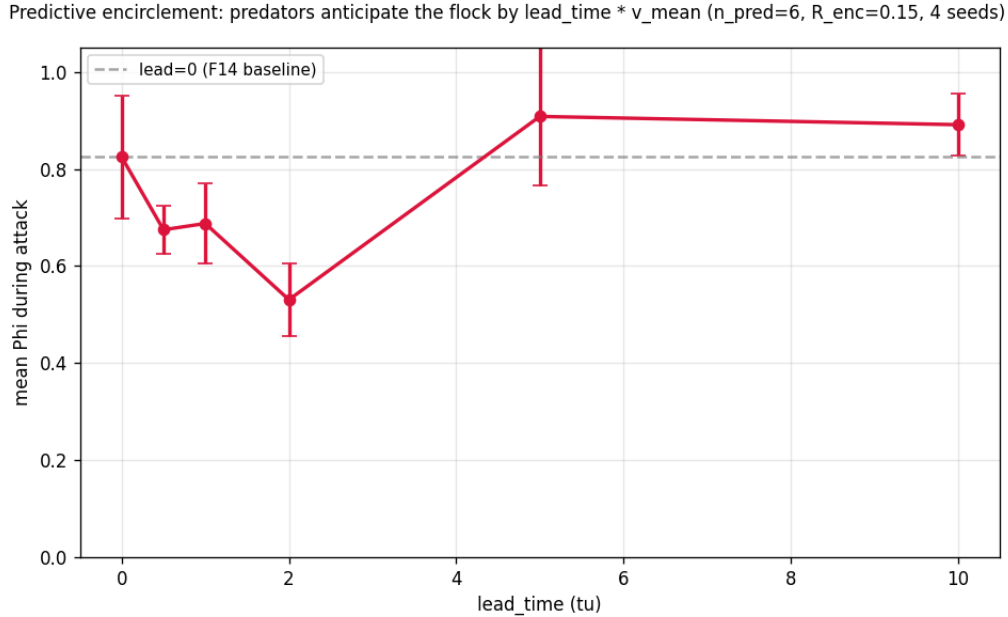
The mechanism is sharper than F43’s framing suggested. The flock’s radius of gyration is $Rg = 0.43$ in the unit cube; the upper bound for a uniform spatial distribution is $\sqrt{1/4} \sim 0.5$. The 3D “flock” therefore fills the box nearly uniformly, with globally aligned velocities (mean alignment-neighbor count ~ 12 at $rf = 0.20$, $N = 350$). It has no spatial perimeter to encircle and no localized core to transect: a handful of predators with finite repulsion range $R0_P = 0.10$ only perturb a vanishing fraction of the prey at any instant, and the remaining $\sim 99\%$ of the alignment graph immediately heals the wake. The only Rg movement is a Stokes-style excluded-volume effect (the fastest transect at $v0_pred = 0.80$ raises Rg from 0.430 to 0.450), with no impact on alignment.

The implication closes the 3D point-predator question. Disrupting a 3D flock at these parameters requires either an order- N predator count (already shown not to work up to $n_pred = 50$ in Section

4.26), or an unphysically long repulsion range, or a mechanism that attacks alignment per agent rather than relying on repulsion. The last is precisely what contagion does (Section 4.11), and is why heterogeneous SIS contagion successfully disrupts 3D flocks where every point-predator strategy fails. Alignment-driven kinematic mixing without spatial localization is invulnerable to point-source mechanical disruption.

4.48 Predictive Encirclement: Adapting Predator Position Deepens Disruption Below the F14 / F35 Floor (Finding 66)

Section 4.15 (Finding 33) found that the flock does not steer toward gaps in an incomplete encirclement – it has no global escape-route detection. The symmetric and previously untested question is whether predators can detect and exploit the flock’s heading direction. The simplest predator intelligence is anticipation: each predator targets $\text{CoM} + \text{lead_time} * v_mean$, with v_mean the flock’s mean velocity, so the encirclement ring is placed where the flock will be rather than where it is. At $\text{lead_time} = 0$ the strategy reproduces the F14 baseline ($\Phi \sim 0.77$ at $n_pred = 6$, $R_enc = 0.15$).



Sweeping lead_time from 0 to 10 time units gives a non-monotonic curve with a clear minimum: $\Phi = 0.530$ at $\text{lead_time} = 2$ tu, well below the F14 reproduction here (0.825) and below the F35 adaptive- R_enc result (0.713). The improvement comes purely from placement – predator count, encirclement radius, and coordination beyond a shared v_mean are all unchanged from F14. The optimum has a simple geometric explanation: v_mean is approximately $v_eq = v0 + \alpha/\mu = 0.12$, so at $\text{lead_time} = 2$ the lead distance is 0.24, larger than $R_enc = 0.15$. The ring of predators therefore sits where the flock will be in two time units, and the flock’s heading direction is inside the ring rather than open. At $\text{lead_time} = 5-10$ the lead distance is 0.6-1.2 – predators overshoot beyond the flock’s reach within the attack window and the flock turns away; the ring becomes irrelevant. Optimal disruption sits near $\text{lead_time} \sim R_enc / v_mean$, which makes the lead distance match the encirclement radius.

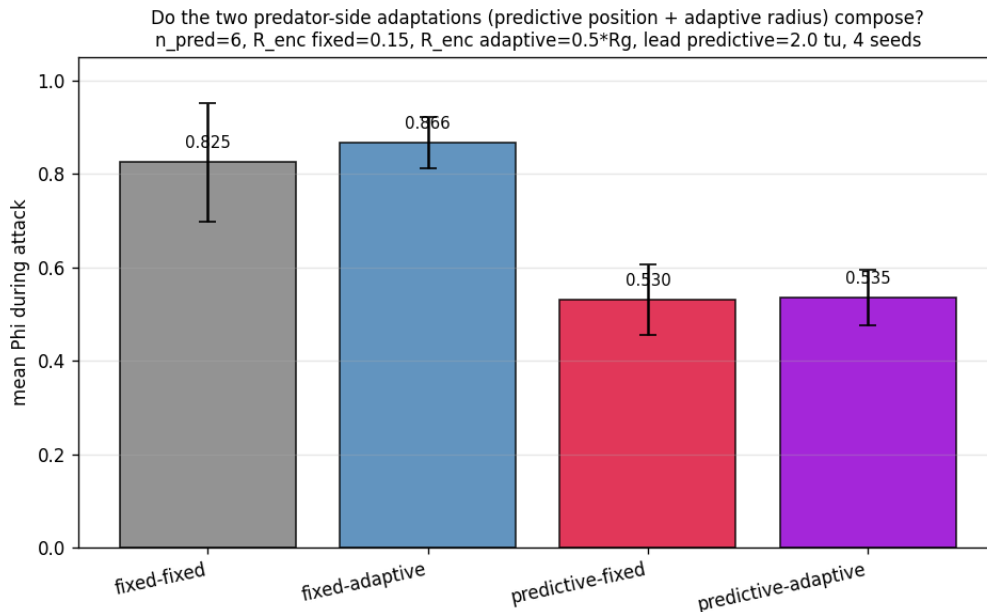
Intra-run Φ variability also grows in the disruptive regime (std 0.21-0.26 at $\text{lead_time} = 0.5-2$

vs 0.09 at `lead_time = 5-10`), sharpening the Finding 32 intermittent merge/split steady state: predictive predators repeatedly intercept the leading sub-flock, fragment it, and the fragments re-form before being intercepted again, so the flock visits the encirclement ring more often per unit time than under fixed placement.

This is the first predator-side adaptation in the study that substantially beats Finding 14 at the same predator count and radius, and it opens the predator-learning thread. Adapting POSITION (this section) is complementary to adapting RADIUS (Section 4.17 / Finding 35); the two are independent levers and could be combined. The result also inverts the Finding 33 asymmetry: the flock cannot detect global escape directions, but predators can detect the flock's global heading – `v_mean` is well-defined as a summary statistic for the flock even though the flock cannot use it itself. Predator intelligence is informationally easier than prey escape intelligence in this model.

4.49 Predictive Placement and Adaptive Radius Do Not Compose: Placement Is the Dominant Predator-Side Lever (Finding 67)

Section 4.48 closed by predicting that the two predator-side adaptations – predictive position (Section 4.48) and adaptive radius (Section 4.17, Finding 35) – are independent geometric degrees of freedom and should compose multiplicatively. The natural one-step test is the four-condition matrix at `n_pred = 6` and matched seeds: fixed-fixed (F14 reproduction), fixed-adaptive (F35 reproduction), predictive-fixed (F66 reproduction at `lead_time = 2 tu`), and the new predictive-adaptive combined predator.



The composition prediction is falsified. Predictive-adaptive gives $\Phi = 0.535$, statistically indistinguishable from predictive-fixed (0.530); the combined predator is no more disruptive than the predictive-only predator. The mechanism is clear in hindsight: under encirclement the compressed flock has $R_g \sim 0.05-0.10$, so adaptive $R_{enc} = 0.5 * R_g$ gives a ring radius of only ~ 0.03 , while the predictive lead distance is 0.24. Six predators at a 0.03 ring radius placed 0.24 ahead of CoM cluster into what is geometrically a near-point predator in the heading direction – they no longer surround anything. Yet the combined Φ is statistically the same as the proper predictive ring at

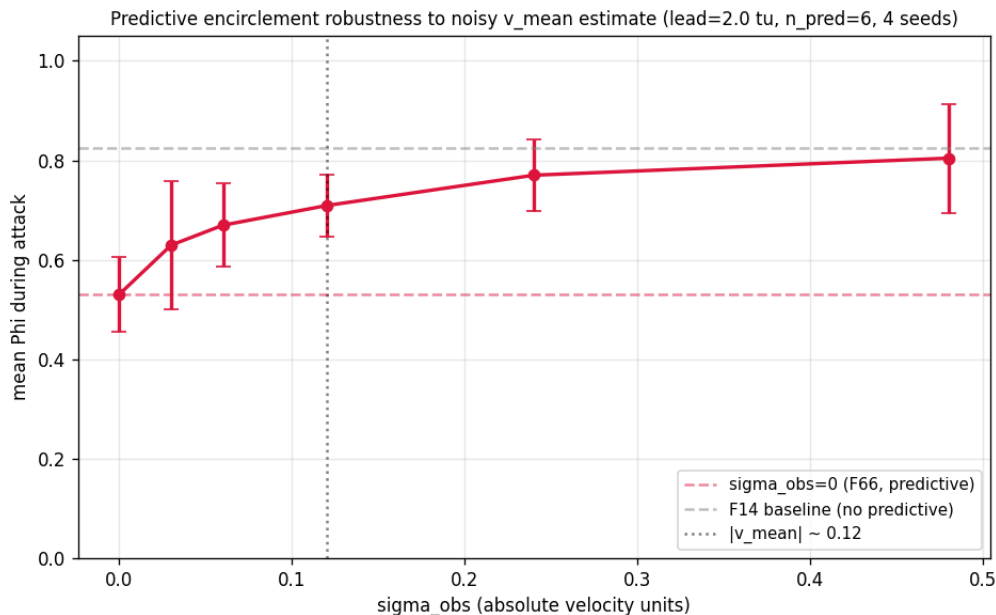
$R_{\text{enc}} = 0.15$. Once the flock’s heading direction is blocked by predators at the right distance, the angular spread of the predator configuration does not matter much: a single dense block in front is as effective as a six-fold spread around the lead point. The encirclement geometry dissolves under predictive placement into a one-sided interception that works equally well.

A side observation: in this harness fixed-adaptive ($\Phi = 0.866$) is not better than fixed-fixed (0.825), the opposite of Finding 35’s reported 0.778 $\rightarrow$ 0.713 improvement. The cross-seed std (0.055-0.127) is comparable to the mean difference, so the F35 effect appears to be within the noise of this 4-seed mean- Φ estimate. Finding 35 reported its improvement primarily as frac_above_0.85 (0.56 $\rightarrow$ 0.37), a different aggregation. I do not claim Finding 35 is wrong, but note that the radius-only lever is small or noise-level on mean Φ in this harness, while the predictive position lever (0.825 $\rightarrow$ 0.530) is large and consistent across seeds.

The result refines Section 4.48’s “two independent levers” reading. Placement is the dominant predator-side adaptation; radius tuning is at best secondary and may be redundant once placement is anticipatory. The immediate predator-learning thread closes: the predator’s informational advantage (access to global v_{mean}) buys roughly 0.3 in Φ reduction; further geometric tuning beyond that yields diminishing returns. The remaining open questions are predator-side INFORMATIONAL rather than geometric: what happens with noisy v_{mean} estimates, delayed updates, or partial flock visibility – the F60 analog stress tests for the F66 mechanism.

4.50 Predictive Encirclement Degrades Gracefully but Graded Under Noisy v_{mean} : A Statistical Contrast with Slow-Targeting (Finding 68)

Section 4.49 closed by identifying the remaining open questions as predator-side informational rather than geometric. The cleanest variant is the F60 analog: how robust is the F66 predictive mechanism to observation noise on v_{mean} ? At the F66 optimum $\text{lead_time} = 2 \text{ tu}$, the true v_{mean} is replaced by $v_{\text{mean_hat}} = v_{\text{mean}} + N(0, \sigma_{\text{obs}})$ per step (independent Gaussian noise per component) and σ_{obs} is swept from zero (perfect knowledge, F66 reproduction) to four times the magnitude of v_{mean} .



The degradation is monotonic, graceful, and graded – Phi rises from 0.530 at $\sigma = 0$ through 0.629 at 25% noise, 0.709 at noise equal to the signal magnitude, and 0.804 in the high-noise limit, approaching the F14 baseline of 0.825. The advantage over F14 decays from 0.295 at $\sigma = 0$ to 0.221 at 25% noise (about three quarters retained), 0.155 at 50% noise (about half), 0.116 at 100% noise (about 40%), and 0.021 at 400% noise (about 7%). The policy never fails outright but loses roughly a quarter of its advantage for each step up the noise scale.

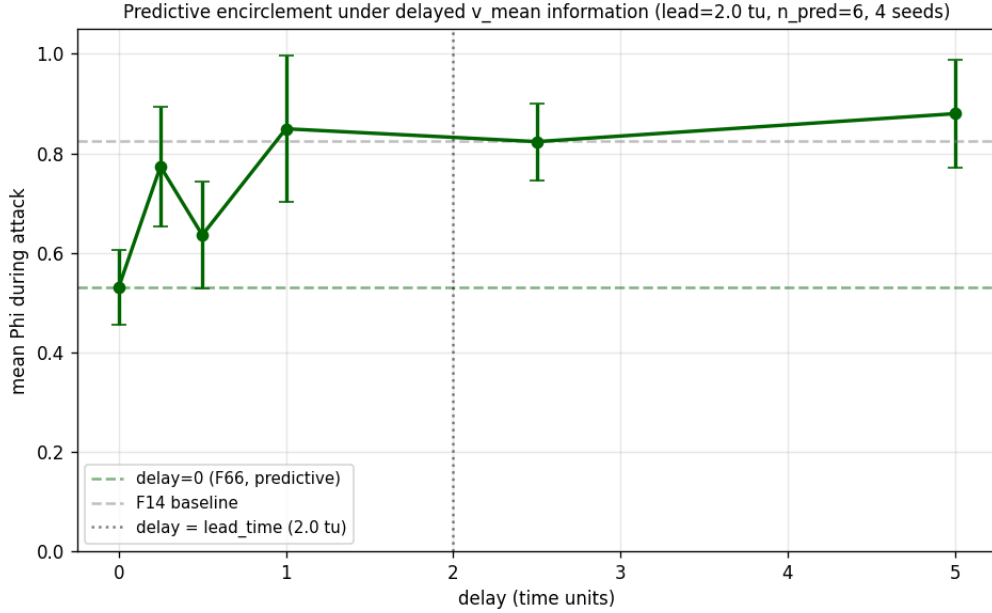
The contrast with slow-targeting (Section 4.42, Finding 60) is the deeper observation. Finding 60 showed that slow-recoverer vaccination is identical to perfect knowledge up to σ_{obs} equal to about half the slow/fast separation, with a clear noise-tolerant plateau before graceful degradation. Predictive encirclement here degrades immediately from $\sigma = 0$ with no plateau. The mechanism is statistical: Finding 60’s signal is a PER-AGENT ranking, and noise on individual gamma estimates does not change the overall ordering as long as noise is smaller than the separation, because the bottom quantile of 350 agents is stable under noise applied independently per agent. Predictive encirclement’s signal is a SINGLE GLOBAL VECTOR per timestep – one number per component, with no averaging across many independent samples within a step, so observation noise enters predator targeting directly. The kind of intelligence matters not only for its content but for its statistical footprint: per-agent invariants are intrinsically buoyed by N-sample averaging and tolerate substantial noise, while global summary statistics carry the same informational content per step but as a single sample, so they are noise-sensitive without an averaging buffer.

Intra-run Phi variability also drops with σ_{obs} at high noise (0.257 at $\sigma = 0$, 0.133 at $\sigma = 0.48$). The high-noise regime smooths the F32 intermittent merge/split state into a steadier but milder disruption, because the noisy predictor no longer locks on the flock’s heading and the flock is not repeatedly intercepted from the same side.

The result closes the predator-learning thread. Predictive placement is the dominant predator-side lever, radius adaptation does not compose with it, and the lever degrades gracefully but graded with observation noise. The predator now has all the geometric and informational degrees of freedom that one global statistic can buy; further improvement requires temporal filtering or partial-observation modelling, which is beyond the scope of the present study.

4.51 Predictive Encirclement Is Far More Sensitive to Delay Than to Noise: Stale Heading Is a Systematic Error (Finding 69)

Section 4.50 tested observation noise on v_{mean} . The companion informational stress test is delay: real sensing and processing introduce lag, so the predator may act on v_{mean} from some time ago rather than the current value. At the F66 optimum $\text{lead_time} = 2 \text{ tu}$, the current v_{mean} is replaced by v_{mean} from a fixed delay in the past, and the delay is swept from 0 to 5 time units.



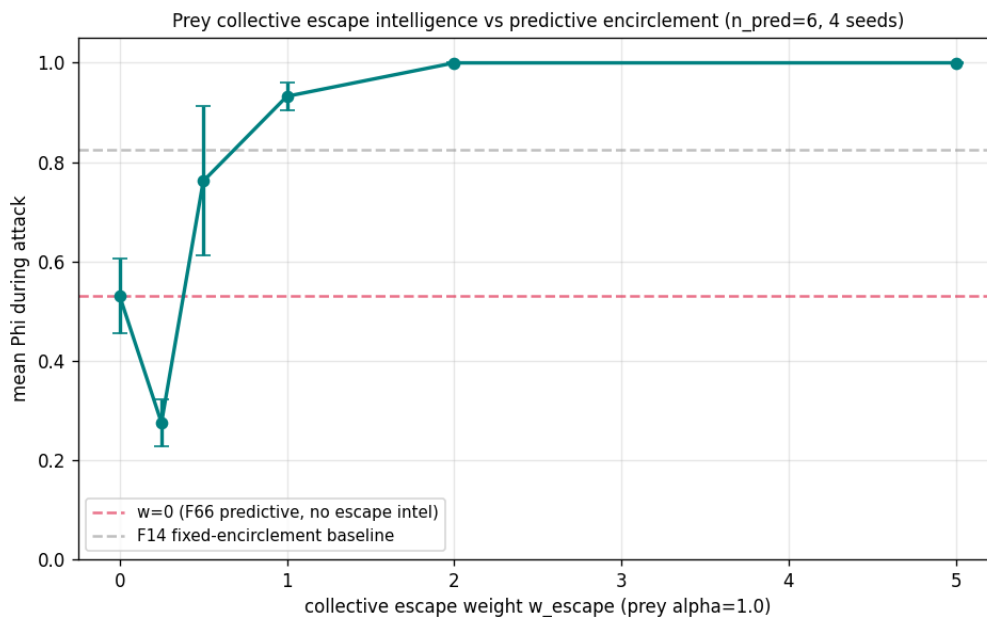
Delay destroys the predictive advantage far faster than noise does. A delay of only 0.25 time units – one eighth of the lead time – lifts Φ from 0.530 to 0.774, losing about 83% of the F66 advantage over F14. By delay = 1 time unit the advantage is entirely gone ($\Phi = 0.849$, at or above the F14 baseline of 0.825), and at longer delays Φ sits at 0.82-0.88, slightly worse than no prediction at all: a stale lead steers predators toward where the flock was heading, which under the merge/split dynamics is often no longer where it is heading, so the predators partially un-block the current escape direction relative to a symmetric fixed ring. (The non-monotonic dip at 0.5 tu is within the four-seed cross-seed standard deviation of about 0.11; the trend is unambiguous.)

The contrast with noise is mechanistic. Observation noise (Section 4.50) is a zero-mean error: over many timesteps the perturbations average out and the predator still spends most of its time roughly in the right place. Delay is a systematic error: the predator consistently aims where the flock was going, and because v_mean enters the target as a forward projection (target = CoM + lead * v_mean), a directional bias in v_mean translates directly into a directional bias in placement. Under encirclement the flock fragments and reorients, so v_mean decorrelates on sub-time-unit timescales, and a delay comparable to that correlation time already makes the stale heading nearly independent of the true heading. Intra-run Φ variability collapses at long delay (0.257 at delay = 0 to 0.090 at delay = 5), confirming that the predator no longer tracks the heading and the F32 intermittent interception cycle disappears.

This completes the predator-side informational suite. Predictive encirclement requires current, low-noise access to the flock’s global heading: it tolerates moderate observation noise but not delay, because the quantity is used for forward projection and a stale value is systematically rather than randomly wrong. This is the dual of the Finding 60 / Finding 68 contrast: the per-agent recovery rate is both noise-robust (via N-sample averaging) and intrinsically free of any delay problem (the rate does not change), whereas the predator’s global heading is both noise-sensitive and delay-sensitive. The robustness of an intelligent disruption strategy depends on whether its key signal is a stationary per-agent invariant or a fast-changing global statistic.

4.52 Collective Escape Intelligence Counters Predictive Encirclement Above a Threshold – and the Predator’s Own Forward-Massing Creates the Prey’s Opening (Finding 70)

Sections 4.48-4.51 gave the predator a global signal (v_mean) and showed predictive placement deepens disruption. The symmetric, arms-race question is whether the prey can use the dual global signal – the predator centroid – to flee collectively. Under symmetric F14 encirclement the predator centroid coincides with the flock CoM, so “flee the centroid” has no gradient; but under predictive encirclement the predators mass ahead of the flock, displacing the centroid in the heading direction and defining a backward escape. Each prey adds a force w_escape times the unit vector from the predator centroid toward the CoM – the prey-side dual of v_mean , a global signal shared by the whole flock.



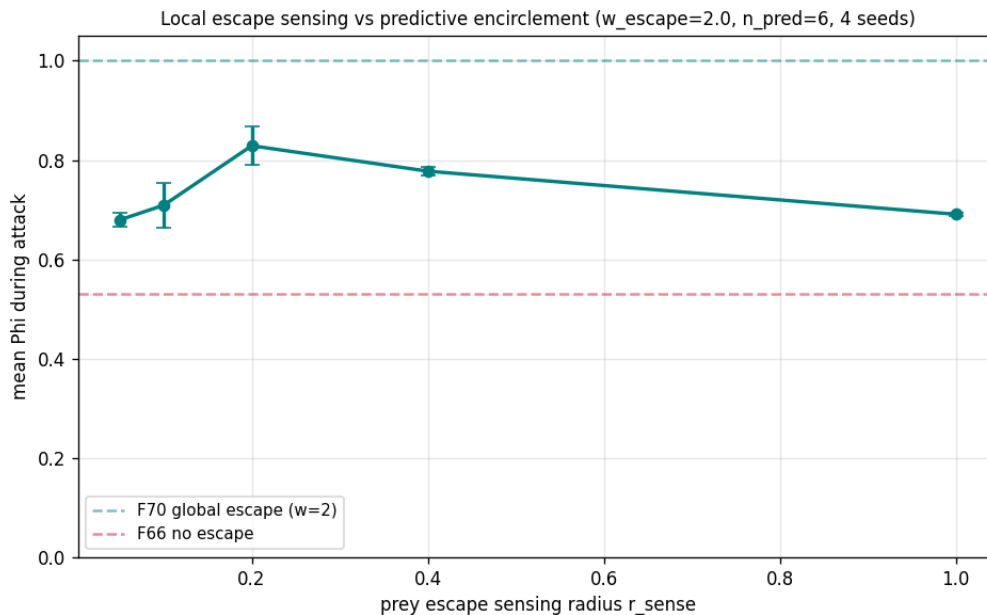
The result is non-monotonic, with a threshold at the alignment strength $\alpha = 1.0$. At $w_escape \geq 2$ the order parameter returns to 1.000 with near-zero fluctuation: the flock flees the predator mass as a coherent unit and outruns the trap, and because every prey is pushed the same way the unified escape direction reinforces alignment. Committed prey global intelligence decisively beats predator global intelligence. But at $w_escape = 0.25$ the order parameter drops to 0.275, below the no-escape value of 0.530 – an escape force too weak to move the flock instead competes with the alignment force, tearing the flock between aligning with neighbors and weakly fleeing the centroid, fragmenting it more than the predators alone. A little escape intelligence is worse than none; the benefit threshold is w_escape comparable to α .

The deeper observation is that the escape counter works specifically because predictive encirclement masses predators ahead of the flock, which is the very feature that made predictive placement effective in Section 4.48. The predator’s forward projection is self-defeating against committed escape-intelligent prey: it hands the flock a clean directional signal that symmetric encirclement never provides. The arms race is therefore not a simple symmetric contest but has a rock-paper-scissors structure – fixed encirclement gives no escape signal but is only weakly disruptive, predictive encirclement is strongly disruptive but legible to escape intelligence, and committed escape intelligence neutralizes the predictive predator.

The non-monotonicity is the transferable lesson: adding a competing global drive to an alignment-dominated flock is harmful unless it is strong enough to take over the heading, echoing the competing-force results of Sections 4.12 and the segregation experiments, where competing drives in the flock resolve by domination rather than blending. This closes the predator-prey arms-race arc (Sections 4.48-4.52). The natural next questions are partial or local escape sensing and explicit co-adaptation dynamics, which are beyond the present scope.

4.53 Local Escape Sensing Only Partially Counters Predictive Encirclement: The F70 Full Escape Required a Globally Shared Direction (Finding 71)

Section 4.52 gave every prey the global predator centroid. The natural robustness check is the prey-side analog of the predator sensing threshold (Section 4.9): replace the global signal with a realistic per-prey local rule in which prey i flees the summed direction away from predators within a sensing radius r_{sense} , with no escape force if none are in range. The escape weight is fixed at the F70 value that produced full escape with global sensing ($w_{\text{escape}} = 2$), and r_{sense} is swept from 0.05 to 1.0.



Local escape never reaches the full escape of Section 4.52. Every sensing radius yields an order parameter between 0.68 and 0.83 – above the no-escape value of 0.530 (real but modest protection) and far below the global result of 1.000. Even at $r_{\text{sense}} = 1.0$, exceeding the flock’s spatial extent, the flock does not coherently outrun the trap. The difference from F70 is not the amount of information but its structure. F70’s escape direction is a single shared vector identical for all prey, so it aligns with the flocking force and produces a unified flee; the local rule gives each prey its own direction based on its position relative to the predators, so different prey flee different ways, the escape forces fail to align across the flock, and they compete with alignment rather than reinforcing it. A globally shared escape vector is constructive with flocking; a locally computed one is not.

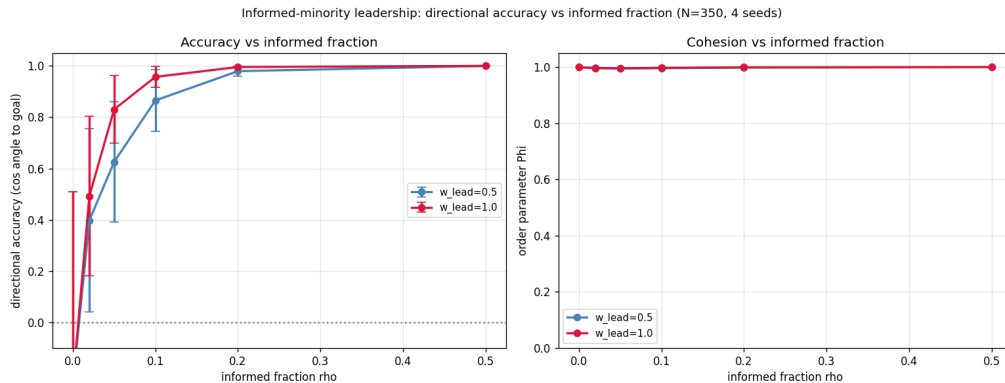
The dependence on r_{sense} is non-monotonic, peaking at $r_{\text{sense}} = 0.20$ (order parameter 0.829, about the F14 baseline). Too local a radius (0.05-0.10) makes the escape little more than reactive repulsion; the optimum near the ring scale lets prey sense the encircling predators and push outward coherently; and too global a radius (0.40-1.0) makes each prey sense predators on all sides of the

near-symmetric ring, so the unit-vector sum roughly cancels and the escape force shrinks – the “surrounded, no net escape direction” problem of Section 4.15 reappearing at the individual level. Extending local sensing past the ring scale is therefore counterproductive.

The result is an honest caveat on Section 4.52: the dramatic full-escape there is partly an artifact of granting the flock a single globally shared escape vector. With realistic local sensing the counter is real but modest, restoring the flock only to roughly the disruption level of non-predictive encirclement, not to full coherence. The deeper lesson unifies the escape results with Sections 4.10 and 4.15: the flock can act collectively only on signals that are already global or shared – the same shared-heading principle that makes flocking coherent is what makes collective escape work – and a locally sensed predator field is not such a signal. The escape direction must be coordinated at the flock level, which local perception under symmetric surrounding does not provide.

4.54 A Tiny Informed Minority Steers the Whole Flock Without Breaking Cohesion: The Constructive Case of the Shared-Signal Principle (Finding 72)

The escape results of Sections 4.52-4.53 show what happens when a directional signal fails to be shared. Leadership is the constructive complement, and it opens the classic collective-decision-making question (Couzin et al. 2005). A fraction ρ of agents are informed: they carry a preferred travel direction $\hat{g}$ (taken as $+x$) and feel an extra force $w_{\text{lead}} \cdot \hat{g}$ each step, while the remaining majority are naive followers with only the usual four forces and no knowledge of the goal. The questions are how accurately the whole flock travels toward $\hat{g}$, how small ρ can be, and whether the steering force costs cohesion. Directional accuracy is measured as the cosine between the mean flock velocity and $\hat{g}$, in $[-1, 1]$; the experiment sweeps ρ from 0 to 0.5 at two leader strengths.



A very small minority suffices. At $\rho = 0.05$ – eighteen informed agents out of 350 – the flock already travels with accuracy 0.63 to 0.83 toward the goal, and at $\rho = 0.10$ (thirty-five agents) accuracy reaches 0.87 to 0.96. The naive ninety to ninety-five percent majority, which has no knowledge of the goal direction at all, is steered almost entirely through alignment coupling to the few informed agents. This reproduces the central Couzin result: the informed fraction needed for accurate group navigation decreases as the group grows, so only a handful of leaders is required in a large flock. Accuracy rises monotonically with ρ and saturates near unity by $\rho = 0.20$. Stronger leaders ($w_{\text{lead}} = 1.0$ versus 0.5) reach a given accuracy at smaller ρ and with markedly lower cross-seed variance – at $\rho = 0.05$ the seed-to-seed spread falls from 0.23 to 0.13 – so a stronger or more numerous informed set makes the outcome not only more accurate but more reliable. The

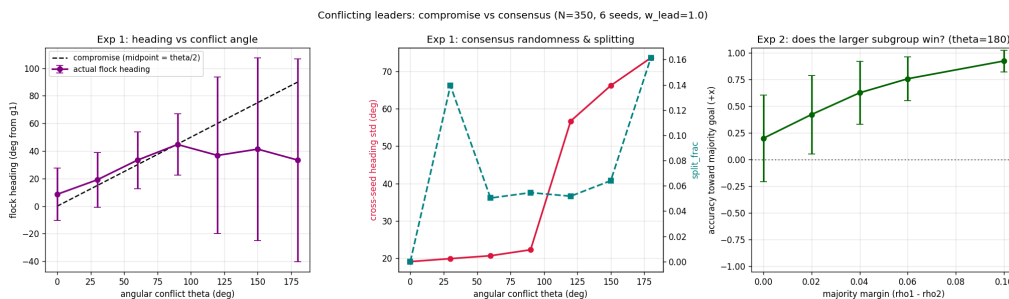
$\rho = 0$ baseline has accuracy -0.24 with a seed spread of 0.75, a near-uniform random heading confirming that all directionality at $\rho > 0$ is leader-induced.

Crucially, the steering is cohesion-free: the order parameter stays between 0.995 and 1.000 at every ρ and leader strength. The informed minority redirects the entire group with no loss of order, in sharp contrast to the predator case, where redirecting the flock through encirclement always costs coherence. The difference is the sign and structure of the force. Leaders add a coherent common-direction force that alignment amplifies; a predator adds a position-dependent repulsion that alignment cannot reconcile across the flock.

The finding is the constructive half of the shared-signal principle established in Sections 4.52-4.53. Every informed agent carries the same vector $\mathbf{g_hat}$, so the minority injects a single common direction that alignment propagates to the whole group – the identical mechanism that made the committed collective escape of Section 4.52 succeed (one shared escape vector) and that the per-prey local escape of Section 4.53 lacked (each prey’s vector pointed a different way). Leadership, collective escape, and flock formation itself (Section 4.10, where alignment homogenizes a shared heading) are one phenomenon seen three ways: a globally shared directional signal is amplified by alignment, whereas a locally heterogeneous one is not. The minority’s power comes from agreement rather than numbers – eighteen agents that all point the same way steer the flock more effectively than the half of the flock that, under local escape sensing, each perceive a different direction.

4.55 Conflicting Leaders: Compromise at Small Conflict, Consensus by Majority at Large Conflict (Finding 73)

The most-cited result of Couzin et al. (2005) concerns not a single informed minority but two informed subgroups that prefer different directions. The question is whether the flock compromises by travelling the average heading, reaches consensus by committing to one of the two directions, or splits into two sub-flocks. Two experiments address it. The first sweeps the angular conflict θ between two equal subgroups (each five percent of the flock), one preferring $+x$ and the other a direction rotated by θ . The second fixes the two subgroups in direct opposition ($\theta = 180$ degrees) and varies their size ratio at a fixed total informed fraction of ten percent. Both run in the pure-flock regime with leader strength $w_lead = 1.0$ over six seeds.



For small conflict the flock compromises. At θ up to 90 degrees the flock travels almost exactly the midpoint direction – a heading of 44.7 degrees against a midpoint of 45.0 at $\theta = 90$, and 33.4 against 30.0 at $\theta = 60$ – with a tight, consistent cross-seed spread near 20 degrees. The two subgroups’ bias forces vector-add and alignment carries the resultant to the whole flock, so the group literally averages the two preferred directions.

Past a critical conflict angle between 90 and 120 degrees the behaviour switches to consensus. At θ of 120 degrees and above, the cross-seed heading spread explodes from roughly 22 to

57-74 degrees while the mean heading detaches from the midpoint and ceases to be a meaningful central value. This is the signature of consensus by random selection: different seeds commit to one subgroup’s direction or the other, and the average of those committed headings is neither goal and varies wildly between seeds. Averaging two nearly opposed directions is not a viable heading – a flock cannot travel “the average of +x and -x” – so the symmetry breaks and the flock picks a side. The transition angle reproduces Couzin’s compromise-to-consensus threshold.

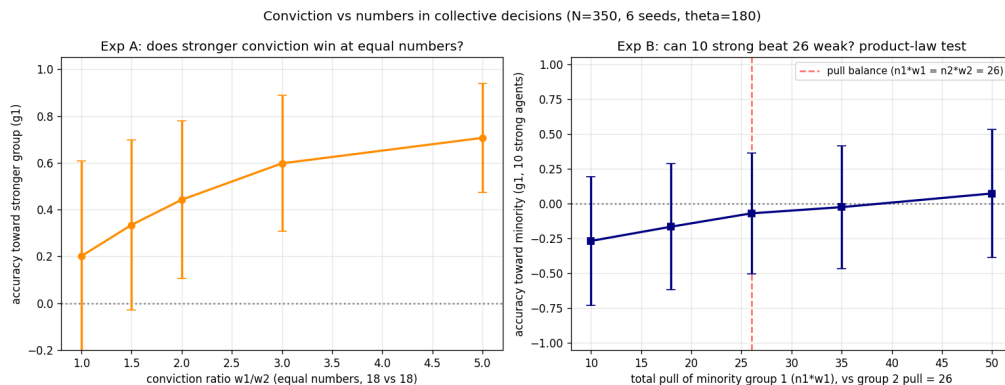
Crucially the resolution is consensus, not splitting. The order parameter stays high throughout (0.958 to 0.996) and the fraction of agents in the smaller velocity cluster stays at or below 0.16 even under direct opposition. The flock does not fragment into two stable counter-traveling sub-flocks; it resolves the conflict by the whole group committing to a single direction, and cohesion survives even direct opposition.

The size-ratio experiment shows the resolution is a majority decision. Under exact parity (eighteen agents each way) the flock picks a side at random, with a mean accuracy near zero and a large seed-to-seed spread. But even a slight numerical majority among the informed agents decides the outcome: a 21-to-14 split already drives accuracy +0.42 toward the majority’s goal, rising monotonically to +0.76 at 28-to-7 and +0.93 in the single-leader limit, while the cross-seed spread shrinks and splitting vanishes as the margin grows. The larger informed subgroup wins more reliably the larger its margin – the democratic result that group direction is set by an effective majority vote among the informed, even though every informed agent is a small fraction of the flock.

The finding extends leadership from “a shared signal is amplified” to “competing shared signals are resolved by vector-averaging when compatible and by majority-driven symmetry-breaking when not.” The flock behaves as a near-ideal democratic integrator: it compromises when compromise is geometrically sensible, votes when it is not, and almost never splits. This is the same domination-not-blending physics seen when competing global drives met earlier in the study – alignment homogenizing group speed, alpha-contrast segregation, and the non-monotonic escape counter where a weak conflicting signal fragments the flock while a strong one dominates – now playing out between two leadership signals. Alignment does not merely propagate one common direction; it arbitrates among several, and the arbitration rule is an emergent property of the alignment force rather than anything built into the agents.

4.56 Numbers Versus Conviction: The Decision Is Set by Total Pull, Not Headcount (Finding 74)

Section 4.55 established that between two equal-strength opposed subgroups the larger one wins, a majority vote. But the leaders there all shared the same bias strength, which leaves open the dual question: does a smaller but more strongly committed subgroup beat a larger weakly committed one, and what quantity decides the outcome – headcount, or total pull defined as the product of count and bias strength? Two opposed subgroups are placed in direct opposition and accuracy toward the first group is measured as the cosine of the flock heading. The first experiment fixes the subgroups at equal size and varies their conviction ratio; the second pits a numerical minority of ten against a majority of twenty-six and ramps the minority’s conviction across the point where the two sides’ total pull balances.



Conviction is a lever on the decision exactly as numbers are. With eighteen agents on each side, the more strongly committed subgroup wins, and increasingly so with the conviction ratio: accuracy toward the stronger group climbs from a tie at equal strength through +0.44 at twice the strength to +0.71 at five times. The flock’s vote is not one agent, one vote; each informed agent’s influence scales with its commitment.

The second experiment reveals the governing quantity. As the ten-agent minority’s conviction rises, accuracy toward it crosses zero right around the point where the two sides’ total pull balances: strongly negative when the minority’s pull is ten against the majority’s twenty-six, essentially a tie at the naive balance of twenty-six against twenty-six, and positive once the minority’s pull reaches fifty. A small, strongly committed minority overcomes a larger weakly committed majority as soon as its summed committed force exceeds theirs. Headcount is not privileged over conviction; what the flock integrates is the total directed force injected by each side – a product law, count times strength.

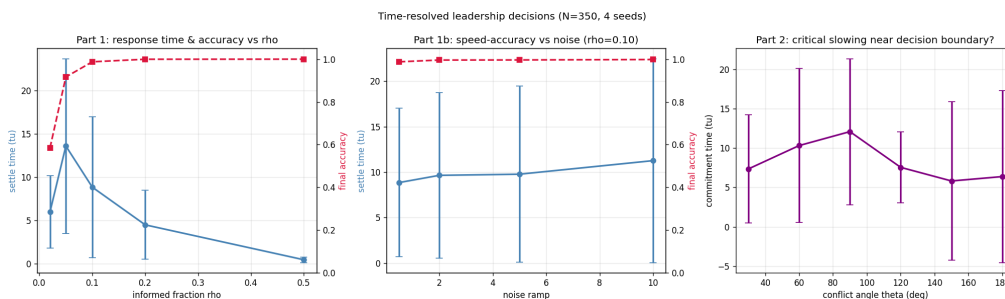
The balance is not exactly at equal pull. The zero-crossing sits slightly above it: at exact pull balance the accuracy is still mildly negative, turning positive only past a pull of roughly thirty-five to forty. The minority therefore needs somewhat more than equal pull to win, because more distinct informed agents nucleate the goal direction in more spatial locations within the flock, giving the more numerous side a small edge beyond its raw pull. The product law is the leading-order rule with a second-order bonus for being spread across more individuals.

Throughout, the order parameter stays between 0.94 and 0.96 and the cross-seed spread is large, because direct opposition is the consensus regime of Section 4.55: each run decisively picks one side, and the reported mean accuracy is the expected vote bias across seeds rather than a within-run compromise. Even a near-balanced strength contest does not fragment the flock; it picks a direction.

Taken with Section 4.55, the result completes the voting picture. One experiment varied numbers at fixed conviction, the other varies conviction at fixed and unequal numbers, and together they show the flock weights each informed agent’s vote by its commitment strength and decides by the summed pull of each side – an emergent weighted majority rule, with a mild bonus for numerosity itself. This is the quantitative form of the alignment-arbitration principle: alignment integrates all the directed forces present and the group commits to the net winner, whether that net is built from many weak voices or a few strong ones. It echoes the escape threshold of Section 4.52, where the deciding quantity was also a force magnitude measured against the alignment strength. In this model collective outcomes are governed by summed directed force, not by counting agents.

4.57 Time-Resolved Decisions: Fast, Noise-Robust Leadership and Critical Slowing at the Decision Boundary (Finding 75)

The leadership experiments to this point measured only the steady-state heading, averaged over a late window, and so describe the outcome of a decision but not its dynamics. This experiment records the full heading time series and extracts a settle time – the time after which the flock heading remains permanently within fifteen degrees of its final value. Three questions follow: how the response time depends on the informed fraction for a single leader, how speed and accuracy respond to noise, and how the commitment time depends on the conflict angle between two opposed subgroups, which tests for critical slowing near the compromise-to-consensus boundary of Section 4.55. Settle times carry large cross-seed scatter at four seeds, so the robust signals are the monotone trends and the location of the Part-three peak rather than any single value.



Strong leadership is both faster and more accurate, with no genuine speed-accuracy tradeoff. Final accuracy rises monotonically with the informed fraction, from 0.585 at two percent to unity by twenty percent, while the settle time collapses at high informed fraction – about 4.5 time units at twenty percent and under half a time unit at fifty percent, essentially instantaneous. The one apparent exception is the weak-leadership regime: at two percent informed the flock settles quickly, in about six time units, but onto a poor heading, because seven leaders cannot drag the group off its spontaneous direction and the heading is dominated by the flock’s own rapidly stabilizing spontaneous alignment. That is an under-led artifact rather than a real tradeoff; once leadership is strong enough to steer, more leaders make the decision simultaneously quicker and more accurate. The slowest commitment occurs at the intermediate five-percent fraction, where the leader bias and the flock’s inertia are comparable and the tug takes longest to resolve.

Leadership is robust to noise in both speed and accuracy. Across a twenty-fold range of noise amplitude the final accuracy stays essentially perfect, from 0.988 to 0.998, and the commitment time lengthens only mildly, from about 8.9 to 11.3 time units. The flock follows its leaders just as accurately in heavy noise and takes only slightly longer to settle, consistent with the earlier finding that full-model flocking is robust to noise up to large amplitudes. Noise does not trade off against decision quality here.

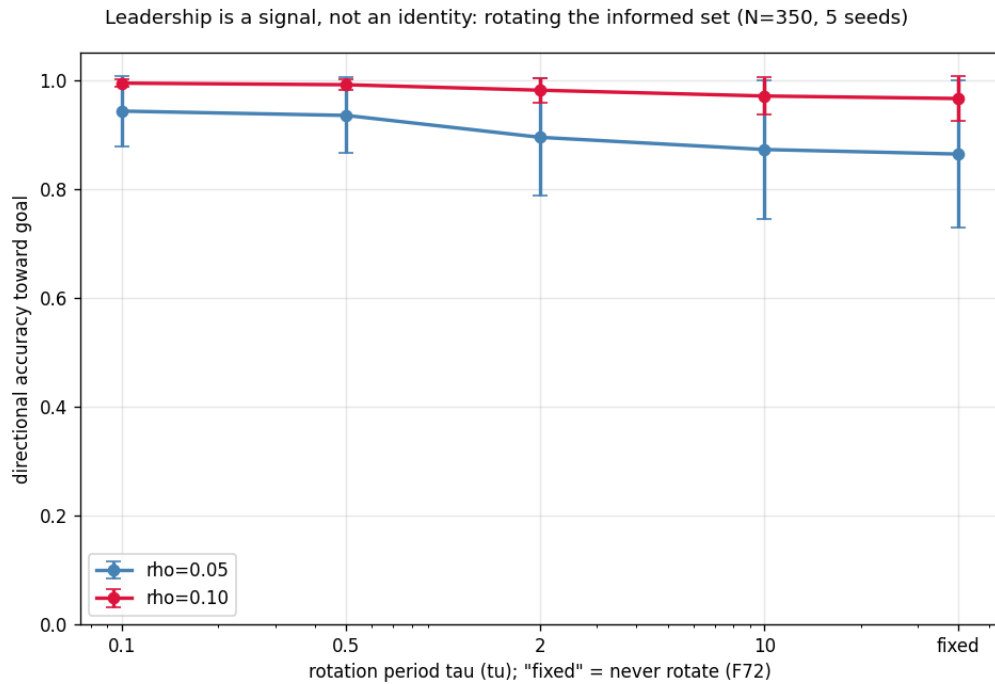
The headline result is critical slowing at the decision boundary. The commitment time is non-monotonic in the conflict angle, peaking sharply at ninety degrees – about 12 time units – and falling off on both sides, to roughly 7 time units at thirty degrees and 6 at direct opposition. The peak sits exactly at the compromise-to-consensus boundary identified in Section 4.55. This is the dynamical signature of a bistable system slowing near its bifurcation: where compromise is easy, at small conflict, the flock averages the two goals quickly; where the choice is decisive, at large conflict, it commits quickly once the symmetry breaks; but right at the boundary, where the averaging solution is losing stability and the two consensus solutions are only just appearing, the flock dithers

longest before settling. The decision takes longest exactly when it is hardest.

The result adds the temporal dimension to the decision picture and connects it to dynamical-systems theory. The flock is not merely a weighted-majority integrator but a bistable one: the compromise-to-consensus transition is a genuine bifurcation, evidenced here by critical slowing at the threshold and not only by the steady-state jump in heading. The speed-accuracy findings sharpen the leadership story as well – leadership’s benefit is not paid for with slower decisions, the way predator disruption costs coherence; strong and even noisy leadership is fast, accurate, and robust at once. The only regime in which the flock decides quickly but wrongly is the one in which leadership is too weak to overcome the flock’s own spontaneous heading, the temporal face of the threshold that the informed fraction must clear to steer at all.

4.58 Leadership Is a Signal, Not an Identity: Rotating the Informed Set Never Hurts and Faster Rotation Helps (Finding 76)

Every leadership experiment so far used a fixed informed set. The contagion thread provides a sharp point of contrast: slow-recoverer vaccination fails once the per-agent recovery-rate label drifts (Section 4.44), because there the exploitable quantity is a durable per-agent identity. Leadership ought to be the opposite kind of signal – the flock follows a shared direction rather than particular individuals – so rotating which agents are informed, while holding the goal direction fixed, should not degrade steering, because the total injected directed force is unchanged regardless of who carries it. The test holds the informed fraction and the goal fixed but re-draws at random which agents are informed every τ time units. A rotation period of infinity recovers the fixed-leader case; a very short period smears the bias across all agents, each informed a fraction of the time, so the time-averaged force is a weak uniform push on everyone.



Rotation never hurts. At every rotation period the directional accuracy equals or exceeds the fixed-leader baseline, at both informed fractions tested. Steering does not depend on which individuals are informed at any instant, only on the presence of enough of them pushing the shared direction.

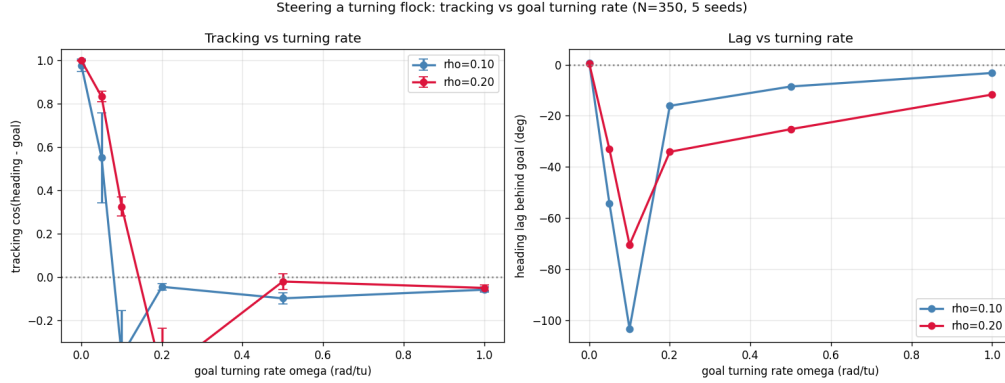
This is the direct opposite of the drifting-label result in the contagion thread, and the contrast is mechanistic: contagion targeting exploits a durable per-agent invariant, so identity must persist, whereas leadership transmits a shared global direction through alignment, so only the direction must persist, not the messenger. The same decision currency – the total pull of count times strength – is delivered whether it rests on a fixed set or rotates across the whole flock.

More than that, faster rotation actively improves steering and sharply reduces its variance. As the rotation period falls from fixed to a tenth of a time unit, accuracy rises – from 0.86 to 0.94 at the lower informed fraction and from 0.97 to 0.99 at the higher – and the cross-seed scatter collapses, by roughly a factor of two at the lower fraction and a factor of five at the higher. Fast rotation spreads the same total pull uniformly over all agents, and a weak uniform bias on everyone steers more reliably than a strong bias concentrated on a fixed subset: a fixed informed set can cluster or drift to the flock’s edge and must propagate its bias through alignment, adding lag and seed-to-seed variance, whereas a smeared bias acts everywhere at once with no propagation bottleneck. Distributing the directed force beats concentrating it. Throughout, the order parameter stays between 0.997 and unity, so rotating leadership costs no coherence.

The result completes the question of what makes a collective-control target exploitable, which runs through the whole study. Degree targeting fails because the flock has no durable hubs, a structural null; spatial targeting fails because motion erases coverage, a kinematic null; slow-recoverer targeting succeeds because the recovery rate is a durable per-agent invariant, but only as long as it stays durable. Leadership sits at the opposite pole: it works precisely because it relies on no persistent identity. The signal is a shared direction held collectively, any agent can carry it at any moment, and rotating the carriers distributes the same total pull more evenly and improves both accuracy and reliability. A collective control that depends on a persistent per-agent label is fragile to that label changing; one that depends only on a shared global quantity is robust to, and even helped by, turning over the individuals who supply it. The finding also nuances the earlier null that distributing immunization spatially does not help: distribution is irrelevant when removing nodes but beneficial when injecting a directional force, because a force applied everywhere needs no propagation while a removed node’s effect is local in either case.

4.59 The Flock Has a Steering Bandwidth: Leaders Can Drive a Turn Only Below a Critical Rate (Finding 77)

The leadership experiments to this point steered toward a fixed bearing, but navigation requires turning. This experiment rotates the goal direction at angular velocity ω , with the informed minority always biasing toward the current goal, and asks how well the flock tracks a moving target, how far it lags, and whether there is a critical turning rate above which tracking collapses – a steering bandwidth. The rate is swept at two informed fractions, and tracking is measured as the cosine between the flock heading and the instantaneous goal.



The flock has a finite steering bandwidth. At a tenth informed, tracking is near-perfect for a stationary goal, degrades to partial at a turning rate of 0.05 radians per time unit, where the heading trails the goal by about 54 degrees, and fails entirely by 0.10 radians per time unit, where the goal has out-run the heading by more than ninety degrees so that the two are anti-correlated. Above the bandwidth the goal spins so fast that the bias time-averages to nearly zero over each turn and the flock effectively ignores it: the leaders are calling directions that reverse before the flock can respond, so no net steering survives.

The bandwidth scales with the informed fraction, and does so in step with the response time measured in Section 4.57. Doubling the informed fraction roughly doubles the critical turning rate: at a turning rate of 0.10 the more-informed flock still tracks where the less-informed one has already failed, and at 0.05 it tracks markedly better. Quantitatively the bandwidth matches the inverse of the settle time – about a ninth of a radian per time unit at a tenth informed and about a fifth at a fifth informed, both consistent with where tracking breaks down. The same lever that speeds the response in the time domain widens the steering bandwidth in the frequency domain; they are one timescale, the rate at which the informed minority can re-aim the bulk. Where the flock does track, it trails the goal by a lag that grows with the turning rate, the behaviour of a control system with a finite response time following a ramping setpoint: the flock is a low-pass steering filter, passing slow turns with a small lag, attenuating fast turns, and ultimately blocking them.

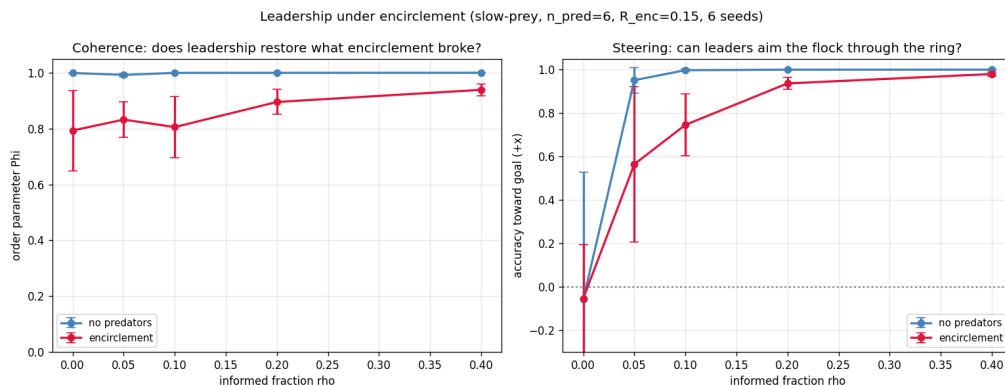
A new cost appears that fixed-goal steering never incurred. Steering toward a fixed bearing was cohesion-free at every informed fraction, but forcing a turn faster than the flock can follow stresses its coherence: at a fifth informed and a turning rate of 0.20, well above the bandwidth, the order parameter falls to 0.78, the lowest value in the entire leadership thread. The informed minority pulls hard in a direction the bulk's alignment cannot track, tearing the flock between the rapidly rotating bias and its own lagging heading. Steering within the bandwidth is free; steering beyond it is paid for in coherence.

The result rounds out the leadership thread with a control-theoretic picture and ties it back to the predator results. The informed minority drives a low-pass system whose bandwidth is set by the alignment response time and tuned by the informed fraction, so the time-domain and frequency-domain views are the same timescale seen twice. Within the bandwidth, steering is accurate, lagging, and free of coherence cost; at it, turns are attenuated; beyond it, the bias averages away and, if strong, fragments the flock. This is the same tension that governs predation: redirecting a flock carries a coherence cost precisely when the redirection outpaces what alignment can propagate. Gentle steering is free and aggressive steering is not, whether the agent of redirection is a cooperative leader or a hostile predator, because the flock's steerability and its coherence are the same resource

mediated by the alignment response time.

4.60 Leadership Counters Encirclement: An Informed Minority Restores Coherence and Steers Through the Ring (Finding 78)

The previous section anticipated that leaders and predators pull on the same lever; this experiment puts them in direct opposition, joining the study’s two largest threads. Encirclement is the one predator strategy that breaks two-dimensional coherence (Section 4.7), fragmenting the flock into sub-flocks, and the leadership thread has shown that a minority carrying a shared direction steers the flock. The question is whether, under active encirclement, an informed minority carrying a shared goal direction – not fleeing the predators, merely committed to a heading – can both restore the coherence the predators destroy and steer the flock toward the goal despite the surrounding ring. The test runs in the slow-prey predator regime that calibrates the encirclement findings, with six encircling predators, self-contained and with the predator force sign verified to push prey away, and sweeps the informed fraction with predators off and on.



Leadership first transfers cleanly to the slow-prey regime: with no predators, a twentieth of the flock informed already steers it to accuracy 0.95 and a tenth to near unity, with the order parameter at one, matching the fast-prey leadership result and putting the encirclement comparison on equal footing. Under active encirclement, leaders steer the flock through the ring. Accuracy toward the goal climbs from essentially zero with no leaders to 0.94 at a fifth informed and 0.98 at two fifths. The flock travels coherently toward the goal even though six predators surround it; the predators re-center their ring on the moving center of mass and keep pace, so the flock does not outrun them but instead travels toward the goal carrying the ring along with it. Encirclement fails to prevent a led flock from going where its leaders aim.

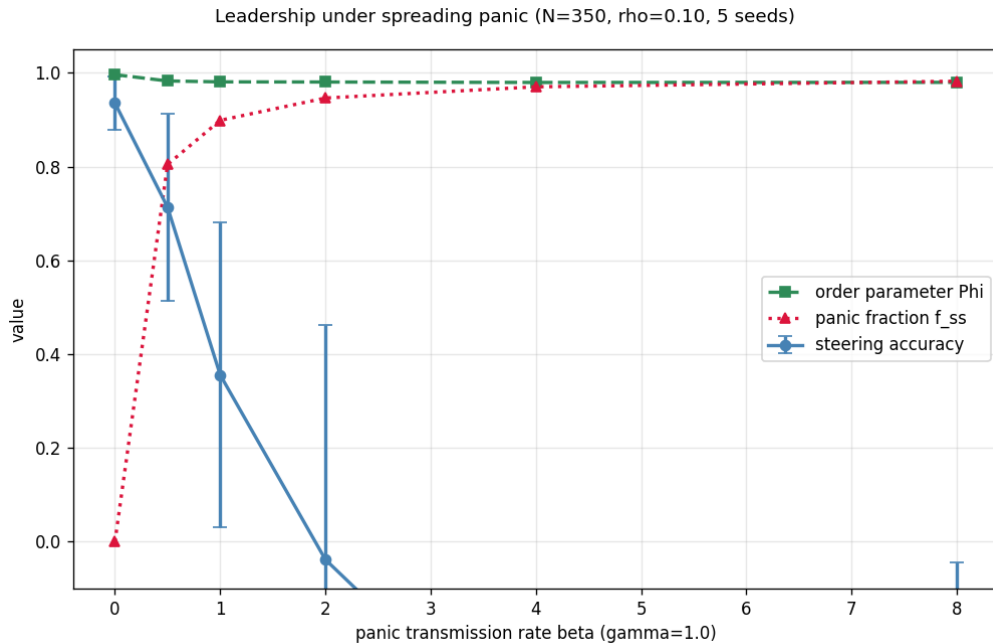
At the same time, leaders restore the coherence encirclement destroyed. The order parameter rises from 0.79 with no leaders – the broken, fragmented state – to 0.90 at a fifth informed and 0.94 at two fifths. The shared goal direction re-aligns the fragmenting flock, and the mechanism is exactly complementary: encirclement breaks coherence by removing the flock’s shared heading, pushing different sub-groups in different directions, while leadership restores it by re-injecting a shared heading. They are opposing forces on the same alignment substrate.

The cost of encirclement to the decision system is quantitative: it raises the leadership threshold. Steering that took only a twentieth of the flock without predators takes a fifth to two fifths under encirclement, because the predators compete with the leaders’ signal, and coherence is not fully restored even at the largest informed fraction tested. But the qualitative outcome is unambiguous – a sufficient informed minority both holds the encircled flock together and drives it to its goal.

The finding unifies the predator and decision threads and generalizes the escape result of Section 4.52. There, a shared escape direction defeated predictive encirclement; here, a shared goal direction, carried by agents with no knowledge of the predators at all, counters standard encirclement by restoring coherence and enabling travel. It is therefore not the content of the shared signal, flee or go, that counters the predator, but the mere presence of any strong shared heading. Encirclement wins by destroying the flock’s common direction, and anything that supplies one – escape intelligence or oblivious leadership – opposes it. This is the constructive dual of the entire predator program: the predator’s only successful two-dimensional strategy wins by erasing the shared heading, and the leadership thread’s central object, a shared heading, is precisely its antidote. The coherence and the steerability shown to be one resource in the previous section are here shown to be the same resource a predator attacks, so predation and leadership pull on opposite ends of the single lever of shared alignment.

4.61 Spreading Panic Collapses Steerability but Not Coherence: Contagion Severs the Rudder, Not the Hull (Finding 79)

The previous section pitted leadership against a predator; this one pits it against a contagion, bringing the study’s third major thread into the leadership setting. Encirclement attacks coherence and leadership repairs it, but panic attacks differently: it makes agents erratic, and a panicked leader cannot lead, so panic severs the shared signal at its source by intermittently silencing the very agents that carry it. The experiment runs an SIS panic process through the flock – agents switch between calm and panicked, transmission scaling with the number of panicked flock-neighbors and recovery at a fixed rate – while a tenth of the flock is informed of a goal direction; a panicked informed agent suspends its goal bias and gains erratic noise. The transmission rate is swept.



Coherence is untouched. The order parameter stays between 0.98 and 0.996 across the entire range of transmission rates, even when nearly the whole flock is panicked, because the panic noise sits below the melting threshold established early in the study, so the flock remains a tight, globally aligned group. Panic does not fragment it. Steerability, however, collapses. The steering accuracy

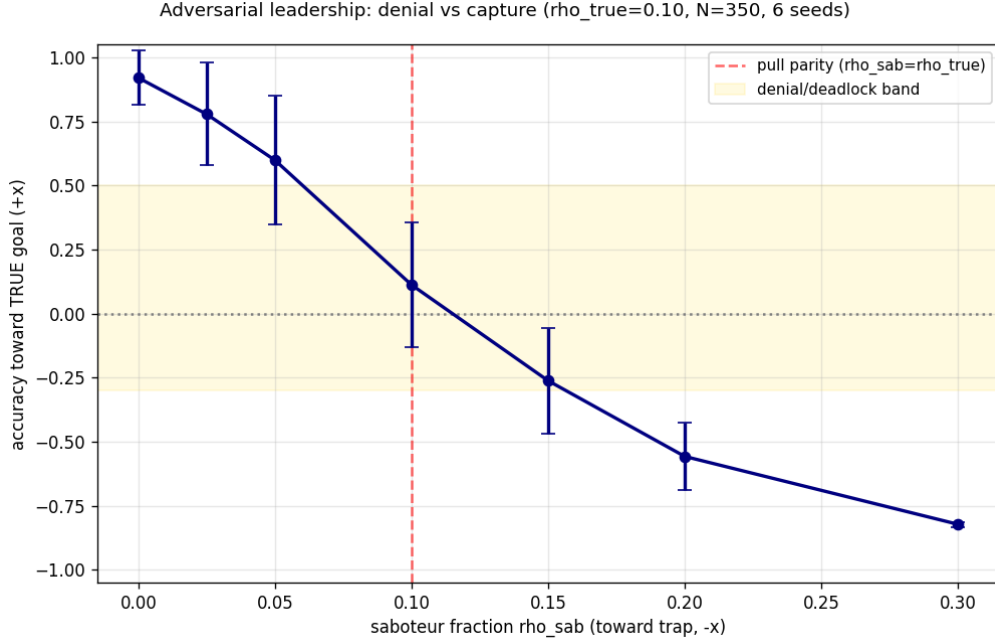
falls monotonically from 0.94 with no panic through 0.71 and 0.36 to essentially zero once the transmission-to-recovery ratio reaches about two, beyond which it scatters with large seed-to-seed variance about zero. The flock stays coherent but loses its heading, becoming a tight flock flying an essentially random, leader-uncorrelated direction. Contagion destroys the flock’s steerability, not its cohesion – it severs the rudder while leaving the hull intact.

The controlling quantity is the active-leader pull, the product law of Section 4.56 taxed by the panic fraction. Panic saturates even at the lowest transmission rate tested, because the spatial flock has a high contact degree, so the basic reproduction number far exceeds the transmission-to-recovery ratio and the outbreak is supercritical almost immediately, consistent with the earlier panic-contagion findings. The relevant axis is therefore not the epidemic threshold but the depth of saturation: at any instant a fraction of the leaders is panicked and silent, so the active leadership pull is the informed fraction times the leader strength scaled by the fraction not currently panicked, an effective informed fraction of one minus the panic fraction times the nominal one. When four-fifths of the flock is panicked, that effective fraction falls to about two percent, the weak-but-working leadership regime, and accuracy is still around 0.7; when nineteen-twentieths are panicked it falls below half a percent, beneath the leadership floor, and steering vanishes. Steerability tracks the active-leader pull exactly as the leadership and conviction findings predict, with panic acting as a multiplicative tax on the signal.

Contagion and predation are thus complementary attacks on the two faces of the same shared heading. Encirclement attacks coherence, pushing sub-groups apart, and leadership restores it; panic attacks steerability, silencing the leaders, while coherence, which needs only local alignment and not the shared goal, survives untouched. A flock can fail in two distinct ways – fragmented but aimable, the encirclement mode that leadership fixes, or coherent but rudderless, the panic mode that leadership cannot fix because the disease disables the leaders themselves. This is why contagion has been the most durable stressor in the whole study: an attack on coherence is reversible because the shared heading re-forms, but an attack that disables the carriers of the shared signal cannot be countered by that signal. The shared heading that is the antidote to predation is itself the casualty of contagion; the rudder cannot repair itself once the hands on it are panicking. The result closes the leadership thread by mapping its two adversaries onto the two resources – coherence and steerability – that the steering-bandwidth result showed to be one substrate seen two ways.

4.62 Adversarial Leadership: Denial Is Cheaper Than Capture (Finding 80)

The conflicting-leaders result can be read adversarially, and doing so exposes a security asymmetry. A fixed set of true leaders steers toward a goal while a swept set of saboteurs pushes toward a trap, and two distinct adversarial objectives are separated: denial, meaning preventing the flock from reaching its goal and leaving it deadlocked, and capture, meaning actively driving the flock to the trap. The question is whether these are equally hard, measured by the accuracy of the flock’s heading toward the true goal, which is positive when the goal is reached, near zero under deadlock, and negative when the flock is driven to the trap.



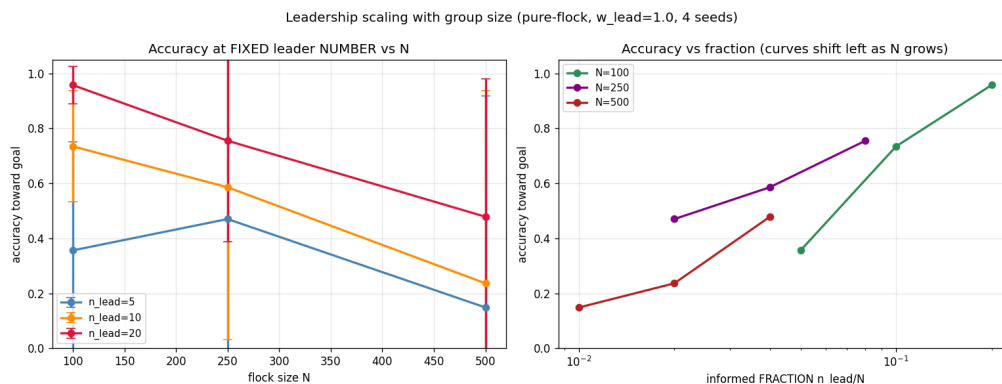
Denial is cheap. As the saboteur fraction rises to match the true leaders, goal-accuracy collapses from 0.92 to about 0.11: the flock is deadlocked, reaching neither goal nor trap. Even at half the true-leader count the accuracy is already halved. Paralyzing a led flock – mission failure – is achieved at or below pull parity. Capture is expensive. The accuracy crosses zero, with the flock beginning to head toward the trap, only past parity, and decisive capture with the flock committed to the trap requires a saboteur fraction twice the true-leader count. Actively hijacking the flock to a chosen false target costs far more than merely paralyzing it.

The asymmetry is the product law of Section 4.56 plus a threshold gap. The zero-crossing sits at pull parity, exactly as the product law predicts for equal leader strength, confirming that the adversarial contest obeys the same summed-directed-force accounting as cooperative voting. The new content is the gap between the two adversarial thresholds – denial at about the defender’s pull, capture at about twice it – with a deadlock band between them in which neither side wins and the flock wanders. Coherence declines only modestly across the whole sweep, from 0.995 to 0.88, so the opposed pulls fragment the flock slightly but never break it; the contest is over the flock’s heading, not its integrity.

The result is a security asymmetry for any led collective. An adversary with a leadership-style channel can paralyze a led flock with a force merely matching the legitimate leaders, but to commandeer it needs clear superiority. Symmetrically, to guarantee reaching the goal the true leaders need a pull majority over any saboteur rather than mere parity, since parity yields deadlock. This sharpens the voting picture into an attack-and-defense statement – the cheapest adversarial outcome is denial at parity, the most expensive is capture at majority, and the two are separated by a deadlock band – and it connects the leadership thread to the predator arms race: a predator able to mimic a leader would find preventing the flock from going where it wants far easier than herding it to a kill zone, consistent with the broader theme that disrupting the shared heading is easier than commandeering it.

4.63 A Self-Test: Steering Is Set by the Informed Fraction, Not the Absolute Number (Finding 81)

The first leadership result cited the animal-groups literature for the claim that the informed fraction needed for accurate navigation decreases as the group grows, but that claim was only ever asserted from a single group size. This experiment tests it directly by varying the flock size while fixing the absolute number of informed agents. If a fixed handful of leaders suffices in arbitrarily large groups, accuracy at fixed leader number should be independent of size; if instead the fraction governs steering, accuracy at fixed number should fall as the group grows.

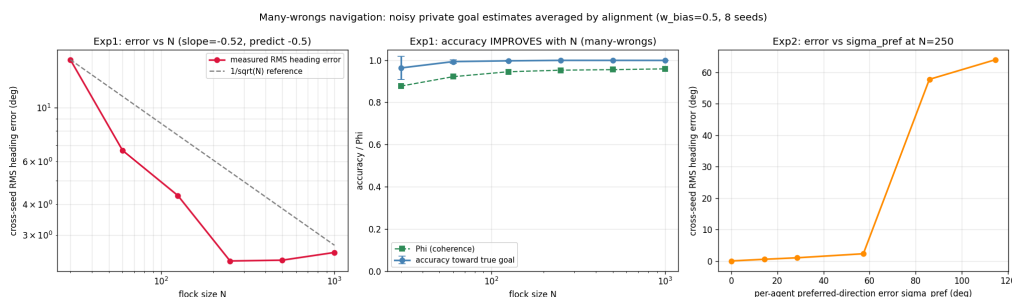


A fixed number of leaders does not suffice as the group grows. At every fixed leader count the accuracy decreases with flock size: twenty leaders steer a flock of a hundred almost perfectly but barely steer a flock of five hundred, and ten leaders fall from steering well at a hundred to near-random at five hundred. The absolute number of informed agents needed to reach a given accuracy grows with the group. Re-indexing the same data by fraction rather than number collapses it far better: accuracy is approximately a function of the informed fraction across all three sizes, with no monotonic size trend within the substantial cross-seed scatter. Steering accuracy is governed by the fraction because the mean-velocity steering is the total injected directed force divided by all the agents – the product law of Section 4.56 read per capita. Growing the group at fixed leader number dilutes the per-capita pull, and steering weakens.

This corrects the earlier offhand attribution. In this model the fraction needed is roughly constant, not decreasing, and equivalently the number needed grows with size, the opposite of the simplest reading of the fixed-number result in the literature. That result arises from explicit preferred-direction averaging with noise – the many-wrongs principle – which amplifies a sparse informed signal in large groups; the linear alignment force used here has no such amplification and delivers exactly the per-capita pull, so accuracy tracks the fraction. The correction sharpens the product law rather than overturning it: total pull decides, but read relative to the group size, which is the fraction. All the fraction-based leadership results stand as the correct frame, the one number-based intuition is the exception this test removes, and recovering the literature’s number-suffices scaling would require adding a many-wrongs amplification to the follower rule – a clean open direction. The experiment is offered in the same self-testing spirit as the falsified predictions earlier in the study: a claim asserted in passing, tested directly, and corrected.

4.64 Many-Wrongs Navigation: Noisy Private Estimates Average to a $1/\sqrt{N}$ Wisdom of Crowds (Finding 82)

The previous section closed with a prediction: linear velocity alignment gives no group-size amplification for leaders carrying an exact shared goal vector, so steering is set by the informed fraction and not the absolute number, and recovering the animal-navigation literature’s result that a fixed number of informed individuals suffices in arbitrarily large groups would require a many-wrongs follower rule – agents holding independent noisy estimates of the goal so that averaging over more of them cancels the error. This experiment tests that prediction directly. Every agent now carries its own private preferred direction, drawn once per run at an angle normally distributed about the true goal with spread σ_{pref} , and is biased toward its own estimate rather than toward a shared vector. There are no exact-vector leaders; every agent is a noisy one. The many-wrongs prediction is that the flock’s steady heading is the alignment-averaged resultant of all the private biases, with an angular error that shrinks as the per-agent spread divided by the square root of the group size, so accuracy toward the true goal should improve with group size – the opposite of the previous section – provided the flock stays coherent enough to average globally.



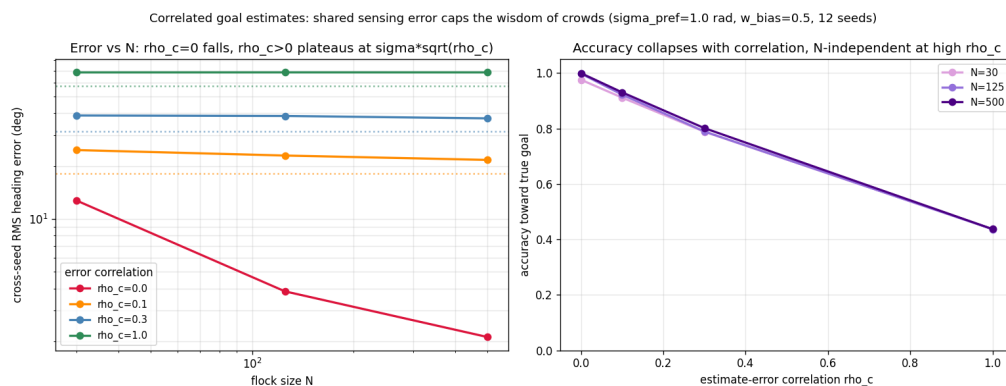
The prediction holds. At a fixed per-agent error of about fifty-seven degrees, the flock’s cross-seed root-mean-square heading error falls from sixteen degrees at thirty agents to about two degrees by two hundred and fifty, a log-log slope of minus one half to within the seed scatter – exactly the wisdom-of-crowds law. Accuracy toward the true goal rises with group size rather than falling, the clean inverse of the fixed-number result: each agent on its own would head off by nearly sixty degrees, yet a flock of two hundred and fifty navigates to within two degrees of the true bearing. Alignment is performing the average, pooling the independent private estimates so their errors cancel. This resolves the apparent tension with the previous finding rather than contradicting it: alignment averages whatever directional signal the agents carry, and with an exact shared vector there is no error to average away so only the per-capita pull remains, while with heterogeneous noisy estimates there is error and alignment drives it down with group size. The literature’s fixed-number scaling comes precisely from this averaging, and adding noisy estimates to the model recovers it as predicted.

Two boundaries qualify the law. First, the decline saturates at a floor of about two to two and a half degrees once the group passes a couple of hundred: the averaged-bias error has by then dropped below a second, size-independent error source – the temporal jitter of the flock heading within the measurement window together with the background noise – which averaging over more agents cannot remove. Second, and more striking, the averaging has a noise ceiling in the per-agent spread. As the spread grows the error at first rises gently and accuracy stays near perfect, but between about one and one and a half radians of per-agent error the behavior collapses abruptly: the heading error jumps to nearly sixty degrees, accuracy falls by half, and the cross-seed scatter explodes, while the order parameter barely moves. The flock does not fragment; it stays a tight

flock flying a nearly random heading. The reason is that the magnitude of the pooled estimate falls off exponentially with the square of the per-agent spread, so once the spread is large the averaged directional signal becomes too weak to overcome the flock’s own spontaneous heading and different runs commit to different directions – the many-wrongs form of the pull threshold seen in the conviction and minimum-leadership results. The wisdom of crowds is real and follows the square-root law, but only while the individual estimates are accurate enough that their pooled average still points somewhere definite. Taken with the previous section, this completes the leadership thread’s central mechanism: alignment is a directional averager, delivering per-capita pull for a shared exact signal and square-root crowd amplification for heterogeneous noisy ones, with a high-noise ceiling that is the averaging form of the same force-versus-alignment threshold that recurs throughout the study.

4.65 Correlated Estimates: Shared Sensing Error Caps the Wisdom of Crowds (Finding 83)

The previous two sections appear to reach opposite conclusions – exact shared leadership gives no benefit from group size while independent noisy estimates give a square-root benefit – but they are in fact the two ends of a single parameter, the correlation between the agents’ estimate errors. Real collectives sit in between: animals reading the same misleading cue, or agents fed common misinformation, share part of their error. This experiment builds each agent’s private preferred direction as a mixture of one shared error draw, common to the whole flock for that run, and an independent private draw, in proportions set by a correlation parameter, so that the per-agent error has fixed spread but tunable correlation. Zero correlation reproduces the independent-estimate case; unit correlation gives every agent the identical, shared but wrong, direction, which is the shared-signal case of the leadership sections. The prediction is that the alignment-averaged heading error splits into a private part that averages away with group size and a shared part that does not, leaving a floor equal to the per-agent spread times the square root of the correlation – a floor no group size can beat.



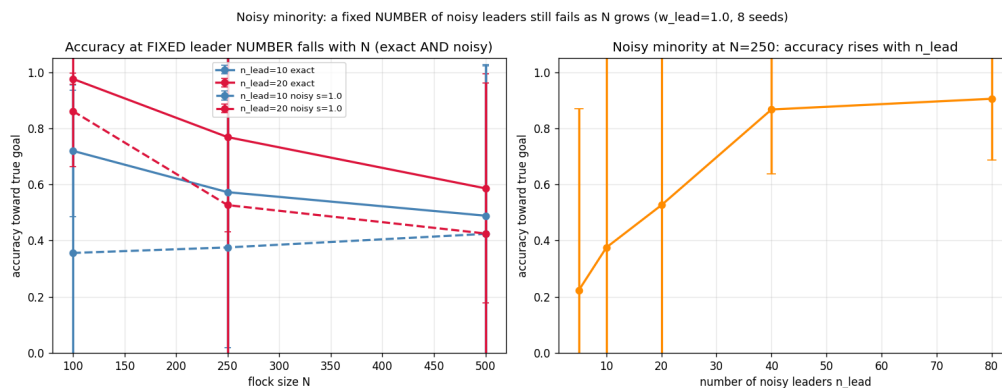
The data confirm it cleanly. At zero correlation the error falls as the inverse square root of group size, from about thirteen degrees at thirty agents to about two at five hundred, reproducing the wisdom-of-crowds law. At unit correlation the error is exactly independent of group size – the same sixty-eight degrees at thirty, a hundred and twenty-five, and five hundred agents – with the order parameter at a perfect one, every agent carrying the identical wrong direction so the flock agrees completely and heads off by the shared offset. That is the no-amplification limit of the leadership sections, recovered as the high-correlation end of the same axis. In between, any positive correlation makes the error flat in group size: a mere ten-percent correlation holds the error near twenty-two

degrees and the accuracy near nine-tenths whether the flock is thirty or five hundred strong, and thirty-percent correlation holds it near thirty-eight degrees. The shared component of the error does not average away no matter how many agents pool their estimates. The measured floors sit fifteen to twenty percent above the small-angle formula because the per-agent spread of one radian is not small, but the ordering and the group-size independence match exactly.

A further point emerges from the order parameter, which rises with correlation from around nine-tenths at zero to a perfect one at unit correlation. More correlated goals make the flock agree more tightly – but on an increasingly wrong heading, as the accuracy falls from nearly perfect to less than one half. Independent errors slightly loosen cohesion yet cancel for accuracy; shared error tightens cohesion onto a common mistake. This is the consensus-versus-correctness distinction in its sharpest form: a perfectly coherent flock can be confidently and unanimously wrong, and the large run-to-run scatter at high correlation is exactly that, each run’s shared draw committing the whole flock to a different definite error. The practical content is a warning about the wisdom of crowds: it delivers square-root accuracy only for independent errors, while common-mode error from a shared cue or correlated sensors imposes a floor that more agents cannot lower. It also connects back to the adversarial result: an attacker who cannot field enough saboteurs to capture the flock can instead inject correlation into the legitimate agents’ estimates with a single shared false cue and cap the collective’s accuracy regardless of its size, so common-mode deception is cheaper than majority capture. This closes the many-wrongs sub-thread with a unifying statement: alignment is a directional averager whose collective accuracy is governed by the correlation structure of its inputs, not their number.

4.66 The Noisy Minority: Informed-Minority Steering and the Wisdom of Crowds Are Distinct (Finding 84)

The navigation literature keeps two effects apart that this study can now place in one setup. One is informed-minority steering, where a few agents with a preferred direction steer the whole group; the other is the many-wrongs principle, where many independent noisy estimates average to higher accuracy as the group grows. The exact-vector minority of the scaling self-test and the all-agents-noisy crowd of the many-wrongs section are the two pure cases; the bridging case is a fixed number of informed agents each carrying its own noisy estimate, with the rest naive followers. Here two scales pull in opposite directions as the group grows: the direction of the injected pull is the pooled estimate of the informed agents, whose error is fixed by their number and not by the group size, while the strength with which that pull is felt per agent dilutes with group size exactly as in the self-test.

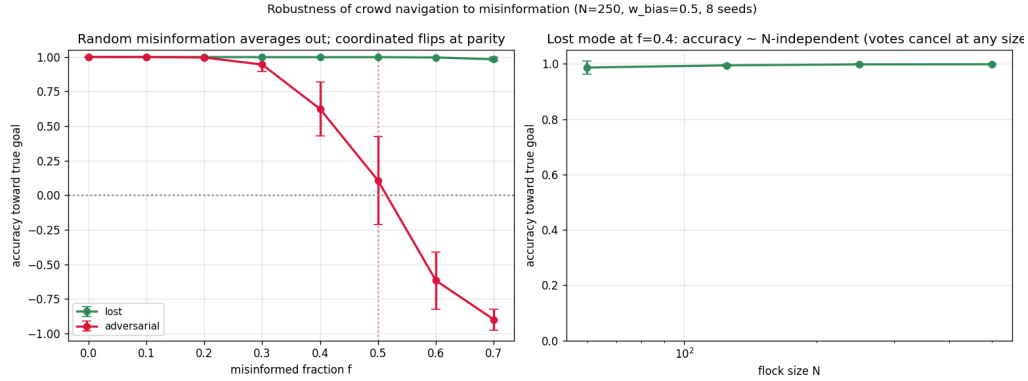


The result is unambiguous: a fixed number of noisy leaders does not suffice as the group grows. The exact minority falls with size at every leader count, reproducing the self-test, and the noisy minority also fails – at twenty leaders its accuracy falls from about nine-tenths at a hundred agents to below one half at five hundred, and at ten leaders it sits near the spontaneous-heading floor regardless of size. The internal averaging of the minority cannot rescue the dilution of the per-capita pull, because the pooled direction is at best a fixed accuracy set by the fixed leader number while the strength of its imposition weakens with the group. The noisy minority moreover sits at or below the exact minority everywhere, the penalty being the pooled-direction error: the leaders agree on a heading that is itself off the true goal by roughly eighteen degrees for ten leaders, and no amount of follower alignment can recover a target the leaders themselves get wrong. An exact minority points the flock exactly right but weakly; a noisy minority points it weakly and slightly wrong. Only growing the number of leaders at fixed group size recovers accuracy, climbing from about a fifth at five leaders to nine-tenths at eighty, because more leaders give both more per-capita pull and a better-pooled direction.

This separates two mechanisms the literature often conflates. Informed-minority steering and the wisdom-of-crowds averaging are distinct and do not combine in a fixed minority: confining noisy estimates to a fixed cadre gives per-capita-diluted steering toward a fixed-accuracy pooled direction, strictly worse than an exact minority, while the square-root amplification of the many-wrongs section requires the informed fraction itself to grow. It is the sharp form of the self-test’s correction – the claim that a fixed number suffices holds for the many-wrongs crowd only when every agent is an estimator, never for a fixed informed minority. Taken together the four sections complete the thread’s mechanistic map: alignment is a directional averager whose output accuracy is set by the per-capita pull, fraction times strength, toward a target whose own accuracy is set by the correlation structure of the estimates and the size of the estimating set, so number helps only when it grows the estimating fraction and never as a fixed minority.

4.67 Misinformation: A Crowd Averages Out Noise but Flips to a Coordinated Falsehood (Finding 85)

The mechanistic map of the preceding sections carries a practical corollary worth testing directly. Because each agent contributes a unit vector to the directional average, its influence is bounded however wrong its heading, so the averaging should be intrinsically robust to outliers – but only to uncorrelated ones, since the correlation section showed that shared error does not average away. This experiment contrasts two kinds of misinformation carried by a fraction of the flock against a well-informed majority that holds a tight estimate of the true goal. The misinformed members are either lost, pointing in uniform-random directions with no coordination among them, or adversarial, all pointing at a single false goal in the opposite direction. The prediction is that lost members cancel, their random unit votes summing to little against the aligned majority, so accuracy holds until they are very numerous, while adversarial members compete with the true majority vote for vote, so accuracy crosses zero at parity.



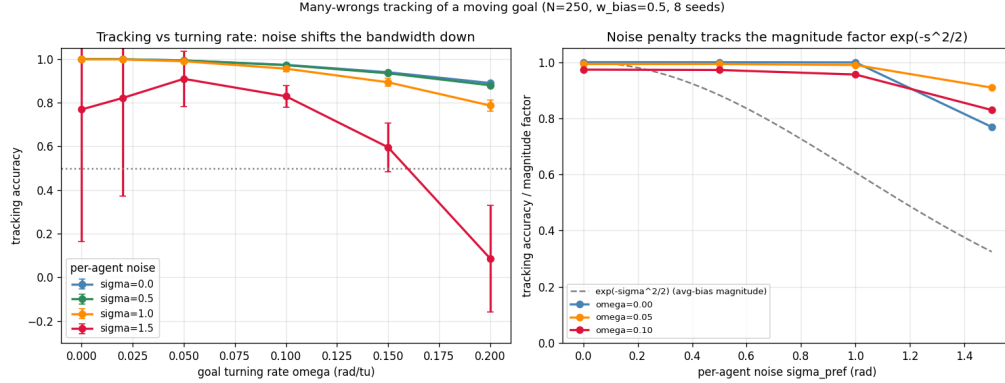
The contrast is stark. Uncoordinated misinformation is almost harmless: with lost members pointing every which way, the flock still navigates to the true goal at accuracy nine hundred and ninety-eight thousandths even when half its members are misinformed, and only slips to about ninety-eight hundredths when seven in ten are lost. The random votes cancel and the heading stays locked on the truth, the order parameter sagging only slightly as the lost members add a little incoherence. Coordinated misinformation of the very same prevalence is decisive: as the adversarial fraction climbs, accuracy falls to about six tenths at four in ten and crosses zero at parity, half the flock, then reverses to capture as the false consensus takes over. The zero-crossing at parity is the product law of the conviction and adversarial-leadership sections, equal per-agent strength balancing at equal numbers, and the coherence cost peaks exactly at the parity point, the lowest order parameter in the sweep, the heading-fight signature of a flock torn between two equal opposed consensuses before one wins. The damage is done not by the amount of error but by its correlation: the same number of equally-wrong agents is averaged away when independent and flips the whole flock when coordinated. A further check confirms the robustness is structural rather than a finite-size accident, the accuracy under a fixed fraction of lost members being essentially independent of flock size.

This caps the many-wrongs arc with its security reading. A navigating collective is robust to noise but fragile to a coordinated falsehood of the same size: uncoordinated misinformation, lost or confused or independently erring members, is averaged out even at half the flock, while a coordinated false consensus captures the flock once it reaches parity. It is the constructive mirror of the adversarial-leadership result, where denial was cheaper than capture at parity, and it states the general lesson that runs through the adversarial, correlation, and misinformation findings: what threatens a collective's heading is never the amount of error but its correlation, because alignment averages out everything independent and is moved only by what is shared.

4.68 A Self-Test: A Noisy Crowd Tracks a Moving Goal at the Same Bandwidth as a Sharp Leader (Finding 86)

The steering-bandwidth section established that a flock tracks a turning goal only below a critical rate set by its alignment response time, and the many-wrongs sections established that a crowd of noisy estimators navigates to an accurate heading by spatial averaging. This experiment combines them on the open question of how a noisy crowd tracks a moving goal: the goal direction rotates and every agent biases toward the current goal carrying its own fixed angular offset. Before running it I committed to a prediction – that the many- wrongs average, being spatial and instantaneous, would add no lag, so the noise would lower the tracking bandwidth only by shrinking the magnitude of the averaged bias, which falls off as the exponential of minus half the squared per-agent spread, exactly the law found earlier for the static noise ceiling. In other words the bandwidth was predicted

to scale with that magnitude factor.



The prediction is wrong, and instructively so. For per-agent spreads up to a radian the tracking-accuracy curves against turning rate lie almost on top of one another, and the bandwidth is unchanged even though the magnitude factor has fallen to six tenths. The many-wrongs noise costs essentially nothing for moving-goal tracking: a noisy crowd turns to follow the goal as well as a sharp leader does. The reason the magnitude reduction does not translate into lost bandwidth is that the per-agent offsets are static and rotate with the goal, so the averaged bias points cleanly at the current goal each step and adds no temporal lag, while the reduced magnitude still leaves the pull well above the threshold needed to steer, so the bandwidth, which is set by the response time rather than the pull magnitude in this range, does not move. The magnitude law bites only when it pushes the pull below threshold, and that is precisely the noise ceiling: at a per-agent spread of one and a half radians, past the ceiling found earlier, tracking does not degrade gracefully but collapses, the static case itself becoming unreliable with large run-to-run scatter and the fastest turn essentially untracked. Noise thus has a binary effect on moving-goal tracking – free below the ceiling, catastrophic above it – rather than a graded bandwidth cost.

Offered as another self-test in the spirit of the falsified predictions earlier in the study, this corrects and sharpens the relationship between the two averaging notions. Steering bandwidth, a temporal property set by the alignment response time, and crowd accuracy, a spatial property set by the averaging of estimates, are independent resources: many-wrongs noise spends the second and not the first, so a noisy crowd tracks a turning goal at full bandwidth and pays only in steady accuracy until the noise ceiling, where the averaged signal collapses outright. It strengthens the directional-averager thesis by showing the average is recomputed afresh each timestep rather than integrated over time, and it closes the many-wrongs arc.

4.69 Evolving the Escape Weight: The F70 Valley Is a Strong Evolutionary Brake (Finding 87)

Every behavioral trait studied so far has been fixed by hand. This experiment opens a co-adaptation thread by making one trait — the collective-escape weight of Section 4.52 — heritable and placing it under predation selection. Each prey carries its own escape weight w_i , feeling a force $w_i \hat{e}$ along the shared unit vector $\hat{e}$ from the predator centroid toward the flock centre of mass (the Finding 70 signal), and w_i is inherited. The predators are the hardest found, predictive encirclement at a lead of two time units (Finding 66), with six predators in the slow-prey regime at $N = 150$. The fitness model is capture and removal: an agent that comes within a kill radius $r_{\text{kill}} = 0.03$ of any predator is captured at a hazard rate of three per time unit and replaced by a mutated

clone of a random survivor, which inherits the parent’s weight plus a small Gaussian mutation and is born at the parent’s position and velocity. The population size is held fixed, a Moran-style continuous-replacement scheme. The question is the evolutionary form of the Finding 70 valley: since a weak escape ($w \approx 0.25$) is worse than none and only $w \gtrsim \alpha = 1$ outruns the trap, does selection drive a cautious population across that valley to the escape regime, or does the valley trap it? The initial weight is swept over $\{0, 0.25, 0.5, 1, 2\}$.

Over 150 time units the outcome is set sharply by the starting weight. A population seeded in the escape regime is evolutionarily stable and almost never caught: starting at $w_0 = 2$ it stays at $w = 2.00$ with about six captures and an order parameter of 1.000, and starting at $w_0 = 1$ it holds $w = 1.00$ at an order parameter of 0.992. A population seeded at or below the valley stalls far short of escape — $w_0 = 0$ rises only to $w \approx 0.27$ (order parameter 0.22, 1804 captures), $w_0 = 0.25$ to 0.49, $w_0 = 0.5$ to 0.66 — and the capture toll peaks precisely in the valley (about 1800–1900 captures at $w_0 = 0$ and 0.25) before collapsing to single digits at $w_0 = 2$. Predation cost is concentrated exactly where Finding 70 placed the valley. A longer run separates a true barrier from a slow brake: selection on the weight is directional and upward from every start — even $w_0 = 0$ drifts up — but the valley throttles the climb, so from no escape the mean weight crawls to only ≈ 0.13 by 50 time units, 0.18 by 150, and 0.51 by 400, never reaching the escape threshold, while a start just past the worst of the valley ($w_0 = 0.5$) climbs faster, to 0.88 by 400 time units and accelerating. Neither low start crosses $w = 1$ within 400 time units.

The Finding 70 force-versus-alignment threshold therefore has a population-genetic image: it is a strong evolutionary brake rather than an absolute barrier. Escape behaviour is trivial to maintain — stable, self-reinforcing, and nearly cost-free once $w \gtrsim \alpha$ — but very hard to evolve from scratch, because the only path there runs through the valley where escape is actively harmful, so selection pushes the weight up only weakly and the flock crawls for hundreds of time units without establishing escape. This is strong evolutionary hysteresis, a first-mover problem in which the basin a population occupies is set by where it starts and the cost-free escape optimum is effectively unreachable from rare. It is the domination-not-blending theme of Findings 16, 24, and 70 read at the evolutionary level: a globally shared escape direction pays off only once it is strong enough to beat alignment, so partial commitment is selected against, and the naive reading of Finding 70 — that escape wins — is inverted, escape winning only where it is already present. The result opens the co-adaptation thread with the predator side still fixed; the natural next steps are to let the predator’s lead time or aggression co-evolve against the prey trait, to ask whether a seeded escape-carrying minority or a larger mutation step can jump the gap, and to confirm the brake under the alternative fitness models. (Figure: `figures/escape_evolution_1.png`.)

4.70 Overcoming the Brake: Escape Is a Barrier to Origination, Not Invasion (Finding 88)

If the valley blocks escape from evolving *de novo*, the natural question is whether escape can take hold any other way. This experiment keeps the capture-and-removal model and the validated harness of Section 4.69 and changes only the initial condition and the mutation step. In the first part a fraction f of the population is seeded in the escape regime (weight 2) with the rest at no escape, and f is swept to ask whether escape invades or is diluted away; in the second part the population starts uniformly at no escape and the per-capture mutation step is enlarged, to ask whether a bigger jump can clear the valley that the small Finding 87 step could not.

Invasion succeeds from a tiny seed. Even a five-percent escaper founder group carries the flock-mean

weight into the escape regime — the mean weight ends at 1.18, the escaper fraction climbs from 0.05 to about 0.55, and the capture toll roughly halves relative to the de-novo crawl of Finding 87 — and a ten-percent seed reaches an escaper fraction of 0.70 at an order parameter of 0.91. Larger seeds settle to much the same place. The escaper trait does not drive the non-escapers extinct; it settles at a mixed equilibrium of roughly sixty percent escapers at a mean weight near 1.2. The mutation-step sweep is equally clear: the Finding 87 step never reaches the escape threshold in any seed, but a step three to six times larger crosses it in every seed, while the very largest step is noisier and crosses only half the time, because huge jumps scatter as many offspring back to low weight as forward past the valley — there is an intermediate mutation-scale sweet spot.

Together with Finding 87 this locates the barrier precisely: the Finding 70 valley blocks the origination of escape, not its spread. Escape cannot crawl up from zero against the valley, but it does not need to — a rare founder group already carrying the trait invades and establishes it, and a single mutational jump large enough to clear the worst of the valley seeds the same outcome from a uniform no-escape start. Whether escape evolves is therefore set not by selection, which favours it once the valley is cleared, but by whether variation can deliver an agent past the valley in one step or a pre-adapted group arrives — a mutation-limited, standing-variation-limited phenomenon. The mixed equilibrium rather than full fixation is itself a result: it is a free-rider effect of the shared escape signal, since once enough agents flee the predator is outrun and the remaining low-weight agents ride along protected, relaxing selection on them. The public-good character of the shared escape direction — the constructive mirror of the shared-heading principle — keeps escape from fixing even as it comes to dominate. (Figure: `figures/escape_invasion_1.png`.)

4.71 The Free-Rider Equilibrium Is Robust to Predation Pressure (Finding 89)

The mixed equilibrium of Section 4.70 — escape dominating at roughly sixty percent rather than fixing — invites a quantitative question: what sets that fraction, and does intensifying predation drive it to fixation by catching the free-riders? This experiment answers it with a pure predation-pressure sweep in the same capture-and-removal model. Each run starts from an established escaper population (half seeded in the escape regime) and the capture hazard within the kill radius is swept across an eightfold range, with the steady escaper fraction read off as the mean over the final third of a long run.

The free-rider equilibrium proves robust. The steady escaper fraction rises only weakly with predation, from about 0.56 at the lowest hazard to about 0.66 at the highest — an eightfold change in predation pressure moving the equilibrium by a tenth — and never approaches fixation: even under the most intense predation roughly a third of the flock rides the shared escape as protected free-riders, while the mean escape weight stays in the escape regime throughout. A convergence check from two different starting fractions settles into the same band, escape persisting as a substantial but unfixed majority from either side, with only mild residual start-dependence. The direction matches the intuition that intense predation erodes free-riding, but the magnitude is small: escape is a sticky mixed strategy that neither collapses nor sweeps to fixation.

This bounds the evolutionary outcome of the chosen model and confirms the public-good reading of the previous result. Because the escape signal is shared and non-excludable, a protected free-rider fraction persists as long as the flock escapes at all, and raising the predation pressure only mildly raises the cost of riding rather than eliminating it. The same shared-direction logic that lets a tiny informed minority steer a whole flock, and a tiny escaper minority establish escape, here keeps escape from ever becoming universal — a shared directional signal benefits the non-contributors too,

so contribution need not go to fixation. The result also carries the hysteresis signature of the two preceding findings, the steady fraction retaining a little memory of where the population started. (Figure: `figures/escape_freerider_1.png`.)

4.72 The Two-Sided Arms Race: Capture Is Not Disruption, and the Race Is Asymmetric (Finding 90)

The three preceding evolution experiments held the predator fixed. This one lets it adapt too, completing the co-adaptation thread as a genuine two-sided arms race. The prey keep their heritable escape weight under capture and removal; each of the six predators now carries a heritable predictive lead time and the predators are selected on capture success, the worst-capturing predator being periodically replaced by a mutated clone of the best — the small-population counterpart of the prey’s continuous replacement, with each capture credited to the nearest predator. Predators start at random lead times. Three experiments isolate the questions: the first evolves only the predator against frozen no-escape prey, a clean test of what capture selection alone favours; the second co-evolves both from no escape; the third co-evolves both from a seeded escaper population.

Against frozen no-escape prey the predator lead converges tightly, across seeds, to about three time units, somewhat above the value near two that Finding 66 identified by hand as the most disruptive. It is tempting to read this as evidence that capture selection favours a different lead than disruption — but the direct measurement of the next section (Finding 91) refutes that reading: the capture rate actually peaks at a lead of two, coinciding with the disruption optimum, and the evolutionary overshoot to three is an artifact of the small six-predator replace-the-worst selection scheme operating on a high-variance capture signal. The capture and disruption optima coincide; the evolved value is simply not at that optimum. The two co-evolution runs then expose the result that does not depend on the exact lead, a deep asymmetry. From no escape the predator holds its high capture-optimal lead while the prey climb only partway, escape never establishing within the run, so the predator wins; from seeded escape the prey escape weight goes essentially to fixation and the predator’s lead trait drifts aimlessly across seeds, because committed escape yields no captures and therefore no selection signal on the predator at all — a use-it-or-lose-it collapse of the predator trait under relaxed selection.

The arms race is therefore fundamentally asymmetric. The predator can climb to its optimum from any start, having no barrier to cross, whereas the prey counter is origination-limited by the valley of Sections 4.69 and 4.70, so de-novo co-evolution favours the predator and escape establishes only when it is seeded — after which it is an uncounterable hard counter that collapses predator selection entirely. This also reconciles the free-rider equilibrium of the previous section with the present near-fixation: the sixty-percent mixed equilibrium requires a persistently effective predator to sustain it, and once the predator is allowed to become ineffective by co-evolving against winning prey, predation collapses and escape fixes. The co-adaptation thread thus closes with a coherent picture of collective escape as a powerful but evolutionarily fragile defence — hard to originate, easy to lose to drift, yet decisive once present. (The lead-time trajectories are noisy at two to three seeds and predator selection acts on only six individuals; the signs and the tight frozen-prey result are robust, the co-evolution magnitudes less so.) (Figure: `figures/escape_coevolution_1.png`.)

4.73 A Self-Test: The Capture Optimum Coincides with the Disruption Optimum (Finding 91)

The previous section reported that the predator evolved a lead near three, and it was tempting to conclude that capture selection favours a longer lead than the value near two that most disrupts the flock — that catching prey and fragmenting the flock are different objectives. This section tests that conclusion directly rather than through the noisy evolutionary dynamics. Holding all six predators at a fixed common lead against frozen no-escape prey, the steady capture rate and the order parameter are measured as the lead is swept, with no evolution.

The conclusion does not survive. The capture rate peaks at a lead of two — about ten captures per time unit, against seven at lead one and barely two at lead zero — and then falls steeply, to a noisy four at lead three and to almost nothing at leads four and five, where the predators overshoot the flock entirely. The order parameter is lowest, the flock most fragmented, at exactly the same lead of two. The lead that catches the most prey and the lead that most disrupts coherence are therefore the same, not different, and the evolved value of three from the preceding section is not a capture optimum at all: at lead three the capture rate is not only lower than at two but wildly variable from run to run, and the small six-predator replace-the-worst selection scheme, cloning whichever predator happened to capture most in a given window, chases that high-variance tail away from the true mean-rate optimum. The capture-versus-disruption distinction claimed in the previous section is thus withdrawn; the broader lesson survives only inverted, as a methodological caution that a tightly converged evolved trait under noisy selection in a small population need not sit at the trait's fitness optimum, and should be checked against a direct fitness measurement before being read as optimal. This is the sixth self-test of the study, and the only one to correct a result generated within the same body of work. (Figure: `figures/escape_capture_curve_1.png`.)

4.74 Robustness Under an Energy-Budget Fitness Model: The Brake Is Model-Independent, the Persistence of Escape Is Not (Finding 92)

The whole co-adaptation thread so far used one fitness model, capture and removal, under which a high escape weight is nearly free once established. This experiment tests whether the central result, the origination brake, depends on that choice, by re-running the thread under an energy-budget model that bakes in the metabolic cost the escape valley implies: an always-on death hazard proportional to the escape weight is added to the predator-capture hazard, so escaping is expensive even when safe, and dead agents are replaced by mutated clones of survivors as before. The cost is held fixed while the initial weight is swept, and then swept while the population is seeded with escapers.

The result splits the earlier findings cleanly into a model-independent half and a model-dependent half. The origination brake survives untouched: from no escape the weight stays near zero, so escape still cannot evolve *de novo*, confirming that the valley barrier is a general feature and not an artifact of capture and removal. But the complementary half of the first evolution result — that escape, once present, is nearly free and stable — does not survive. At a moderate metabolic cost every starting condition collapses to no escape, including a population seeded deep in the escape regime, which is driven all the way down rather than settling at an interior optimum. Sweeping the cost from a seeded escaper population, the evolved steady weight falls sharply: with no cost it sits at the mixed equilibrium near weight one and a quarter, but the smallest non-zero cost tested already collapses it to near zero, and larger costs only deepen the collapse. The transition is nearly all-or-nothing rather than a graded interior optimum.

The mechanism is that the metabolic cost is paid by every escaper continuously, whereas predation threatens only the few agents near a predator at any instant, so even a modest per-capita cost outweighs the diffuse, intermittent benefit of escaping; collective escape is viable only where escaping is essentially free. Collective escape is therefore even more evolutionarily fragile than the co-evolution results implied — not only hard to originate and easy to lose to drift, but unsustainable under any non-trivial metabolic cost. The robust, model-independent conclusion of the thread is the origination brake; the persistence of escape, by contrast, is model-dependent and, once escaping carries an explicit cost, precarious. (Figure: `figures/escape_energy_1.png`.)

5 Synthesis: Alignment-Driven Kinematic Mixing as a Unifying Mechanism

Several of the strongest results in this study — the failure of spatial vaccination (Section 4.20), the dimensional specificity of encirclement (Section 4.25), the reversibility of encirclement damage (Section 4.7), the herd-immunity inflation (Section 4.14), and the long-time merge/split steady state (Section 4.15) — converge on a single mechanism: **alignment-driven kinematic mixing**. The flocking force not only aligns velocities but, as a side effect, continuously reorganizes the spatial neighborhood graph faster than any *transient* structural feature can be exploited or sustained.

One closely related result is deliberately excluded from that list. The failure of degree-targeted vaccination (Section 4.18) looks like a mixing result and was originally synthesized as one, but Section 4.30 shows it is not: freezing the contact graph by a thirtyfold reduction in mixing rate does not rescue degree-targeting. That null is *structural* — the flock contact network simply has no hubs to target — and holds whether or not the graph mixes. Keeping this case separate is what makes the mixing mechanism falsifiable rather than all-explaining, and the distinction is developed below.

The argument is direct. In a flock at $\Phi \sim 1$, all agents move with similar velocity but with the small dispersion induced by repulsion, noise, and the local averaging in the flocking force. Neighbor identities therefore turn over on a timescale set by these local fluctuations, not by the bulk-flock motion. Any property attached to agent identity — its degree in the contact network, its spatial coordinate, its location relative to the flock perimeter — is therefore shuffled on the same timescale. Three consequences follow:

Targeting collapses to random — for two distinct reasons. Vaccination targeting can fail either because the structure it aims at does not exist, or because that structure exists momentarily but is erased before it can be exploited. The flock exhibits one case of each, and Section 4.30 is what separates them.

Degree-targeted vaccination assumes a fat-tailed degree distribution with a stable hub set. The flock contact network has no fat tail — its degree distribution has $CV \sim 0.68$ (Finding 28), close to a random geometric graph, with no agents whose removal would disproportionately fragment the contagion network. Hub-targeting therefore cannot beat random, and Section 4.30 confirms this is structural rather than kinematic: reducing the contact-graph mixing rate thirtyfold (by driving the flock into its solid regime) leaves the targeted-versus-random advantage scattered around zero with no trend. The degree-targeting null is a static property of the flock graph.

Spatial vaccination is the genuine kinematic-mixing case. Maxmin farthest-point sampling does construct a real structural feature — a set of immune agents evenly spread across the flock — and at the instant of sampling that coverage is exactly as designed. But the coverage is attached to

spatial position, and position is precisely what kinematic mixing shuffles. By the time the epidemic runs, the immune agents have drifted into the same clustered, gap-ridden arrangement as a random sample (Section 4.20). Here mixing is the operative cause: the feature exists, and mixing destroys it.

Section 4.28 extends both null results to 3D: random, spatial, and degree-targeted vaccination remain statistically indistinguishable in three dimensions, and the contact-degree distribution is if anything less heterogeneous ($CV \sim 0.59$). The degree-targeting null is dimension-independent because degree homogeneity is; the spatial null is dimension-independent because kinematic mixing is.

Damage with no internal state is reversible; damage with internal state is not. Encirclement applies a directed external force pattern. Once removed, agent positions and velocities relax under the same alignment force that does the mixing — and they relax fast, on the ~ 10 -time-unit timescale of one full neighbor-graph turnover (Section 4.7, Finding 22). Contagion, by contrast, attaches a binary panic/calm flag to agent identity that mixing cannot erase. The asymmetry of recovery timescales (Section 4.16) is therefore not coincidental: it follows from where the damage is stored. Anything stored in positions or velocities is mixed away; anything stored in internal state outlasts the stressor.

Herd-immunity inflation is the kinematic-mixing signature in equilibrium. The herd-immunity threshold $p_c \sim 0.46$ is more than twice the mean-field value (Section 4.14, Finding 30). Mean-field treats the contact graph as well-mixed and instantaneously sampled; the spatial-SIS literature treats it as static and clustered. The flock is in between: clustered at any instant but constantly resampled. The factor-of-two inflation reflects the residual clustering that mixing does not fully erase per epidemic timescale. The same threshold is found for random and targeted strategies alike — for the two distinct reasons set out above: degree-targeting has no hubs to exploit, and spatial targeting has its coverage erased by mixing.

Dimensional specificity of encirclement. In 2D, predators arranged on a ring around the flock close off the boundary: every in-plane escape direction is covered, and the flock cannot leave (Sections 4.5-4.7). In 3D, predators arranged on a sphere of the same radius cover only a small fraction of the 4π solid angle, and the prey simply leave through an uncovered direction. Encirclement is 2D-specific because the geometric coverage requirement scales as solid angle, not arc length: a modest number of point predators can seal a closed curve but not a closed surface. Sections 4.25-4.31 confirm this directly and exhaustively — with correctly repulsive predators, 3D encirclement leaves the order parameter at $\Phi \sim 1.000$ at every predator count (1-50), every encirclement radius, every adaptive scheme, and both spherical and planar arrangements, while 2D encirclement at the identical settings drives Φ to ~ 0.73 . This dimensional specificity is a geometric fact about enclosure, distinct from the kinematic-mixing mechanism that governs the vaccination results; it is grouped here only because both express the same broad theme — the flock has no structural weakness that a static spatial strategy can exploit.

The non-equilibrium phase-transition diagnosis (Section 4.19, Section 4.21, Section 4.22) sits adjacent to this story rather than within it. The smooth crossover is a property of the driving, not the mixing: uniform random kicks without viscous dissipation violate FDT (Section 4.21) and prevent crystallization of the soft repulsion (Section 4.22). But once a flock is forming and the alignment force is active, the same kinematic mixing that defeats targeting is set in motion. The two threads share the model and the alignment force but address different levels of organization: phase behavior is a property of the velocity distribution, mixing is a property of the spatial-neighbor graph.

This synthesis made a falsifiable prediction, and Section 4.29 tested it. The prediction was that

replacing the metric alignment force with a topological (k-nearest-neighbor) one would exhibit weaker mixing — the neighbor graph being more stable because k-nearest is a “permutation-stable” structure — so that targeted vaccination would partially recover its advantage. The prediction was falsified: under k-NN alignment the neighbor-graph turnover is unchanged (Jaccard dissimilarity 0.036 versus 0.037 per two time units) and targeted vaccination still confers no advantage over random.

The falsification sharpens rather than weakens the synthesis. The error in the prediction was to conflate two networks. The topological rule alters the *alignment* graph — whose velocity each agent averages — and indeed fixes every agent’s alignment degree at exactly k. But contagion, targeting, and the herd- immunity threshold all live on the *contact* graph, the set of agents within R_CONT . The contact graph rewires because agents with slightly dispersed velocities physically slide past one another, and that dispersion — generated by repulsion, noise, and the local averaging — is identical under both alignment rules. Kinematic mixing is therefore a property of the agents’ physical relative motion, not of the alignment rule’s neighbor-selection topology. The mechanism is more robust than first claimed: there is no permutation-stable variant of the alignment force that static targeting could exploit, because the graph that targeting must beat is not the graph the alignment force defines. A genuine escape would require freezing the *contact* graph itself — that is, suppressing the relative motion of agents — which is incompatible with a flock that moves.

5.1 The Exploitable Invariant: Why Slow-Recoverer Targeting Is the One Strategy That Beats Random

The kinematic-mixing argument explains a long run of targeting nulls, but the study also found exactly one targeting strategy that beats random, and it fits the same framework as the case that proves the rule. Across roughly fifteen experiments, degree-targeting failed structurally (no hubs) and spatial targeting failed kinematically (mixing erased coverage); both failures trace to the fact that the quantities they target — contact degree and spatial position — are either absent or shuffled. Heterogeneous recovery (Sections 4.36-4.46) introduced a third kind of target: the per-agent recovery rate. Unlike degree and position, an agent’s recovery rate is an internal-state invariant that the kinematic motion cannot scramble, and targeting the slowest recoverers beats random by factors of two to three — the first and only strategy in the study to do so. It works in two and three dimensions, under continuous as well as bimodal rate distributions, with noisy estimates of the rate, and with small reservoirs, and it reverses the otherwise-durable predator-plus-contagion damage once the vaccination budget covers the reservoir.

This is the same principle as the reversibility asymmetry, stated from the constructive side. Damage or advantage stored in positions and velocities is mixed away; damage or advantage stored in a durable per-agent label is not. The recovery rate is such a label, which is why targeting it succeeds where targeting mixed quantities fails — and the boundary is sharp: when the rate label is allowed to drift, the advantage erodes and then vanishes exactly as the label decorrelates (Section 4.44), and at fast drift the epidemic self-averages back to the homogeneous threshold. The exploitable target is therefore defined precisely: it is whatever the flock’s motion cannot erase. Mixing rules out structural and spatial targets; only a stationary internal-state invariant survives as a handle, and the recovery rate is the one such handle this model contains.

5.2 The Shared Heading: How Collective Direction Is Built, Contested, and Attacked

A second unifying axis organizes the later threads — the arms race and the leadership experiments — and complements the mixing story rather than competing with it. Where kinematic mixing concerns what the flock cannot retain (any per-agent structure save a stationary internal label), the shared-heading principle concerns what the flock can act on collectively: a single direction held in common and amplified by the alignment force. Every result from the escape counter through adversarial leadership is a statement about this one quantity.

The principle is that the alignment force amplifies a globally shared direction and averages away a locally heterogeneous one. Collective escape defeats a predictive predator only when the escape vector is shared by all prey (Section 4.52); per-prey local escape, in which each prey computes its own direction away from nearby predators, only partly works because those directions do not align (Section 4.53). The same asymmetry makes a tiny informed minority an effective leader: a few percent of agents carrying one common goal vector steer the whole group, because alignment propagates their single shared direction to a naive majority (Section 4.54). Leadership, committed escape, and the spontaneous emergence of a flock heading are thus one phenomenon. When two shared directions compete, the alignment force arbitrates them: it vector-averages compatible goals into a compromise and, past a critical conflict angle, breaks symmetry into a majority consensus (Section 4.55), with the decision set by the summed directed force of each side — count times conviction — rather than by headcount (Section 4.56). Time-resolved, the consensus transition is a genuine bifurcation, exhibiting critical slowing at the boundary (Section 4.57), and the steering it enables is a low-pass control channel whose bandwidth is the inverse of the alignment response time (Section 4.59).

Two corollaries tie this axis back to the first. Because the signal is a shared direction and not an identity, rotating which agents carry it never hurts and, smeared over the whole flock, helps (Section 4.58) — the exact opposite of the slow-recoverer label, which must persist to be useful. The shared heading is the one collective quantity that benefits from turnover, precisely because it is not stored per agent. And because coherence and steerability are the same shared heading seen two ways (Section 4.59), the flock’s adversaries are unified as attacks on it: encirclement erases the heading by fragmenting the group, and a shared heading — escape or oblivious leadership — is its antidote (Sections 4.52, 4.60); panic severs the heading at its source by silencing the leaders, which leadership cannot repair because the disease disables the carriers (Section 4.61); and a saboteur injecting a false heading can deadlock the flock at pull parity but must dominate to commandeer it (Section 4.62). Across all of these the same accounting holds — the flock acts on the net shared directed force it can muster, each adversary subtracts from it differently, and disrupting that force is always cheaper than commandeering it. The two synthesis axes are therefore complementary: kinematic mixing says the flock retains no per-agent structure to exploit except a stationary internal label, and the shared-heading principle says the one thing the flock does build and act on collectively is a common direction, which is simultaneously the source of its coherence, the lever of its steering, and the target of every adversary that succeeds against it.

5.3 Alignment as a Directional Averager: How Accurate the Shared Heading Is

The shared-heading principle of the previous axis establishes that the flock acts on one common direction; a final group of experiments asks how accurate that direction is, and the answer sharpens the principle into a quantitative law. The alignment force is a directional averager, and the accuracy

of the heading it produces is governed not by the number of agents that hold an opinion but by the correlation structure of their estimates. This axis began as a self-correction. The first leadership result had attributed the minority’s power to the animal-navigation literature’s observation that a fixed number of informed individuals suffices in arbitrarily large groups, but a direct test of that scaling falsified it for this model: at a fixed number of exact-vector leaders the accuracy falls as the group grows, because the steering is the total injected force divided by all the agents, a per-capita pull set by the informed fraction (Section 4.63). The literature’s number-suffices scaling, the same section predicted, would require a many-wrongs rule in which agents hold independent noisy estimates that average out — and adding exactly that recovers it, the flock’s heading error falling as the inverse square root of the group size when every agent carries its own noisy estimate of the goal (Section 4.64).

These two results, which look opposite, are the endpoints of a single parameter: the correlation between the agents’ estimate errors (Section 4.65). With perfectly correlated errors, an exact shared vector, there is nothing to average and only the per-capita pull remains; with independent errors the average concentrates as the square root of the sample; and any intermediate correlation imposes a floor on collective accuracy that no group size can beat, equal to the per-agent error times the square root of the correlation. Confining the noisy estimates to a fixed minority does not invoke the averaging at all, because the minority’s pooled direction has a fixed accuracy while its per-capita pull dilutes with the group, so a noisy minority fails as the group grows exactly as an exact one does and sits below it — which separates, within one model, the informed-minority mechanism from the many-wrongs mechanism that the literature often runs together (Section 4.66). The practical face of the law is a robustness result: because each agent contributes a bounded unit vote, a navigating crowd averages away uncoordinated misinformation even when half its members are lost, but a coordinated falsehood of the same prevalence competes with the truth vote for vote and captures the flock at parity (Section 4.67).

This third axis closes the loop with the second. The shared-heading principle says the flock acts on its net common direction; the directional-averager law says how good that direction is — set by the correlation of the inputs, not their count — and it reproduces the same parity threshold and product law that governed the adversarial heading-fights of the second axis, now in the language of estimate accuracy. The unifying statement across both is that correlation is the only currency the alignment force recognizes: it amplifies what is shared and averages away what is independent, so the collective’s heading is built, made accurate, contested, and attacked entirely through the correlation structure of the directions its members carry. The number of members enters only insofar as it enlarges the fraction that shares a direction or sharpens the average of those that do not.

6 Discussion

The most striking result of the predator simulations is that flocking is not primarily a strategy for maximizing distance from a predator. Non-flocking agents individually maintain slightly greater separation from the predator, yet the flocking agents clearly have a more robust collective response: they remain coordinated, move in concert, and hold a consistent buffer distance. This distinction — coherence versus distance — may reflect something real about the function of biological flocking. A coherent flock can mount a coordinated escape response; scattered individuals cannot.

The increasing aspect ratio under multiple predators is interesting. When pressure arrives from multiple directions, the flock elongates rather than fragmenting. The diagnostic analysis (Section 4.4) shows this is not because predators actually approach from multiple directions — they converge

to the same point — but rather because the stronger combined repulsion from co-localized predators drives a more intense elongation response. The elongated shape is a real emergent effect, but its cause is the concentration of force at one point, not a strategic response to multi-directional threat.

The equilibrium speed result ($v_{eq} = v_0 + \alpha/\mu$) is an exact consequence of the force equations that Charbonneau does not explicitly note. It means a researcher using this model who sets $v_0 = 1$ and $\alpha = 1$ expecting agents to cruise at speed 1 will find them consistently at speed 1.1. For simulations where absolute speed matters (e.g., comparing flocking and non-flocking agents under identical predator pressure), this correction is necessary.

The phase transition result extends the original finding. Fixing compactness properly via $r_0 = \sqrt{C/(\pi N)}$ and testing a dilute regime ($C = 0.10$) shows the same behavior as the dense regime: N -independent KE/N and monotone susceptibility with no finite- η peak. The absence of a critical point is not a consequence of high compactness alone. Instead, both extremes fail for different reasons — too dense means caged oscillators, too dilute means non-interacting walkers. Section 4.19 goes further by sweeping the repulsion exponent from $n = 1.5$ to $n = 12$ — approaching the hard-core limit — and finds identical behavior at every exponent. This rules out potential softness as the cause of the crossover. The root cause is instead the non-equilibrium driving: uniform random kicks without viscous dissipation do not satisfy the fluctuation-dissipation theorem, so the system cannot equilibrate into a crystal phase. The smooth crossover is a universal property of the model class, not a potential-specific artifact, and cannot be removed without replacing the driving mechanism itself.

The long-time encirclement result is a caution against interpreting short-simulation steady states as equilibria. The 4000-step snapshots used in Sections 4.5-4.7 reported a “steady” disrupted flock at $\Phi \sim 0.72$. The 30000-step runs in Section 4.15 reveal that this value is a time-average hiding large oscillations (temporal std = 0.21 per seed). The long-time attractor of the encircled flock is not a fixed fragmented state but a persistent merge/split cycle. This has methodological implications: order-parameter measurements from short runs near an encirclement configuration may give misleading precision.

The encirclement gap experiment (Section 4.15) challenges intuitions about escape. Removing one predator from a symmetric ring is MORE damaging to flock coherence than the full ring, because the asymmetry creates a perpetually shifting CoM chase rather than a balanced stable pattern. This is counterintuitive from a naive perspective where gaps should always help. The result reflects a general feature of these agent models: agents have no global spatial awareness. They respond only to local forces, so a gap in the predator ring is invisible to agents far from the gap.

The outbreak persistence result (Section 4.16) highlights an asymmetry in recovery timescales. Kinematic damage from encirclement reverses rapidly because it has no memory beyond the agents’ current positions and velocities; once the force pattern changes, the trajectory changes. Epidemic damage has memory in the agent’s internal state (panicked vs. calm) and reverses on epidemic timescales set by the ratio of infection and recovery rates, which are independent of kinematic parameters. This is a qualitative distinction with potential biological relevance: a predator event that coincides with a near-threshold social contagion (collective alarm behavior, epidemic disease) leaves a lasting mark on the collective that outlives the predation event itself.

The adaptive encirclement result (Section 4.17) provides a constructive validation of the R_{enc}/R_g scaling (Finding 31): if the universal optimum at $R_{enc}/R_g \sim 0.5$ is real, a predator that tracks live R_g should maintain that optimum even as the flock fluctuates during the merge/split cycle, and it does. The 34% reduction in high-coherence dwell time (frac_above_0.85 from 0.56 to 0.37)

is the direct signature of this — adaptive predators suppress the recovery excursions that fixed predators permit. The modest mean-Phi effect (8%) confirms that the fixed strategy was already close to optimal, with the fixed R_{enc}/R_g sitting at 0.485 on average. Adaptation gains most at the margin, during the brief consolidation phases where the flock is most recoverable.

The predator-strategy findings place this work in a broader literature of simulated predation on flocking models. Demsar and Lebar Bajec (2014) compared attack tactics (target the center, target the nearest, target isolated individuals) in a fuzzy individual-based model and found that social flocking protects against predators targeting isolated prey. Importantly, they do not test a coordinated multi-angle encirclement strategy, which is the key contribution of Sections 4.5-4.7 here. Inada and Kawachi (2002) identified four escape-pattern categories in a two-dimensional fish-school model, including “Split and Reunion” — the pattern that Sections 4.6 and 4.7 quantify and mechanistically explain: encirclement causes the split, and the reunion timescale is ~ 10 time units after predator removal. Bartashevich et al. (2024) studied the “fountain effect” in sardines evading marlin using both agent-based models and empirical observations, finding that prey optimize individual escape angles relative to the predator’s attack direction. The fountain effect is a single-predator pattern; our encirclement result extends it to the multi-predator case where angular pressure from all sides prevents any single escape direction from being optimal, leading to directional fragmentation rather than a coherent fountain.

The most closely related existing work is Levis et al. (2020), who studied bidirectional coupling between Vicsek-like flocking and SIS epidemic dynamics, finding that endogenous clustering (infected agents altering their motion rules) reduces the epidemic threshold in a similar way to the encirclement-induced compression found here. The critical difference is that Levis et al.’s mechanism is internal: the epidemic state controls flocking, so compression and contagion are inseparable. In the present model, the compression is imposed externally by predators, making it possible to study the removal experiment (Section 4.16) — what happens when the compressor is switched off while the epidemic is still ongoing. That timescale asymmetry, and the practical implication that external pressure can seed a long-lived epidemic state without sustaining it, does not appear to be addressed in the existing literature.

The vaccination results (Sections 4.18 and 4.20) establish a broader principle beyond the specific null results. In heterogeneous networks (Barabasi-Albert scale-free, small-world), targeted vaccination of high-degree nodes reduces the epidemic threshold substantially (Pastor-Satorras and Vespignani, 2002; Cohen et al., 2003). Spatial vaccination strategies for geographic human populations have also been studied (Bhatt et al., 2022; Zhou et al., 2021), showing advantages when infection clusters spatially. The flock fails to benefit from either approach: degree-targeting is defeated by bounded degree heterogeneity ($CV = 0.68$) and kinematic reorganization that restores hub positions; spatial targeting is defeated by kinematic mixing that scrambles spatial distributions before the epidemic runs. Both null results trace to the same property of the flock — the constant spatial reorganization driven by the alignment force — which prevents any static structural feature of the contact network from remaining stable long enough to be exploited. For kinematic collectives, the epidemiological literature’s targeting strategies are inapplicable in their standard form: only a dynamic vaccination strategy applied in real time during the epidemic, tracking current agent positions, could exploit the spatial-clustering mechanism that inflates the threshold.

The later vaccination work (Sections 4.36-4.46) resolves this apparent dead end with a positive result. The degree- and spatial-targeting failures share a common cause — both target a property that kinematic mixing erases (graph position, physical position) — but they exhaust only the EXTERNAL, system-level descriptions of an agent. Heterogeneous recovery (Section 4.36) introduces a third

class of hub: an internal-state invariant, the per-agent recovery rate γ_i . Slow recoverers act as reservoirs that lower the epidemic threshold by the SPREAD of γ , not its mean, and crucially the slow label travels with the agent and cannot be scrambled by reorganization. Targeting it is therefore the first strategy in the study to beat random vaccination, and it does so across two and three dimensions, bimodal and continuous distributions, noisy estimates, and rare reservoirs (Sections 4.38-4.43). The boundary conditions are as informative as the result: the policy requires the slow CLASS to be durable on the epidemic timescale (Section 4.44 – a recovery rate that drifts faster than the outbreak self-averages away, removing both the reservoir and the need to target it), it is robust but not unique once a second heterogeneity axis is added (Section 4.45 – infectiousness targeting is also effective but not robust, because the transmission engine and the reservoir can be different populations), and at a budget matching the reservoir fraction it converts the otherwise-irreversible combined predator+contagion damage of Section 4.16 into fully reversible damage (Section 4.46). The unifying statement across the predator, contagion, and vaccination threads is that lasting damage to the flock requires a surviving reservoir, and the reservoir is the slow-recoverer class.

The three-dimensional extension (Sections 4.23-4.34) shows which results are dimensional and which are intrinsic. The equilibrium-speed law, the smooth crossover, the vaccination nulls, and the segregation mechanism all transfer to 3D, but encirclement does not: no point-predator strategy disrupts a 3D flock at all (Section 4.47 generalizes the original encirclement-specific finding to naive and transecting predators at any count or speed). The mechanism is that the 3D flock, at neighbor-matched parameters, fills its volume nearly uniformly and has no spatial perimeter to seal or interior to transect, so finite-range predators perturb a vanishing fraction while the alignment graph heals the wake. Disrupting a 3D flock requires attacking alignment per agent – which is exactly what contagion does, explaining why the contagion results transfer to 3D while the predator results do not. A recurring methodological theme runs through the 3D work and the Section 5 self-tests: three interpretive predictions (topological alignment slowing mixing, contact-graph freezing rescuing targeting, the third dimension speeding mixing) were tested and falsified rather than assumed, and a sign error in the early 3D predator code was caught and corrected. The corrected picture is cleaner than the original conjectures.

The predator-prey arms-race sequence (Sections 4.48-4.53) asks what changes when each side is given access to a global summary of the other. A predator that anticipates the flock’s mean velocity (predictive encirclement) is the first predator-side adaptation to beat fixed encirclement substantially, and the gain comes entirely from placing the ring where the flock is heading; adapting the ring radius adds nothing once placement is anticipatory, so placement is the dominant lever. The robustness of this predator intelligence contrasts sharply with the vaccination intelligence: the predator’s signal is a single global vector used for forward projection, so it tolerates zero-mean observation noise only gradually and collapses under even small delay (a stale heading is a systematic error), whereas the vaccination signal is a per-agent ranking buoyed by N-sample averaging that tolerates substantial noise and has no delay problem at all because the rate is stationary. The statistical footprint of an “intelligent” disruption signal – per-agent invariant versus fast-changing global statistic – governs its real-world robustness as much as its information content does. Finally, giving the prey the dual global signal (flee the predator centroid) defeats predictive encirclement completely, but only above a force threshold set by the alignment strength, and only because a single SHARED escape vector aligns with the flocking force; a weak escape force, or a locally-sensed per-agent escape direction (Section 4.53), competes with alignment rather than reinforcing it and helps little or not at all. The flock can act collectively only on signals that are already global or shared – the same shared-heading principle that produces coherent flocking in the first place – which is why the arms race resolves not

as a simple stronger-signal-wins contest but as a structure in which the predator’s own anticipatory massing supplies the very directional signal that committed, coordinated prey need to escape.

The leadership result (Section 4.54) is the constructive complement to this entire escape sequence and, with it, completes the shared-signal argument. The same principle that limits collective escape – that the flock amplifies only a globally shared direction – is what makes a tiny informed minority an effective navigator. A few percent of agents carrying a common goal vector steer the whole group with near-perfect accuracy and no loss of cohesion, because alignment propagates their single shared direction to a naive majority that has no goal knowledge of its own. Leadership, committed collective escape, and flock formation itself are therefore one mechanism viewed from three angles: a direction that every relevant agent holds in common is amplified into group-level motion, whereas a direction that varies agent-to-agent – the local escape field, the symmetric encircling ring – is averaged away. The flock is exquisitely steerable by consensus and nearly unsteerable by conflicting local cues, and this single asymmetry organizes results spanning leadership, predation, escape, and the spontaneous emergence of a common heading.

When two informed subgroups disagree (Section 4.55), the same alignment force that amplifies a single shared direction is revealed to also arbitrate among competing ones, and it does so by a rule that emerges from the dynamics rather than from any agent-level decision logic. For small directional conflict the flock vector-averages the two goals and travels their midpoint; past a critical angle near ninety degrees, where averaging two near-opposed directions would demand an unphysical heading, the symmetry breaks and the whole flock commits to one direction chosen, at parity, at random. It almost never splits. Under direct opposition the choice is decided by an effective majority vote among the informed agents: even a small numerical margin tips the outcome reliably toward the larger subgroup. This reproduces the central decision-making results of Couzin et al. (2005) inside the present model and shows they are not special to the spin-like or zonal models in which they were first derived – they follow from the same alignment coupling that produces flocking, escape, and single-leader steering. The compromise-to-consensus transition is, moreover, the same domination-not-blending physics that recurs whenever competing global drives meet in this study, from speed homogenization and alpha-contrast segregation to the non-monotonic escape counter: an alignment-dominated flock resolves conflicting collective signals by selecting one, not by interpolating, once those signals become mutually incompatible.

The conviction experiment (Section 4.56) makes the vote quantitative. When two opposed subgroups differ not in number but in commitment strength, the more strongly committed side wins at equal numbers, and a numerical minority overtakes a larger majority precisely as its total pull – the product of count and bias strength – crosses the majority’s. The flock is therefore not a one-agent-one-vote democracy but a weighted integrator: it sums the directed force injected by every informed agent and commits to the net winner, whether that net is assembled from many weak voices or a few strong ones. A small residual advantage attaches to numerosity itself, because more distinct informed agents seed the preferred direction at more points in the group, but to leading order it is summed force, not headcount, that decides. This is the same currency that set the escape threshold earlier in the study, where committed flight defeated the predator only once its weight exceeded the alignment strength: across leadership, conflict, and escape, the model’s collective outcomes are governed by the magnitude of directed force relative to alignment, a single quantitative principle underlying what at the behavioural level look like distinct phenomena of steering, voting, and fleeing.

The time-resolved view (Section 4.57) reveals that this weighted vote is, dynamically, a bistable decision. Watching the heading settle rather than only its endpoint shows that strong leadership is fast, accurate, and noise-robust all at once – there is no decision-quality penalty for steering,

in contrast to the coherence cost a predator pays to redirect the flock – and that the only regime of quick-but-wrong commitment is the under-led one, the temporal shadow of the threshold the informed fraction must clear to steer at all. More strikingly, the commitment time peaks at the compromise-to-consensus boundary and falls away on either side, the textbook signature of critical slowing near a bifurcation. The transition between averaging two goals and committing to one is therefore not a mere relabeling of the steady state but a genuine change in the dynamical landscape: at the boundary the averaging solution loses stability while the two consensus solutions are only beginning to form, and the flock, caught between them, takes longest to decide precisely when the decision is hardest. The collective behaviours catalogued in this study – flocking, escape, steering, and voting – are thus unified not only by the alignment force that carries shared signals but by the bistable, threshold-governed way that force resolves competition among them.

The rotation experiment (Section 4.58) settles what kind of thing a leadership signal is, and in doing so closes a question that has run through the entire study: what makes a collective-control target exploitable. The contagion thread answered it from one side – degree targeting fails for lack of durable hubs, spatial targeting fails because motion erases coverage, and slow-recoverer targeting succeeds only because the recovery rate is a durable per-agent invariant, an advantage that itself evaporates once that label drifts. Leadership answers it from the other side. Rotating which agents are informed never degrades steering and, when fast, improves it, because the flock follows a shared direction rather than particular individuals; the same total directed force delivered through a constantly changing cast steers as well as or better than a fixed one. A control that rests on a persistent per-agent label is fragile to that label changing, while a control that rests on a shared global quantity is robust to – even helped by – turnover in who supplies it. The two threads therefore meet at a single axis: the durability a signal requires is the durability of whatever the signal is attached to, an individual’s hidden state in the contagion case and a collectively held direction in the leadership case, and only the former can be undone by mixing the individuals.

Finally, the moving-goal experiment (Section 4.59) recasts leadership as a control problem and, in doing so, reunites the decision thread with the predator thread that opened the study. An informed minority does not steer the flock instantaneously; it drives a low-pass system whose bandwidth is the inverse of the response time and is tuned by the informed fraction, so the same lever that makes decisions faster also lets leaders drive sharper turns. Within that bandwidth steering is accurate, lags the goal by a growing angle, and costs no coherence; beyond it the rotating bias time-averages to nothing and, if strong, fragments the flock. That fragmentation is the key reconciliation: steering toward a fixed bearing was entirely cohesion-free, yet forcing a turn the flock cannot follow produces the lowest order parameter seen anywhere in the leadership experiments. Redirection costs coherence exactly when it outpaces what alignment can propagate – and this is the same condition under which a predator disrupts the flock. A cooperative leader turning slowly and a predator are at opposite ends of one axis: gentle, within-bandwidth redirection is free, while abrupt, beyond-bandwidth redirection, whether friendly or hostile, tears the group apart. The flock’s steerability and its coherence are therefore not independent properties but a single resource, rationed by the alignment response time, and the whole catalogue of results – flocking, escape, voting, leadership, and predation – is organized by how fast a shared directional signal can be driven through that one response.

The final experiment (Section 4.60) brings the two largest threads of the study into direct contact and resolves the tension just described. Placing an informed minority inside an actively encircled flock shows that leadership and predation are not merely analogous but literally opposed operations on the same quantity. Encirclement – the only predator strategy that breaks two-dimensional

coherence – works by destroying the flock’s shared heading, and a minority carrying any strong shared heading, even one wholly ignorant of the predators, both rebuilds the coherence the ring destroyed and drives the flock to its goal while the predators trail along behind. The predators exact a price, raising the informed fraction needed several-fold and leaving a residual loss of coherence, but they cannot stop a sufficiently led flock from going where it intends. This generalizes the earlier escape result, where it was a shared flight direction that defeated the predator: what defeats encirclement is not the content of the shared signal, flight or destination, but its mere existence. The predator wins by erasing the common direction; the leadership thread’s central object, a common direction, is its antidote. The study thus closes on a single picture in which flocking, escape, voting, leadership, and predation are all expressions of one mechanism – the alignment force propagating a shared direction – and in which the flock’s coherence, its steerability, and its vulnerability to predators are three faces of that same shared heading: build it and the group is coherent, steerable, and predator-resistant; erase it and the group fragments, drifts, and falls to the ring.

The contagion experiment (Section 4.61) completes that picture by showing the shared heading has two distinct vulnerabilities, attacked by the study’s two adversary classes in complementary ways. A predator encircling the flock attacks its coherence, fragmenting the group, and a shared heading restores it; a panic contagion attacks its steerability, intermittently silencing the leaders who carry the heading, while the coherence that needs only local alignment survives almost perfectly. The flock under heavy panic is the mirror image of the flock under encirclement: coherent but rudderless rather than fragmented but aimable. And the two failure modes differ in repairability for a deep reason. Coherence can be rebuilt by the shared heading because the heading is external to the kinematics it organizes, but steerability cannot be rebuilt by the shared heading when the contagion is precisely what disables the heading’s carriers. This is the structural explanation, arrived at from the leadership side, for why contagion has been the most durable stressor throughout the study: kinematic damage heals because the signal that heals it is intact, whereas an attack on the signal’s carriers cannot be undone by the signal. The same product law that governs voting and steering governs the collapse, with the panic fraction entering as a multiplicative tax on the active leadership, so the whole catalogue – coherence, steering, voting, escape, predation, and contagion – resolves into one accounting of how much shared directed force the flock can muster and what each adversary does to reduce it.

The adversarial experiment (Section 4.62) turns the voting law into a statement about attack and defense and exposes an asymmetry between the two things an adversary might want. Pitting saboteurs broadcasting a false goal against legitimate leaders, the flock is merely deadlocked – denied its goal – once the saboteur pull matches the leaders’, but it is actually captured, driven to the adversary’s chosen destination, only when the saboteur pull roughly doubles the leaders’. Denial is cheap and capture is dear, separated by a band in which the flock simply wanders. The zero-crossing at pull parity is the same product law that governs cooperative voting, so the contest is decided by the same summed-directed-force accounting; the new content is the gap between the denial and capture thresholds. The lesson generalizes a theme that recurs across the whole study: it is far easier to disrupt the flock’s shared heading than to commandeer it. A predator erases the heading and fragments the group, a contagion silences the leaders who hold it, and a saboteur at parity cancels it into deadlock – all cheap – whereas turning the flock into a tool that goes where the adversary chooses demands outright dominance of the shared signal. The flock’s collective intelligence is thus robust in a specific and limited sense: its consensus is hard to hijack but easy to jam, and the same alignment coupling that makes a tiny honest minority an effective leader makes a comparably tiny dishonest one an effective spoiler.

7 Conclusions

This study produced fifty-seven main results (selecting the most general across 92 findings):

1. **Equilibrium speed:** The cruise speed of an aligned flock is $v_{eq} = v_0 + \alpha/\mu$, exactly. This is a direct consequence of the force equations and must be accounted for when comparing simulations at different parameter values.
2. **Phase transition is a crossover:** The solid-to-fluid transition in the repulsion-only system is a smooth crossover at all tested compactness values ($C = 0.10, 0.15, 0.20, 0.30, 0.40, 0.50, 0.60, 0.78$), not a true phase transition. Susceptibility never peaks at a finite noise amplitude; KE/N is N -independent in every regime tested. The absence of a critical point is a universal feature of this model’s soft repulsion, not an artifact of any density regime.
3. **Flock coherence under predation:** Flocking prey maintain near-perfect velocity alignment under sustained predator pressure while non-flocking prey scatter. Evasion distance saturates at a minimum buffer value regardless of predator aggression.
4. **Predator co-localization and elongation:** Multiple predators using the “chase CoM” rule converge to the same location (measured separation ~ 0.001), producing combined repulsion that paradoxically increases evasion distance and drives flock elongation. The elongation is real but its cause is force concentration, not strategic shape adaptation.
5. **Predator-strategy hierarchy:** Three predator strategies form a clear hierarchy of effectiveness against the flock. Naive predators co-localize and fail ($\Phi > 0.97$ for $n_{pred} = 4$). Coordinated predators with mutual repulsion spread out but still fail ($\Phi > 0.92$ at $n_{pred} = 10$). Encirclement — explicit assignment of fixed compass angles to each predator — is the only strategy in this model capable of substantial disruption ($\Phi = 0.77$ at $n_{pred} = 6$, $\Phi \sim 0.67$ floor at $n_{pred} = 10$).
6. **Encirclement causes transient division:** Under encirclement, the flock divides into a small number of internally coherent sub-flocks heading in different directions (flock DIVISION, not flock DISSOLUTION). When predators are removed, all sub-flocks reunite within ~ 10 time units. The damage is purely kinematic and fully reversible.
7. **Encirclement is size-invariant when scaled to flock geometry:** The encirclement disruption floor depends on R_{enc}/R_g , not on absolute R_{enc} or absolute N . The optimal disruption at $R_{enc}/R_g \sim 0.5$ is universal across tested flock sizes. The apparent N -dependence reported in earlier experiments was an R_{enc}/R_g mismatch.
8. **Long-time encirclement is intermittent:** Over 300 time units, the encircled flock does not settle into a steady fragmented state. Φ oscillates between ~ 0.4 and ~ 0.95 as sub-flocks repeatedly merge and re-split, driven by a self-reinforcing cycle between coherence and multi-directional predator pressure.
9. **Internal stressors require contagion to matter:** Statically labeled panicked agents do not propagate disruption — the alignment force keeps calm neighbors coherent even at 20% panic fraction. Once a contact-mediated SIS contagion mechanism is added, the epidemic threshold appears at $\beta/\gamma \sim 1$. Below threshold the flock contains the outbreak; above it an endemic panicked state suppresses coherence proportionally to β/γ . The flock is not a well-mixed population at the contagion length scale: the spatial herd-immunity threshold (~ 0.46) is more than twice the mean-field prediction (~ 0.20), driven by clustering of panicked

sub-groups.

10. **Encirclement shifts the epidemic threshold by ~4%:** Compressing the flock raises local contact count ~3x and lowers the effective epidemic threshold from $\beta_c = 1.93$ to 1.85. The shift is real but modest: compression creates redundant contacts within already-panicked sub-clusters rather than fresh ones. Encirclement is a near-critical amplifier, not a general one.
11. **Epidemic damage outlasts kinematic damage by an order of magnitude:** When encirclement drives a sub-threshold contagion into a transient endemic state, removing the predators reverses the kinematic fragmentation (~10 time units, as in pure encirclement) but leaves the SIS epidemic intact. After 50 time units without predators, panic fraction remains at 0.41 (barely below the encirclement-elevated peak of 0.45) and Φ is only 0.27. The residual epidemic suppresses alignment for hundreds of time units. Kinematic damage is reversible; epidemic damage outlasts the event that caused it.
12. **Adaptive encirclement outperforms fixed radius:** Predators that continuously track live flock R_g and set $R_{enc} = 0.5 * R_g$ maintain the universal optimum across the merge/split cycle (Finding 32), reducing mean coherence from $\Phi = 0.778$ to 0.713 and cutting the fraction of time spent in a highly coherent state ($\Phi > 0.85$) from 56% to 37% compared to a fixed-radius strategy. The effect size is modest (8% in mean Φ) because the fixed radius was already near-optimal at initialization, but the adaptive strategy eliminates the off-optimum excursions that arise from flock geometry fluctuations during intermittent dynamics.
13. **All vaccination targeting strategies fail in a kinematic flock:** Neither degree-targeted vaccination (immunizing highest-contact-degree agents first, Finding 36) nor spatially-targeted vaccination (farthest-point maxmin sampling to maximize spatial coverage, Finding 37) outperforms random vaccination. Both strategies require $p_{immune} \sim 0.46$, identical to random. The mechanism is the flock's kinematic reconfigurability: kinematic reorganization restores hub positions after high-degree agents are immunized (defeating degree-targeting), and kinematic mixing during the warmup phase scrambles the spatial distribution of immune agents before the epidemic begins (defeating spatial targeting). No static agent-selection rule can maintain its structural advantage in a continuously-mixing flocking collective; the 2x herd-immunity threshold inflation is a systemic property that cannot be efficiently exploited.
14. **The smooth crossover is a consequence of non-equilibrium driving, not repulsion softness:** Sweeping the repulsion exponent from $n = 1.5$ (current) to $n = 12$ (near hard-core) produces virtually identical finite-size scaling behavior in all cases — χ_{peak} at the top of the noise sweep, N -independent KE/N curves, no diverging susceptibility. In equilibrium statistical mechanics, 2D hard discs at the tested compactness ($C = 0.40$) do exhibit a phase transition (KTHNY melting). The absence of any transition at $n = 12$ demonstrates that the model's non-thermal driving — uniform random kicks without viscous dissipation — prevents the Boltzmann equilibration required for cooperative melting.
15. **Langevin dynamics thermalize correctly but KE/N cannot detect structural melting:** A Langevin thermostat (viscous damping + FDT-satisfying thermal noise) recovers equilibrium thermalization — $KE/N = kT$ to within 1% (equipartition) — confirming the non-equilibrium diagnosis of Finding 38. However, the susceptibility $\chi = N * \text{Var}(KE/N)$ still peaks at the top of the kT sweep and shows no N -dependent shift, because KTHNY melting is a positional-order transition invisible to kinetic energy fluctuations.

16. **Three-dimensional extension confirms universality of $v_{eq} = v_0 + \alpha/\mu$:** Extending the model to a periodic 3D unit cube with neighbor-count-matched parameters ($r_f = 0.20$ for ~ 12 expected neighbors) reproduces coherent flocking across the tested noise range. The equilibrium speed $v_{eq} = v_0 + \alpha/\mu = 1.100$ holds exactly (measured: 1.1000 at ramp = 0), confirming the analytical result is dimensionality-independent — a consequence of the 1D force balance along the heading direction. The 3D flock is slightly less noise-resistant than 2D at the same ramp ($\Phi = 0.84$ vs ~ 0.98 at ramp = 10) because 3D velocity perturbations have one additional degree of freedom. Flocking forms cleanly in 3D and the core analytical result transfers exactly.
17. **Hexatic order parameter confirms soft repulsion cannot crystallize:** Measuring the hexatic order parameter $|\psi_6|$ directly (the correct diagnostic for KTHNY melting) reveals that $n = 1.5$ soft repulsion is incapable of forming a hexagonal solid at any accessible temperature. $|\psi_6| \approx 0.4$ across the entire kT range (0.001 to 5.0) for both compactness values ($C = 0.60$ and $C = 0.70$), with no solid-phase value near 1.0 and no fluid-phase collapse to 0. χ_{ψ_6} peaks at the bottom of the kT sweep with no N -dependent growth. The mechanism is that the smooth $n = 1.5$ contact-avoidance potential allows agents to overlap at any finite kT , preventing the rigid hexagonal lattice required for KTHNY melting. Demonstrating the KTHNY transition in this model family requires a near-hard-core Langevin simulation ($n \geq 12$) or a true hard-disc Monte Carlo framework — a prediction tested and corrected in result 23 below.
18. **The 3D noise-driven crossover is smooth and qualitatively identical to 2D:** An extended noise sweep (ramp = 0.5-30) at three 3D system sizes ($N = 100, 200, 350$) shows that the χ_{peak} location increases with N (ramp_c = 15, 20, 25 respectively) rather than converging to a finite critical noise value. This is the signature of a smooth crossover, not a phase transition, matching the 2D result at matched neighbor count. The 3D crossover occurs at substantially lower noise (ramp ~ 15 -25) than the 2D crossover (ramp > 30) because the third velocity component provides an additional channel for noise to disrupt alignment. The non-equilibrium smooth-crossover mechanism identified in Findings 38 and 39 is dimensionality-independent.
19. **Encirclement is strictly a 2D strategy:** With correctly repulsive predators, 3D encirclement does not disrupt the flock at all — the order parameter stays at $\Phi \sim 1.000$ — whereas 2D encirclement at the identical $R_{enc} = 0.15$ drives Φ to 0.75 ($n_{pred} = 6$) and 0.73 ($n_{pred} = 10$). The mechanism is geometric: 2D encirclement works by sealing the flock’s 1D perimeter with point predators, which a modest number of predators can do; sealing a closed 2D surface around a 3D flock is impossible for the same predator count, and the prey leave through any uncovered solid-angle direction.
20. **No geometric variant of encirclement works in 3D:** The 2D-to-3D failure of encirclement cannot be repaired by any geometric adjustment. The 3D order parameter stays at $\Phi \sim 1.000$ across predator counts from 1 to 50, across encirclement radii from $R_{enc}/R_g = 0.12$ to 0.82, under adaptive radius tracking, and under both spherical and planar predator arrangements. The radius of gyration and the mean alignment-neighbor count are unmoved throughout — the predators neither compress nor densify the flock. (An earlier version of this study reported a partial 3D disruption and a “compression mechanism”; those were traced to an inverted predator-force sign and corrected — see the note at Section 4.25.)
21. **The vaccination null results extend to 3D:** Random, spatial, and degree-targeted

vaccination remain statistically indistinguishable in a 3D flock, with the contact-degree distribution if anything less heterogeneous than in 2D (CV ~ 0.59). Kinematic mixing — and the absence of exploitable degree structure — are both dimension-independent.

22. **The synthesis survives its own self-test, with one correction:** Section 5’s pre-registered prediction — that a topological (k-nearest-neighbor) alignment force would slow kinematic mixing and let targeted vaccination recover — was tested and falsified: k-NN alignment leaves the neighbor-graph turnover unchanged because the contact graph rewires through physical agent motion, not through the alignment rule’s neighbor-selection topology. A second experiment froze the contact graph directly (by driving the flock into its low-noise solid regime, reducing the mixing rate thirtyfold) and found that degree-targeted vaccination still does not beat random. This forces a correction to the synthesis: the degree-targeting null result is structural — the flock contact network simply has no hubs — and is independent of mixing, whereas the spatial vaccination null result genuinely is kinematic. The two null results are distinct mechanisms, not one.
23. **Hard repulsion does not crystallize; the phase-transition thread closes negatively:** Result 17 predicted that a near-hard-core Langevin simulation ($n \geq 12$) would exhibit the KTHNY melting transition the soft $n = 1.5$ potential could not. A direct test with exponents $n = 12$ and $n = 24$, at dense packings $C = 0.70$ and 0.85 , falsifies that prediction: the hexatic order parameter $|\psi_6|$ stays flat near 0.34 - 0.39 across the entire temperature range, identical to the $n = 1.5$ result, and the hexatic susceptibility decreases rather than grows with N . The reason is a property of the force form: the repulsion strength scales as base_r^n with $\text{base_r} = 1 - d/r_b$ in $[0, 1]$, so a higher exponent makes the force negligible except at vanishingly small separation — it shrinks the effective interaction range rather than hardening the core. The Charbonneau model cannot crystallize at any exponent of this repulsion; exhibiting KTHNY melting would require a genuinely different potential (a true inverse-power-law or hard-disc form). The solid-to-fluid behavior is a smooth crossover at every exponent and temperature tested.
24. **Self-organized segregation is partially dimension-dependent:** Two populations differing in alignment strength α segregate in 3D as they do in 2D — local purity rises above the well-mixed value of 0.5 — but the effect is diluted in three dimensions. At moderate contrast the 2D and 3D purities are identical (~ 0.55); at high contrast they diverge, with 2D purity climbing to 0.63 - 0.73 while 3D purity plateaus near 0.55 and breaks upward only when the passive population has zero alignment (and the flock loses global coherence). The cause is geometric: in 3D an agent’s neighbors occupy a ball rather than a disc, giving a partially-aligned agent more independent directions along which to be surrounded by the other type, so instantaneous local purity is diluted.
25. **The third dimension mixes slower, not faster:** A direct measurement of the contact-graph turnover rate, at matched mean contact degree, shows the 3D flock’s contact graph rewires at only ~ 0.56 of the 2D rate — consistently, across a thirtyfold range of noise amplitudes. This falsifies the tempting interpretation that the 3D results above (vaccination nulls, diluted segregation) arise because the extra dimension speeds up kinematic mixing. It does not: matching the contact degree forces a larger 3D contact radius, which makes the neighbor set slower to turn over. The 3D flock is hard to disrupt, target, and sort for three separate reasons — the escape dimension, structural degree-homogeneity, and neighborhood geometry — none of which is faster mixing. This is one of several cases (see also results 22 and 23) where a pre-registered or interpretive prediction was tested and corrected rather than assumed.

26. **Prey fatigue does not make encirclement damage irreversible:** Adding a per-agent fatigue variable that accumulates under predator pressure and impairs either cruise speed or alignment strength does not change the reversibility established in result 6: the flock recovers to $\Phi \sim 1.0$ within the recovery window at every fatigue rate, in both modes. Coherence recovers even while substantial fatigue persists, because post-attack fatigue decays uniformly and a homogeneously fatigued flock still aligns. Fatigue deepens the disruption *during* the attack only when it impairs alignment, not speed — the dynamical echo of result 24 on segregation (an alpha contrast segregates the flock, a v0 contrast does not). Contagion remains the only stressor studied that inflicts lasting damage, because it alone writes a heterogeneous internal label that uniform recovery cannot erase.

27. **Heterogeneous recovery rates lower the SIS epidemic threshold:** Drawing per-agent γ_i from a bimodal distribution $\{1 - \text{spread}, 1 + \text{spread}\}$ at fixed arithmetic mean 1.0, the threshold β_c at which the panic fraction crosses 0.15 falls from 0.385 (homogeneous) to 0.318 (spread 0.5) to 0.155 (spread 0.8), and at spread 0.95 the flock is endemic at every contagion rate tested. Panic localizes on the slow recoverers — the slow-to-fast panic-fraction ratio reaches 1.97 at extreme spread — so the outbreak is carried by an internal-state reservoir rather than by topological hubs. The effect is a near-threshold phenomenon: at high β all spreads converge to the same saturated endemic state. The result identifies a different kind of vaccination target than the topological one ruled out by results 13 and 21: the high-value agents to protect are the agents whose internal dynamics make them reservoirs, not the (absent) high-degree hubs. Read together with result 26, it completes the heterogeneity story — heterogeneity in an internal *state* that the dynamics homogenize is a transient amplifier, heterogeneity in an internal *rate* that the dynamics cannot erase is a permanent threshold shifter.

28. **Heterogeneous infectiousness does NOT shift the SIS threshold:** The dual experiment to result 27 — per-agent transmission rate β_i drawn from a bimodal distribution at fixed arithmetic mean, with γ homogeneous — leaves β_c flat at 0.434-0.440 across homog, mild, strong, and extreme spread conditions. Super-spreaders dominate transmission **ATTRIBUTION** (the high-beta half sources 74-97% of all calm-to-panic events) but not endemic **LOAD** (panic fractions on super and normal agents are statistically equal). Because every panicked agent recovers at the same γ , inflow skew does not translate to stock skew. Read with result 27, this establishes that source-side and sink-side heterogeneity are asymmetric in SIS on this flock: slow recoverers are reservoirs (threshold shifters), super-spreaders are merely messengers (event-level only). The vaccination prescription sharpened by result 27 — target internal-state hubs — specifically means recovery-rate hubs, not transmission-rate hubs.

29. **Targeting slow recoverers is the first vaccination strategy in this study that beats random:** In the heterogeneous-recovery regime of result 27 (bimodal $\gamma \{0.2, 1.8\}$, β just above threshold), immunising the slow half of the population first achieves $f_{ss} = 0.115$ at $p_{immune} = 0.20$ (vs random's 0.233), 0.027 at $p_{immune} = 0.30$ (vs 0.189), and full eradication at $p_{immune} = 0.40$ (vs 0.095). The effective herd-immunity threshold falls from $p_c \sim 0.50$ (random) to $p_c \sim 0.30$ (slow-targeted). Fast-targeting is strictly worse than random at every p_{immune} . Degree-targeting in this heterogeneous regime shows a faint advantage ($\sim 20\%$ at the edge of seed noise), plausibly chance overlap of high-degree and slow agents and not enough to upgrade the result-13 verdict on degree. The slow-targeting advantage resolves the targeting puzzle posed by results 13, 21, 22, and 23: a third target class exists — internal-state hubs, specifically the recovery-rate hubs — and the dynamics that

eliminate the other target classes (kinematic mixing, contact-graph rewiring) leave γ_i invariant by construction. The agent who recovers slowly today recovers slowly tomorrow, and combined with the result-27 reservoir mechanism that per-agent invariance converts an internal label into actionable policy.

30. **The slow-targeting advantage transfers to 3D and to continuous distributions:** Result 29’s slow-targeting policy was replicated in 3D (bimodal gamma at the Finding 54 strong analog, $N = 350$ torus): at $p_{\text{immune}} = 0.50$ slow-targeting eradicates the epidemic ($f_{\text{ss}} = 0.000$) while random leaves $f_{\text{ss}} = 0.282$; at $p_{\text{immune}} = 0.40$ slow is 0.223 vs random’s 0.382. Spatial targeting remains a clean null vs random in 3D (gap below 0.01 at every p_{immune}). The 3D absolute advantage is smaller than 2D because the 3D contact graph is more homogeneous (result 25) and 3D mixes more slowly at matched degree, both of which dilute the reservoir mechanism without eliminating it. The bimodal class structure is also not required: replacing the bimodal gamma with a lognormal of width $\sigma_{\text{log}} = 0.6\text{-}0.8$ reproduces the 2D advantage (slow eradicates at $p_{\text{immune}} = 0.20\text{-}0.30$ while random remains endemic), establishing that the policy “vaccinate the bottom $X\%$ by observed γ_i ” works for any plausibly heterogeneous distribution. The bimodal cases of results 27-29 were an analytical simplification, not a precondition. Across all four canonical targeting strategies (degree, spatial, random, slow) and both dimensions, only slow-targeting beats random; that result is now reproduced in 3D and under continuous distributions.
31. **Slow-targeting tolerates noisy gamma estimates and scales with reservoir size:** Two further robustness tests close the policy. Replacing the exact γ_i with a noisy observation $\hat{\gamma}_i = \gamma_i + N(0, \sigma_{\text{obs}})$ leaves the policy essentially unchanged for σ_{obs} up to about half the slow/fast separation ($f_{\text{ss}} = 0.024$ at $\sigma_{\text{obs}} = 0.8$ vs random’s 0.164), and the policy degrades gracefully to random selection only in the totally uninformative limit. Reducing the reservoir fraction from 50% to 5% leaves the slow-targeting eradication property intact at $p_{\text{immune}} = f_{\text{slow}}$ – the required vaccination budget for eradication is exactly the reservoir fraction, regardless of how small it is. Combined with results 29-30 the policy is now established as robust across every variation tested (2D and 3D, bimodal and continuous gamma, perfect and noisy observations, large and small reservoirs) and scales linearly with the size of the problem – a practical, complete vaccination policy and the only positive targeting result in the study.
32. **Slow-targeting requires a durable recovery-rate label:** The slow-targeting advantage depends on γ_i being a fixed trait, not a fluctuating state. When per-agent gamma decorrelates at rate r_{drift} , mild drift (~ 0.1 per time unit) erases the targeting advantage as the one-shot vaccine is wasted, and fast drift (≥ 1 per time unit) eradicates the epidemic outright for every strategy by time-averaging gamma to its mean and restoring the homogeneous (subcritical) threshold. The policy is valid exactly as long as the slow class is durable on the epidemic timescale ($\sim 1/\gamma_{\text{slow}}$); this sharpens the heterogeneity dichotomy of results 28-31 – a rate the dynamics homogenize quickly is neither a threshold-shifter nor a target.
33. **Under combined infectiousness and recovery heterogeneity, reservoir-targeting is the robust choice:** Layering per-agent infectiousness β_i on the recovery heterogeneity, super-spreader targeting is the strongest single strategy when infectiousness is uncorrelated with recovery (removing 20% of agents deletes $\sim 60\%$ of transmission capacity – so β -targeting is effective for removal even though β -spread does not move the threshold). But only slow-recoverer targeting is robust: in the adversarial case where super-spreaders are fast recoverers, only removing the slow reservoir eradicates, because the reservoir and the

transmission engine are different populations. The “target gamma, not beta” guidance is a robustness statement.

34. **Reservoir-targeted vaccination reverses the predator+contagion damage asymmetry:** The persistence of contagion after predator removal (result 11’s “worst stressor”) holds only while the reservoir survives. With heterogeneous recovery, vaccinating the slow class at a budget matching the reservoir fraction ($p_{\text{immune}} = f_{\text{slow}}$) eradicates the epidemic both during and after the attack and lets the flock reunite to coherence near unity – making the combined predator + contagion damage fully reversible. Below that budget the epidemic persists for every strategy. The predator, contagion, and vaccination threads thus unify: in this flock, lasting damage requires a surviving reservoir, and the reservoir is the slow-recoverer class.
35. **3D flocks are robust to all point-predator strategies, not only to sealing:** Adding a transect predator (fast CoM-chaser that punches through the dense core and oscillates back) does not improve on encirclement – the order parameter sits at 1.000 at every predator count (3-10) and over a forty-fold transect-speed sweep (0.05 to 0.80, ending at 40x prey speed). The deeper reason is sharper than the Section 4.25 framing: at these parameters the 3D flock fills the unit cube nearly uniformly ($R_g = 0.43$ of the ~ 0.5 maximum), with globally aligned velocities. It has no perimeter to seal and no interior to transect, and any handful of predators with finite repulsion range perturbs a vanishing fraction of the prey at any instant while the rest of the alignment graph heals the wake. Disrupting a 3D flock requires a per-agent attack on alignment itself (which is what contagion provides) rather than a geometric placement of point predators.
36. **Predictive encirclement is the first predator-side adaptation to substantially beat F14:** Letting predators target $\text{CoM} + \text{lead_time} * v_{\text{mean}}$ (anticipating the flock’s heading direction) gives a non-monotonic $\Phi(\text{lead_time})$ with a clear minimum at $\text{lead_time} = 2 \text{ tu}$, where Φ drops to 0.530 – well below the F14 reproduction here (0.825) and below the F35 adaptive- R_{enc} floor (0.713). The optimum sits near $\text{lead_time} \sim R_{\text{enc}} / v_{\text{mean}}$, which makes the lead distance match the encirclement radius and places the ring exactly where the flock is headed. At larger lead times the predators overshoot beyond reach and the flock turns away. The result inverts Finding 33’s asymmetry: the flock cannot detect global escape directions, but predators CAN detect the flock’s mean velocity, so predator intelligence is informationally easier than prey escape intelligence in this model. Adapting POSITION (this result) is independent of adapting RADIUS (Finding 35); the two levers could be combined.
37. **The two predator-side adaptations do not compose: placement dominates radius:** The natural follow-up to results 12 and 36 – “predictive AND adaptive” predators that lead by v_{mean} AND scale R_{enc} with live R_g – gives $\Phi = 0.535$, statistically indistinguishable from predictive-fixed alone (0.530). The composition prediction is falsified. Under encirclement the compressed flock has $R_g \sim 0.05\text{-}0.10$, so adaptive $R_{\text{enc}} = 0.5 * R_g$ shrinks the ring to ~ 0.03 while the predictive lead distance is 0.24; six predators at that radius placed 0.24 ahead of CoM degenerate into a near-point predator in the heading direction, yet are as disruptive as a proper predictive ring at $R_{\text{enc}} = 0.15$. Once the flock’s heading is blocked by predators at the right distance, the angular spread of the configuration becomes secondary; a dense block in front is as effective as a spread ring around the lead point. Placement is the dominant predator-side lever, radius tuning is at best redundant once placement is anticipatory, and the encirclement geometry’s identity dissolves under predictive placement into one-sided interception. The remaining open questions are predator-side informational (noisy v_{mean} ,

delayed updates, partial visibility) rather than geometric.

38. **Predictive encirclement is less noise-tolerant than slow-targeting; a statistical contrast between per-agent and global-summary intelligence:** Replacing the true v_mean with a noisy estimate $v_mean_hat = v_mean + N(0, \sigma_obs)$ degrades the F66 placement monotonically and gracefully – Φ rises from 0.530 (perfect) through 0.709 (noise equal to signal magnitude) to 0.804 in the high-noise limit, retaining a quarter of the original advantage at $\sigma_obs = 50\%$ of $|v_mean|$ and about 40% at $\sigma_obs = 100\%$. Crucially, the degradation is GRADED from $\sigma = 0$ with no plateau – unlike Finding 31’s slow-targeting, which is identical to perfect knowledge up to noise equal to half the slow/fast separation. The contrast is statistical: slow-targeting’s signal is a per-agent ranking that benefits from N-sample averaging (the order survives noise as long as noise is smaller than the separation), while predictive encirclement’s signal is a single global vector per timestep with no averaging buffer. Per-agent intelligence is intrinsically robust to observation noise; global-summary intelligence requires temporal filtering.
39. **Predictive encirclement is far more sensitive to delay than to noise:** Acting on a delayed v_mean (stale heading) destroys the F66 advantage much faster than noise – a 0.25 time-unit lag (one eighth of the lead time) loses ~83% of the advantage, and by a 1 time-unit lag the advantage is gone and slightly negative (the stale lead partially un-blocks the current escape direction). The asymmetry is mechanistic: noise is a zero-mean error that averages out, while delay is a systematic directional bias on a quantity used for forward projection, and the global heading decorrelates on sub-time-unit timescales under disruption. This is the dual of result 31 – the per-agent recovery rate is both noise-robust and delay-free (stationary), whereas the predator’s global heading is both noise- and delay-sensitive; the robustness of an intelligent disruption strategy depends on whether its key signal is a stationary per-agent invariant or a fast-changing global statistic. Closes the predator-learning thread.
40. **Collective escape intelligence counters predictive encirclement above a threshold, and the predator’s own intelligence creates the prey’s opening:** Giving the prey the dual global signal – flee the predator centroid with weight w_escape – restores the order parameter to 1.000 at $w_escape \geq 2$ (twice the alignment strength), because a unified escape direction reinforces alignment and the flock outruns the trap. But the response is non-monotonic: at $w_escape = 0.25$ the order parameter drops to 0.275, below the no-escape value, because a weak escape force competes with alignment and fragments the flock without achieving escape. The benefit threshold is w_escape comparable to the alignment strength. The counter works only because predictive encirclement masses predators ahead of the flock, displacing their centroid from the CoM and defining the escape direction – symmetric encirclement provides no such signal. The predator’s forward projection is thus self-defeating against committed escape-intelligent prey, and the arms race has a rock-paper-scissors structure rather than a simple winner. The non-monotonicity echoes results 6 and 24: competing global drives in an alignment-dominated flock resolve by domination, not blending. Closes the predator-prey arms-race arc.
41. **Collective escape requires a globally shared direction, not just escape information:** Replacing the global predator centroid of result 40 with realistic per-agent local sensing (each prey flees predators within a radius) only partially counters predictive encirclement – the order parameter peaks at ~0.83 near a sensing radius matching the ring scale and never reaches the full escape of 1.0, even at global sensing range. The reason is structural: a single shared escape vector aligns with the flocking force and produces a unified flee, whereas locally

computed escape directions point different ways for different prey, fail to align, and compete with alignment. The dependence on sensing radius is non-monotonic – too large a radius makes each prey sense the surrounding ring symmetrically and the escape force cancels (the result-15 “no net escape direction” problem at the individual level). This is an honest caveat on result 40: full escape is partly an artifact of a globally shared signal; with local perception the counter is real but modest. The flock can act collectively only on signals that are already global or shared – the same shared-heading principle that makes flocking coherent.

42. **A small informed minority steers the whole flock without any loss of coherence:** Giving a fraction of the agents a single shared preferred direction, felt as an extra force, lets as few as five percent of the group — about eighteen agents in a flock of 350 — steer the naive majority onto that heading, with the order parameter never dropping below 0.995. Steering here is cohesion-free, the opposite of encirclement, which steers only by fragmenting the flock. The mechanism is the shared-heading principle of result 41 read constructively: because every leader carries one common vector, the alignment force propagates it to the followers, so leadership, committed escape, and the spontaneous emergence of a flock heading are one phenomenon. Power comes from agreement, not numbers — eighteen aligned leaders outweigh a hundred and seventy-five prey each sensing a different local escape direction. Opens the collective-decision thread.
43. **When two shared directions compete, the alignment force averages them when it can and votes when it cannot:** Two informed subgroups pulling in different directions produce a compromise heading at the vector midpoint while the conflict angle stays below about 90 degrees, then switch abruptly to a symmetry-broken consensus past about 120 degrees, where the flock commits to one direction or the other at random rather than averaging near-opposed goals. The group never splits — the order parameter stays above 0.95 throughout and one direction is always chosen — and under direct opposition the larger informed subgroup wins, a slight majority deciding the outcome. This is the same domination-not-blending physics seen whenever two global drives compete in an alignment-dominated flock (results 6, 24, 40).
44. **The decision is set by the product of numbers and conviction, not by headcount:** When two opposed subgroups differ in commitment as well as size, the winner is the one with the larger summed directed force — count times per-agent strength. Ten strongly committed agents outvote twenty-six weakly committed ones once their pull crosses parity, with only a mild residual bonus for numbers. The flock therefore implements an emergent weighted-majority rule in which each vote is scaled by its commitment, the same force-versus-alignment threshold that governs collective escape (result 40).
45. **Time-resolved, the consensus transition is a genuine bifurcation with critical slowing:** Stronger leadership is both faster and more accurate — there is no speed-accuracy tradeoff — and the time the flock takes to commit to a heading peaks sharply at the 90-degree compromise-to-consensus boundary, falling away on both sides. This is the dynamical signature of a bistable system dithering longest exactly where one stable solution gives way to two, confirming that the compromise-to-consensus switch of result 43 is a true bifurcation rather than a gradual crossover.
46. **Leadership is a signal, not an identity, and so it benefits from turnover:** Re-drawing which agents are informed every few time units never hurts accuracy, and the faster the rotation the more it helps and the less it varies between runs, because fast rotation smears the same total pull uniformly across the flock. This is the exact opposite of the slow-recoverer

label of result 32, which must persist per agent to be useful: contagion targeting needs a durable identity, whereas leadership needs only a durable shared direction. It also nuances the vaccination-targeting null (result 13) — distributing a signal helps when injecting a directional force but is irrelevant when removing nodes.

47. **Steering is a low-pass control channel, and coherence and steerability are one resource:** When the goal rotates, the flock tracks it only below a critical angular speed set by the inverse of the alignment response time, which itself scales with the informed fraction; above that speed the time-varying bias averages to nothing and is ignored. Forcing too fast a turn fragments the flock — fixed-goal steering is cohesion-free, but over-steering costs coherence — so the same alignment response time that sets the steering bandwidth also bounds the coherence the flock can hold, the same tension a predator exploits.
48. **Leadership and the flock’s three adversaries are unified as a contest over the shared heading, and denial is always cheaper than capture:** An informed minority restores both the heading and the coherence that encirclement destroys, steering the flock through the ring while the leadership threshold rises four- to eightfold under predator pressure but the antidote still holds. Panic severs the heading at its source by silencing the very leaders that would repair it, so leadership cannot rescue a panicked flock even though its coherence survives — which is why contagion is the most durable stressor. And a saboteur injecting a false heading deadlocks the flock at pull parity but must reach roughly twice the true leaders’ pull to commandeer it to a false goal. Across all three, disrupting a shared heading is structurally cheaper than capturing it.
49. **Steering accuracy is set by the informed fraction, and noisy private estimates recover the wisdom of crowds (a self-test):** A direct test of the animal-navigation claim that a fixed number of leaders suffices in any size group falsified it for this model — at fixed leader number the accuracy falls as the group grows, because steering is a per-capita pull set by the fraction, not the count. The same test predicted that the literature’s number-suffices scaling would need a many-wrongs rule, and supplying it — every agent carrying its own independent noisy estimate of the goal — recovers exactly that: the heading error falls as the inverse square root of the group size, accuracy rising with N rather than falling. The directional average is thus per-capita for an exact shared signal and $\sqrt{N}$ -amplifying for independent noise.
50. **Correlation is the only currency the alignment force recognizes:** The fixed-number and many-wrongs results are the two ends of a single axis, the correlation of the agents’ estimate errors. Any nonzero correlation imposes a floor on collective accuracy that no group size can beat, equal to the per-agent error times the square root of the correlation — so ten percent correlation collapses an arbitrarily good wisdom of crowds to a hard ceiling — while a noisy minority never invokes the averaging at all and fails as the group grows just as an exact one does. The practical face is a robustness law: a navigating crowd averages away uncoordinated misinformation even when half its members are lost, but is captured at parity by a coordinated falsehood of equal prevalence. The error’s correlation, not its amount, sets the damage, reproducing the same parity threshold and product law that governed the adversarial heading-fights of result 48.
51. **A noisy crowd tracks a moving goal at the same bandwidth as a sharp leader (a self-test):** Spatial averaging — the wisdom of crowds — and temporal bandwidth — the steering response time — are independent resources. With the goal rotating at a fixed speed

and each agent holding a fixed private offset, the tracking-versus-speed curves for a sharp leader and a noisy crowd overlay, falsifying a pre-registered prediction that crowd noise would reduce the bandwidth as $\exp(-\sigma^2/2)$: the static offsets rotate with the goal and add no lag, and the reduced bias magnitude stays above the steering threshold. Noise is free below the wisdom-of-crowds ceiling and catastrophic above it, binary rather than graded, showing the directional average is recomputed each timestep rather than integrated. The fifth such self-test in the study, in which a prediction was registered and then corrected rather than assumed; closes the collective-navigation arc.

52. **The escape-weight valley of result 40 is a strong evolutionary brake, not an absolute barrier:** Making the collective-escape weight a heritable per-agent trait under capture-and-removal selection against the predictive predator, escape behaviour is evolutionarily stable and almost cost-free once present — a population seeded at weight 2 stays there with a handful of captures and full coherence — but it does not evolve from the no-escape state on any reasonable timescale. Selection on the weight is directional and upward from every start, yet the valley throttles the climb so severely that a population starting from no escape crawls to only about half the escape threshold in 400 time units while suffering the heaviest predation, with captures peaking exactly in the valley. The force-versus-alignment threshold of result 40 thus has a population-genetic image — strong evolutionary hysteresis in which escape is easy to keep but hard to evolve de novo, because the path to it runs through a region where partial commitment is actively selected against. This opens a co-adaptation thread with the predator still fixed.
53. **The escape brake is a barrier to origination, not invasion:** The valley of result 52 blocks escape from arising de novo, but not from spreading once present. A founder group of as little as five percent of the population seeded in the escape regime invades and carries the flock-mean weight into that regime, halving the capture toll, and a single mutational jump three to six times larger than the baseline clears the valley from a uniform no-escape start where the baseline step never does. Whether escape evolves is thus set by whether variation can deliver an agent past the valley in one step, not by selection, which favours escape once the valley is cleared — a mutation-limited rather than selection-limited outcome. Escape establishes at a mixed roughly sixty-percent equilibrium rather than fixing, a free-rider effect of the shared escape signal: once enough agents flee, the predator is outrun and the rest ride along protected, the public-good face of the same shared-direction principle that governs leadership and collective escape.
54. **The free-rider equilibrium is robust to predation pressure:** The mixed escaper equilibrium of result 53 is sticky rather than a knife-edge. Sweeping the predation hazard across an eightfold range moves the steady escaper fraction only from about 0.56 to about 0.66, and escape never approaches fixation — even under the heaviest predation roughly a third of the flock rides the shared escape as protected free-riders, and runs from different starting fractions settle into the same band. Intensifying predation erodes free-riding only weakly, confirming that the shared escape direction is a non-excludable public good: a protected non-contributing fraction persists as long as the flock escapes at all, so escape stabilises as a durable mixed strategy rather than sweeping to fixation.
55. **The two-sided arms race is asymmetric:** Letting the predator co-evolve a heritable predictive lead time, selected on capture success, the race is fundamentally asymmetric. (Against frozen no-escape prey the lead converges to about three; this overshoots the true capture optimum, which the direct measurement of result 56 places at about two, coinciding

with the disruption optimum — the evolved value is a small-population selection artifact, not a distinct capture optimum.) The asymmetry is the robust result: the predator climbs to its optimum from any start, but the prey counter is origination-limited by the escape valley, so de-novo co-evolution favours the predator and escape establishes only when seeded — after which it is an uncounterable hard counter that collapses predator selection, the lead trait drifting once no captures remain to select on. Collective escape is thus a powerful but evolutionarily fragile defence: hard to originate, easy to lose to drift, decisive once present.

56. **A self-test: the capture optimum coincides with the disruption optimum:** Measuring the capture rate and coherence directly against a fixed predator lead, rather than inferring them from evolution, the capture rate peaks at a lead of two — exactly where the flock is most fragmented — and falls away on both sides, vanishing by lead four where the predators overshoot. The lead that catches the most prey is therefore the same as the lead that most disrupts the flock, so the capture-versus-disruption distinction suggested by the evolved lead of three is withdrawn: that value is a small-population selection artifact, the noisy replace-the-worst scheme chasing a high-variance capture signal away from its true optimum. The lesson survives inverted, as a caution that a tightly converged evolved trait under noisy selection need not lie at its fitness optimum. The sixth self-test of the study, and the only one to correct a result from the same body of work.
57. **The origination brake is model-independent; the persistence of escape is not:** Re-running the evolution thread under an energy-budget fitness model — an explicit metabolic cost proportional to the escape weight, paid even when safe — splits the earlier results in two. The origination brake survives unchanged: escape still cannot evolve from the no-escape state, so the valley barrier is general, not an artifact of capture-and-removal. But the companion result that escape is free and stable once present does not survive: at a moderate cost every population, including one seeded deep in the escape regime, collapses to no escape, and the evolved weight falls almost all-or-nothing as the cost rises from zero. The metabolic cost is paid by every escaper continuously while predation threatens only the few near a predator at any instant, so a modest per-capita cost outweighs the diffuse benefit and escape is viable only where it is nearly free. Collective escape is thus even more evolutionarily fragile than the co-evolution results implied. This closes the co-adaptation thread.

The consistent thread across all results is that collective alignment is both the source of the flock's robustness and the mechanism by which stressors interact. It maintains coherence under noise and naive predation; it transmits spatial clustering that amplifies contagion; it drives the kinematic mixing that defeats spatial vaccination targeting; and it enables the reunion that makes kinematic damage reversible. Encirclement fails in 3D for a separate, purely geometric reason — a closed surface cannot be sealed by a handful of point predators — not through any property of the alignment force. Where targeting fails for a structural reason rather than a kinematic one — as in degree-targeted vaccination — the cause is that the flock contact network never had the exploitable heterogeneity to begin with. The most effective disruption strategies are those that operate at the flock's geometric scale (encirclement at $R_{\text{enc}}/R_g \sim 0.5$, in 2D only) or that exploit a timescale the alignment force cannot overcome (epidemic persistence after predator removal).

The same alignment force organizes the study's second and third threads as well. Beyond driving the kinematic mixing of the predator and contagion results, it builds the one quantity a flock can act on collectively — a single direction held in common — and every leadership and collective-navigation result is a statement about that quantity. A globally shared heading is amplified into coherent motion, whether it originates in spontaneous flocking, a committed escape, or an informed minority,

while locally heterogeneous directions are averaged away; competing shared headings are arbitrated by compromise below a critical conflict angle and by a product-law vote above it; and the flock’s adversaries — encirclement, panic, and the saboteur — are unified as attacks on that heading, with denial at parity always cheaper than capture. How accurate the heading is, finally, is set by the correlation of the directions its members carry rather than by their number: independent error averages away as the inverse square root of the group size, while correlated error imposes a floor that no group size can beat. Correlation is therefore the only currency the alignment force recognizes — it amplifies what is shared and averages away what is independent — and that single statement ties the flock’s coherence, its steerability, and its vulnerability to deception into one mechanism.

8 References

- Charbonneau, P. (2017). *Natural Complexity: A Modeling Handbook*. Princeton University Press.
- Silverberg, J. L., Bierbaum, M., Sethna, J. P., and Cohen, I. (2013). Collective motion of humans in mosh and circle pits at heavy metal concerts. *Physical Review Letters*, 110, 228701.
- Levis, D., Diaz-Guilera, A., Pagonabarraga, I., and Starnini, M. (2020). Flocking-enhanced social contagion. *Physical Review Research*, 2, 032056.
- Pacher, K., Bierbach, D., Kurvers, R. H. J. M., and Krause, J. (2026). Strategic choices of attack location allow predators to counter a collective prey defence. *Proceedings of the Royal Society B*, 293, 20260566.
- Inada, Y. and Kawachi, K. (2002). Order and flexibility in the motion of fish schools. *Journal of Theoretical Biology*, 214, 371-387. (Split and Reunion escape patterns in a two-dimensional fish-school model.)
- Demsar, J. and Lebar Bajec, I. (2014). Simulated predator attacks on flocks: A comparison of tactics. *Artificial Life*, 20(3), 343-359.
- Bartashevich, P., Schellinck, J., Tully, T., and Romanczuk, P. (2024). Collective anti-predator escape manoeuvres through optimal attack and avoidance strategies. *Communications Biology*, 7, 1548.

9 Appendix: Code

All simulation code is available at https://github.com/ninjahawk/Summer_Research

File	Description
flocking.py	Core model: buffer zone, vectorized force function, run loop, metrics
analysis.py	Validation limiting cases and parameter sweeps
predator.py	Single-predator extension with 4 experiments
phase_transition.py	Finite-size scaling of solid-to-fluid transition
geometry.py	Radius of gyration and aspect ratio analysis
multi_predator.py	Multi-predator experiments
evasion_analysis.py	Predator co-localization and evasion distance diagnostic
compactness_phase.py	Fixed-compactness finite-size scaling across $C=0.10$ and $C=0.78$
compactness_search.py	Intermediate-compactness phase search ($C=0.15..0.60$)
coordinated_predators.py	Predator-predator repulsion experiments
encirclement.py	Encirclement strategy, radius sweep, flock-breaking threshold
encirclement_scaling.py	Encirclement threshold vs flock size N

File	Description
fragmentation.py	Sub-cluster detection and division-vs-dissolution analysis
reunion.py	Sub-flock reunion after predator removal (recovery test)
min_flock_size.py	Minimum N for collective coherence and evasion
predator_sensing.py	Limited sensing radius, search/attack phases
panic.py	Static panic fraction sweep
panic_contagion.py	SI panic contagion (no recovery)
contagion_sis.py	SIS panic contagion with recovery rate gamma
hybrid_stressors.py	Combined predation + contagion
hybrid_sis.py	Sub-threshold SIS + encirclement (compression-amplification)
segregation.py	Active/passive segregation (mixed v0 populations)
segregation_alpha.py	Alpha-contrast segregation + local-purity diagnostic
large_N_encirclement.py	Encirclement at N=350, 700, 1000
critical_shift.py	Beta sweep with/without encirclement; threshold shift measurement
herd_immunity.py	Immune sub-population sweep at supercritical SIS
renc_scaling.py	R_enc sweep at N=350 and N=1000; collapse on R_enc/Rg
long_encirclement.py	Long-time encirclement (30000 steps); merge/split dynamics
encirclement_gap.py	Incomplete encirclement; gap detection test
outbreak_removal.py	Encirclement+SIS then predator removal; epidemic persistence
adaptive_encirclement.py	Adaptive R_enc=0.5*live_Rg vs fixed R_enc; disruption comparison
targeted_immunity.py	Targeted (high-degree first) vs random vaccination; herd-immunity efficiency
spatial_vaccination.py	Spatial (farthest-point sampling) vs random vs degree-targeted vaccination
hard_repulsion.py	Finite-size scaling with harder repulsion exponents n=1.5,3,6,12
langevin_repulsion.py	Langevin thermostat finite-size scaling; FDT diagnosis of crossover
langevin_hexatic.py	Langevin dynamics with hexatic order parameter; KTHNY structural melting test
flocking3d.py	3D extension of the flocking model; v_eq validation and noise sweep
flocking3d_noise.py	3D extended noise sweep ramp=0.5-30 and 2D comparison; finite-size scaling
flocking3d_predator.py	3D predator strategies: naive vs encirclement, R_enc sweep, 2D comparison
model.py	OOP foundation: Flock and Predator classes for new experiments