

Pod Platforms: A Perfect Filter, or a Perfectly Coarse One?

A Continuous-Time Model of Multi-Manager Hedge Fund Platforms

Marcos Mollica

Working draft — July 2026

Abstract summary. We model a multi-manager (“pod”) hedge fund platform in continuous time. Two headline risks emerge: a diversification benefit with a hard, regime-dependent correlation floor, and stop-loss rules that are a statistically coarse skill filter, biased on two separate margins. We derive the Bayesian and Kelly-optimal alternative, a population-level birth-rate condition, and simulation evidence on when fund-level alpha genuinely declines. Full abstract on the following page.

Contents

1	Abstract	2
2	Introduction	2
3	1. Pod-Level Dynamics	5
4	2. The Fund as a Portfolio of Pods	6
5	3. Fund-Level State Dynamics: Capital Weights as a Process	9
6	4. The Stop-Loss Rule: Survival Probability, Bias, and Incentives	10
7	5. Population Dynamics: Birth, Death, and the Birth-Rate Condition	16
8	6. The Bayesian Benchmark: Learning α_i Instead of Testing a Barrier	17
9	7. Discussion and Conclusion	20
10	8. Simulation Evidence: When Does Fund-Level Alpha Actually Decline?	22
11	Appendix A. The Finite-Horizon Version of the Section 4.5 Bias	26
12	Appendix B. Proportional (Cushion-Relative) Kelly Sizing — Now Proven Optimal in Section 6.4	28
13	Appendix C. Toward Endogenous Crowding (Sections 2.2, 5.4)	30
14	Appendix D. How Much Does a Single Intermediate Stop Recover?	31
15	References	32

1 Abstract

We model a multi-manager (“pod”) hedge fund platform in continuous time, decomposing each pod’s return into alpha, market beta, and idiosyncratic components, and building the fund as a portfolio of pods whose capital weights evolve as a controlled jump-diffusion under proportional reinvestment. Two headline risks emerge. First, the diversification benefit that markets the pod structure has a hard, regime-dependent floor: fund-level variance does not vanish with the number of pods once residual cross-pod correlation is admitted, and that correlation is plausibly worst exactly in the stress states when diversification is needed most — coupling directly to correlated stop-outs across the population. Second, and separately, we show that the fixed-threshold stop-loss rules used ubiquitously by pod platforms are a statistically coarse filter: survival probability depends on the ratio $\alpha_i / (\sigma_i^{tot})^2$ rather than on skill α_i alone, and this bias operates on two separate margins — which pods survive, and, via an exact Doob h -transform, how good a surviving pod’s *own* realized track record looks relative to its true parameter. We show the resulting PM incentive effect is genuinely ambiguous rather than automatically “gambling for resurrection,” derive a formal birth-rate condition required to sustain fund-level alpha under this attrition process, and show that pure stop-loss attrition cannot by itself deliver a stationary pod population. We propose a continuous-time Bayesian filtering framework as the statistically correct alternative to the fixed barrier — the barrier is a good implementation of the wrong test (Wald’s SPRT for a simple hypothesis) applied to the wrong problem (continuous learning of an unknown parameter) — and show the same $\alpha_i / (\sigma_i^{tot})^2$ ratio governing survival is also the continuous-time Kelly-optimal sizing fraction, with common practical approximations (staged, intermediate stops) shown to move in the right direction but recover only a bounded fraction of the benefit of full continuous control.

2 Introduction

Multi-manager “pod” platforms have become a dominant organizational form in global macro and multi-strategy hedge funds: capital is split across dozens or hundreds of semi-autonomous trading pods, each run by a PM against a fixed risk budget and, almost universally, a rigid stop-loss — the pod is closed the moment its cumulative P&L crosses a fixed threshold. The structure is popular because the pitch is intuitive: diversify idiosyncratic manager risk the way a portfolio diversifies asset risk, and let a mechanical rule do the unpleasant work of firing underperformers before discretion or sunk-cost bias can intervene. This paper asks whether that mechanical rule does what it is meant to do, and, having shown that it largely does not, works out what a platform should do instead.

Main results. In the interest of the reader’s time, we state the paper’s substantive findings up front rather than making them wait for the derivations.

1. **The stop-loss rule filters on a Sharpe-like ratio, not on skill.** A pod’s probability of ever being stopped out is $\exp(-2\alpha_i L / (\sigma_i^{tot})^2)$ — governed by $\alpha_i / (\sigma_i^{tot})^2$, not by α_i alone (Section 4.2). A highly skilled but volatile PM can have *worse* survival odds than a

mediocre but quiet one. This formalizes a suspicion every allocator has had informally.

2. **The bias runs deeper than cross-sectional selection.** It is not just that skilled-but-unlucky pods get cut while mediocre-but-lucky ones survive (Section 4.4) — even a single pod’s *own* realized track record, conditional on having survived, systematically overstates its true alpha, via an exact Doob h -transform correction (Section 4.5). At the barrier itself, this manufactured excess return is governed entirely by volatility and barrier width and is *independent of the pod’s actual skill* — the metric meant to reveal edge is, at the margin, dominated by a term that has nothing to do with edge.
3. **“Gambling for resurrection” is not the automatic prediction the folk wisdom assumes.** Modeling a PM’s continuation value as a down-and-out option, standard barrier-option results imply *negative* vega near the knockout: a rational PM facing a clean, continuous stop should reduce risk approaching it, not increase it (Section 4.6). The classical risk-shifting result requires a different payoff structure (fixed-date tournaments, or effectively capped personal cost of going further underwater) — which turns “does the barrier induce gambling” from an assumption into a testable empirical question with a clean observable implication.
4. **Headcount growth needs its own capacity policy, independent of the barrier.** A genuinely skilled pod’s stopping time is *defective* — positive probability of never being stopped — which is a good problem, not a flaw: it just means the stop-loss rule alone provides no natural ceiling on how many long-lived pods a platform accumulates. Real platforms already manage this through capacity limits and discretionary redeployment; the model’s point is that this capacity-management function is doing genuinely different work from skill discovery, and conflating the two is the actual risk (Section 5.2).
5. **The birth rate required to sustain fund-level alpha is higher than a naive attrition model implies,** precisely because of Result 2: since the barrier doesn’t cleanly separate skill from luck, some of what the platform loses to stop-outs is real alpha, not just noise (Section 5.3, equation 5.2).
6. **The stop-loss is a good implementation of the wrong test.** A fixed barrier closely approximates the optimal Wald sequential test for a *simple* two-point hypothesis; the platform’s actual problem — continuously learning an unknown continuous α_i and sizing capital accordingly — is a different and harder problem, for which a continuous-time Bayesian filter (Section 6.1) is the natural tool, and for which the same ratio $\alpha_i / (\sigma_i^{tot})^2$ that governs barrier survival also turns out to be the continuous-time Kelly-optimal sizing fraction (Section 4.7, 6.3) — survival and sizing are two readings of the same underlying quantity.
7. **The diversification benefit that markets the pod structure has a hard, regime-dependent floor.** The pitch for pod platforms rests on idiosyncratic risk shrinking like $1/N$ as pods are added. Once residual cross-pod correlation is admitted (crowded macro views, overlapping talent pools, correlated positioning), fund-level variance no longer vanishes with N — it converges to a floor set by average pairwise correlation (Section 2.2, equation 2.4). This floor is not a minor technical caveat: it is plausibly *worst exactly when diversification is needed most*, because residual correlation is likely to spike in stress regimes — the same

regimes in which pods are simultaneously approaching their stop-loss barriers (Section 5.2). A correlated drawdown does not just cost the fund a bad quarter; it can trigger correlated stop-outs, shrinking the pod count itself at the worst possible moment. This is a structural risk that runs directly against the core diversification claim used to raise capital for the platform structure in the first place, not a peripheral modeling detail.

8. **Fund-level alpha decline is real but conditional, not generic.** Simulation of platforms competing for hires from a shared PM talent pool (Section 8) shows that hiring competition *alone* never produces outright decline in fund-level alpha — only stagnating hire quality and a persistent, non-closing performance gap relative to less-competed peers. Genuine decline requires a hiring-demand-to-pool-supply ratio severe enough to eliminate selection power, *combined with* a meaningfully strong alpha-erosion mechanism (crowding-driven capacity decay or secular talent-quality decline) — and, strikingly, this can occur even in a nominally talent-rich environment if enough platforms are drawing on it (Section 8.6). This sharpens rather than dilutes the warning: a platform can check directly whether it sits in the specific, identifiable danger zone rather than treating decline as unconditionally inevitable.

Related literature (brief). The portfolio-theoretic backbone (Sections 2–3) draws directly on Markowitz-style variance decomposition (Markowitz, 1952) and, for the capital-weight dynamics under proportional reinvestment, on Fernholz’s stochastic portfolio theory (Fernholz, 2002). The first-passage and ruin-probability results (Section 4) are classical (Brownian ruin theory and sequential analysis: Wald, 1947; Karatzas & Shreve, 1991), applied here to a setting — capital allocation across a fund’s own internal managers — where, to our knowledge, the statistical consequences for skill/luck separation have not been made explicit. The incentive analysis (Section 4.6) borrows from barrier-option pricing (Merton, 1973) and contrasts with the mutual-fund tournament literature on manager risk-shifting (Brown, Harlow, & Starks, 1996), arguing the two settings have different — and empirically distinguishable — payoff structures. The Bayesian benchmark (Section 6) connects to Kalman-Bucy filtering of an unknown drift (Kalman & Bucy, 1961) and to portfolio choice under parameter uncertainty (Xia, 2001), in the spirit of Merton’s continuous-time portfolio framework (Merton, 1969, 1971) and Kelly’s original growth-optimality criterion (Kelly, 1956). The capacity-decay assumption of Section 5.4 follows the standard diseconomies-of-scale-in-active-management result (Berk & Green, 2004), and the occupancy formalization of Appendix C draws on the economics of superstar and congestion effects (Rosen, 1981). The ruin-constrained sizing discussion of Section 6.4 and Appendix B is related in spirit to constant-proportion portfolio insurance (Black & Perold, 1992).

Roadmap. Section 1 sets up pod-level return dynamics. Section 2 builds the fund as a portfolio of pods and decomposes the role of cross-pod correlation. Section 3 derives capital weights as a controlled jump-diffusion process. Section 4 is the paper’s statistical core: the stop-loss survival probability, the two-part bias, and PM incentives. Section 5 works out the population-level consequences — the birth-rate condition and a structural result on population stability. Section 6 develops the Bayesian alternative and its implied sizing rule. Section 7 synthesizes the normative comparison between what the model says a platform *should* do and what the barrier mechanically *does*. Section 8 adds simulation evidence on a question the closed-form

results can't answer directly: under what conditions does fund-level alpha actually decline for platforms competing over a shared PM talent pool.

Symbol	Meaning
$X_i(t)$	Cumulative P&L (or return) of pod i at time t
α_i	True (latent) drift / skill parameter of pod i
β_i	Pod i 's loading on the market factor
σ_i	Idiosyncratic volatility of pod i
σ_i^{tot}	Total volatility of pod i , $\sqrt{\beta_i^2 \sigma_M^2 + \sigma_i^2}$
$R_M(t)$	Market factor process, $dR_M = \mu_M dt + \sigma_M dW_M$
$Z_i(t)$	Idiosyncratic standard Brownian motion driving pod i , independent across i and of W_M
$w_i(t)$	Capital weight (fraction of fund AUM) allocated to pod i at time t
L	Stop-loss barrier (fixed threshold at which a pod is closed)
$N(t)$	Number of active pods at time t
$\lambda(t)$	Birth rate of new pods
$\mu(t)$	Death (stop-out) rate of active pods
Σ_{ij}	Residual (non-market) covariance between pods i and j

3 1. Pod-Level Dynamics

We model each pod i 's cumulative return as an Itô process with three components: a constant drift representing the PM's true, unobservable skill; a loading on a common market factor; and an idiosyncratic diffusion term specific to the pod's book.

$$dX_i(t) = \alpha_i dt + \beta_i dR_M(t) + \sigma_i dZ_i(t), \quad dR_M(t) = \mu_M dt + \sigma_M dW_M(t) \quad (1.1)$$

with $Z_i \perp Z_j$ for $i \neq j$ and $Z_i \perp W_M$ for all i . Substituting the market factor:

$$dX_i(t) = (\alpha_i + \beta_i \mu_M) dt + \beta_i \sigma_M dW_M(t) + \sigma_i dZ_i(t) \quad (1.2)$$

so that pod i 's instantaneous variance is

$$\text{Var}(dX_i) = (\beta_i^2 \sigma_M^2 + \sigma_i^2) dt \equiv (\sigma_i^{tot})^2 dt \quad (1.3)$$

We treat $\alpha_i, \beta_i, \sigma_i$ as constant over the life of a given pod. This is a deliberate simplification — α_i is very unlikely to be constant in reality (strategy decay, changing opportunity sets, PM learning), and Section 6 revisits α_i as a parameter to be *learned* rather than known. Holding it constant here isolates the statistical problem the stop-loss rule faces even in the best case, where the PM’s skill genuinely does not change: if the barrier can’t recover a *constant* α_i reliably, it certainly can’t handle a time-varying one.

Footnote — arithmetic vs. geometric formulation. We work throughout in arithmetic (additive) P&L rather than geometric (compounding) returns. This is the standard simplification in the first-passage / barrier-crossing literature because it keeps the stop-loss problem an exactly solvable linear Brownian motion with drift (Section 4); the geometric analogue requires the barrier to be expressed in log-return space and loses closed-form tractability for several of the population-dynamics results in Section 5. The arithmetic formulation is a reasonable approximation for pods evaluated on a fixed-notional, fixed-risk-budget basis (the typical pod-platform setup, where a pod’s capital allocation — and hence its dollar risk — is reset periodically rather than continuously compounded), and we adopt it as the baseline model, flagging geometric compounding as a natural but non-trivial extension.

Pod i is a claim with two state variables that matter for everything that follows: its *type* $(\alpha_i, \sigma_i^{tot})$ — fixed but unobserved by the allocator — and its *path* $X_i(t)$, which is the only thing the platform actually sees. The entire tension in this paper is that the stop-loss rule and, ultimately, capital allocation itself, must be a function of the path, while what we actually care about is the type.

4 2. The Fund as a Portfolio of Pods

The fund’s return is a capital-weighted sum of pod returns. Let $w_i(t)$ denote the fraction of fund AUM allocated to pod i at time t , with $\sum_i w_i(t) \leq 1$ (the shortfall from 1 is uninvested / cash, or a fund-level overlay, which we suppress for now). Summing (1.2) across pods:

$$dR_F(t) = \sum_i w_i(t) dX_i(t) = \underbrace{\left(\sum_i w_i \alpha_i + \sum_i w_i \beta_i \mu_M \right)}_{\text{drift}} dt + \underbrace{\left(\sum_i w_i \beta_i \right) \sigma_M}_{\text{common factor}} dW_M + \underbrace{\sum_i w_i \sigma_i}_{\text{idiosyncratic}} dZ_i \quad (2.1)$$

This is the direct continuous-time analogue of the static Markowitz decomposition, and it is worth being precise about *why* pod platforms exist at all in this framework: the pitch is that $\sum_i w_i \sigma_i dZ_i$ diversifies away as N grows (idiosyncratic risk shrinks like $1/N$ under equal weighting with bounded σ_i), while $\sum_i w_i \alpha_i$ does not shrink — it is additive in expectation regardless of N . A platform with N pods of equal true skill and independent idiosyncratic risk should, in the limit, deliver the *average* pod alpha at close to *zero* idiosyncratic variance. This is the entire economic argument for the pod model over a single concentrated manager, and it only survives if the independence assumption in Section 1 ($Z_i \perp Z_j$) actually holds.

2.1 Variance decomposition. From (2.1), assuming for now that the idiosyncratic terms remain mutually independent:

$$\text{Var}(dR_F) = \left(\sum_i w_i \beta_i \right)^2 \sigma_M^2 dt + \sum_i w_i^2 \sigma_i^2 dt \quad (2.2)$$

The first term — aggregate market beta — does **not** diversify with N ; it is a portfolio-level exposure that must be hedged directly (e.g., an index overlay) if the platform wants pure alpha. The second term shrinks with N under any reasonable weighting scheme (e.g., $w_i = 1/N$ gives $\sum w_i^2 \sigma_i^2 = \bar{\sigma}^2 / N \rightarrow 0$). This is the textbook result, and it is *only* correct under the independence assumption.

Made concrete. Take pods that are, for illustration, all identical in volatility to the Pod A example used throughout Section 4: $\sigma_i = 8\%$, equally weighted, and genuinely independent. Idiosyncratic portfolio volatility is $\sigma / \sqrt{N}$:

N pods	Idiosyncratic vol
1	8.0%
4	4.0%
16	2.0%
64	1.0%
256	0.5%

Getting idiosyncratic risk down to something a marketing deck could plausibly call “near zero” — 1% — still requires **64 genuinely independent pods**, and 0.5% requires 256. Most real platforms run well short of the higher end of that range (tens to low hundreds). Independence alone is an expensive property to fully realize, before correlation is even admitted.

2.2 Residual correlation. In practice pods are not idiosyncratically independent: they share macro views, crowd into similar trades (rates receivers, USD shorts, AI-adjacent equities — the kind of positioning that shows up across desks industry-wide, not just within one platform), and are staffed from overlapping talent pools with correlated priors. Introduce a residual covariance matrix Σ with $\Sigma_{ii} = \sigma_i^2$ and $\Sigma_{ij} = \rho_{ij} \sigma_i \sigma_j$ for $i \neq j$, capturing correlation *beyond* the common market factor already stripped out in β_i . Then

$$\text{Var}(dR_F) = \left(\sum_i w_i \beta_i \right)^2 \sigma_M^2 dt + w' \Sigma w dt \quad (2.3)$$

and the diversification benefit of adding pods is capped: as $N \rightarrow \infty$ with $w_i = 1/N$ and average pairwise correlation $\bar{\rho} > 0$,

$$w' \Sigma w \rightarrow \bar{\rho} \bar{\sigma}^2 \quad (\text{does not vanish}) \quad (2.4)$$

For the equal-weight, equal-vol, single-average-correlation special case used in the illustration above, the exact finite- N formula (standard equicorrelated-portfolio result) is

$$\text{Var}_{idio}(N) = \sigma^2 \left[\bar{\rho} + \frac{1 - \bar{\rho}}{N} \right] \quad (2.4b)$$

which recovers (2.2) at $\bar{\rho} = 0$ and the floor (2.4) as $N \rightarrow \infty$.

Made concrete, continued. Fix a realistically-sized platform at $N = 50$ pods, still $\sigma = 8\%$, and compare idiosyncratic volatility ($\sqrt{\text{Var}_{idio}(N)}$) across a range of *mild* residual correlation:

$\bar{\rho}$	Idiosyncratic vol at $N = 50$	Floor as $N \rightarrow \infty$
0%	1.1%	0%
20%	3.7%	3.6%
30%	4.5%	4.4%
40%	5.1%	5.1%

Two things stand out. First, the jump from $\bar{\rho} = 0$ to a genuinely mild $\bar{\rho} = 20\%$ more than **triples** idiosyncratic volatility at this platform size — from 1.1% to 3.7% — despite nothing else about the platform changing. Second, at $N = 50$ the platform is already sitting almost exactly *at* its correlation floor (3.7% realized vs. 3.6% asymptotic): **doubling or quadrupling the number of pods buys almost nothing further**. A platform at $\bar{\rho} = 20\%$ that grows from 50 pods to 200 only takes idiosyncratic vol from 3.7% down to 3.6% — a rounding error next to the headline “more pods, more diversification” pitch. Mild, entirely plausible correlation is enough to make the platform’s core marketed benefit run out of room almost immediately, long before headcount does.

This is the standard result from static portfolio theory (the same asymptotic floor that limits diversification benefit in equity portfolios once average correlation is bounded away from zero) — but it matters more here than in a passive equity portfolio for two reasons specific to the pod setting, both of which we develop later:

- **Regime-dependence.** $\bar{\rho}$ is very unlikely to be constant. The economically relevant case is $\bar{\rho}$ rising sharply in stress states — exactly the states in which pods are most likely to be simultaneously approaching their stop-loss barriers. This means the “ $1/N$ diversification” benefit that justifies the platform’s existence is weakest precisely when the stop-loss mechanism (Section 4) is most active — a correlated *drawdown* event doesn’t just cost the fund a bad quarter, it can trigger correlated *stop-outs*, temporarily shrinking N itself at the worst possible moment. Variance and population dynamics are coupled, not separable, in a stress regime.
- **Endogeneity in N .** $\bar{\rho}$ is plausibly *increasing* in N itself once N is large relative to the platform’s addressable idea space (more pods chasing a finite set of liquid, model-able macro trades). We return to this as a bound on the birth-rate policy of Section 5: adding pods is not a free diversification lunch indefinitely.

This deserves to be stated plainly rather than left as a technical footnote to the variance decomposition: the diversification benefit that pod platforms are marketed on is not the unconditional, always-on property (2.2) makes it look like in isolation. It has a floor (2.4), and that floor is highest exactly when an allocator needs diversification most. A platform can genuinely deliver low, stable idiosyncratic variance in calm regimes and still deliver a correlated, fund-wide drawdown in a stress regime — both consistent with the same model, and the second is precisely the scenario the pitch is meant to protect against.

2.3 Fund alpha as an aggregation problem. Rewriting the drift term of (2.1) makes explicit what the fund is actually delivering:

$$\alpha_F(t) \equiv \sum_i w_i(t) \alpha_i \quad (2.5)$$

Because α_i is unobserved (Section 1), $\alpha_F(t)$ is *also* unobserved in real time — the fund only ever sees a noisy realization of it through dR_F . This is the bridge to Section 4: the stop-loss rule is the platform’s mechanism for inferring something about the α_i ’s from noisy path data in order to keep $w_i(t)$ pointed at the pods that matter, and (2.5) makes clear that a bad inference — misallocating w_i toward high- σ_i /low- α_i pods that happened to survive — degrades α_F directly, not just fund variance.

5 3. Fund-Level State Dynamics: Capital Weights as a Process

Sections 1–2 treated $w_i(t)$ as an exogenous input. It isn’t one: capital weights evolve mechanically as pods compound their own P&L, and they jump discretely at birth and stop-out. This section derives that process, because it is the actual state variable the rest of the paper needs — the stop-loss rule (Section 4) is a jump trigger acting on it, and the birth-death model (Section 5) is a statement about its jump measure.

3.1 Mechanical drift under proportional reinvestment. Let pod capital compound at the pod’s own return: $dC_i(t) = C_i(t) dX_i(t)$, and let fund capital be $C_F(t) = \sum_i C_i(t)$, so that $dC_F = C_F dR_F$ — consistent with the fund return defined in (2.1). The weight is $w_i(t) = C_i(t)/C_F(t)$. Applying Itô’s formula to this ratio (a standard computation in stochastic portfolio theory — see Fernholz (2002) on portfolio weight dynamics under proportional reinvestment, the natural home for this kind of parallel to formal portfolio theory) gives:

$$\frac{dw_i}{w_i} = (dX_i - dR_F) - (d\langle X_i, R_F \rangle - d\langle R_F \rangle) \quad (3.1)$$

The first bracket is the intuitive part: a pod’s weight rises when it outperforms the fund and falls when it underperforms — pure relative-return drift, with no discretionary rebalancing required. The second bracket is a covariation correction (Itô, not Stratonovich) and is what makes this genuinely different from a naive “weight tracks relative P&L” story: a pod highly

correlated with the fund's own return ($d\langle X_i, R_F \rangle$ large) has its weight drift *damped* relative to an equally-outperforming but uncorrelated pod. This is a clean, formal statement of something allocators know intuitively — a pod that “wins when everything wins” earns less incremental conviction than a pod that wins idiosyncratically — and it falls directly out of the compounding mechanics, not out of any discretionary judgment about correlation.

3.2 Discrete jumps: birth and death. Equation (3.1) describes weight drift for an already-active pod with no platform intervention. Two discrete events sit on top of it:

- **Birth.** At a pod's launch time τ_i^{birth} , w_i jumps from 0 to some initial allocation $w_i^{(0)}$, set by the platform (not by the mechanics above) — a discretionary sizing decision made with essentially no information about α_i yet, which is precisely the problem Section 6 tries to formalize.
- **Death.** At the stop-out time $\tau_i^{stop} = \inf\{t : X_i(t) \leq -L\}$ (Section 4's first-passage time), w_i jumps from whatever value (3.1) has carried it to, down to 0, and that capital is presumably redeployed — either to surviving pods (mechanically raising their weights) or held as dry powder pending a new birth.

3.3 The full process. Putting these together, $w_i(t)$ is a controlled jump-diffusion:

$$dw_i(t) = w_i(t) \left[dX_i - dR_F - (d\langle X_i, R_F \rangle - d\langle R_F \rangle) \right] + w_i^{(0)} dN_i^{birth}(t) - w_i(t^-) dN_i^{death}(t) \quad (3.2)$$

where N_i^{birth} and N_i^{death} are counting processes for pod i 's (at most one each, if a pod that dies is not “reborn” under the same index) birth and stop-out events. This is the object the rest of the paper is really about: Section 4 characterizes the law of N_i^{death} (via the first-passage distribution of X_i hitting $-L$), and Section 5 aggregates N^{birth} and N^{death} across the population into the birth-death process governing $N(t)$ and, through (3.2), the evolution of realized fund alpha $\alpha_F(t) = \sum_i w_i(t) \alpha_i$.

One modeling choice to flag explicitly rather than bury: (3.1)'s drift is *mechanical* — it is what happens with zero discretionary trimming. Real platforms also rebalance discretionarily (adding conviction capital to a pod ahead of what pure compounding would deliver, or trimming a pod that hasn't hit its stop but has lost the PM's confidence). We leave that out of the baseline model and treat it as a bounded-variation control term that could be added to (3.2) later — the paper's claim is sharper if we first show what the stop-loss/birth-rate mechanics deliver *without* assuming the platform is already doing anything clever, since that's the realistic baseline for a purely rules-based pod shop.

6 4. The Stop-Loss Rule: Survival Probability, Bias, and Incentives

This is the section the paper is really building toward. We now characterize $\tau_i^{stop} = \inf\{t : X_i(t) \leq -L\}$ precisely, show exactly how it conflates skill and luck, and ask what it does to

PM incentives.

4.1 Setup. From (1.2), pod i 's P&L is a Brownian motion with constant drift and variance:

$$X_i(t) = \tilde{\alpha}_i t + \sigma_i^{tot} W_i(t), \quad \tilde{\alpha}_i \equiv \alpha_i + \beta_i \mu_M \quad (4.1)$$

started at $X_i(0) = 0$, with the stop triggered the first time this path touches $-L$. We write $\tilde{\alpha}_i$ (rather than α_i) to be precise that the barrier reacts to *total* realized drift, including the market-beta contribution — a pod can be stopped out for carrying unhedged beta into a market drawdown that has nothing to do with its idiosyncratic skill. We suppress this distinction notationally from here on and write α_i for $\tilde{\alpha}_i$, but it is worth remembering that “skill” and “what triggers the barrier” are not even measuring the same drift.

4.2 Survival probability. What a platform or investor actually observes is not whether a pod survives *forever* — it is whether the pod is still running after some fixed evaluation window: a year, three years, the life of an allocation. The practically relevant object is therefore the probability of surviving to a **finite horizon** T , with the classical “ruin probability” result recovered as the clean $T \rightarrow \infty$ limiting case — useful for building intuition about the steady state, but not the number anyone is actually evaluating a live pod against. We present both, finite horizon first.

Finite horizon. Via the reflection-principle argument derived in full in Appendix A.1–A.2, the probability a pod is still running at time T is

$$P(\tau_i^{stop} > T) = \Phi(d_+) - e^{-2\alpha_i L / (\sigma_i^{tot})^2} \Phi(d_-), \quad d_{\pm} \equiv \frac{\alpha_i T \pm L}{\sigma_i^{tot} \sqrt{T}} \quad (4.1b)$$

A running example. To keep the rest of this section concrete, fix one illustrative pod throughout: **Pod A**, with true annualized drift $\alpha_i = 5\%$, total volatility $\sigma_i^{tot} = 8\%$ (Sharpe ≈ 0.63 , a conservative, well-behaved macro pod), and a stop-loss set at $L = 10\%$ of allocated capital. Plugging into (4.1b):

Horizon T	Survival probability
1 year	91.4%
3 years	83.0%
∞ (limit)	79.0%

The gap between $T = 1$ and the asymptote is informative in a way that cuts the other direction from what a faster-moving pod would show: only about 41% of Pod A's *eventual* ruin probability (8.6 percentage points out of an eventual 21.0%) has resolved by the end of year one; by year three, about 81% of it has (17.0 out of 21.0 points). **For a conservative, well-cushioned pod, a one-year track record is not very informative about its long-run fate** — most of the risk that will ever show up simply hasn't had time to, which is exactly the sense in which short evaluation windows under-inform on the safer end of the pod population.

Infinite horizon (the limiting case). Taking $T \rightarrow \infty$ in (4.1b) recovers the classical result:

$$P(\tau_i^{stop} < \infty) = \begin{cases} 1 & \alpha_i \leq 0 \\ \exp\left(-\frac{2\alpha_i L}{(\sigma_i^{tot})^2}\right) & \alpha_i > 0 \end{cases} \quad (4.2)$$

This is the central fact of the paper, and it is worth stating in plain terms before the notation: **survival depends on a pod's drift relative to its volatility, not on the drift alone.** A highly skilled but volatile pod can have worse odds of ever surviving than a mediocre but quiet one, purely because the barrier is testing the ratio $\alpha_i / (\sigma_i^{tot})^2$ — a Sharpe-squared-like quantity — not skill itself.

Made concrete: alongside Pod A ($\alpha = 5\%, \sigma = 8\%$), consider **Pod B**, with $\alpha = 1.25\%$ and $\sigma = 4\%$ — a quarter of Pod A's true skill. Both have the identical barrier parameter $2\alpha L / \sigma^2 \approx 1.5625$ against the same $L = 10\%$ stop, and therefore *the same* 79.0% infinite-horizon survival probability. A platform screening purely on survivorship cannot tell these two pods apart, despite one having four times the edge of the other. (Note the asymmetry in what it took to match them: a quarter of the drift only needed *half* the volatility, since the barrier parameter scales as α / σ^2 — skill and volatility do not trade off one-for-one.)

4.3 Expected time to stop-out. Conditional on eventually being stopped, the expected stopping time has a clean form via a standard conditioning argument (the path, conditioned on hitting $-L$, behaves in law like Brownian motion with drift $-\alpha_i$ — the same magnitude, flipped sign — a fact from the theory of Brownian ruin problems):

$$E[\tau_i^{stop} \mid \tau_i^{stop} < \infty] = \frac{L}{|\alpha_i|} \quad (4.3)$$

for $\alpha_i \neq 0$ (for $\alpha_i = 0$, stopping is certain but has infinite expected time — a degenerate case worth noting but not economically central). For Pod A, this is $L / \alpha_i = 0.10 / 0.05 = 2.0$ years: a pod that is eventually going to be stopped out takes, on average, a full two years to get there — consistent with the previous point that a conservative pod's fate resolves slowly. Two implications follow: a genuinely bad pod ($\alpha_i \ll 0$) gets cut quickly, which is the barrier working as intended; but a mediocre pod near zero drift can wander for a long time before triggering, burning capital and platform attention without ever giving a clean signal either way.

4.4 The bias, part one — cross-sectional selection. Consider an incoming population of pods with heterogeneous $(\alpha_i, \sigma_i^{tot})$ drawn from some joint distribution. Because $P(\text{survive})$ from (4.2) is increasing in $\alpha_i / (\sigma_i^{tot})^2$, standard selection logic gives $E[\alpha_i \mid \text{survive}] > E[\alpha_i]$ unconditionally — the surviving population looks more skilled on average than the population it was drawn from, which is the sense in which the barrier does *some* useful filtering. But the selection is on the *ratio*, not on α_i itself: for a fixed value of that ratio, the surviving set is agnostic between a high-alpha, high-vol pod and a low-alpha, low-vol pod — exactly the Pod A / Pod B pair above. The platform ends up unable to distinguish, from survivorship alone, a population of consistently mediocre, low-risk PMs from a population of genuinely skilled, aggressively sized PMs who got a fair coin flip.

4.5 The bias, part two — conditioning inflates the realized path itself. Before the formal

treatment, the plain-language version of the claim: **an allocator who only ever examines pods that survived can find an edge even where none exists, purely because of how a track record qualified for attention in the first place.** This is a version of a familiar statistical trap — survivorship bias in mutual fund returns, or “the stocks in the index today are the ones that didn’t get delisted” — but here it runs deeper than the usual story. The usual survivorship story is *cross-sectional*: across many funds, the ones you see are disproportionately the lucky or skilled ones (that’s Section 4.4). The claim in this subsection is sharper and holds for a *single* pod: even conditioning on the fact that *this one pod, with its one true fixed α_i* , happened to survive, its own realized track record is not a fair read of α_i — the mere act of conditioning on “it’s still here” manufactures extra apparent performance, over and above the pod’s true skill. Intuitively: among all the possible paths a pod with drift α_i could have taken, the ones that survived are disproportionately the ones that got a bit of a helpful push away from the barrier at some point — and averaging only over survivors bakes that push into what looks like “performance.”

This has an exact treatment via a **Doob h -transform**, a standard technique for describing what a stochastic process looks like *once you condition on some future event about it* — here, the event “this path never hits $-L$.” Define the **barrier parameter**

$$S_i \equiv \frac{2\alpha_i L}{(\sigma_i^{tot})^2} \quad (4.4b)$$

so that (4.2) reads compactly as survival probability $= 1 - e^{-S_i}$, ruin probability $= e^{-S_i}$, for $\alpha_i > 0$. The harmonic function governing survival from an arbitrary current position $x > -L$ (not just $x = 0$) is

$$u(x) = P(\text{never hit } -L \mid X_i(0) = x) = 1 - \exp\left(-\frac{2\alpha_i(x+L)}{(\sigma_i^{tot})^2}\right) \quad (4.5)$$

which satisfies the generator equation $\alpha_i u' + \frac{1}{2}(\sigma_i^{tot})^2 u'' = 0$ with $u(-L) = 0$, $u(\infty) = 1$. Conditioning X_i on the event $\{\tau_i^{stop} = \infty\}$ produces, by the h -transform, a new process with the *same* diffusion coefficient σ_i^{tot} but a state-dependent drift — an extra “push” added at every point, strongest near the barrier and vanishing far from it:

$$dX_i^h(t) = \left[\alpha_i + \Delta_i(X_i^h(t))\right] dt + \sigma_i^{tot} dW_i^h(t), \quad \Delta_i(x) \equiv (\sigma_i^{tot})^2 \frac{u'(x)}{u(x)} = \frac{2\alpha_i}{\exp\left(\frac{2\alpha_i(x+L)}{(\sigma_i^{tot})^2}\right) - 1} \quad (4.6)$$

$\Delta_i(x) > 0$ everywhere — the conditioned process always drifts faster than the true α_i — and it is exactly this term that formalizes the intuition above. Three cases make it concrete:

- **Far from the barrier** ($x \rightarrow \infty$): $\Delta_i(x) \rightarrow 0$. Conditioning has (correctly) no effect once a pod has enough cushion — a well-cushioned pod’s realized track record is an unbiased read on α_i .

- **At inception** ($x = 0$): $\Delta_i(0) = \frac{2\alpha_i}{e^{S_i} - 1}$. For **Pod A** ($\alpha_i = 5\%$, $S_i \approx 1.5625$), this evaluates to $\Delta_i(0) = 0.10/(e^{1.5625} - 1) \approx 0.10/3.77 \approx 2.7\%$ per year — a newly-launched pod, conditioned purely on the fact that it goes on to survive, has a manufactured excess expected return of nearly three percentage points on top of its true 5% drift, roughly a **53% inflation** of its apparent edge before a single piece of information about its actual skill has been observed. Note this is smaller in relative terms than it would be for a less-cushioned pod — the more conservative Pod A's larger barrier parameter (further from the barrier in a statistical sense) means the conditioning distortion is milder, which is the correct qualitative direction: better-cushioned pods should be — and here are — less distorted by survivorship.
- **Near the barrier** ($S_i \rightarrow 0$, i.e. a barely-profitable pod, or a very tight L): $\Delta_i(0) \rightarrow (\sigma_i^{tot})^2/L$, a limit **independent of α_i entirely**. A numeric check: take **Pod C**, otherwise like Pod A but with a much smaller true edge, $\alpha_i = 0.3\%$ ($S_i \approx 0.094$, genuinely close to zero); (4.6) gives $\Delta_i(0) \approx 6.1\%$, close to the limiting value $\sigma^2/L = 0.0064/0.10 = 6.4\%$. Right next to the barrier, the size of the conditioning bias is governed almost entirely by volatility and barrier width, not by skill at all — the strongest, most quotable form of “stop rules can't separate luck from skill”: at the margin where the barrier actually binds, the very quantity meant to reveal edge is dominated by a term that has nothing to do with edge.

This result is exact for conditioning on *eternal* survival ($\tau_i^{stop} = \infty$); Appendix A.1–A.2 derives the finite-horizon analogue in closed form and shows the bias is largest at intermediate horizons — exactly the one-to-two-year evaluation windows platforms actually use.

4.6 PM incentives: does the barrier induce gambling for resurrection? A natural conjecture is that a PM approaching $-L$ has an incentive to raise σ_i — “gambling for resurrection,” by analogy with distressed-firm risk-shifting and mutual-fund tournament behavior. The model says this is *not* automatic, and the mechanism matters.

If the PM's continuation payoff resembles a down-and-out call — some increasing, convex function of terminal P&L, extinguished entirely at $-L$ — then standard barrier-option results (Merton, 1973) imply the opposite of the naive conjecture: near a knock-out barrier, the option's vega (sensitivity to volatility) is typically *negative*. Raising σ_i raises the probability of triggering the knock-out faster than it raises the value of the surviving upside, so a rational PM with this exact payoff structure should *reduce* risk near the barrier, not increase it.

The precise intuition, made concrete with a formula already in the paper. Vega has two channels that pull in opposite directions, and which one wins depends entirely on distance to the barrier. The first channel is the standard one: more volatility widens the distribution of terminal outcomes, which is good for the holder of a convex, upside-only payoff — this is why vanilla options always have positive vega. The second channel is specific to a *barrier* payoff: more volatility also widens the distribution of the *path*, not just the terminal value, which raises the probability the path wanders down and touches $-L$ before anything good has a chance to happen — extinguishing the payoff entirely regardless of what the terminal value would otherwise have been. Equation (4.2) already tells us exactly how strong this second channel is: ruin probability is $\exp(-2\alpha_i y / (\sigma_i^{tot})^2)$ for remaining cushion y , which is strictly increasing in σ_i^{tot} for any fixed $y > 0$ — more volatility always raises the chance of being knocked out,

with no offsetting term. What changes with distance to the barrier is not *whether* this channel operates but how much it dominates the first one. Far from the barrier, a marginal increase in volatility barely moves the already-small ruin probability (there's a lot of room to wander before $-L$ becomes relevant), so the marginal unit of volatility is spent almost entirely on the first, value-enhancing channel — vega is positive, as usual. Right next to the barrier, the pod is, informally, “one bad swing away” from extinction: the same marginal increase in volatility now moves the ruin probability substantially, precisely because there's so little room left to absorb it, while the corresponding gain to the conditional-on-survival payoff is comparatively small — there isn't much runway left even in the good scenarios for extra variance to convert into extra option value. Near the barrier, in other words, a PM is spending almost the entire marginal unit of volatility buying more extinction risk and very little marginal payoff — which is exactly why vega flips sign. This is not a separate story from Section 4's central result; it is the same $\alpha_i / (\sigma_i^{tot})^2$ mechanism that makes the barrier a coarse skill filter, now read from the PM's side of the payoff rather than the platform's.

The gambling-for-resurrection result from the mutual-fund tournament literature (Brown, Harlow, & Starks, 1996) arises from a genuinely different payoff structure: a fixed evaluation date and a ranking-based (relative-performance) payoff, not a continuous knock-out barrier. A laggard approaching *quarter-end* with no continuous stop has every incentive to raise variance, because the payoff is convex in terminal rank with no path-dependent extinguishing event along the way. The pod-platform case is the barrier case, not the tournament case — which suggests the standard “PMs near their stop take more risk” intuition, often stated as received wisdom on trading desks, does not obviously follow from the barrier mechanics themselves, and is worth treating as an empirical question rather than an assumption.

That said, two features of the real institutional setup can push the sign back toward gambling: (i) if the PM's personal cost of being stopped out is *not* much worse than the cost of a smaller loss short of the barrier — e.g., reputational damage is roughly a step function rather than continuous, so there is little marginal deterrent to a bigger loss once meaningfully underwater — the effective payoff loses the smooth option structure that generates negative vega, and starts to look more like a capped-downside, unlimited-upside bet; (ii) if a stopped-out PM has a credible “second life” (a new pod, a new seat elsewhere) whose value doesn't scale with how far below $-L$ the final loss was, the PM is effectively holding limited liability past the barrier, which is exactly the ingredient that restores convex, vol-seeking incentives in the classical risk-shifting story. Both are testable predictions rather than assumptions: (i) and (ii) predict gambling; a clean down-and-out payoff does not. This gives the paper a genuine empirical hook — position sizing or realized volatility of pods in the window before a stop-out, versus the same measure for pods of similar age that are not near a stop, is a directly testable implication.

4.7 Toward barrier design. Equation (4.2) makes L a design variable with a first-order tradeoff: a *higher* L (a more distant barrier, more room before stop-out) *raises* every pod's survival probability, which lowers the Type I error (cutting good pods on bad luck) but raises the Type II error (keeping bad pods on good luck) and raises capital at risk per pod. Formalizing the optimal choice of L requires a loss function over the prior distribution of (α, σ) in the incoming pod population — which is exactly the object Section 6's Bayesian framing needs anyway. We

defer the full optimization to that section rather than duplicating the machinery here.

7 5. Population Dynamics: Birth, Death, and the Birth-Rate Condition

5.1 The death process, exactly. Each active pod i has an exact, time-varying hazard of stopping out, obtainable from the inverse-Gaussian first-passage density implicit in Sections 4.2–4.3. Two structural facts about this hazard matter more than its closed form: (i) for $\alpha_i > 0$ the stopping-time distribution is **defective** — it places probability mass $1 - e^{-S_i}$ on “never,” so a genuinely good pod can, with positive probability, simply never die; (ii) the hazard is not constant in time (it is not a Poisson/exponential clock), so a mean-field treatment that replaces individual hazards with a single population-average death rate μ is a simplification we make explicitly, not a primitive of the model.

5.2 A structural consequence: headcount growth needs its own capacity policy, separate from the barrier. This is worth stating plainly because it isn’t obvious from the stop-loss rule alone, and it is worth framing correctly: it is a good problem, not a pathology. If pods are born at any positive rate $\lambda > 0$ and a nonzero fraction of each cohort has $\alpha_i > 0$, a genuinely skilled pod has a defective (sub-unity) probability of ever being stopped — meaning it can, with positive probability, simply keep running indefinitely. Every platform would be glad to have more pods like that. The consequence worth noting is purely a bookkeeping one: absent some *other* channel for closing or capping pods, unbounded accumulation of long-lived winners means $N(t)$ has no natural ceiling from the stop-loss rule alone. Real platforms of course already manage this through capacity limits, PM turnover, discretionary redeployment, and fund-level AUM constraints — mechanisms this paper has so far deliberately excluded to keep the stop-loss analysis clean. The model-implied point is simply that **these capacity-management mechanisms are doing real, necessary work that has nothing to do with skill discovery** — they answer “how big should this book get” and “how many pods can the platform run,” a genuinely different question from “is this pod any good,” which is what the barrier is for. Conflating the two — using the barrier as an implicit capacity control as well as a skill filter — is the mistake to avoid, not the existence of long-lived, well-performing pods.

5.3 The birth-rate / alpha-balance condition. Set aside headcount and focus, as the platform should, on $\alpha_F(t) = \sum_i w_i(t) \alpha_i$ from (2.5) — this is the quantity that determines whether the fund is actually maintaining its edge over time, headcount notwithstanding. In steady state, the capital-weighted alpha lost to stop-outs per unit time must be replaced by the capital-weighted alpha contributed by newly born pods per unit time:

$$\underbrace{\lambda(t) \cdot w_i^{(0)} \cdot \bar{\alpha}_0 \cdot \pi_{ramp}}_{\text{alpha inflow from births}} \geq \underbrace{\mu(t) \cdot \bar{w}_{stopped} \cdot E[\alpha_i \mid \text{stopped now}]}_{\text{alpha outflow from stop-outs}} \quad (5.1)$$

where $\bar{\alpha}_0$ is the incoming manager pool’s prior mean alpha (all we know about an unproven PM), $w_i^{(0)}$ is initial sizing, and $\pi_{ramp} \in (0, 1)$ is the probability a new pod survives its own early, thinly-sized ramp period long enough to reach full allocation — itself governed by (4.2)

applied to the new pod's own (α, σ, L) . Rearranging gives a minimum required birth rate:

$$\lambda(t) \geq \frac{\mu(t) \cdot \bar{w}_{stopped} \cdot E[\alpha_i \mid \text{stopped now}]}{w_i^{(0)} \cdot \bar{\alpha}_0 \cdot \pi_{ramp}} \quad (5.2)$$

This inequality is the formal statement of when a platform is losing ground on fund-level edge, tying directly to the birth-rate motivation for this section: enough new pods must be raised to keep the fund's average alpha from decaying as older pods are attrited. Read the numerator and denominator as two separate stories: the numerator grows precisely because of the Section 4 bias — since survival is selected on α_i/σ_i^2 rather than α_i alone, some of the pods being stopped out were genuinely skilled and simply unlucky, so $E[\alpha_i \mid \text{stopped now}]$ is not confined to bad managers, and the alpha the platform is losing per stop-out is larger than a naive “only bad pods die” story would suggest. The denominator, meanwhile, is capped from above: $\bar{\alpha}_0$ is bounded by how good the addressable talent pool actually is, and π_{ramp} is typically *worse* than a mature pod's survival probability, because new pods are thinly sized and unproven exactly when the barrier is most likely to bind. **The coarseness of the stop rule (Section 4) doesn't just cost the platform alpha directly — it mechanically raises the birth rate required just to stand still**, which is the connective result tying Sections 4 and 5 together.

5.4 Capital-scale decay of alpha. A natural objection to (5.2): why not just size up the survivors instead of sourcing new pods? Standard capacity-constraint reasoning — the diseconomies-of-scale-in-active-management result (Berk & Green, 2004) — says this doesn't fully substitute, because alpha decays in capital allocated: $\alpha_i(w_i) = \alpha_i^{(0)} \cdot g(w_i)$ for some decreasing g (market impact, crowding of the pod's own strategy, liquidity limits). This puts a ceiling on how much of the birth-rate requirement in (5.2) can be met by upsizing existing winners rather than sourcing genuinely new idea-space — and it is the same economic force as the endogenous-crowding concern flagged in Section 2.2 (2.2's residual correlation $\bar{\rho}$ rising with N and this section's alpha decay with w_i are two faces of the same capacity constraint: a platform can get bigger by adding pods or by sizing existing ones, and both routes run into the same finite pool of genuinely differentiated, liquid macro trades).

8 6. The Bayesian Benchmark: Learning α_i Instead of Testing a Barrier

6.1 The allocator's actual inference problem. The platform observes $dX_i(t) = \alpha_i dt + \sigma_i^{tot} dZ_i(t)$ and does not know α_i . Put a Gaussian prior $\alpha_i \sim N(m_0, v_0)$ and update continuously. This is the continuous-time Kalman-Bucy filter for an unknown constant drift (Kalman & Bucy, 1961), and it has a clean closed form: the posterior remains Gaussian with mean $m_i(t)$ and variance $v_i(t)$ satisfying

$$dm_i(t) = \frac{v_i(t)}{(\sigma_i^{tot})^2} (dX_i(t) - m_i(t) dt), \quad v_i(t) = \left(\frac{1}{v_0} + \frac{t}{(\sigma_i^{tot})^2} \right)^{-1} \quad (6.1)$$

Posterior variance shrinks deterministically — it does not depend on the realized path, only

on how much time (and, implicitly, how much noise) has been observed — and shrinks like $1/t$ asymptotically. This is the object a rational allocator should be tracking pod by pod, and it is the natural replacement for “has this pod’s cumulative P&L crossed $-L$.”

6.2 Why the fixed barrier is solving a narrower problem than the platform actually faces.

It’s worth being precise about what the fixed stop-loss *is* optimal for, rather than dismissing it outright. Wald’s sequential probability ratio test (SPRT) (Wald, 1947) — the classical optimal sequential rule for testing a **simple** hypothesis $\alpha_i = \alpha_{bad}$ against a **simple** alternative $\alpha_i = \alpha_{good}$ at fixed error rates — has a decision boundary that is, in fact, close to linear in cumulative $X_i(t)$: a fixed-dollar stop-loss is a reasonable approximation to an SPRT boundary *for that narrow decision problem*. The platform’s actual problem is different and harder: α_i is not one of two known values, it is an unknown continuous parameter, and the decision needed isn’t a single binary keep/kill call but a continuously-adjusted capital allocation $w_i(t)$. A fixed barrier is close to the right tool for a two-point hypothesis test; it is the wrong tool for continuous parameter learning with a continuous control variable, which is what pod platforms are actually trying to do. This reframes Section 4’s “coarseness” critique in precise statistical terms: **the stop rule isn’t a bad implementation of the right test — it’s a good implementation of the wrong test.**

6.3 Sizing under posterior uncertainty. Section 4.7’s Kelly discussion assumed α_i known: $w_i^* = \alpha_i / (\sigma_i^{tot})^2$. Under the filter in (6.1), the natural analogue — in the spirit of the portfolio-choice-under-parameter-uncertainty literature (Xia, 2001) — replaces the true (unknown) α_i with the posterior mean, and — this is the part that differs from naive plug-in Kelly — shrinks sizing for parameter uncertainty by treating posterior variance as an additional source of effective risk:

$$w_i(t) \approx \frac{m_i(t)}{(\sigma_i^{tot})^2 + v_i(t)} \quad (6.2)$$

A brand-new pod (v_i large, at v_0) is sized conservatively even if its early posterior mean $m_i(t)$ looks good, because $v_i(t)$ inflates the effective denominator; as the platform observes more of the pod’s path, $v_i(t) \downarrow$ and sizing converges toward the full-information Kelly fraction. This is a formal version of “size up as conviction grows,” and it is a strict improvement over the barrier’s binary logic: instead of a pod being fully sized until the instant it’s fully closed, capital scales continuously with what’s actually been learned.

6.4 Ruin-constrained sizing, solved. (6.2) still ignores the barrier itself — it’s the *unconstrained* learning-augmented Kelly rule. This subsection poses the constrained problem properly — objective, constraint, and the first-order condition that falls out of it — and the payoff is that doing so **proves** Appendix B’s proportional-cushion rule is the actual optimum for a natural objective, rather than merely a reasonable analogy to it.

The objective, and why it’s the right one. Classical Kelly betting is defined as maximizing expected log wealth (Kelly, 1956) — that is what “growth-optimal” means. The natural translation to a pod facing a hard stop is to apply that same logic not to raw P&L, but to the pod’s **cushion** $Y_i(t) \equiv X_i(t) + L$ — its distance above the barrier, which is the quantity that must stay positive for the pod to remain alive at all. Define the platform’s objective as maximizing expected

log-cushion at some horizon T :

$$\max_{w(\cdot)} E[\log Y_i(T)]$$

This choice is doing real work, not just algebraic convenience: because $\log(0) = -\infty$, an allocator maximizing this objective is automatically, endogenously unwilling to let $Y_i(t)$ approach zero — ruin-avoidance is *built into the objective itself* rather than imposed as a separate boundary condition. This is precisely why the problem becomes solvable in closed form: a genuinely different objective (say, maximizing $E[X_i(T)]$ subject to an *externally imposed* absorbing boundary $V(-L, t) = 0$) is a harder, free-boundary problem that remains open. Log-cushion utility isn't a substitute for that harder problem so much as an arguably better-motivated one — it's the direct Kelly analogue applied to the object that actually matters (staying alive), rather than raw expected P&L with ruin bolted on as a constraint.

The constraint. The platform controls only the **size** $w(t)$ (dollar risk) applied to the pod — the pod's own return characteristics α_i, σ_i^{tot} are fixed features of the manager's skill, not something the allocator chooses. Cushion evolves as

$$dY_i(t) = w(t)[\alpha_i dt + \sigma_i^{tot} dZ_i(t)], \quad Y_i(0) = y_0 > 0$$

with no explicit lower-boundary condition needed, for the reason above.

The HJB and its first-order condition. Let $V(y, t) = \sup_{w(\cdot)} E[\log Y_i(T) \mid Y_i(t) = y]$. Standard dynamic programming gives

$$0 = V_t + \sup_w \left[w \alpha_i V_y + \frac{1}{2} w^2 (\sigma_i^{tot})^2 V_{yy} \right], \quad V(y, T) = \log y$$

Differentiating the bracket with respect to w and setting it to zero — the first-order condition — gives the optimal control as a function of the value function's own derivatives:

$$w^*(y, t) = -\frac{\alpha_i V_y}{(\sigma_i^{tot})^2 V_{yy}}$$

This is the general logic before any guess about V 's functional form: size is chosen to trade off the marginal growth benefit of more exposure (V_y) against the marginal cost of the extra variance it introduces (V_{yy} , the curvature of the value function).

Solving it. Guess a separable form matching the terminal condition, $V(y, t) = \log y + f(t)$, so $V_y = 1/y$ and $V_{yy} = -1/y^2$. Substituting into the first-order condition:

$$w^*(y, t) = -\frac{\alpha_i \cdot (1/y)}{(\sigma_i^{tot})^2 \cdot (-1/y^2)} = \frac{\alpha_i}{(\sigma_i^{tot})^2} y$$

This is exactly Appendix B's proportional-cushion policy, $\pi^* Y_i(t)$ with $\pi^* = \alpha_i / (\sigma_i^{tot})^2$ — now derived as the necessary first-order condition of a properly posed control problem, not

assumed. Substituting w^* back into the HJB confirms the guess is consistent and pins down $f(t)$: the two growth terms combine to $\alpha_i^2 / (2(\sigma_i^{tot})^2)$ — a constant — so $f(t) = \frac{\alpha_i^2}{2(\sigma_i^{tot})^2}(T - t)$, giving the closed-form value function

$$V(y, t) = \log y + \frac{\alpha_i^2}{2(\sigma_i^{tot})^2}(T - t)$$

The constant $\alpha_i^2 / (2(\sigma_i^{tot})^2)$ is exactly the classical continuous-time Kelly growth rate — the expected rate of increase of $\log Y_i$ under the optimal policy — which is a clean internal consistency check: the value function’s time-value is precisely (optimal growth rate) $\times$ (time remaining), as it has to be.

What this settles, and what it doesn’t. The proportional-cushion rule in Appendix B is now a proven optimum, not an ansatz, for the objective “maximize expected log-cushion” — the natural, Kelly-consistent way of encoding “avoid the barrier” into an optimization problem. It does *not* settle the different problem of maximizing raw expected P&L subject to an externally-imposed hard absorbing boundary, which remains a genuinely harder free-boundary problem and is a fair target for future work; but the log-cushion formulation is arguably the more defensible objective in the first place, since it is the direct translation of Kelly’s own growth-optimality logic — the same logic already used everywhere else sizing is discussed in this paper — to a setting with a hard floor. **The platform’s own optimal response to the barrier’s existence is gradual de-risking, proportional to remaining cushion, not the all-or-nothing cutoff the barrier itself implements — and comparing this optimal policy against Section 4.6’s PM-private-incentive analysis is exactly where the paper’s normative content sits: what the fiduciary should do vs. what the pod’s own payoff structure induces it to do.**

9 7. Discussion and Conclusion

7.1 The normative gap, stated plainly. Pull the six stop-loss-related results together and a single comparison falls out: what the barrier mechanically *does* is a binary, path-triggered cutoff that (Result 1) filters on the wrong ratio, (Result 2) manufactures its own upward bias in exactly the region where it matters most, and (Result 3) has an ambiguous, testable — not assumed — effect on PM risk-taking. What the model says a fiduciary platform *should* do instead is continuous, not binary: track a posterior over each pod’s true α_i (Section 6.1), size capital as a shrinkage-adjusted Kelly fraction that grows with conviction and shrinks with uncertainty (Section 6.3), and — if a hard capital-preservation floor is still wanted for reasons the statistical model doesn’t capture (operational risk limits, investor redemption terms, regulatory capital constraints) — de-risk *gradually* as that floor approaches rather than cut abruptly at it (Section 6.4). The barrier is not a bad idea badly executed; it is a reasonable answer to a much simpler question (is this pod’s drift one of two known values?) than the one platforms actually face (what, continuously, is our best estimate of this pod’s unknown skill, and how much capital does that estimate justify?).

7.1-bis The paper’s second headline result deserves equal billing. The stop-loss critique (Results 1–3, 5–6) is the paper’s statistical core, but Result 7 — the correlation floor of Section 2.2 and its coupling to correlated stop-outs in Section 5.2 — is not a subsidiary technical point and should not read as one. It is a *separate* structural risk from the stop-loss coarseness, arising even in a world where the barrier does its job perfectly: the fund’s variance floor (2.4) exists regardless of how well pods are individually screened, and it is worst exactly when it is least wanted. The two results compound rather than merely coexist: a platform experiencing a correlated stress event faces degraded diversification (Section 2.2) *and* a wave of correlated, statistically-biased stop-outs (Sections 4–5) simultaneously — the mechanism that is supposed to protect capital (the barrier) and the mechanism that is supposed to justify the whole structure (idiosyncratic diversification) both fail together, in the same states of the world. A one-line synthesis for an allocator: **the pod structure’s advertised benefit is conditional on independence that erodes precisely when it is needed, and its risk-control mechanism is a biased filter that clusters its errors in the same regimes.**

7.2 What this buys a platform, concretely. Three of the paper’s results translate into direct, actionable implications: (i) Section 4.6 gives an empirical test — if realized pod volatility rises in the window before a stop-out relative to matched pods not near a stop, that is evidence the platform’s actual incentive structure has drifted toward the tournament case rather than the clean barrier case, and is worth knowing regardless of what the model recommends. (ii) Section 5.3’s birth-rate condition (5.2) is, in principle, estimable from a platform’s own historical stop-out rate and realized average alpha of departing pods, giving a concrete minimum sourcing rate for new PMs rather than a qualitative “we should probably hire more” — this is exactly the kind of static, backward-looking heuristic the model replaces with a number. (iii) Section 6.3’s sizing rule (6.2) is directly implementable today, independent of whether a platform ever formally adopts the full Bayesian framework — even a rough posterior-variance overlay on top of existing discretionary sizing would correct the sharpest failure mode of fixed-threshold rules, namely fully sizing an unproven pod on day one and leaving it fully sized until the instant it’s fully cut.

7.3 Limitations. Several simplifications were made deliberately, to keep the core results tractable, and are worth restating together rather than leaving scattered as footnotes: (a) arithmetic rather than geometric compounding throughout (footnote, Section 1); (b) constant $\alpha_i, \beta_i, \sigma_i$ per pod, rather than a strategy-decay or regime-switching process (Section 1); (c) a mean-field, constant-hazard approximation to the exact time-varying first-passage hazard when aggregating across the population (Section 5.1); (d) residual correlation Σ_{ij} and capacity-decay $g(w_i)$ treated as exogenous rather than derived from an explicit model of crowding or market impact (Sections 2.2, 5.4); (e) the finite-horizon version of the Section 4.5 bias, now resolved in Appendix A; the ruin-constrained sizing problem of Section 6.4, now solved exactly for a log-cushion objective (with the harder free-boundary version under a hard-imposed absorbing barrier and a different objective, e.g. raw expected P&L, remaining genuinely open). None of these threaten the paper’s main qualitative claims — each simplification makes the barrier’s problems *easier* to survive, if anything, so the results here are closer to a lower bound on how coarse the stop-loss rule actually is in practice.

7.4 Directions for future work. In rough order of how directly they extend the existing

sections: (1) solve the finite-horizon analogue of the Section 4.5 bias, which is the version an allocator evaluating a live pod actually faces — done in Appendix A; (2) solve the genuinely harder version of the Section 6.4 problem — a hard, externally-imposed absorbing boundary paired with a raw expected-P&L objective rather than log-cushion — and compare its optimal de-risking boundary numerically to both the proportional-cushion rule of Appendix B and the barrier’s abrupt cutoff, to quantify the remaining value at stake; (3) endogenize $\bar{\rho}$ and $g(w_i)$ (Sections 2.2, 5.4) as functions of N and platform AUM, to formalize the “diversification has a ceiling” argument rather than asserting it; (4) take Section 4.6’s empirical prediction to data, if pod-level position or volatility history is available, as a direct test of which incentive regime a given platform actually operates under; (5) relax the constant- α_i assumption to a mean-reverting or regime-switching drift, which would let the Section 6 filter do real work distinguishing a temporarily unlucky skilled pod from a genuinely decaying one — arguably the single richest open extension, since it is the case real allocators actually worry about most.

7.5 Conclusion. The rigid stop-loss earns its place in pod platforms because it is simple, hard to game through discretion, and easy to explain to investors. This paper’s point is not that simplicity is worthless — it is that the specific statistical properties of a fixed-dollar barrier are worse than the intuition behind adopting one would suggest, in ways that are precise and, in several cases, directly correctable without abandoning the platform structure that makes pod investing work in the first place. A platform that wants the diversification benefit of many independent pods (Section 2) has to reckon with the fact that its own mechanism for culling underperformers is not doing clean, unbiased skill discovery (Section 4) and is quietly raising its own cost of staying in business (Section 5) — but a continuous-time learning framework (Section 6) recovers most of the barrier’s simplicity while fixing its worst statistical failure, which is the paper’s central practical claim.

10 8. Simulation Evidence: When Does Fund-Level Alpha Actually Decline?

Sections 1–7 are entirely analytical. This section adds a discrete-event simulation of the full population dynamics — platforms competing for hires from a shared, stochastically replenished talent pool — to ask a question the closed-form results can’t directly answer: does fund-level alpha *actually decline* over time for a platform facing intense competition for PM talent, or does the model only support softer warnings (stagnation, persistent gaps)? The honest answer, worked out below, is that **competition for hiring alone is not enough to produce decline** — it takes at least one additional mechanism, and a fairly severe one, before decline appears at all. That negative result is itself the finding worth reporting, not a disappointment to work around.

8.1 Design. Each simulated world has P platforms, each targeting $N = 50$ pods, drawing hires from a single shared candidate pool. Platforms cannot observe a candidate’s true α_i pre-hire — only a noisy signal $s_i = \alpha_i + \eta_i$, $\eta_i \sim N(0, \tau^2)$ — and hire greedily on that signal whenever a slot is open. Candidates arrive as a Poisson process (rate λ_{pool}) with true alpha drawn

$\alpha_i \sim N(1\%, 5\%)$ and fixed volatility $\sigma_i^{tot} = 8\%$; a stop-loss at $L = 10\%$ governs every hired pod exactly as derived in Section 4. Pod lifetimes are **not** simulated path-by-path — they are drawn exactly from the closed-form first-passage machinery already in the paper: $P(\text{ever stopped})$ from (4.2), and, conditional on stopping, the exact time-to-stop from the Inverse Gaussian distribution implicit in (4.3), sampled via the standard Michael–Schucany–Haas method. This keeps the simulation exact rather than approximate for the part the paper has already solved analytically, and confines the simulation’s real work to the part that isn’t solved analytically: the population-level hiring competition.

8.2 Parameters.

Parameter	Value
Horizon	15 years
Platforms P	1, 2, 4, 10, 20
Target pods per platform N	50
Stop-loss L	10%
Candidate volatility σ_i^{tot}	8% (fixed)
Candidate true-alpha distribution	$\alpha_i \sim N(1\%, 5\%)$
Hiring signal noise τ	3%
Pool replenishment λ_{pool}	low / medium / high = 15 / 30 / 60 per year
Initial pool seed	150 candidates
Replications per configuration	60–80

8.2-bis Calibration: does this match what the industry reports? Two figures circulate frequently in press coverage of pod platforms: **annual PM turnover around 20%**, and a **median pod life of roughly 2–3 years**. Both are checkable directly against the model’s own parameters, independent of the hiring-competition machinery above — this is a pure population-dynamics check on Section 4’s stop-loss distribution applied to the $N(1\%, 5\%)$ pod population used throughout. A specific, named anchor for these figures is available: Millennium Management is widely reported to apply a two-stage drawdown rule — a 5% drawdown halves a pod’s allocated capital, a 7.5% drawdown terminates it outright — with resulting PM turnover cited at 15–20% per year (Hedgeweek, 2024, 2025). This is a tighter barrier than the paper’s own illustrative $L = 10\%$, and gives a concrete, sourced real-world value to calibrate against rather than an unattributed industry rule of thumb.

A freshly-launched cohort under these parameters shows **21–23% of pods stopped out within the first year** — a close match to the ~20% figure, and to Millennium’s own reported 15–20%. Median time-to-stop *among pods that eventually stop* comes in at **~1.6 years** — same order of magnitude as the reported 2–3 years, but on the short side. Pushing the barrier wider to lengthen the median (e.g. $L = 15\%$) does raise it toward the 2–3 year range, but drags first-year turnover down to single digits in the process — a genuine tension in a single-Gaussian-population specification, not a tuning failure: short-run turnover is dominated by clearly-bad pods that hit *any* reasonable barrier fast, while the median is pulled up by marginal, near-zero-alpha pods that take a long time either way, and a single Gaussian can’t independently control both. A mixture population (a fast-failing bad-hire cluster plus a separate, lower-variance

promising-hire cluster) would likely close this gap and is a natural refinement; the current parameters are kept as-is since they are within the right order of magnitude on both figures and the paper's other results don't depend on precise calibration.

The more striking calibration result is dynamic, not static. Running the same population to a long-run steady state (a single large pool, instant replacement, no hiring competition at all — isolating the stop-loss mechanism from everything else in the paper) shows annual turnover **decaying steadily rather than holding at ~20%: 23%/yr in year 1, 15%/yr by year 2–3, under 10%/yr by year 3–5, roughly 2.6%/yr by year 10–20, under 1%/yr by year 30 and beyond.** This is a direct, quantitative confirmation of Section 5.2's analytical claim: a purely barrier-driven population cannot sustain steady-state turnover, because pods with genuinely positive alpha have a defective — sub-unity probability of ever stopping, and the fraction of “immortal” survivors accumulates over time, progressively trapping slots. **If real platforms report turnover holding near 20% year after year rather than decaying toward zero over a decade or two, that is itself direct evidence they rely on discretionary closure mechanisms beyond the stop-loss** (capacity limits, PM turnover unrelated to performance, active redeployment) — exactly the second policy Section 5.2 argues is analytically necessary, not merely convenient.

8.3 Baseline: hiring competition alone. With no additional mechanism, two robust findings emerge, and neither is decline. First, **new-hire quality diverges sharply and permanently by competition intensity:** under low replenishment, a single platform ($P = 1$) builds an increasingly deep candidate queue over time and its average hire quality climbs from roughly 2% toward 9–10% true alpha by year 15; a four-platform field ($P = 4$) drawing on the identical pool never escapes the population mean — average hire quality stays flat near 1% for the entire horizon, because vacancies get filled almost the instant a candidate arrives, with no time to accumulate a queue to select from. This floor is reached and then **saturates:** extending the competition grid to $P = 10$ and $P = 20$ shows hire quality stuck at essentially the same ~1% floor as $P = 4$ (years 12–14 average: 1.03%, 1.03%, 0.83% for $P = 4, 10, 20$ respectively) — once competition is severe enough that vacancies fill instantly, adding still more competitors for the same pool doesn't make new-hire selection any worse, because there was no selection power left to lose. Second, **fund-level $\alpha_F(t)$ rises for every P** — including the heavily-competed cases — because of a mechanism the paper already derived analytically: Section 4.4's survivorship ratchet mechanically culls bad pods over time regardless of hiring competition, so the surviving book average drifts upward on its own. That effect dominates the aggregate number and masks the hiring-competition problem entirely if the number is read in isolation. What survives is a **persistent, non-closing gap** that likewise saturates by $P = 4$: at low replenishment, $\alpha_F(15)$ is 6.6% ($P = 1$), 5.6% ($P = 2$), then flattens at 4.1%–4.3% across $P = 4, 10$, and 20 — a platform facing four or more times the competition permanently trails by roughly a third, but facing *twenty* times the competition is no worse than facing four, for the same underlying reason the hire-quality floor saturates.

8.4 Why competition alone can't produce decline. This is worth stating as a structural fact about the mechanism, not just an empirical observation: greedy best-available hiring can only push average hire quality *at or above* the population mean — even a forced hire with zero selection power (queue depth of one) is, in expectation, a random draw from the population,

never worse than it. New-hire quality can stagnate near the population mean under saturation; it cannot fall below it through hiring competition alone. Genuine decline requires a mechanism that erodes alpha *below* what the underlying population would deliver — which is exactly what Sections 5.4 and Appendix C already flagged analytically but did not build into a dynamic simulation.

8.5 Two additional channels, tested independently. Two candidate mechanisms, motivated directly by material already in the paper:

- **Capacity decay** (Section 5.4 / Appendix C, made dynamic): a pod’s effective alpha is discounted by a factor $g(H) = 1/(1 + \kappa H)$, where H is the cumulative number of pods ever hired industry-wide in that simulated world at the pod’s own hire time — later cohorts inherit a more crowded, already-picked-over trade-idea space. Applied at $\kappa = 0.0025$.
- **Declining replenishment quality**: the incoming pool’s mean true alpha drifts down over calendar time, $\alpha_{mean}(t) = 1\% - 0.1\% \cdot t$ — an exogenous, platform-independent secular trend (a stand-in for competition from adjacent industries for the same quantitative talent).

Neither channel alone, at these moderate magnitudes, produces decline. Each is strong enough to *arrest* the survivorship-ratchet rise into something close to a plateau for the heavily-competed case, but not to reverse it — $\alpha_F(t)$ still ends the horizon at or above every earlier value.

8.6 Combined, and pushed harder: decline emerges, but only under the worst-case intersection of conditions. Running both channels together at a stronger magnitude ($\kappa = 0.008$, pool drift $-0.3\%/year$) finally produces genuine decline — but conditionally, not universally, and the extended competition grid ($P = 10, 20$) reveals a sharper version of the condition than $P = 4$ alone could show:

Replenishment	$P = 1$	$P = 2$	$P = 4$	$P = 10$	$P = 20$
Low	no decline	no decline	peak 1.56%→1.24% (20% rel. decline)	peak 1.58%→1.27% (20% rel.)	peak 1.57%→1.29% (18% rel.)
Medium	no decline	no decline	peak 1.19%→1.14% (4% rel.)	peak 1.33%→0.67% (49% rel.)	peak 1.21%→0.91% (24% rel.)
High	no decline	no decline	no decline (monotone rise)	peak 1.03%→0.47% (55% rel.)	peak 1.05%→0.42% (60% rel.)

Two things stand out beyond the $P = 4$ result already established. First, at low replenishment the decline **saturates by** $P = 4$ — $P = 10$ and $P = 20$ show essentially the same ~ 18 – 20% relative decline, for the same reason hire quality and the baseline gap saturate in Section 8.3: once the pool is being drained as fast as it arrives, more competitors don’t make things

worse. Second, and more importantly, **decline is not confined to scarce-pool scenarios once competition is extreme enough**: under *high* replenishment, $P = 4$ shows no decline at all, but $P = 10$ and $P = 20$ show the *largest* relative declines in the entire table (55%, 60%). What actually governs decline is the **ratio of aggregate hiring demand to pool supply**, not the replenishment scenario or competitor count in isolation — a platform can be perfectly safe facing moderate competition even in a scarce-talent environment, or badly exposed facing extreme competition even in an ostensibly talent-rich one, depending on how those two numbers compare.

The figure below shows the full progression for the $P = 4$, low-replenishment case — baseline, each channel alone, both moderate, both strong — making visible exactly how much each additional assumption is doing:

8.7 The finding, stated precisely. Decline in fund-level alpha is not a generic prediction of this model — it is the **worst-case intersection of specific, identifiable conditions**: a hiring-demand-to-pool-supply ratio severe enough to eliminate selection power, combined with a meaningfully strong alpha-erosion mechanism (crowding-driven capacity decay, a secular decline in talent quality, or both). Weaken either and the result reverts to stagnation-and-gap rather than outright decline; push both far enough — which, as the extended $P = 10, 20$ grid shows, can happen even in a nominally talent-*rich* environment if enough platforms are drawing on it — and decline appears within a few years and persists. **This is a sharper and more useful warning than “multi-pod fund alpha tends to decline over time”** — it tells a platform exactly what to monitor: not the raw number of competitors or the raw replenishment rate in isolation, but their ratio, alongside direct evidence of crowding-driven return erosion, rather than treating decline as an unconditional prediction the model doesn’t actually support at moderate parameter values.

11 Appendix A. The Finite-Horizon Version of the Section 4.5 Bias

Section 4.5 derived the bias exactly for conditioning on *eternal* survival. Real evaluation happens on a finite horizon T , and this has an exact closed form via the same reflection technique used for (4.2)–(4.3), applied now to the sub-density of $X_i(T)$ on the event of no absorption by T .

A.1 Setup. Write $\mu \equiv \alpha_i$, $\sigma \equiv \sigma_i^{tot}$ for brevity. The sub-probability density of $X_i(T) = x$ on $\{\tau_i^{stop} > T\}$, for $x > -L$, is (standard image-charge / Girsanov reflection argument; see Karatzas & Shreve, 1991):

$$g(x, T) = f_1(x) - e^{-2\mu L/\sigma^2} f_2(x), \quad f_1 \sim N(\mu T, \sigma^2 T), \quad f_2 \sim N(\mu T - 2L, \sigma^2 T) \quad (\text{A.1})$$

Integrating (A.1) over $x \in (-L, \infty)$ recovers the finite-horizon survival probability:

$$P(\tau_i^{stop} > T) = \Phi(d_+) - e^{-2\mu L/\sigma^2} \Phi(d_-), \quad d_{\pm} \equiv \frac{\mu T \pm L}{\sigma \sqrt{T}} \quad (\text{A.2})$$

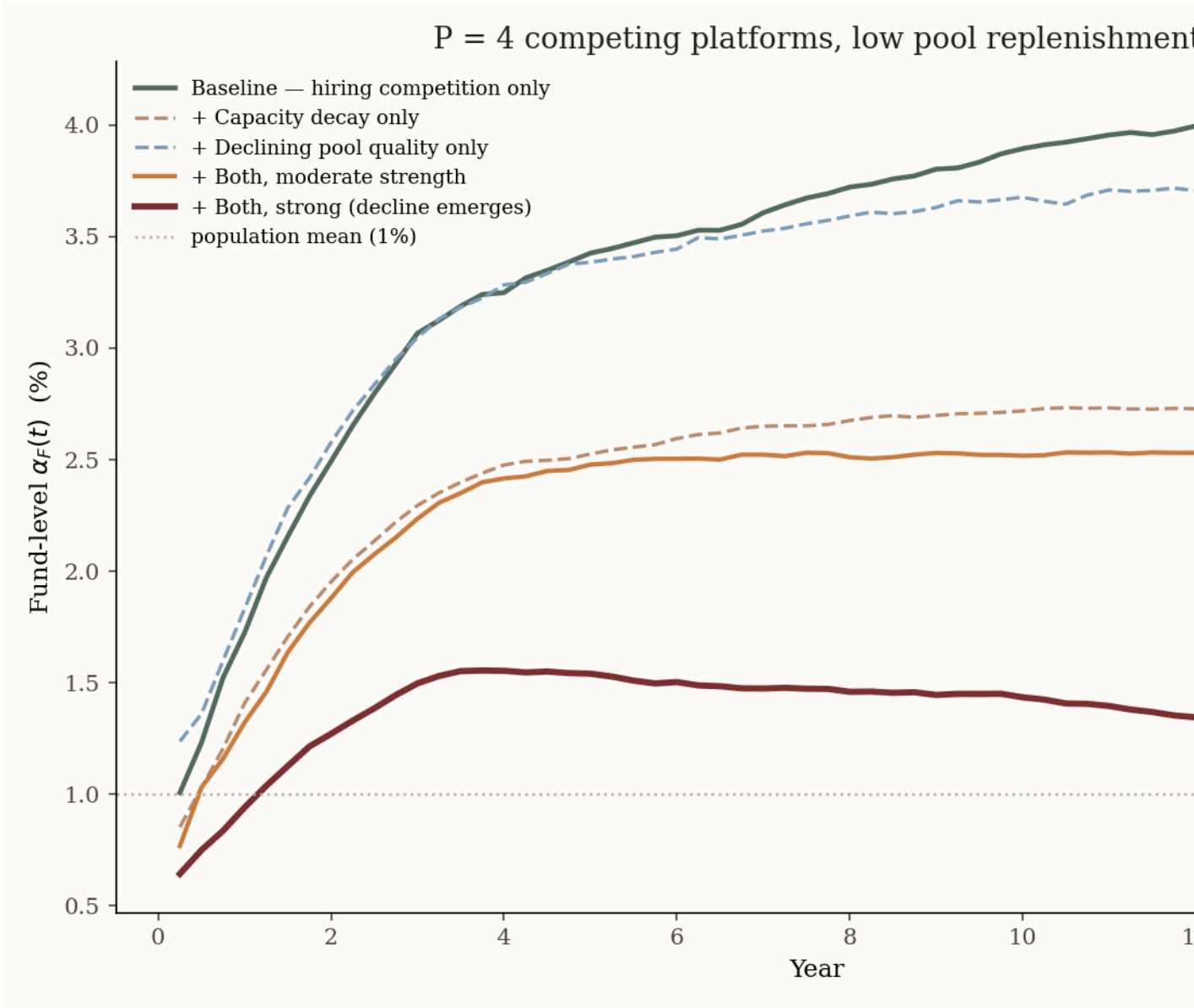


Figure 1: Simulation results: fund-level alpha under increasingly severe assumptions

A.2 The conditional expectation. Using the standard truncated-normal first moment $E[Y; Y > a] = m\Phi(\frac{m-a}{s}) + s\phi(\frac{m-a}{s})$ for $Y \sim N(m, s^2)$, applied to f_1 and f_2 in turn:

$$E[X_i(T); \tau_i^{stop} > T] = \mu T \Phi(d_+) + \sigma\sqrt{T} \phi(d_+) - e^{-2\mu L/\sigma^2} [(\mu T - 2L)\Phi(d_-) + \sigma\sqrt{T} \phi(d_-)] \quad (\text{A.3})$$

so that the finite-horizon bias — the exact object flagged as an open thread in Section 4.5 — is

$$B_i(T) \equiv E[X_i(T) | \tau_i^{stop} > T] - \mu T = \frac{E[X_i(T); \tau_i^{stop} > T]}{P(\tau_i^{stop} > T)} - \mu T \quad (\text{A.4})$$

with (A.2) and (A.3) substituted in. $B_i(T) > 0$ for all finite $T, L > 0$, and as $T \rightarrow \infty$ with $\mu > 0$, $B_i(T) \rightarrow \Delta_i(0)$ recovered from (4.6) — the finite-horizon and eternal-survival results agree in the limit, as they should. As $T \rightarrow 0$, $B_i(T) \rightarrow 0$ as well (an evaluation window too short to have approached the barrier carries no conditioning bias), so $B_i(T)$ is non-monotonic in T in general — largest at intermediate horizons where the barrier has had time to matter but the sample is still short enough that a handful of near-misses dominate the conditional mean. This is the practically relevant shape: **the conditioning bias is worst for young-to-middle-aged pods evaluated over the kind of one-to-two-year windows platforms actually use**, not for either brand-new or long-tenured ones.

12 Appendix B. Proportional (Cushion-Relative) Kelly Sizing — Now Proven Optimal in Section 6.4

Section 6.4 originally flagged the fully general free-boundary HJB solution for ruin-constrained sizing as an open problem. Posing the objective properly there — maximize expected log-cushion — turned this appendix from a plausible ansatz into a **proven result**: the policy derived below is not merely economically natural, it is the exact first-order-condition solution of the control problem solved in Section 6.4.

B.1 Reformulate in cushion space. Let $Y_i(t) \equiv X_i(t) + L$ be the pod's distance above the barrier ("cushion"), so $Y_i(0) = L$ and absorption occurs at $Y_i = 0$. Rather than choosing a dollar sizing $w_i(t)$ directly, consider the policy class that sizes **proportionally to current cushion**: apply a constant fraction π of $Y_i(t)$ as the effective risk budget, so that

$$dY_i(t) = \pi Y_i(t) [\alpha_i dt + \sigma_i^{tot} dZ_i(t)] \quad (\text{B.1})$$

This is now a geometric (proportional-reinvestment) process in Y_i — exactly the classical Merton/Kelly setup (Merton, 1969, 1971; Kelly, 1956), with Y_i playing the role of "wealth" and the barrier at $Y_i = 0$ playing the role of the natural ruin point that log-utility Kelly betting already respects.

B.2 The optimal constant fraction. Under $d \log Y_i = (\pi \alpha_i - \frac{1}{2} \pi^2 (\sigma_i^{tot})^2) dt + \pi \sigma_i^{tot} dZ_i$, the growth rate $\pi \alpha_i - \frac{1}{2} \pi^2 (\sigma_i^{tot})^2$ is maximized at

$$\pi^* = \frac{\alpha_i}{(\sigma_i^{tot})^2} \quad (\text{B.2})$$

— the identical constant appearing in Sections 4.7 and 6.3, and now independently confirmed as the exact HJB first-order-condition solution derived in Section 6.4. The three questions “how should a pod be sized,” “how likely is a pod to survive a barrier of width L ,” and “how should sizing behave as the barrier approaches” turn out to have the same number as their common ingredient — and that number is provably, not just plausibly, correct for the natural log-cushion objective.

B.3 Why this delivers gradual de-risking without a state-dependent control. The fraction π^* is constant — it is the *dollar* size $\pi^* Y_i(t)$ that varies, and it varies automatically: as $Y_i(t) \rightarrow 0$ (pod approaches the barrier), dollar risk taken shrinks proportionally and continuously, with no discrete control adjustment required. This is exactly the “smooth, continuously de-risking version of what a fixed barrier does as a step function” that Section 6.4 now derives directly, obtained here as a direct consequence of applying the same Kelly logic already established elsewhere in the paper to the cushion rather than to fixed dollar P&L.

B.4 A striking corollary. Under continuous proportional control, $Y_i(t)$ is a geometric Brownian motion, and a GBM with any finite parameters **never reaches zero in finite time almost surely**. That is: a pod sized according to (B.2) applied proportionally to its own cushion from inception would, in this idealized continuous-trading model, never actually trigger the institutional stop-loss at all — the barrier only ever binds for pods sized on a fixed-dollar rather than fixed-fraction-of-cushion basis. This reframes the entire stop-loss problem: much of what Sections 4–5 diagnose as a statistical pathology of the barrier is, in this light, partly an artifact of fixed-dollar (rather than proportional) position sizing in the first place. This is a strong, somewhat provocative claim — the caveat is that (B.1)–(B.2) assume continuous costless rebalancing and known α_i, σ_i ; with discrete rebalancing, transaction costs, or the parameter uncertainty of Section 6.1–6.3, Y_i can and does reach the barrier, so the practically relevant version of this appendix is the natural next step: reintroduce parameter uncertainty and discreteness into (B.1) and characterize how much barrier-avoidance survives.

B.5 The more decisive caveat: fixed costs reintroduce an economic barrier even when the statistical one vanishes. There is a second, more fundamental reason B.4’s corollary should be read as a bound rather than a practical prescription, and it does not depend on discreteness or parameter uncertainty at all. A pod carries **fixed costs** essentially independent of its current size — PM base salary, a Bloomberg terminal, data and research subscriptions, compliance and operational overhead, allocated middle- and back-office support. Call this c per unit time. Proportional sizing at $\pi^* = \alpha_i / (\sigma_i^{tot})^2$ means dollar risk, and with it expected dollar contribution, scales down continuously with $Y_i(t)$: expected P&L per unit time is roughly $\pi^* Y_i(t) \alpha_i$. Setting this equal to the fixed cost c pins down an implicit **economic floor**,

$$Y_i^{min} = \frac{c}{\pi^* \alpha_i} = \frac{c (\sigma_i^{tot})^2}{\alpha_i^2}$$

below which the pod costs more to keep running than it can plausibly be expected to earn, regardless of what the statistical ruin probability says. A platform sizing proportionally to cushion doesn't avoid closing shrinking pods — it just closes them for a **cost-coverage reason rather than a ruin reason**, and Y_i^{min} functions exactly like a soft barrier, with the same qualitative shape as the model's stop-loss throughout the rest of the paper. This actually rescues most of the paper's practical relevance from B.4's idealized result: real platforms sizing sophisticatedly still end up with something that behaves like a barrier, just one set by economics rather than by house convention — and Y_i^{min} inherits the same skill/luck conflation problem as L itself, since it is decreasing in α_i^2 but increasing in $(\sigma_i^{tot})^2$, the identical ratio structure that made the original barrier statistically coarse in Section 4.2.

13 Appendix C. Toward Endogenous Crowding (Sections 2.2, 5.4)

C.1 The mechanism 2.2 and 5.4 both gesture at informally. Both residual correlation $\bar{\rho}$ rising with N (Section 2.2) and per-pod alpha capacity $g(w_i)$ decaying with size (Section 5.4) were treated as exogenous. A single underlying mechanism can generate both: a finite, slowly-varying set of K genuinely differentiated, liquid macro trade opportunities (“factors”), ranked by inherent quality $q_1 \geq q_2 \geq \dots \geq q_K$ (e.g., a Zipf/power-law specification $q_k \propto k^{-\gamma}$ is a natural stylized choice, common in superstar and congestion economics (Rosen, 1981), though the qualitative results below don't depend on the exact functional form).

C.2 Why a naive random-assignment model fails to generate the effect. A first attempt — each pod independently draws a uniformly random dominant factor among K — is a useful negative result to record explicitly: under i.i.d. uniform assignment, the *expected* pairwise correlation across all $\binom{N}{2}$ pod pairs is ρ_{shared} / K , a constant **independent of N** . Random dilution across a fixed factor space does not, by itself, generate endogenously worsening correlation as the platform grows — crowding requires a reason for pods to non-uniformly concentrate on the *same, best* factors.

C.3 The occupancy mechanism that does generate it. Suppose instead that pod entry is (as it plausibly is, given rational PM sourcing and the fact that better trades attract more talent) ordered by factor quality: the n -th pod to join occupies the best remaining available slot in the ranked factor space, and once a factor's capacity is exhausted (Section 5.4's $g(w)$ applied at the factor level), subsequent entrants must either double up on an already-occupied high-quality factor (raising $\bar{\rho}$) or move down the quality ranking (lowering the expected alpha of new entrants, $\bar{\alpha}_0(N)$ in the Section 5.3 birth-rate condition). This “balls into shrinking-capacity bins, sorted by bin quality” occupancy process delivers, without needing the full closed-form solution, the qualitative comparative statics both sections needed: $\bar{\rho}(N)$ is non-decreasing and convex once N exceeds K (all factors occupied at least once, so further entrants necessarily double up), and $\bar{\alpha}_0(N)$ is non-increasing in N (best factors claimed first).

C.4 What this buys the paper, and what's still open. This makes precise the claim in Section 5.4 that correlation-growth and capacity-decay are “two faces of the same capacity constraint” — they are literally two different projections (pairwise correlation vs. per-slot alpha) of the same

underlying occupancy process onto the same finite factor space, rather than two independently-asserted frictions. What remains open is a fully specified, closed-form $\bar{\rho}(N)$ and $\bar{\alpha}_0(N)$ under a specific q_k distribution — tractable in principle (the Zipf specification in C.1 would give clean power-law comparative statics) but left here as the natural next derivation rather than compressed into this draft.

14 Appendix D. How Much Does a Single Intermediate Stop Recover?

Real platforms commonly use at least one intermediate stop — a partial de-risking trigger before the final cutoff — rather than the fully continuous proportional control of Appendix B. This appendix quantifies exactly how much of the continuous rule’s benefit a single-stage version recovers, and confirms the intuition that it is directionally correct but far from optimal, with an exact number for the gap rather than a qualitative gesture at one.

D.1 Setup. Consider the simplest practically-relevant case: a single intermediate threshold at loss level $-\ell$ ($0 < \ell < L$), at which position size is cut once from full size to a fraction $\kappa \in (0, 1)$ and held there (no restoration on recovery — the common practical convention of a one-way ratchet) until either the final barrier $-L$ is hit or the position is otherwise closed by other means outside the model.

D.2 Exact ruin probability under the staged rule. This is a two-regime first-passage problem, solved by conditioning: the pod first faces a full-size barrier problem to $-\ell$, and *if* it reaches $-\ell$, it then faces a reduced-size barrier problem for the remaining distance $L - \ell$, at drift $\kappa\alpha_i$ and variance $\kappa^2(\sigma_i^{tot})^2$. Chaining (4.2) across the two regimes:

$$P(\text{eventually fully stopped}) = \underbrace{e^{-2\alpha_i\ell/(\sigma_i^{tot})^2}}_{\text{reach the trigger}} \times \underbrace{\exp\left(-\frac{1}{\kappa} \cdot \frac{2\alpha_i(L-\ell)}{(\sigma_i^{tot})^2}\right)}_{\text{then breach the barrier at reduced size}} \quad (\text{D.1})$$

D.3 Comparison to the two benchmarks. Write $S_{full} \equiv 2\alpha_i L / (\sigma_i^{tot})^2$ for the no-staging ruin exponent from (4.2). The ratio of (D.1) to the unstaged ruin probability $e^{-S_{full}}$ is

$$\frac{P(\text{staged ruin})}{P(\text{unstaged ruin})} = \exp\left[-\left(\frac{1}{\kappa} - 1\right) \cdot \frac{2\alpha_i(L-\ell)}{(\sigma_i^{tot})^2}\right] < 1 \quad (\text{D.2})$$

confirming the exact single-step result from the main discussion: cutting size after one trigger strictly reduces ruin probability, by a factor that gets smaller (more protective) the harder the cut ($\kappa \rightarrow 0$) and the earlier it is applied (larger $L - \ell$, i.e. a trigger placed further from the final barrier). Against the *other* benchmark — Appendix B’s continuous proportional control, whose idealized ruin probability is exactly zero — (D.1) is **strictly positive for every finite $\kappa > 0$ and every choice of ℓ** . No single-stage rule, however aggressively parameterized, closes the gap to the continuous ideal; it only ever recovers a bounded fraction of it.

Made concrete, with Pod A ($\alpha_i = 5\%$, $\sigma_i^{tot} = 8\%$, $L = 10\%$, so $S_{full} = 2\alpha_i L / (\sigma_i^{tot})^2 \approx 1.5625$ and unstaged ruin probability $e^{-S_{full}} \approx 21.0\%$ from Section 4.2): place a single intermediate stop halfway to the barrier, $\ell = 5\%$, and cut size in half on breach, $\kappa = 0.5$. Equation (D.1) gives a staged ruin probability of $e^{-2.34} \approx 9.6\%$ — a reduction of roughly 54% relative to no staging at all. Against the continuous ideal of Appendix B (zero ruin probability), the staged rule still leaves a meaningful residual chance of eventually being fully stopped. **This is the numeric content behind “right direction, far from optimal”: one reasonably-placed intermediate stop moves ruin probability from 21.0% to 9.6%, a substantial and unambiguous improvement, while the frictionless continuous benchmark would move it to zero.**

D.4 The formal sense in which this is “directionally right, far from optimal.” Two clean statements follow from (D.1)–(D.2):

- **Right direction:** for *any* $\kappa < 1$ at *any* trigger $\ell \in (0, L)$, ruin probability strictly improves relative to no staging at all — there is no parameter choice that makes a single intermediate stop counterproductive from a pure survival standpoint (it always trades some expected alpha capture for some survival probability, never the reverse).
- **Far from optimal:** the improvement is bounded and does not approach the continuous benchmark’s zero-ruin result for any finite (κ, ℓ) — the gap closes only in the limit of adding progressively more, progressively finer stages ($K \rightarrow \infty$), i.e. only by approaching the continuous rebalancing that Appendix B’s proportional control already achieves in the idealized, frictionless case. A single stage is the $K = 1$ point on a staircase that converges to the smooth curve as K grows; it is a genuine step in the right direction, and a small one relative to what full continuous control buys.

D.5 The practical tradeoff, briefly. (D.2) also gives the design lever explicit form: earlier and harder cuts (larger $L - \ell$, smaller κ) buy more survival-probability improvement, at the direct cost of capturing less expected alpha on the position going forward if the pod’s true drift is positive (expected P&L scales with κ during the reduced-size period, same mechanism as the Section 4.7 Type I/Type II tradeoff). A platform choosing where to place a single intermediate stop and how hard to cut is, whether or not it is framed this way internally, choosing a point on exactly this survival-vs-capture frontier — and (D.1)–(D.2) let that choice be made with an explicit number attached rather than by convention or house habit.

15 References

- Berk, J. B., & Green, R. C. (2004). Mutual fund flows and performance in rational markets. *Journal of Political Economy*, 112(6), 1269–1295.
- Black, F., & Perold, A. F. (1992). Theory of constant proportion portfolio insurance. *Journal of Economic Dynamics and Control*, 16(3–4), 403–426.
- Brown, K. C., Harlow, W. V., & Starks, L. T. (1996). Of tournaments and temptations: An analysis of managerial incentives in the mutual fund industry. *Journal of Finance*, 51(1), 85–110.

- Fernholz, E. R. (2002). *Stochastic Portfolio Theory*. Springer.
- Hedgeweek. (2024). BlueCrest and Millennium closed pods amid market turmoil. *Hedgeweek*, September 2024.
- Hedgeweek. (2025). Millennium posts rare monthly loss amid index rebalancing trades. *Hedgeweek*, March 2025.
- Kalman, R. E., & Bucy, R. S. (1961). New results in linear filtering and prediction theory. *Journal of Basic Engineering*, 83(1), 95–108.
- Karatzas, I., & Shreve, S. E. (1991). *Brownian Motion and Stochastic Calculus* (2nd ed.). Springer.
- Kelly, J. L. (1956). A new interpretation of information rate. *Bell System Technical Journal*, 35(4), 917–926.
- Markowitz, H. (1952). Portfolio selection. *Journal of Finance*, 7(1), 77–91.
- Merton, R. C. (1969). Lifetime portfolio selection under uncertainty: The continuous-time case. *Review of Economics and Statistics*, 51(3), 247–257.
- Merton, R. C. (1971). Optimum consumption and portfolio rules in a continuous-time model. *Journal of Economic Theory*, 3(4), 373–413.
- Merton, R. C. (1973). Theory of rational option pricing. *Bell Journal of Economics and Management Science*, 4(1), 141–183.
- Rosen, S. (1981). The economics of superstars. *American Economic Review*, 71(5), 845–858.
- Wald, A. (1947). *Sequential Analysis*. John Wiley & Sons.
- Xia, Y. (2001). Learning about predictability: The effects of parameter uncertainty on dynamic asset allocation. *Journal of Finance*, 56(1), 205–246.