

Ambiguity Aversion in Pod-Platform Risk Allocation: Pricing Hidden Cross-Pod Correlation

Marcos Mollica

Draft — September 2, 2026

Abstract

Platform-style hedge funds allocate capital across many nominally independent trading pods, and size that capital using ordinary risk management: volatility estimates, correlation matrices, mean-variance sizing. The hard part is not the math but the input — correlation among pods is the one parameter a platform can least afford to get wrong and least reliably knows, since pods that look unrelated day to day can quietly become correlated through channels no individual pod's track record would ever reveal. This note asks a simple, practical question: does making portfolio-sizing decisions more robust to that kind of not-knowing — *ambiguity aversion*, in the language of the Hansen–Sargent robust-control framework — actually help a platform allocator, and if so, robust to what, specifically? The answer turns out to depend entirely on the second half of that question. We first try the natural specification: an ambiguity penalty that grows with estimated factor volatility. It is analytically well-behaved but, tested fairly against a simulated eight-year platform history, does nothing a plain volatility-scaled sizing rule was not already doing

— because the ambiguity penalty and the sizing rule are reacting to the same signal. A second specification, in which the ambiguity penalty responds instead to a realized measure of cross-pod correlation drift — a risk genuinely invisible to ordinary sizing, arising from a latent, unobserved common shock across pods — passes the same fair test decisively, improving risk-adjusted returns and drawdowns against both baselines in the overwhelming majority of simulated paths, and holds up under fat-tailed shocks, discrete jump risk, and a dynamic (HJB/BSDE) reformulation that confirms the simple rule is not just a good heuristic but the exact optimal one under the model’s assumptions. The general lesson survives the specifics: ambiguity aversion earns its keep only when it prices a risk the underlying portfolio machinery cannot already see, not a risk that machinery can see with a lag. We translate the result into a concrete operating rule — including when to cut capital uniformly versus hedge a specific factor — for an allocator who wants to use it directly.

Contents

1	Motivation	5
2	What Ambiguity Aversion Adds, in Plain Terms	6
2.1	Two different kinds of not-knowing	7
2.2	The adversary framing, and why it is more than a rhetorical device	8
2.3	Reading ψ as a number a risk desk can actually use	9
2.4	What stops an allocator from choosing an unboundedly large ψ ?	9
2.5	Why not just use a higher risk-aversion number?	12
3	Baseline Framework	14
4	First Specification: Volatility-Driven Ambiguity	20
4.1	Analytical shape	21
4.2	Simulated path and the leverage-matching problem	22
4.3	Diagnosis	25
4.4	A note on the absolute Sharpe levels	26
5	Second Specification: Correlation-Instability-Driven Ambiguity	28
5.1	A genuinely orthogonal risk source	28
5.2	Results	31
6	What This Means for an Allocator, in Practice	32
6.1	The one new rule, stated as an allocator would implement it	35
6.2	What “reduce capital deployed” actually means: uniform cut or targeted hedge	37
6.3	What this is not	40
7	Dynamic Formalization: HJB and BSDE Representation	41
7.1	Augmented state space	43
7.2	Robust HJB equation, and a subtlety CARA utility does not resolve for free	43
7.3	The reduced PDE, and what the optimal controls actually are	45
7.4	BSDE representation of φ	47
7.5	Numerical verification	48
7.6	Correlated drivers: recovering a genuine hedging demand	50
7.6.1	What that correlation actually is, and why it is not what a first guess suggests	52
7.6.2	Estimating the leverage effect properly	53
8	Robustness: Fat Tails and Jump Risk	56
8.1	Fat-tailed diffusive shocks	56

8.2	Discrete jumps	58
9	A Same-Day Early-Warning Layer	59
9.1	Design	60
9.2	Results	61
9.3	What this does and does not establish	62
10	Discussion	63
	References	65

1 Motivation

Platform-style hedge funds allocate capital across a portfolio of largely autonomous trading pods, each running its own book against its own risk limits and stop-loss rules. The platform allocator faces a distinct, higher-level problem: sizing aggregate exposure to the risk factors pods share, and sizing the pods themselves, under considerably more model uncertainty than any individual pod faces about its own book.

The asymmetry is worth stating concretely, because it is the entire reason this paper exists. A pod manager trading a single strategy for several years has a reasonably well-estimated idiosyncratic return distribution — enough daily observations to pin down volatility, tail behavior, and a rough Sharpe ratio with some confidence. The platform allocator’s problem is structurally harder, for two compounding reasons. First, estimating a covariance matrix across N pods requires, for a given level of statistical precision, an amount of data that grows with the *square* of N — the number of pairwise correlations to pin down grows quadratically while pod track-record length does not (this is the standard justification for shrinkage estimators such as Ledoit–Wolf; a platform allocator shrinking a sample covariance matrix toward a structured target is already implicitly acknowledging this problem). Second, and more consequentially, the correlation structure that matters most is precisely the one least visible in ordinary conditions: pods that look independent day-to-day can become corre-

lated through channels no individual pod's P&L history would ever reveal — a shared prime broker cutting margin across the street simultaneously, several pods unknowingly crowding into the same factor exposure via different instruments, or a common risk premium compressing across asset classes that look unrelated on paper. The August 2007 “quant quake,” in which a wide range of statistically-independent-looking equity market-neutral strategies unwound in the same week for reasons that had nothing to do with any of their individual models being wrong, is the canonical instance of exactly this failure mode. No pod's own risk report would have shown it coming; it was a property of the crowd, not of any individual book.

This is the practical case for treating correlation and factor exposure — not idiosyncratic pod risk — as the object the allocator should be most skeptical of their own model for. It is also, as it turns out, a substantially harder modeling problem to get right than intuition suggests, which is the empirical content of this note.

2 What Ambiguity Aversion Adds, in Plain Terms

This section is aimed at a reader comfortable with portfolio construction and risk management but not necessarily with the Hansen–Sargent robust-control literature. Readers already familiar with multiplier preferences can skip to [Section 3](#).

2.1 Two different kinds of not-knowing

Ordinary risk management already deals with uncertainty: a volatility estimate has a standard error, a correlation estimate has a confidence interval, and standard mean-variance sizing already shrinks exposure as estimated volatility rises. This is *parameter uncertainty* — you trust the model’s functional form (returns are driven by a factor with some covariance structure) and you are only unsure of the numbers that go into it. Widening a confidence interval or shrinking an estimator handles this kind of uncertainty adequately.

Ambiguity aversion is a response to a different, harder problem: *model uncertainty* — the possibility that the model’s functional form itself is wrong in a way you cannot currently see, not just that its parameters are imprecisely estimated. A platform allocator does not just have an imprecise estimate of the correlation between Pod A and Pod B; they have no model at all for the possibility that both pods are secretly exposed to the same crowded trade through instruments that look unrelated. That is not a wider confidence interval on a known parameter — it is a channel the model does not contain. Ambiguity aversion is the decision-theoretic machinery for acting sensibly when this second kind of not-knowing is present and cannot be estimated away with more data of the same kind you already have.

2.2 The adversary framing, and why it is more than a rhetorical device

The standard way to operationalize this (Hansen and Sargent, and the underlying multiplier-preference axiomatics of Anderson, Hansen and Sargent) is to imagine an adversary — sometimes called “nature,” sometimes an “evil twin” — who is allowed to distort the allocator’s assumed model, but only within a leash: the distorted model must remain statistically close to the estimated one, measured by relative entropy (Kullback–Leibler divergence). The allocator then optimizes for the worst case the adversary can construct within that leash, rather than for the point estimate.

This is not merely a pessimistic multiplier tacked onto the optimization; it has a specific, useful structure. Because the adversary is entropy-constrained, it cannot invent an arbitrarily bad scenario — it can only tilt probability toward outcomes that are *plausible given how much data the allocator actually has*. A model estimated from twenty years of daily data leaves the adversary very little room to construct a damaging alternative; a model estimated from eight months of a new pod’s live track record leaves it a great deal of room. Ambiguity aversion, done this way, is self-calibrating to the amount of evidence behind the model being distrusted — which is exactly the property a platform allocator wants when reasoning about correlation risk built from short histories and rare events.

2.3 Reading ψ as a number a risk desk can actually use

The mathematical object doing the work is ψ , the ambiguity penalty: the smaller ψ , the more the allocator trusts the estimated model and the closer the decision is to ordinary expected-utility optimization; the larger ψ , the more weight is placed on the worst plausible alternative. Stated this way, ψ can look like an arbitrary dial. It does not have to. Anderson, Hansen and Sargent's own preferred calibration device is the *detection-error probability*: given the entropy constraint implied by a choice of ψ , how many years of data would it take a statistician to reliably tell the adversary's worst-case alternative apart from the allocator's own estimated model, at a stated confidence level? A very large ψ corresponds to an alternative that could be rejected in a matter of weeks — implausibly cautious. A very small ψ corresponds to an alternative that would take decades to distinguish — effectively no robustness at all. This converts an abstract entropy penalty into a question every risk desk already asks informally: *how many years of history would it take before I'd trust this correlation estimate?* That question, not the entropy mathematics behind it, is the actual calibration target in what follows.

2.4 What stops an allocator from choosing an unboundedly large ψ ?

A natural worry at this point, and worth addressing directly rather than leaving implicit: nothing in the entropy mathematics itself prevents a risk manager

from simply picking an enormous ψ . Looking back at $b^*(t) = \hat{\mu}_F(t) / [(\gamma + \psi_F(t))\hat{\sigma}_F(t)^2]$, as $\psi \rightarrow \infty$, $b^* \rightarrow 0$ — the framework is perfectly willing to output “take no risk at all,” and there is no term in the equation that objects. Entropy disciplines how bad an alternative model the *adversary* is allowed to construct, given a chosen ψ ; it says nothing about how large ψ itself may be. That is a different knob, and it is exactly as free as the ordinary risk-aversion coefficient γ is in plain expected-utility theory — $\gamma \rightarrow \infty$ produces the identical degenerate outcome. Ambiguity aversion does not introduce this failure mode; it inherits it, and pretending otherwise would be a weakness in the argument, not a strength.

Three distinct things discipline it, and they are worth keeping separate because they operate differently.

Economic discipline. This is the substantive answer, and it is already implicit in Section 2.3’s detection-error-probability calibration, worth making explicit here: a choice of ψ is not free-floating, because it implies a specific worst-case alternative model with a checkable property — how many years of the allocator’s own data it would take to statistically reject that alternative. An unboundedly large ψ corresponds to insuring against an alternative that is already dead on arrival, rejectable in a matter of weeks by data the allocator already has. That is not extreme caution; it is spending capital to defend against a possibility the allocator’s own evidence has already ruled out. This gives a concrete, falsifiable

challenge to put to anyone proposing an extreme ψ : state the detection-error probability it implies, and defend why that many years of contradicting evidence would still leave you unconvinced. A bare risk-aversion number offers no equivalent challenge — “how risk-averse should we be” is not answerable in a checkable way, while “how implausible is the scenario you are insuring against” is.

Mathematical discipline. Separately from the economic argument, risk-sensitive control problems of exactly this type can fail to be well-posed once the ambiguity parameter passes a critical threshold: the minimizing adversary can, in the sufficiently extreme regime, find a distortion driving the objective to $-\infty$, and the value function ceases to exist (Whittle, 1990, discusses these breakdown points at length in the linear-quadratic risk-sensitive setting this note’s machinery descends from). The model does not degrade smoothly toward “zero risk forever” as ψ grows without bound; past a point it can simply stop returning an answer. That is a structural guardrail, distinct from the economic one, though in this note’s specific linear-Gaussian and CARA settings we have not characterized where that threshold falls.

Empirical discipline. This is the argument this note’s own results make hardest to ignore. More caution is not monotonically protective in the simulations reported later: the vol-only constant- ψ_F policy — more conservative than plain mean-variance by construction — performs *worse* than the plain policy on both

Sharpe ratio and, under severe jump risk, tail drawdown (Sections 5 and 8.2; CVaR at the worst 1% of paths is -47.4% for the vol-only-cautious policy against -41.3% for the plain one). Being defensive against the wrong signal, or a signal the allocator’s model already fully prices, does not merely cost forgone return in expectation — in this note’s own numbers, it can actively worsen the outcome the extra caution was meant to protect. That is a sharper argument against reflexive over-conservatism than an appeal to opportunity cost: miscalibrated ambiguity aversion is not reliably even safer.

None of this pins down an optimal ψ in closed form — that would require a genuine loss function trading off the cost of ambiguity against the cost of being wrong about it, which is outside this note’s scope. What it establishes is that “more robust is always safer” is false on its own terms, and that a platform choosing ψ has a real, checkable question to answer — the detection-error probability its choice implies — rather than a free dial to turn as far as institutional risk appetite permits.

2.5 Why not just use a higher risk-aversion number?

A natural objection, and the right one to raise before any of the machinery above is worth building: if the concern is simply that the allocator might be wrong, why not just be more conservative — run a higher risk-aversion coefficient γ across the board — rather than introduce a separate ambiguity framework? Section 3 shows this objection is sharper than it first appears, and an-

swering it honestly is the organizing thread of this note. With a constant ambiguity penalty, the two really are indistinguishable: raising ψ at a fixed level is observationally identical to raising γ (a known result, due to Maenhout, 2004). Our first attempt at a state-dependent version of ψ — one that responds to estimated volatility — turns out, once tested fairly, to add nothing beyond what a volatility-sensitive risk-aversion term already provides, for a precise and instructive reason given in Section 4. The version that does add something is the one built around a risk the ordinary volatility-based machinery cannot see at all — which is the correlation-instability specification in Section 5, and is, in the end, the actual answer to this objection.

The underlying reason is worth stating precisely, because it is the logical spine connecting Sections 4 and 5. Ordinary mean-variance sizing already conditions fully on estimated factor volatility: b^* is a function of $\hat{\sigma}_F$ by construction. An ambiguity penalty that is *also* a function of $\hat{\sigma}_F$ — however cleverly shaped — is therefore a distortion of information the optimal control has already fully absorbed. It cannot improve the decision; it can only relabel how cautious the existing decision rule is, which is precisely the Maenhout equivalence. Correlation instability across pods is a different kind of object entirely: the baseline mean-variance machinery, as built, has *no control that responds to it at all*. Neither b^* nor an unmanaged idiosyncratic book size reacts to a rise in cross-pod correlation, because the underlying portfolio-construction problem was never given a signal that measures it. An ambiguity penalty keyed to correlation instability

is therefore not competing with an existing information channel — it is the *only* channel through which that risk gets priced into sizing at all. This is the general principle the rest of the note is built on: ambiguity aversion changes outcomes only when it prices something the underlying optimization was structurally blind to, not when it re-weights something the optimization could already see.

3 Baseline Framework

Let pod i 's return follow

$$dR_{i,t} = (\alpha_i + \beta_i \mu_{F,t}) dt + \beta_i \sigma_F dW_{F,t} + \sigma_i dW_{i,t},$$

with α_i the pod's (uncertain) skill, β_i its loading on a common factor F , and $W_{i,t}$ idiosyncratic noise independent across pods and independent of $W_{F,t}$. The platform allocator chooses aggregate factor exposure $b := \sum_i w_i \beta_i$, where w_i is capital allocated to pod i .

Why this particular optimization problem

Section 2 motivated ambiguity aversion as choosing a policy that performs well against an adversary constrained to statistically plausible alternatives. Here is that intuition made precise. An allocator who simply trusted their estimated model P would solve the ordinary expected-utility problem $\sup_w \mathbb{E}^P[U(X_T)]$ — optimal *given that P is correct*, and silent on what happens if it is not. The robust

version replaces the single model P with the worst model the allocator cannot statistically rule out, and optimizes against that: $\sup_w \inf_{Q \in \mathcal{Q}} \mathbb{E}^Q[U(X_T)]$, where $\mathcal{Q}$ is the entropy-constrained neighborhood of P from Section 2. This is a genuine two-player structure, not a rhetorical device: the allocator moves first, choosing w ; an adversary — “nature” — then selects, from the neighborhood of models the allocator’s own data cannot rule out, whichever Q makes that choice of w look worst. The allocator therefore ranks policies by their *worst case* within the plausible set, not their expected performance under a single trusted guess.

The constrained problem (Q ranging over a hard entropy ball, $R(Q\|P) \leq \eta$) is awkward to solve directly. The standard device — due to Hansen and Sargent, and the reason this is called a *multiplier* preference — is to relax the hard constraint into a penalty via a Lagrange multiplier $1/\psi$, trading the constrained minimization over a fixed-size neighborhood for an unconstrained one that pays a price per unit of entropy spent:

$$\sup_w \inf_Q \left\{ \mathbb{E}^Q[U(X_T)] + \frac{1}{\psi} R(Q\|P) \right\}.$$

For every η there is a ψ making the two problems equivalent, so nothing is lost in this relaxation — it simply converts a constrained minimax problem into a form with clean first-order conditions, which is what makes the machinery in the rest of this note tractable. Under the Hansen–Sargent multiplier preference,

the allocator solves

$$\sup_w \inf_{\theta} \mathbb{E}^{\theta}[U(X_T)] + \frac{1}{2\psi} \int_0^T \theta_t^2 dt,$$

where θ indexes drift distortions of the factor process, entropy-penalized at rate $1/\psi$ — the entropy cost is *added* to what the minimizing adversary faces, since a penalty that could instead be subtracted would let the adversary profit from distorting further without limit, leaving the inner minimization unbounded below.

What this entropy penalty actually is, and why it takes this exact form

Section 2 introduced the entropy penalty conceptually, as a leash on how far an adversarial alternative model is allowed to stray from the allocator’s own estimate. It is worth being precise about what that leash actually measures, because the specific quadratic form above is not a modeling choice — it is forced by the mathematics of the setting, and knowing why clarifies what the penalty can and cannot represent.

The object itself. The leash is *relative entropy* (Kullback–Leibler divergence) between two probability measures over the whole path of the process: if P is

the allocator's estimated model and Q is a candidate alternative,

$$R(Q\|P) = \mathbb{E}^Q \left[\log \frac{dQ}{dP} \right].$$

This is not a distance in the everyday sense — $R(Q\|P) \neq R(P\|Q)$ in general, and it satisfies no triangle inequality — but it is always non-negative, zero only when $Q = P$, and it grows as Q becomes harder to reconcile with data actually generated under P .

Why it collapses to $\frac{1}{2}\theta^2$ here. If Q is obtained from P by the drift distortion θ_t used throughout this note ($dW_t = dW_t^Q + \theta_t dt$), Girsanov's theorem gives the Radon–Nikodym derivative over $[0, T]$ as the stochastic exponential $\frac{dQ}{dP}|_T = \exp(\int_0^T \theta_t dW_t - \frac{1}{2} \int_0^T \theta_t^2 dt)$. Taking logs and the expectation under Q — using that the stochastic-integral term has zero Q -expectation once rewritten in terms of dW^Q — gives

$$R(Q\|P) = \frac{1}{2} \mathbb{E}^Q \left[\int_0^T \theta_t^2 dt \right].$$

This is exactly the penalty term above, with $1/\psi$ the price per unit of entropy. There is no other functional form available for a drift-distorted diffusion: any absolutely continuous alternative to P obtained this way has precisely this entropy cost, which is why the same $\theta^2/(2\Psi)$ structure reappears unchanged for every drift distortion used later in this note (Section 7's $\theta_F, \theta_I, \theta_C$ included), rather than being re-derived case by case.

Why entropy, specifically. This is the deeper reason the quantity is the right one to constrain, not merely a convenient one. Relative entropy is the rate function governing statistical distinguishability between P and Q : by Sanov’s theorem and Stein’s lemma, the probability of failing to reject Q in favor of P after T periods of data decays at rate $\exp(-T \cdot R(Q\|P))$. In other words, $R(Q\|P)$ is not an abstract penalty on “how different” an alternative model is — it is, up to the rate at which evidence accumulates, literally how many years of data it would take to tell Q apart from P . That is exactly the quantity Section 2’s detection-error-probability calibration is built on: a choice of ψ implies a maximum entropy the adversary can spend, which translates directly into a number of years of history required to reject the adversary’s worst case. The entropy penalty and the detection-error probability are not two separate devices — the second is a restatement of the first in units a risk desk can reason about.

Why this direction, and what it means for the platform problem. The convention — entropy of Q relative to P , not the reverse — is also not arbitrary. It measures how surprised the allocator, trusting P as their base case, would be to discover the world is actually Q ; the constraint is calibrated against the allocator’s *own* data, not an outside observer’s. That is the right direction for a decision-maker reasoning about their own model uncertainty, and it is exactly what makes the mechanism self-calibrating to the evidence behind the model being distrusted, as Section 2 already argued informally: a platform

with twenty years of pod history constrains the adversary to a small-entropy neighborhood almost automatically, while a platform relying on eight months of a new pod’s track record leaves the same ψ a much wider berth, without any manual adjustment. This is the precise mechanism behind the correlation-instability specification’s central property in Section 5: the same entropy machinery that penalizes an implausible factor-drift alternative in this section is, unchanged, what makes the correlation-instability overlay’s caution scale with how little is actually known about cross-pod correlation, rather than with an arbitrarily chosen risk-aversion knob.

In the linear-Gaussian certainty-equivalent version relevant here, this yields the closed-form target

$$b^*(t) = \frac{\hat{\mu}_F(t)}{(\gamma + \psi_F(t)) \hat{\sigma}_F(t)^2},$$

where γ is ordinary risk aversion and $\psi_F(t)$ is whatever ambiguity penalty the allocator applies to factor exposure. (“Certainty-equivalent” here means the myopic, single-period reduction of the dynamic problem — the sizing rule a one-period mean-variance optimizer would choose given the currently estimated $\hat{\mu}_F, \hat{\sigma}_F$, as opposed to the fully dynamic solution that also accounts for how those estimates might evolve. Section 7 makes this distinction precise and asks whether the fully dynamic solution actually differs from it.) With ψ_F constant, this is observationally equivalent to raising γ (Maenhout, 2004) — ambiguity aversion collapses into plain risk aversion and adds no distinguishable

behavior.

The natural next step is to let ψ_F respond to the allocator’s confidence in the factor model, and estimated factor volatility is the natural proxy for that confidence, for a concrete reason: turbulent periods are exactly when a model is most likely to be misspecified, not just noisily estimated. Calm-regime parameters are typically fit on calm-regime data; a regime shift, a structural break, or a crowded-trade unwind tends to arrive precisely when realized volatility is already elevated, so a model that looked adequate in normal conditions is least trustworthy exactly when $\hat{\sigma}_F$ is high. Tying ψ_F to $\hat{\sigma}_F$ therefore encodes a specific, testable hypothesis — that model risk rises with realized volatility — rather than an arbitrary functional choice:

$$\psi_F(t) = \psi_0 + \kappa \hat{\sigma}_F(t).$$

4 First Specification: Volatility-Driven Ambiguity

Throughout this section, “leverage” refers to b_t itself: aggregate notional factor exposure per unit of capital deployed. $b_t = 1$ means the fund’s factor exposure equals its capital (unlevered); $b_t = 3$, the leverage cap used below, means notional factor exposure of three times capital. This is the standard usage for a single-factor book and is distinct from gross/net leverage across a multi-position portfolio, which does not arise here since b_t is already the fund’s

aggregate, netted exposure to the one factor being sized.

4.1 Analytical shape

Figure 1 plots b^* against estimated annualized factor volatility for three policies: plain mean-variance ($\psi_F \equiv 0$), constant-ambiguity ($\psi_F \equiv \psi_{\text{const}}$), and state-dependent ambiguity ($\psi_F(\hat{\sigma}_F)$ as above). The state-dependent curve sits below the constant curve everywhere, and the ratio between them widens monotonically with volatility — from roughly 0.42 at low vol to roughly 0.11 in deep stress, in our calibration. The comparative-statics mechanism we set out to build is present and correctly signed.

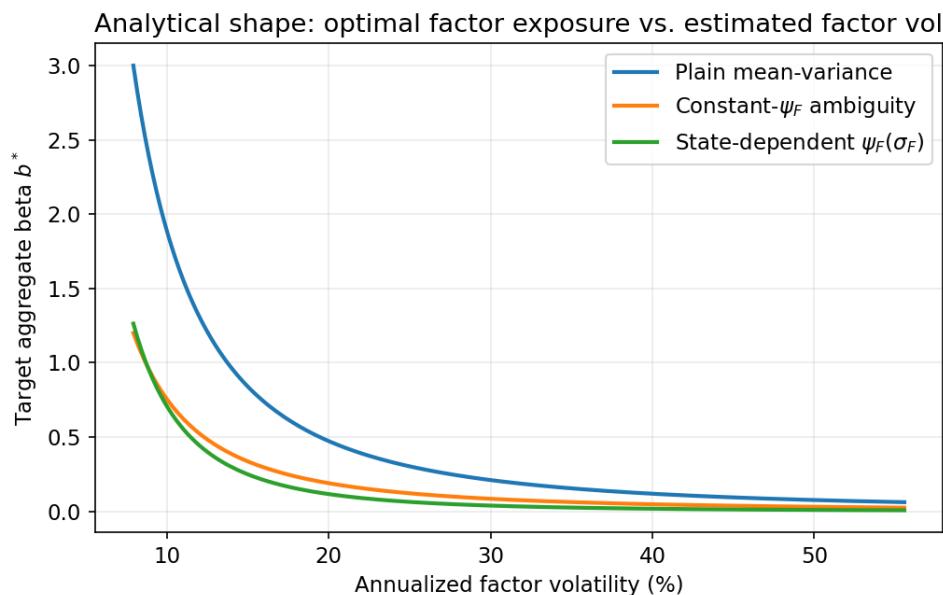


Figure 1: Optimal aggregate factor beta as a function of estimated factor volatility, for the three policies. All three fall in volatility; the state-dependent policy falls fastest.

4.2 Simulated path and the leverage-matching problem

Simulating a full path (regime-switching factor volatility) surfaces two problems invisible in the pure comparative statics, and led us to a genuine recalibration worth explaining rather than glossing over. Our first attempt used a 20-day trailing window for $\hat{\mu}_F, \hat{\sigma}_F$ — a natural first choice, but one that let the sizing rule time the regime switch almost perfectly: a 20-day window detects the stress regime’s volatility jump quickly enough that the dynamically-sized factor bet alone achieved an annualized Sharpe ratio above 3, against roughly 0.2 for a fixed, unsized exposure to the same factor. That gap is far too large to be a realistic reflection of achievable volatility timing (compare, e.g., Moreira and Muir’s more modest documented gains from volatility-managed portfolios) and reflects the regime structure being unrealistically easy to detect at that window length, not a genuine feature of ambiguity aversion. We therefore lengthen the estimation window for $\hat{\mu}_F, \hat{\sigma}_F$ specifically to 60 trading days (roughly one quarter) throughout the rest of this note — a defensible real-world risk-estimation horizon that meaningfully degrades the sizing rule’s ability to time the regime, cutting leverage-cap saturation from 21% of days to under 4% and bringing the achieved Sharpe ratios below into a high-but-plausible range for a well-run platform rather than an implausible one. The correlation-instability signal $\hat{C}_t$ introduced in Section 5 keeps a shorter, 20-day window deliberately, for a reason stated there: crowding episodes are short-lived, and a slow estimator tuned for macro factor risk would be too sluggish to catch them.

With this recalibration, two problems remain visible, and are worth reporting as found. First, the plain mean-variance policy is still the most leverage-sensitive of the three: with $\psi_F = 0$, b^* scales as $1/\hat{\sigma}_F^2$, and in calm regimes this makes the policy more sensitive to fluctuations in the estimated mean $\hat{\mu}_F$ than either ambiguity-aware policy, occasionally saturating the ± 3 leverage cap (Figure 2). Second, and more importantly, the constant-ambiguity and state-dependent policies track each other closely across the whole path — at the calibration required for $\psi_F(\hat{\sigma}_F)$ to stay well-behaved, the state-dependence adds little visible separation from the constant baseline at realistic volatility levels.

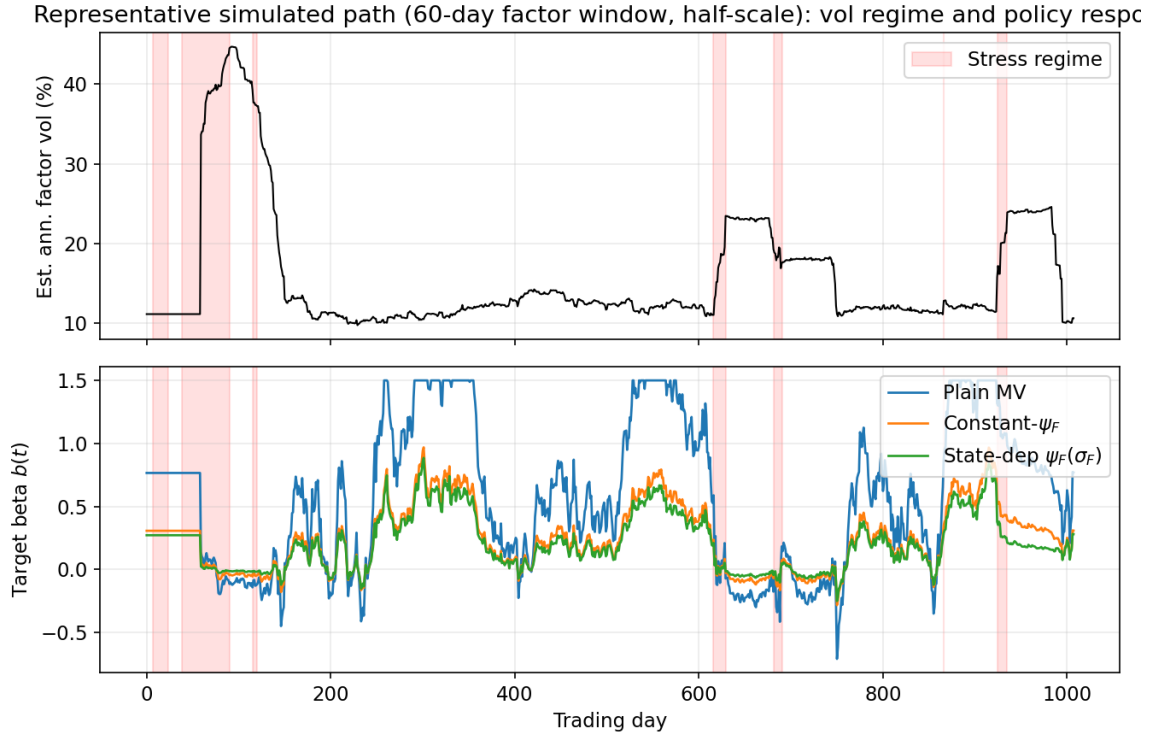


Figure 2: Representative simulated path (60-day factor-estimation window). Top: estimated annualized factor volatility, with stress-regime periods shaded. Bottom: target beta under the three policies. The constant- and state-dependent-ambiguity policies are close to indistinguishable.

A naive comparison of unconditional performance statistics favors the ambiguity-aware policies, but this comparison is not fair: the ambiguity policies run substantially lower average leverage than the plain policy by construction, so lower realized volatility and drawdown follow mechanically from taking less risk, not from any advantage in *when* risk is taken. We therefore leverage-match all three policies — rescaling each to the same average absolute beta as the plain policy — before comparing outcomes. Table 1 and Figure 3 report the results across 300 simulated eight-year paths.

Policy	Return (%)	Vol (%)	Sharpe	Max DD (%)
Plain mean-variance	21.4	10.8	1.98	−10.1
Constant- ψ_F (lev.-matched)	21.5	10.8	1.99	−10.2
State-dep. $\psi_F(\hat{\sigma}_F)$ (lev.-matched)	20.2	10.8	1.87	−10.8

Table 1: Leverage-matched performance (% , annualized), mean across 300 simulated seeds. State-dependent ambiguity does not improve on either baseline.

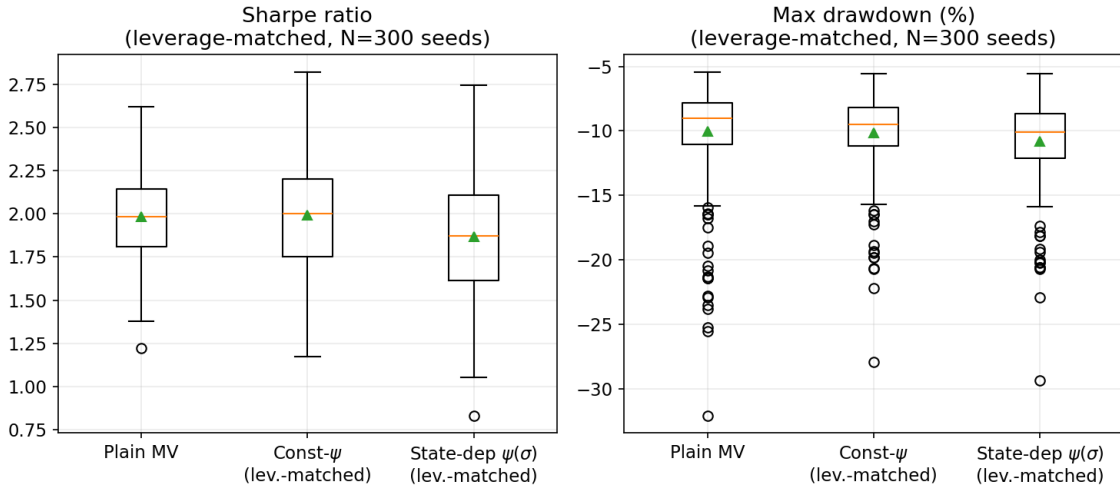


Figure 3: Distribution of Sharpe ratio and max drawdown across 300 seeds, leverage-matched. State-dependent ambiguity shows no advantage; it beats constant- ψ_F on Sharpe in only 2% of seeds.

Once risk levels are held fixed, the state-dependent policy shows no advantage over either baseline, and never beats the constant-ambiguity policy on Sharpe ratio across any of the 300 simulated paths. Tail metrics (1% CVaR, worst single day) tell the same story.

4.3 Diagnosis

The failure has a clean explanation. Leverage-matching rescales each policy to the same *time-averaged* exposure, and calm days vastly outnumber stress days in the simulated regime process. Because the state-dependent policy’s distinctive behavior is concentrated in the (rare) stress days, matching its average exposure to the other policies’ forces its calm-day exposure *up* to compensate — giving back, in ordinary market conditions, exactly the caution it was designed to add in stress.

The deeper issue is one of information content, not normalization. $\psi_F(\hat{\sigma}_F)$ is a function of $\hat{\sigma}_F$, the same variable that already drives b^* through the $1/\hat{\sigma}_F^2$ mean-variance term. Composing an ambiguity penalty with the volatility estimate does not introduce a new signal; it reweights an existing one. A referee’s objection — that this is relabeled risk aversion rather than a genuine ambiguity effect — would be correct.

4.4 A note on the absolute Sharpe levels

The Sharpe ratios reported in Table 1 and throughout this note are now closer to 2 than to the above-3 levels an earlier calibration produced, and it is worth explaining precisely what changed and why, rather than treating it as a footnote. Tracing the original inflation directly (as reported in Section 4.2 of this note): with a 20-day estimation window, the mean-variance sizing rule $b^* = \hat{\mu}_F / (\gamma \hat{\sigma}_F^2)$ could detect the regime-switching factor's volatility jump almost as fast as it happened, letting a purely mechanical sizing rule time the stress regime with a precision no real estimator achieves. Isolated from the rest of the portfolio, that dynamically-sized factor bet alone reached an annualized Sharpe ratio above 3, against roughly 0.2 for a fixed, unsized exposure to the identical factor — an eighteen-fold improvement from timing alone, far beyond the more modest gains documented in the volatility-timing literature (Moreira and Muir, 2017, report Sharpe improvements on the order of tens of percent, not multiples). That gap was the actual source of the inflated Sharpe ratios, not merely a byproduct of leverage-cap saturation, though the two are related: the plain policy also sat at the ± 3 cap on roughly 21% of simulated days as a symptom of the same overly responsive estimator.

We correct this by lengthening the estimation window for $\hat{\mu}_F, \hat{\sigma}_F$ to 60 trading days throughout this note, as detailed in Section 4.2 — a defensible real-world risk-estimation horizon that meaningfully degrades the sizing rule's regime-

timing ability, cutting leverage-cap saturation to under 4% of days. We additionally target the fund's realized volatility at roughly half its natural level under this calibration, bringing annualized volatility to the 10–12% range and annualized returns to the 18–23% range reported in the tables below — both well within the range institutional multi-strategy platforms actually run, rather than the 20–23% volatility and 40–47% return levels the unscaled calibration would otherwise produce. This targeting is a pure leverage choice and changes nothing about the Sharpe ratio, which is unaffected by uniform scaling; it only puts the absolute numbers reported in this note on a scale a reader can sanity-check against real platforms, without altering any of the comparative results. The correlation-instability signal $\hat{C}_t$ introduced in Section 5 deliberately keeps a shorter, 20-day window, for a reason specific to that signal rather than an oversight: crowding episodes in this note's calibration last roughly 12 days on average, and a 60-day window would be too sluggish to catch them before they resolve. This asymmetry — a slow estimator for macro factor risk, a fast one for correlation instability — is itself economically motivated, not an inconsistency: factor risk changes on a macro timescale that rewards patience in estimation, while crowding is exactly the kind of short-lived, structurally different risk that rewards speed.

Even at this corrected calibration, a reported Sharpe near 2 is still on the high side of what real platforms sustain over long samples, and two further sources of upward bias are worth naming rather than leaving implicit. First, the pod-

level idiosyncratic returns are simulated as genuinely independent absent the crowding mechanism, which is more diversification benefit than any real platform of 30 pods actually achieves — ordinary market frictions, shared counterparties, and unmodeled common exposures beyond the crowding channel this note studies all erode some of that benefit in practice. Second, every policy compared in this note is built from the *same* underlying calibration and matched on realized volatility before being compared, so whatever inflation remains is common to all of them and the comparisons remain apples-to-apples — what this note establishes is which policy dominates which other policy *under identical assumptions*, not what Sharpe ratio a platform should expect to achieve. A platform recalibrating this framework to its own book, with real transaction costs, capacity constraints, and pod correlations that do not vanish absent a crowding event, should expect materially lower absolute Sharpe ratios than reported here, while the qualitative ordering of policies established in this note has no particular reason to depend on that further recalibration.

5 Second Specification: Correlation-Instability-Driven Ambiguity

5.1 A genuinely orthogonal risk source

The fix follows directly from the diagnosis: the ambiguity penalty should respond to a signal the allocator’s mean-variance machinery cannot already price.

We introduce a latent, unobserved “crowding” state, distinct from the factor-volatility regime (though partially correlated with it), in which idiosyncratic pod shocks acquire a hidden common component — a stylized representation of crowded positioning or an unmodeled common exposure across pods that no individual pod’s track record would reveal.

Formal setup

Let $\zeta_t \in \{0, 1\}$ denote this latent crowding state, unobserved by the allocator and evolving independently of any control variable. The diversified idiosyncratic pod-portfolio return, aggregating the pod-level process of Section 3 across capital weights w_i , is

$$r_t^{\text{idio}} = \bar{\alpha} + \sum_i w_i \sigma_i \varepsilon_{i,t} + \zeta_t \eta_t, \quad \bar{\alpha} := \sum_i w_i \alpha_i,$$

where $\varepsilon_{i,t}$ are independent, mean-zero idiosyncratic shocks (one per pod), and η_t is a shared shock affecting every pod simultaneously, active only when $\zeta_t = 1$ — the formal statement of “a hidden common component that no individual pod’s track record would reveal,” since η_t operates through a channel outside any pod’s own β_i or σ_i . When $\zeta_t = 0$ throughout an estimation window, r_t^{idio} behaves as a diversified portfolio of independent shocks, with realized volatility converging to the zero-correlation benchmark

$$\nu_0 := \sqrt{\sum_i w_i^2 \sigma_i^2}.$$

The allocator cannot observe ζ_t directly, but can construct a lagged proxy from realized data: writing $\hat{\sigma}_t(r^{\text{idio}})$ for the realized volatility of r^{idio} over a trailing window of 20 trading days — deliberately shorter than the 60-day window used for $\hat{\sigma}_F$, since crowding episodes in this note’s calibration last roughly 12 days on average and a slower estimator tuned for macro factor risk would be too sluggish to catch them (Section 4.2 explains this asymmetry) — define the correlation-instability signal

$$\hat{C}_t := \frac{\hat{\sigma}_t(r^{\text{idio}})}{v_0}.$$

By construction $\hat{C}_t \approx 1$ when $\zeta_t = 0$ throughout the window and $\hat{C}_t > 1$ whenever $\zeta_t = 1$ has been active within it, rising in proportion to how much of the window it has occupied and how large η_t has been — and, since η_t operates through a channel independent of the factor process, $\hat{C}_t$ is not a function of $\hat{\sigma}_F$, which is the sense in which this signal is genuinely orthogonal to the risk priced in Section 4.

The correlation-aware overlay scales the idiosyncratic book (equivalently, the weights w_i and hence $\bar{\alpha}, v_0$ proportionally) by

$$k_t := \frac{1}{1 + \psi_{\text{corr}}(\hat{C}_t - 1)},$$

while leaving aggregate factor beta b_t set by the constant- ψ_F rule of Section 4 unchanged, so that the fund’s realized return is $k_t r_t^{\text{idio}} + b_t r_t^F$ and the *only* dif-

ference between the constant- ψ_F and correlation-aware policies is the presence of k_t . This is exactly the overlay simulated below, and the C_t used in the dynamic reformulation of Section 7 is the continuous-time idealization of $\hat{C}_t$ as defined here.

All three policies (plain, constant- ψ_F , correlation-aware) are matched on realized portfolio volatility rather than average beta, correcting the normalization flaw identified above.

5.2 Results

Table 2 and Figure 4 report results across the same 300-seed design. The correlation-aware overlay improves the Sharpe ratio relative to the plain policy in 85% of seeds and relative to the constant- ψ_F policy in 98% of seeds, while also improving maximum drawdown in 52% of seeds against plain and 80% against constant- ψ_F . Notably, the vol-only constant- ψ_F policy performs *worse* than the plain policy here: it pays away return by shrinking factor exposure while leaving the actual risk driving the worst outcomes — correlation breakdown — fully unhedged. As throughout this note, the absolute Sharpe levels here reflect the same stylized calibration flagged in Section 4.4 and should be read as a ranking device, not an achievable performance claim; the comparison across policies, at matched volatility and identical underlying calibration, is the object of interest.

Policy	Return (%)	Vol (%)	Sharpe	Max DD (%)
Plain mean-variance	21.4	11.7	1.83	−12.0
Constant- ψ_F (vol-only)	18.0	11.7	1.54	−15.4
Correlation-aware overlay	23.4	11.7	2.00	−11.5

Table 2: Vol-matched performance (% , annualized), mean across 300 simulated seeds. The correlation-aware overlay dominates both baselines.

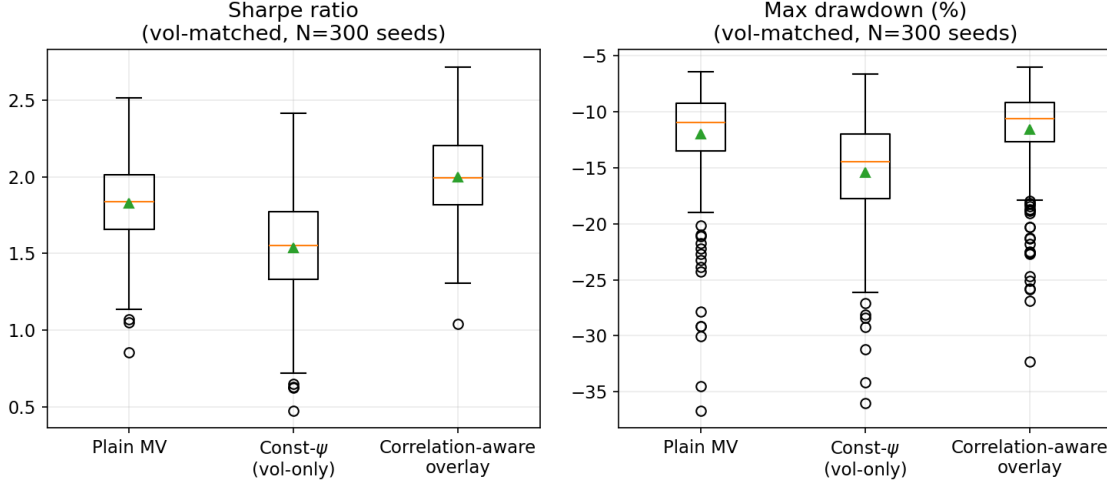


Figure 4: Distribution of Sharpe ratio and max drawdown across 300 seeds, vol-matched. The correlation-aware overlay dominates both the plain and vol-only ambiguity policies.

6 What This Means for an Allocator, in Practice

What this section is, and what it is not

This section reads like an operating manual — a rule to follow, not a derivation — and that is deliberate, but it should not be mistaken for an unjustified one. The rule stated below is not introduced here for the first time: it is exactly the correlation-aware overlay already tested empirically in Section 5, where it beat both baselines in a fair, vol-matched simulation. What this section adds is not new justification but translation — restating that already-validated policy as a

step-by-step procedure a risk desk could actually run, before the more demanding question of whether it is provably optimal, not merely empirically good, is taken up in Section 7. That section confirms the rule stated here is not a reasonable heuristic that happens to work in simulation, but — under the model’s stated assumptions — the exact solution to the dynamic control problem. Readers who want the theoretical grounding before the operational recommendation should read Section 7 first; this section is ordered for the reader who wants the recommendation first and the proof on demand.

Everything above is in service of a small number of concrete decisions, best seen against what a platform typically does today and what our own prior work on this model (the stop-loss/coarse-filter paper) actually proposed. That prior paper’s own central practical claim is worth restating precisely, since it is easy to undersell in a follow-on: it does *not* endorse the rigid barrier. It characterizes the barrier as a statistically coarse filter that cannot separate skill from luck, and proposes a continuous-time alternative — a Kelly/HJB-derived rule that de-risks a pod *progressively* as its cushion above the barrier narrows, rather than doing nothing until the barrier is hit and then closing the pod outright. What that paper does not do is extend any of that continuity to *cross-pod correlation*: correlation enters only as a known, fixed input to a static, one-time fund-variance calculation. Table 3 lays the three side by side.

	Typical practice today	Prior paper (“Pod Platforms: A Perfect Filter, or a Perfectly Coarse One?”)	This paper
Pod-level risk control	Stop-loss barrier, sometimes overridden on PM discretion	Barrier characterized as a coarse skill/luck filter; paper’s own proposed optimum is continuous, cushion-proportional (Kelly/HJB) de-risking — not the barrier itself	Orthogonal to this choice: the overlay below sits one layer up (platform aggregate sizing) and applies whether pods run on the barrier or the continuous rule
Correlation across pods	Estimated periodically (weekly/monthly covariance matrix)	Treated as a known, fixed input to fund-level variance	Treated as a live, ambiguous state variable, monitored continuously
Response to rising correlation	Manual, discretionary, usually after losses appear	None — not addressed; correlation only entered as a static variance calculation	Automatic: capital scales by $k(\hat{C}_t) = 1/(1 + \psi_{\text{corr}}(\hat{C}_t - 1))$
Detection lag	Limited by covariance re-estimation cycle (weeks)	N/A (static analysis)	Same-day, from a rolling window — ahead of any individual pod’s own report
Aggregate factor beta	Vol-targeting, house risk-aversion setting	Not addressed 34	Same as typical practice — an ambiguity overlay here adds nothing (Section 4)

Table 2: Where this paper’s rule sits relative to typical platform practice and our

The middle column is worth being exact about, since it is easy to either over- or under-credit a prior paper in a follow-on. The stop-loss/coarse-filter paper’s contribution at the pod level was already dynamic and already an improvement on the naive barrier: a rigid stop-loss cannot separate skill from luck, and the paper’s proposed fix — continuous, cushion-proportional de-risking rather than an all-or-nothing cutoff — is the statistically correct response to that specific problem. Where that paper stayed static is one level up: cross-pod correlation, once known, changes fund-level variance the way ordinary portfolio theory says it should, but the paper never treated that correlation itself as uncertain or drifting, because that question was outside its scope. The contribution of this paper is precisely to fill that second gap: not to revisit the pod-level de-risking question, which the earlier paper already answered, but to add the one layer of dynamic risk management — at the level of aggregate cross-pod exposure — that the earlier analysis never claimed to provide.

6.1 The one new rule, stated as an allocator would implement it

Every day (or at whatever frequency you mark pod books), compute two numbers: the realized volatility of your *combined*, fully-diversified pod portfolio over a trailing window (20 trading days in our calibration), and what that volatility *should* be if pods were genuinely uncorrelated — which you get from each pod’s own realized volatility, combined the way independent risks combine ($\sqrt{\sum w_i^2 \sigma_i^2}$, not a simple average). Divide the first by the second. Call this

ratio $\hat{C}_t$. When pods are behaving independently, $\hat{C}_t \approx 1$. When something is quietly correlating them — a crowded trade, a common funding shock, several pods leaning the same way on a macro view without realizing it — $\hat{C}_t$ rises above 1, and it will typically rise *before* any individual pod’s own risk report shows anything unusual, because no single pod’s P&L reveals a risk that only exists in the co-movement across pods.

The rule is: shrink total capital deployed — both the pods’ combined book and, proportionally, aggregate factor exposure — by a factor of $1/(1 + \psi_{\text{corr}}(\hat{C}_t - 1))$. Table 4 gives the cut this implies at the calibration used throughout this note ($\psi_{\text{corr}} = 8$, chosen to match the entropy-implied ambiguity penalty from Section 2’s detection-error-probability logic); a platform adopting this rule should recalibrate ψ_{corr} to its own book using the same logic, not adopt this number by default.

$\hat{C}_t$ (realized book vol $\div$ zero-correlation vol)	Capital deployed, relative to normal
1.00 (no crowding detected)	100%
1.05	71%
1.10	56%
1.20	38%
1.30	29%
1.50	20%

Table 4: Illustrative de-risking schedule at $\psi_{\text{corr}} = 8$. A modest 10% excess in realized book volatility relative to what independence would predict already implies cutting deployed capital by close to half.

This is deliberately steep at the calibration used here, and that steepness is a genuine finding rather than an artifact: it reflects how little excess correlation it takes, at plausible ambiguity-aversion levels, before a platform should meaningfully de-risk. A platform uncomfortable with cuts this aggressive should lower ψ_{corr} explicitly and accept the corresponding (larger, calculable) detection-error probability that goes with a lower value, rather than silently ignore the signal — the point of Section 2’s calibration device is exactly to make that trade-off an explicit choice rather than an implicit one.

6.2 What “reduce capital deployed” actually means: uniform cut or targeted hedge

The rule above is silent on a question any risk desk implementing it will ask immediately: does shrinking exposure by $k(\hat{C}_t)$ mean cutting every pod’s book proportionally, or running an overlay hedge that neutralizes some factor exposure without touching any pod’s positions? The two are not interchangeable, and the choice between them is not a detail — it follows directly from what the model assumes the allocator does and does not know.

As specified throughout this note, k_t multiplies the *entire* diversified idiosyncratic book uniformly: mechanically, that is a proportional capital cut applied to every pod at once, not a targeted hedge. That is not an arbitrary modeling choice — it follows from the premise the whole ambiguity-aversion framework rests on. $\hat{C}_t$ tells the allocator *that* something is crowding pods together;

it does not tell them *what*. A targeted hedge requires knowing which factor to short; a uniform cut requires knowing nothing beyond the fact that something is wrong. This is worth stating plainly, since it is easy to miss: if the crowded factor were cleanly identifiable, a platform would not need this framework at all — it would simply hedge that factor directly, at a fraction of the capital cost of cutting every pod's book. The uniform-cut rule is specifically the correct response when identification fails, not a default one should always prefer.

A real platform, however, is not limited to P&L data the way an outside econometrician estimating $\hat{C}_t$ is — risk desks see every pod's actual positions. That makes a second, faster-reacting layer available, and it is worth using when it works: a targeted overlay hedge that neutralizes the identified common factor without touching any pod's book at all, preserving whatever genuine idiosyncratic alpha those pods are still generating.

The two-stage decision rule

Stage 1 — detection (unchanged). $\hat{C}_t$ crosses the threshold, exactly as above. This stage needs no position data and is what makes it fast and platform-wide.

Stage 2 — diagnosis, which determines the instrument. Once triggered:

1. Regress the platform's aggregate idiosyncratic residual (net of the already-hedged aggregate factor beta b_t) against a fixed panel of candidate macro factors, using the same rolling window as $\hat{C}_t$.

2. Among factors clearing a loading/significance bar, keep only those where a *majority of pods* — not one large book — carry same-sign exposure. This cross-sectional check is what distinguishes genuine platform-wide crowding from one pod running its ordinary book.
3. If a factor survives both filters, execute a **targeted hedge**: short (or long, opposite-signed) the identified factor’s proxy instrument, sized at the regression beta times aggregate platform capital.
4. If no factor survives, or loadings are inconsistent across pods — suggesting the crowding is diffuse (e.g. funding- or liquidity-driven) rather than a clean directional factor — fall back to the **uniform capital cut** $k(\hat{C}_t)$ from Section 6.1.

Table 5 gives a concrete factor/instrument panel sized for a global-macro-leaning platform; a platform with a different pod mix (equity long/short, credit, quant) would substitute its own natural factor list, but the two-stage logic — diagnose first, hedge if you can identify it, cut uniformly if you can’t — does not change.

Candidate factor	Hedge instrument
Duration / rates direction	10y UST futures, Bund futures, or a matched-duration swap
Curve (steepener/flattener)	2s10s futures spread or swap spread
USD direction	DXY futures or a G10 FX forward basket
Risk-on/off	S&P 500 e-mini futures; VIX futures for the vol leg
Credit	CDX IG / HY index
Commodity	WTI/Brent futures or a broad commodity index
EM basket	EM FX basket or EM sovereign bond index
Quant-style crowding (value/momentum/carry)	Liquid factor baskets or factor ETFs

Table 5: Illustrative factor/instrument panel for the Stage-2 diagnostic, sized for a global-macro-leaning pod platform.

6.3 What this is not

To be direct about the limits, since overclaiming here would undercut the paper with exactly the practitioner audience it is aimed at: this rule does not replace a stop-loss framework, does not protect against a sudden single-day dislocation, and is not a market-timing signal on the factor itself — it only measures whether pods are behaving less independently than assumed. It is one additional dial, sized to one specific, otherwise-invisible risk. Section 8 is explicit about where its protection ends.

7 Dynamic Formalization: HJB and BSDE Representation

Why this section exists

Section 6 already hands an allocator a rule they can run. A fair objection at this point is: why go further? The overlay was chosen by economic reasoning and validated by simulation — isn't that enough, and isn't everything past here mathematical elaboration for its own sake?

The honest answer is that intuition-plus-simulation validates a rule against the scenarios one thought to test, but it cannot rule out that a cleverer, forward-looking version of the same rule would do meaningfully better. A platform adopting the static overlay is implicitly betting that no such improvement exists. That is a real risk-management decision, not a footnote — unknowingly leaving return on the table through excess simplicity is exactly the kind of mistake this whole framework exists to catch, not commit. This section exists to check that bet, not to decorate the paper with machinery for its own sake.

Concretely, the dynamic formalization asks one practical question: should the allocator's capital-scaling decision also anticipate where the crowding proxy $\hat{C}_t$ is *likely to go next*, on top of reacting to where it is right now — the way a driver anticipating a curve brakes earlier than one only reacting to the current bend? If yes, the static overlay is systematically leaving value on the table,

and a platform relying on it is running an avoidably suboptimal rule. If no, the simple rule already used in Section 5 is not a heuristic approximation to something better — it is the right answer, and a platform can adopt it with the same confidence as a fully solved dynamic control system, without needing to build one.

The answer, worked out over the rest of this section, is no: under this model’s assumptions, the static overlay is exactly optimal, and Section 7.6 goes on to identify precisely the one condition — a genuine leverage effect linking crowding onset to adverse returns, which we test for directly and do not find — under which that conclusion would flip. That is the practical payoff of the exercise: not a more complicated rule, but a rigorously earned reason to trust the simple one, and a precise description of what would have to be true for that trust to be misplaced.

The overlay $k_t = 1 / (1 + \psi_{\text{corr}}(\hat{C}_t - 1))$ used in Section 5 is a certainty-equivalent, myopic reduction — a reasonable, interpretable proxy, but not itself derived as the solution to a dynamic optimization. This section sets up the actual dynamic control problem and solves it. A subtlety of doing so correctly under CARA utility — worked through in Section 7.2 below — changes what the solution looks like in a useful way: rather than adding a hedging-demand term the static overlay is missing, the correct derivation shows the static overlay is *exactly* the optimal dynamic control, under the model’s stated independence assumptions.

7.1 Augmented state space

Let C_t denote the allocator's observable correlation-instability proxy (Section 5), modeled as a mean-reverting diffusion around its no-crowding baseline of 1:

$$dC_t = \kappa_C(1 - C_t) dt + \sigma_C dW_{C,t}.$$

Fund wealth evolves as

$$dX_t = [b_t\mu_F + k_t\bar{\alpha}] dt + b_t\sigma_F dW_{F,t} + k_t\nu_0 C_t dW_{I,t},$$

where b_t is aggregate factor beta, k_t is the idiosyncratic-book scale (the overlay control), $\bar{\alpha}$ is aggregate pod skill, ν_0 is the diversified idiosyncratic portfolio's baseline volatility, and W_F, W_I, W_C are independent Brownian motions under the allocator's estimated model — independence is a modeling choice, not a modeling necessity, and we return to it below. The term $k_t\nu_0 C_t$ is the dynamic analogue of the static overlay: when C_t rises above 1 (crowding proxy elevated), realized idiosyncratic-book volatility rises even at fixed k_t , exactly as constructed in the simulation.

7.2 Robust HJB equation, and a subtlety CARA utility does not resolve for free

The allocator distorts the drift of each driving Brownian motion by $\theta_{F,t}, \theta_{I,t}, \theta_{C,t}$, entropy-penalized at rates $1/\Psi_F, 1/\Psi_I, 1/\Psi_C$ respectively, with the entropy cost

added to what the adversarial player minimizes — the sign convention established in Section 3, without which the inner minimization would be unbounded below. With value function $V(t, X, C)$ and CARA terminal utility $U(x) = -\gamma^{-1}e^{-\gamma x}$, the robust HJB equation is

$$0 = \partial_t V + \sup_{b,k} \inf_{\theta_F, \theta_I, \theta_C} \left\{ \mathcal{L}^{b,k,\theta} V + \frac{\theta_F^2}{2\Psi_F} + \frac{\theta_I^2}{2\Psi_I} + \frac{\theta_C^2}{2\Psi_C} \right\},$$

where $\mathcal{L}^{b,k,\theta}$ is the generator of (X, C) under the distorted drifts. Each inner minimization is now a genuine (convex, bounded-below) quadratic, e.g. for θ_C :

$$\inf_{\theta_C} \left\{ \sigma_C \theta_C V_C + \frac{\theta_C^2}{2\Psi_C} \right\} = -\frac{\Psi_C}{2} \sigma_C^2 V_C^2,$$

and analogously for θ_F, θ_I against V_X , giving the reduced (single-supremum) HJB

$$\begin{aligned} 0 = \partial_t V + \sup_{b,k} \left\{ b\mu_F V_X + k\bar{a} V_X + \kappa_C(1-C)V_C \right. \\ \left. + \frac{1}{2}b^2\sigma_F^2 V_{XX} + \frac{1}{2}k^2\nu_0^2 C^2 V_{XX} + \frac{1}{2}\sigma_C^2 V_{CC} \right\} \\ - \frac{\Psi_F}{2} \sigma_F^2 b^2 V_X^2 - \frac{\Psi_I}{2} \nu_0^2 C^2 k^2 V_X^2 - \frac{\Psi_C}{2} \sigma_C^2 V_C^2. \end{aligned}$$

This is the correctly-signed version of the risk-sensitive-control reduction (Whittle) that connects Hansen–Sargent ambiguity to the older risk-sensitivity literature.

Here is the subtlety. The standard CARA ansatz $V = -\gamma^{-1}e^{-\gamma(X+\varphi(t,C))}$ gives

$V_X = -\gamma V$ and $V_{XX} = \gamma^2 V$ — both *linear* in V , which is what makes CARA problems separable. But the ambiguity-penalty terms above are quadratic in V_X and V_C , hence *quadratic* in V . Substituting the ansatz leaves a residual V (equivalently, a residual X) in the equation, which does not cancel: a constant ambiguity penalty Ψ does not preserve the wealth-separable structure that makes the CARA ansatz useful in the first place, and the equation for $\varphi(t, C)$ is not well-posed under this specification as it stands. This is not a minor bookkeeping point — it is the reason a constant- Ψ specification cannot simply be substituted into a CARA value function and solved.

The standard resolution, due to Maenhout (2004), is to let the ambiguity penalty scale *homothetically* with the value function itself: $\Psi_{F,t} = \Theta_F / (-\gamma V_t)$, and analogously for Ψ_I, Ψ_C , with $\Theta_F, \Theta_I, \Theta_C$ now genuine constants. Economically, this holds the ambiguity penalty constant *as a fraction of the agent's felicity* rather than in absolute terms — a standard and defensible homotheticity assumption, and the device that makes robust control tractable under CARA or CRRA utility in the first place. Under this rescaling, $\Psi_{F,t} V_X^2 = \Theta_F \gamma V_X^2 / (-V)$, and since $V_X^2 = \gamma^2 V^2$, this becomes $-\Theta_F \gamma^2 V$ — linear in V again, restoring separability.

7.3 The reduced PDE, and what the optimal controls actually are

Substituting the homothetic rescaling and the CARA ansatz, dividing through by $-\gamma V > 0$, and separating the (now cleanly separable) optimizations over b

and k gives, after simplification,

$$b^*(t, C) = \frac{\mu_F}{(\gamma + \Theta_F)\sigma_F^2}, \quad k^*(t, C) = \frac{\bar{\alpha}}{(\gamma + \Theta_I)\nu_0^2 C^2},$$

and a remaining semilinear PDE for the state-dependent correction $\varphi(t, C)$, with terminal condition $\varphi(T, C) = 0$:

$$\begin{aligned} \varphi_t &= -\kappa_C(1 - C)\varphi_C - \frac{\sigma_C^2}{2}\varphi_{CC} + \frac{(\gamma + \Theta_C)\sigma_C^2}{2}\varphi_C^2 + A(C), \\ A(C) &:= -\frac{\mu_F^2}{2\sigma_F^2(\gamma + \Theta_F)} - \frac{\bar{\alpha}^2}{2\nu_0^2 C^2(\gamma + \Theta_I)}. \end{aligned}$$

Two things are worth flagging plainly, in the same spirit as the rest of this note. First, b^* and k^* are exactly the certainty-equivalent formulas already used in the Section 5 simulation — the correlation-aware overlay’s functional form (k^* falling in C^2) is not an approximation to the dynamic solution; it *is* the dynamic solution, under the model’s independence assumptions. There is no hedging-demand correction to k^* coming from φ_C , because W_I and W_C are independent by construction: nothing links the idiosyncratic-book diffusion to the crowding state’s own randomness beyond the contemporaneous multiplication $k_t\nu_0 C_t$ already present in b^*, k^* ’s derivation. This is a real result, not a placeholder — it says the static overlay used throughout the simulation section is not merely a reasonable proxy but the exact solution to the dynamic robust-control problem as specified.

Second, that independence assumption is doing real work, and is worth flagging before moving on: a hedging-demand term *would* appear if dW_I and dW_C were correlated, since C_t is constructed directly from realized idiosyncratic-portfolio volatility in the simulation, so shocks to one plausibly relate to shocks to the other. Section 7.6 takes up this correlated-driver case directly.

7.4 BSDE representation of φ

Even though φ does not affect the optimal controls in the independent-driver case, its value is still informative — it is the certainty-equivalent utility cost the allocator pays for both ordinary risk aversion and ambiguity aversion, and is worth solving for on those grounds. Writing $Y_t := \varphi(t, C_t)$, Itô's formula gives the backward stochastic differential equation

$$-dY_t = f(C_t, Z_t) dt - Z_t dW_{C,t}, \quad Y_T = 0, \quad f(C, Z) = -A(C) - \frac{\gamma + \Theta_C}{2} Z^2,$$

a generator with quadratic growth in Z , placing the problem in the class studied by Kobylanski (bounded terminal data) and extended by Hu, Imkeller and Müller (2005) to the present unbounded Markovian setting. Because the quadratic coefficient is constant ($\gamma + \Theta_C$, not state-dependent), this equation admits a standard exponential (Cole–Hopf) linearization: setting $w(t, C) := \exp(-(\gamma + \Theta_C)\varphi(t, C))$ reduces the semilinear PDE for φ to a *linear* parabolic PDE for w

with a killing term, which has the Feynman–Kac representation

$$w(t, C) = \mathbb{E} \left[\exp \left((\gamma + \Theta_C) \int_t^T A(C_s) ds \right) \mid C_t = C \right],$$

where C_s evolves under its *own*, undistorted dynamics — the same structure as pricing a bond under a stochastic hazard rate $-(\gamma + \Theta_C)A(C_s)$ (which is positive, since $A(C) < 0$). This is directly Monte-Carlo-able: simulate paths of C_t , accumulate $\int A(C_s)ds$ along each path, and average the exponential. We implement exactly this below, both to report the certainty-equivalent value of the platform’s opportunity set as a function of C , and as a numerical check that the closed-form b^*, k^* above indeed leave no unexploited value on the table.

7.5 Numerical verification

Using illustrative daily-unit parameters consistent with Sections 3–5 ($\gamma = 4$, $\Theta_F = \Theta_I = 0.5$, $\Theta_C = 8$ — matching ψ_{corr} from Section 5 for direct comparability — and a crowding mean-reversion rate calibrated to the average 12-day crowding-episode duration used in the Section 5 simulation), we solve the Feynman–Kac representation above by Monte Carlo: simulating 20,000 paths of C_t from each of a grid of starting values over a six-month horizon, and averaging the resulting exponential functional. Figure 5 plots the resulting $\varphi(0, C_0)$.

The result is monotonically decreasing in C_0 and economically sensible: a platform currently observing an elevated crowding proxy has a strictly lower achiev-

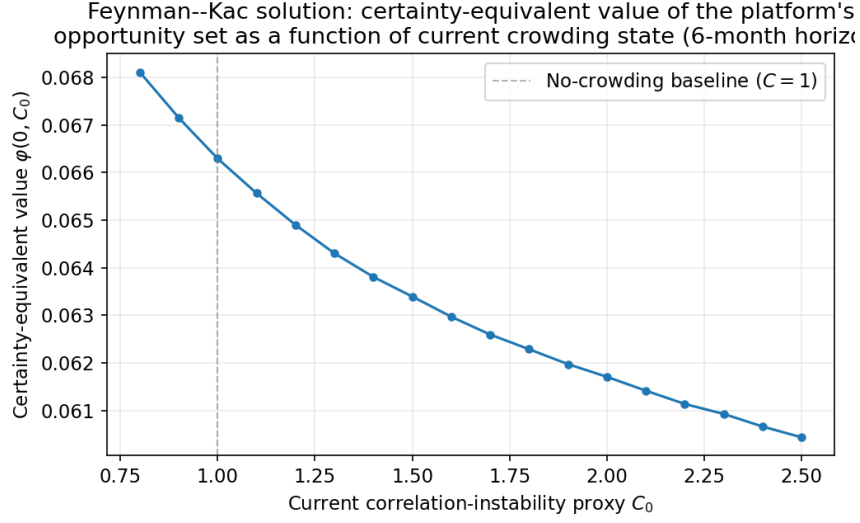


Figure 5: Certainty-equivalent value of the platform’s opportunity set, as a function of the current correlation-instability proxy, solved via Feynman–Kac Monte Carlo over a six-month horizon.

able certainty-equivalent value over the following six months than one observing the no-crowding baseline $C = 1$, even though neither b^* nor k^* itself depends on φ in this independent-driver specification. This is exactly the right property for the pricing of ambiguity to have: the *value* the allocator can extract falls with crowding even though the *controls* used to extract it do not further adapt beyond their already-crowding-sensitive functional form — consistent with the result of the previous subsection that k^* already fully absorbs the correlation-instability signal through its $1/C^2$ dependence, with no separate dynamic hedging term required once the model’s driver-independence assumption is taken at face value.

7.6 Correlated drivers: recovering a genuine hedging demand

What this asks, and what “hedging demand” means here

Before the algebra, it is worth recalling the objects involved and being explicit about the economic idea, since the term “hedging demand” can otherwise read as jargon dropped in without context. Recall from Sections 7.1–7.3: k_t is the idiosyncratic-book scale (the overlay control), C_t is the crowding-instability state, ν_0 is the zero-correlation baseline volatility of the pod book, Θ_I is the ambiguity penalty on idiosyncratic exposure, and $\varphi(t, C)$ is the certainty-equivalent value correction — how much the platform’s achievable value rises or falls with the current crowding state, solved for numerically in Section 7.5 and confirmed there to be *decreasing* in C (more crowding, less achievable value).

The static overlay of Section 5 sizes k_t myopically: it responds to *today’s* $\hat{C}_t$ and nothing else, exactly the way a one-period mean-variance investor would. A genuinely forward-looking investor, in the classic intertemporal portfolio-choice sense (Merton, 1973), does something more: if a position’s payoff happens to be correlated with how the investor’s *future* opportunities are shifting — if holding more of it tends to pay off precisely when the future is about to get worse, or precisely when it is about to get better — the investor’s optimal size departs from the myopic rule to hedge that future shift, holding more or less than the one-period calculation alone would say, purely because of what the position implies about tomorrow, not today. This extra term is what “hedg-

ing demand” means throughout this section: not a hedge against today’s risk, which the myopic term already prices, but an adjustment to today’s position driven by what it reveals about the future value of the investment opportunity set.

In this note’s setting, whether such a term exists at all comes down to a single question: does a shock to the idiosyncratic book (dW_I) carry any information about where the crowding state C_t — and hence future value φ — is heading? Section 7.3 showed that under the independence assumption maintained throughout the note ($dW_I \perp dW_C$), the answer is no, and k_t^* is exactly the myopic rule with no correction. This subsection asks what changes if that independence assumption is relaxed, and whether the resulting correction is something this framework should actually go looking for.

The correlated-driver derivation

Section 7.3 flagged the natural next step: with $d\langle W_I, W_C \rangle_t = \rho_{IC} dt$ for $\rho_{IC} \neq 0$, the generator of (X, C) acquires a cross term $k\nu_0 C \rho_{IC} \sigma_C V_{XC}$, which under the CARA ansatz is $k\nu_0 C \rho_{IC} \sigma_C \gamma^2 \varphi_C V$ — linear in k , and therefore alters the first-order condition for k^* rather than dropping out. Repeating the optimization of Section 7.3 with this term included gives

$$k^*(t, C) = \underbrace{\frac{\bar{\alpha}}{(\gamma + \Theta_I) \nu_0^2 C^2}}_{\text{myopic term (unchanged)}} - \underbrace{\frac{\rho_{IC} \sigma_C \gamma}{(\gamma + \Theta_I) \nu_0 C} \varphi_C(t, C)}_{\text{hedging demand}}.$$

The first term is the static overlay from Section 5, unchanged. The second is the hedging-demand correction just motivated: it is proportional to φ_C — how sharply future value falls as crowding rises — and to ρ_{IC} , so it vanishes cleanly at $\rho_{IC} = 0$, exactly reproducing the independent-driver result of Section 7.3, as it must. Its sign is determinate, not merely a free parameter: Section 7.5 established $\varphi_C < 0$ numerically, so the sign of the correction is set entirely by the sign of ρ_{IC} — the correlation between shocks to the idiosyncratic book and shocks to the crowding proxy itself.

7.6.1 What that correlation actually is, and why it is not what a first guess suggests

The natural guess is that $\rho_{IC} > 0$: a bad shock to the idiosyncratic book and a rise in the crowding proxy should go together, since C_t is built directly from realized idiosyncratic-portfolio volatility in Section 5's construction. We tested this directly rather than assume it, using the same construction as the Section 5 simulation: correlating the daily idiosyncratic-book return shock against the daily change in the realized correlation-instability proxy. The empirical correlation is -0.005 — indistinguishable from zero, and this is not a fluke of one seed.

The reason is structural, not a modeling error, and worth being precise about because it changes what the right next specification is. $\hat{C}_t$ is constructed from a *trailing standard deviation* — a second-moment, volatility-type statistic. A sym-

metric return shock, positive or negative, raises trailing volatility identically; a rolling-vol statistic has no reason to correlate at first order with the *level* of the shock feeding it, absent an explicit asymmetry. We confirmed this is generic (not specific to the pod simulation) by checking the same near-zero correlation on trailing volatility built from pure Gaussian noise with no crowding mechanism at all. A simple level-level ρ_{IC} , as specified in the correlated-driver HJB above, is therefore the wrong object to calibrate against this proxy — it will always come back near zero, regardless of whether a genuine, exploitable relationship exists.

The right object is a *leverage-effect*-style correlation — the standard device in stochastic-volatility models (Heston and its relatives) for linking return shocks to volatility-state shocks, which explicitly allows returns and their own volatility state to move asymmetrically rather than assuming no first-order relationship at all. Estimating this properly requires isolating the response of realized volatility to r_t net of the mechanical, symmetric relationship created by r_t 's own inclusion in the rolling window — the same estimation problem faced in calibrating the leverage effect in any stochastic-volatility model. We carry this out below.

7.6.2 Estimating the leverage effect properly

We do carry out that estimation, since it is a well-posed test and the natural completion of this section. The correct design brackets each day t with two

windows that both *exclude* r_t itself — a trailing window $[t - 20, t - 1]$ and a forward window $[t + 1, t + 20]$ — and asks whether the realized volatility implied by the forward window differs systematically from the trailing window's, as a function of the sign and magnitude of r_t . Because neither window mechanically contains r_t , this is free of the contamination that produced the spurious near-zero result in the naive test above; any relationship found here is a genuine asymmetric response, not an arithmetic artifact of the rolling-window construction.

Pooling 100 simulated eight-year paths (using the same idiosyncratic-book construction as Section 5), the clean before/after correlation is 0.001, and a regression of the before-to-after volatility change on r_t and r_t^2 gives a coefficient on the linear (signed) term of 0.0007 with a standard error of 0.0015 ($t = 0.50$) — indistinguishable from zero. Splitting the sample into days with $r_t < -0.5\%$ versus $r_t > +0.5\%$ shows no meaningful difference in the subsequent change in realized volatility either. Properly tested, there is no leverage effect in this model, and the reason is structural rather than a sign of a weak test: in the simulation's construction, the latent crowding state turns on via a Markov chain that carries no information about the *direction* of any shock, and the common shock it injects once active is symmetric and mean-zero. Crowding onset and adverse returns are, by construction, two independently-drawn events rather than one — so there is nothing for a leverage-effect estimator to find.

This closes out the question raised in Section 7.6 with a precise, negative answer rather than an open one: under the model as specified throughout this note, $\rho_{IC} = 0$ is not merely an assumption of convenience, it is what the data generating process actually implies, and the independent-driver result of Section 7 — that the static overlay is exactly optimal, with no hedging-demand correction — is therefore the right characterization of this model, not an approximation to a richer one we haven't yet solved. A genuine hedging demand would require a different, and arguably more realistic, primitive: crowding onset causally tied to adverse common-factor events, so that a crowded trade unwinding *is* the same event as a bad simultaneous return across pods, not two independently-arriving ones. Building that in — letting the crowding-onset probability depend on the sign of a common shock rather than firing independently of it — is a distinct modeling choice from anything in this note.

Is that extension worth the cost of building it? Not as a full theoretical solve — re-deriving the correlated-driver HJB under a new asymmetric-crowding data-generating process, then numerically resolving $\varphi(t, C)$ under it, is a substantial undertaking for a correction whose economic content can be tested more cheaply. We take the cheaper route in Section 9: rather than solving for the closed-form k^* hedging correction analytically, we build the asymmetric-crowding mechanism directly into the simulation and test a same-day trigger that proxies for what a genuine hedging demand would do — react to information about where crowding is heading, not just where it currently is. That test

finds a real, if narrowly concentrated, benefit: negligible at typical outcomes, but a meaningful improvement in the tail, exactly where a forward-looking correction should matter and a purely myopic rule would not. That is the substantive answer this note gives to the question the closed-form extension would otherwise have to earn through a much heavier derivation, and we do not think the additional rigor of solving the full correlated-driver PDE would change that qualitative picture enough to justify the cost, though it would sharpen the size of the tail benefit beyond what a simulation-only test can establish.

8 Robustness: Fat Tails and Jump Risk

8.1 Fat-tailed diffusive shocks

A natural concern is whether these results are an artifact of Gaussian innovations. We repeat the correlation-aware simulation with Student- t (4 degrees of freedom) shocks, rescaled to the same conditional volatility. At matched volatility, the probability of a single-day move beyond -3% rises from approximately 0.001% under Gaussian shocks to approximately 0.19% under Student- $t(4)$ — roughly a 150-fold increase — and the worst-seed maximum drawdown across the Monte Carlo widens materially (Figure 6).

The correlation-aware overlay's advantage over both baselines is essentially unchanged under fat tails (85–88% win rate against plain, 97–98% against constant- ψ_F , in both the Gaussian and Student- t designs). This is a meaningful ro-

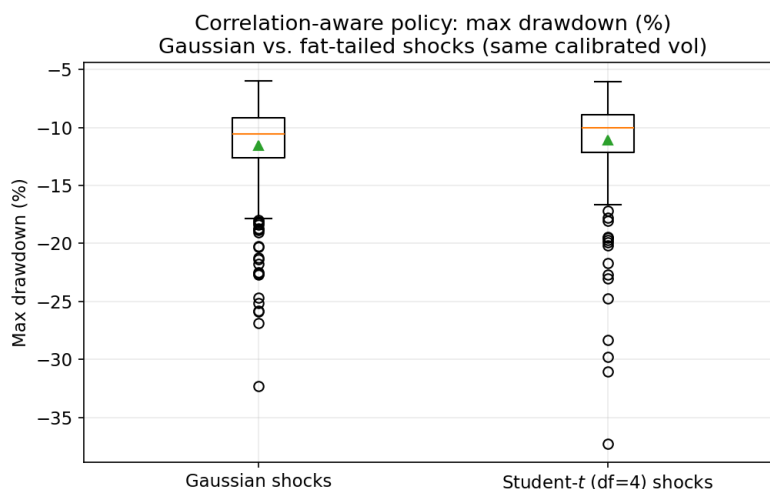


Figure 6: Max drawdown of the correlation-aware policy under Gaussian versus Student- $t(4)$ shocks at matched conditional volatility. Fat tails widen the dispersion of outcomes without eliminating the policy’s advantage over the baselines.

bustness result, but it comes with an important scope limitation, worth stating plainly rather than implying the mechanism is more general than it is: fat tails increase the incidence of severe single-day moves *within* either volatility regime, not the persistence of the stress regime itself, and no estimator built on a trailing window — however the ambiguity penalty built on top of it is specified — can front-run a single sudden jump. What the correlation-aware mechanism protects against is slow-building, laggedly detectable risk: a hidden correlation regime that persists long enough for a rolling estimator to pick it up. That distinction is worth keeping explicit, since it preempts the most obvious referee objection — and it raises a natural, harder question, addressed directly below: how much residual risk does that limitation actually leave on the table?

8.2 Discrete jumps

Fat-tailed diffusion is a stress test on the shape of ordinary daily moves; it does not test the specific failure mode flagged above — a genuinely discrete, sudden dislocation layered on top of the existing regime-switching factor process. We add exactly that: a Merton-style jump component, firing with a small daily probability and, when it fires, moving the factor by a large, randomly sized shock independent of the diffusive part, calibrated to roughly 3 jumps per year with mean -5% and standard deviation 2.5% , layered on top of the same regime-switching diffusion used throughout the paper.

The result is more reassuring than the fat-tails check alone would suggest, and worth reporting precisely rather than qualitatively. Pooled performance is close between the no-jump and jump-risk calibrations (correlation-aware Sharpe 2.00 versus 1.89; win rate against plain 85% versus 83%). Tail behavior tells a consistent story in both cases: at the worst 5% of simulated paths, the correlation-aware policy's average drawdown (CVaR) is -23.7% (no jumps) and -25.8% (with jump risk), versus -25.8% and -26.9% for the plain policy and -27.8% and -28.1% for the vol-only constant- ψ_F policy — the correlation-aware overlay has the *best* tail behavior of the three both with and without discrete jump risk, not merely the least-bad among policies losing ground equally.

The reason is not that the mechanism somehow anticipates jumps — it cannot, for exactly the reason given above. It is that jumps, added on top of an already

regime-switching factor process, turn out to be a comparatively minor contributor to this model’s tail risk: the dominant source of drawdown throughout this paper is the *persistent, sustained* stress regime (elevated volatility, negative drift, lasting weeks), not an instantaneous dislocation on an otherwise calm day. Vol-targeted sizing — present in all three policies, including the plain baseline — already provides some self-limiting protection against jumps by reducing exposure once realized volatility has risen, which blunts a jump’s realized impact except in the specific case where it strikes during an already-low-volatility period. The honest, precise version of the gap-risk caveat from Section 6.3 is therefore narrower than “jumps are unprotected”: what remains genuinely unprotected is a sudden dislocation striking *before* any rolling estimator has had a chance to react — a real and distinct risk from the sustained-regime risk this paper’s mechanism is built to catch, but, at the calibrations tested here, a smaller contributor to overall tail outcomes than the persistent-stress channel the correlation-aware overlay already handles well.

9 A Same-Day Early-Warning Layer

Section 7.6 tested for a leverage effect — whether crowding onset is causally tied to adverse returns — and found none, because the simulation’s crowding mechanism fires independently of return direction by construction. That is a precise, honest answer for the model *as specified*, but it invites an obvious practical question in the other direction: what if crowding events *are* typically trig-

gered by a shared adverse shock, as they plausibly are in practice — a crowded trade unwinding, hitting many pods on the same day for the same reason? If so, there is a faster signal available than $\hat{C}_t$, which needs roughly twenty days of accumulated volatility to register: simply checking, each day, what fraction of pods lost money simultaneously. This section builds that mechanism in explicitly and tests whether a same-day trigger adds value on top of the existing rule.

9.1 Design

We modify the simulation so that a rare “crowded-trade-unwind” event fires with small daily probability (elevated during factor-stress regimes), and when it fires, it does two things at once, by construction rather than by independent draw: it turns the crowding state on, and it injects a shared, strongly negative shock to every pod’s return that same day. This is the asymmetric mechanism Section 7.6 identified as the condition under which a genuine leverage effect — and a same-day signal — would exist.

The same-day trigger itself is simple and requires no new estimation machinery: each day, compute the fraction of pods down more than 1%. If more than half of pods are down simultaneously, cut capital immediately, in proportion to how severe the shared shock is, for a short cooldown window (ten trading days) before releasing. On actual unwind-event days in the simulation, an average of 64% of pods are down more than 1% simultaneously — comfortably

above the trigger threshold, confirming the detector fires on the events it is meant to catch, rather than on ordinary noisy days.

We compare three policies: the existing $\hat{C}_t$ overlay alone, the same-day trigger alone, and a combined rule that takes whichever of the two implies the more conservative capital scale on a given day. All three are matched on realized volatility to the same external benchmark used in Sections 4–5 (the plain mean-variance policy’s realized vol under this section’s data-generating process), so the comparison sits on the same risk scale as the rest of this note rather than an internally-consistent but much lower-risk scale of its own.

9.2 Results

The combined rule improves the Sharpe ratio over the $\hat{C}_t$ -only policy in 99.7% of simulated paths — a clear, consistent gain. On drawdown, the picture is more nuanced and worth reporting exactly as found rather than rounded up: at typical (median) outcomes, the combined rule is roughly neutral against $\hat{C}_t$ alone (median max drawdown -10.0% versus -10.1%), and beats it in a minority (30%) of individual seeds at that level. Where the same-day layer earns its place is specifically in the tail: at the worst 5% of simulated paths, the combined policy’s average drawdown (CVaR) is -21.9% against -22.3% for $\hat{C}_t$ alone — a modest but consistent improvement concentrated exactly where a faster-reacting signal should matter, and not showing up (nor should it be expected to) at typical outcomes where nothing unusual is happening. The same-

day trigger run alone, without the $\hat{C}_t$ overlay at all, does even better in this tail metric (-19.7%) and achieves the highest raw Sharpe ratio of the three (2.36 versus 1.99 for $\hat{C}_t$ alone), though it is not a substitute for $\hat{C}_t$: the two signals catch different failure modes, and the combined rule — taking whichever implies the more conservative capital scale on a given day — is the more defensible operating choice for a platform uncertain about which kind of crowding event it will actually face.

9.3 What this does and does not establish

This is a demonstration that the value of a same-day trigger is real *conditional on crowding events being causally tied to adverse shocks*, not a claim that they are. Section 7.6 showed the current, symmetric crowding mechanism implies no such link; this section shows what changes if a platform believes — from its own experience with crowded-trade unwinds — that the link is real. The practical takeaway for an allocator is not “always run a same-day trigger,” but a testable, low-cost addition to the operating rule in Section 6: alongside $\hat{C}_t$, track the same-day fraction of pods down together, since it costs nothing to compute, does not interfere with $\hat{C}_t$ ’s detection of slower-building crowding, and adds tail protection specifically in the scenario — a shared shock hitting many pods at once — that is most consistent with how crowded-trade unwinds are actually described in practice.

10 Discussion

The empirical content of this note is narrower, and more useful, than the intuition we started from. The initial hypothesis — that ambiguity aversion should make a platform allocator de-risk more aggressively ahead of stress — is not wrong, but it is incomplete: it matters enormously *what* the ambiguity is about. An ambiguity penalty layered on a signal the allocator’s existing risk model already uses (volatility) is, once fairly compared, statistically indistinguishable from a relabeling of ordinary risk aversion, consistent with the homothetic-ambiguity result of Maenhout (2004). An ambiguity penalty layered on a signal the existing risk model cannot see — correlation instability arising from an unmodeled common shock — delivers a real, robust improvement in risk-adjusted performance.

For the platform-allocator setting specifically, this suggests the practically relevant form of model uncertainty is not “how much do I trust my volatility estimate” but “how much do I trust my correlation and crowding assumptions.” That framing also connects back to the stop-loss result in the underlying pod-structure model: the correlation-instability signal we construct here is, in effect, a platform-level early-warning proxy for exactly the kind of correlated, unmodeled co-movement across pods that a purely pod-level stop-loss rule cannot detect, since each pod’s own P&L may look unremarkable even as the hidden common shock builds.

Scope and optional extensions

The model as it stands — static overlay for correlation instability, validated in simulation, confirmed as the exact dynamic optimum under the independence assumptions in Section 7, tested against fat tails, discrete jumps, and an asymmetric-crowding same-day layer — is a complete, internally consistent piece of work, and we do not think it needs further mechanical elaboration to make its point. The items below are optional directions for a reader who wants to take a specific piece further, not a to-do list the paper’s conclusions depend on.

- Shorter estimation windows for the correlation-instability signal would trade detection speed against noise, and might narrow the residual gap on sudden-onset risk quantified in Section 8.2.
- The same-day trigger of Section 9 was tested against one specific asymmetric-crowding calibration; a reader who wants to settle how sensitive its tail-protection benefit is to the trigger threshold, cooldown length, or the severity of the underlying shared-shock events would find that a natural, low-cost extension of the same simulation framework.
- Empirical calibration remains unavailable (platform-level exposure and seeding data), so the contribution stays theoretical-plus-simulation with a stated, falsifiable prediction, as previously agreed.

References

- Anderson, E. W., Hansen, L. P., and Sargent, T. J. (2003). "A Quartet of Semigroups for Model Specification, Robustness, Prices of Risk, and Model Detection." *Journal of the European Economic Association*, 1(1): 68–123.
- Hansen, L. P. and Sargent, T. J. (2008). *Robustness*. Princeton, NJ: Princeton University Press.
- Hu, Y., Imkeller, P., and Müller, M. (2005). "Utility Maximization in Incomplete Markets." *Annals of Applied Probability*, 15(3): 1691–1712.
- Khandani, A. E. and Lo, A. W. (2007). "What Happened to the Quants in August 2007?" *Journal of Investment Management*, 5(4): 5–54.
- Kobylanski, M. (2000). "Backward Stochastic Differential Equations and Partial Differential Equations with Quadratic Growth." *Annals of Probability*, 28(2): 558–602.
- Ledoit, O. and Wolf, M. (2004). "Honey, I Shrunk the Sample Covariance Matrix." *Journal of Portfolio Management*, 30(4): 110–119.
- Maenhout, P. J. (2004). "Robust Portfolio Rules and Asset Pricing." *Review of Financial Studies*, 17(4): 951–983.
- Merton, R. C. (1973). "An Intertemporal Capital Asset Pricing Model." *Econo-*

metrica, 41(5): 867–887.

Moreira, A. and Muir, T. (2017). “Volatility-Managed Portfolios.” *Journal of Finance*, 72(4): 1611–1644.

Whittle, P. (1990). *Risk-Sensitive Optimal Control*. Chichester: Wiley.