

OLD SHOREHAM INVESTMENT MANAGEMENT

Does Behavioural Product Design Pay for Itself?

Assignment Value, Manufacturing Spread And The Boundary
for Behavioural Instruments

Working Paper V.o.1

Oladayo Oduwole
July, 2026

Abstract

Does behavioural product design pay for itself? A payoff feature that keeps an investor invested through a fall has a cost, and the question is whether the implementation benefit exceeds it. Answering it requires an accounting that most behavioural finance does not supply, and this paper builds one and then applies it.

The accounting requires a vector of triggers rather than a single exit threshold. An investor is characterised by thresholds on several path statistics: loss from inception, drawdown from a running peak, shortfall against a reference asset, elapsed inactivity. The coordinates matter because they carry different actions and therefore different mathematics. A depth trigger liquidates, which is optimal stopping and admits a reduction to constrained optimisation. A shortfall trigger reallocates, which is a state-dependent regime switch and admits no pathwise constrained-set representation at all; the recurrent version, in which the investor switches repeatedly, generalises to impulse control and carries the same obstruction. Measuring only the loss arm overstates the benefit a protective feature can recover, because it credits the feature with retention it does not deliver against the other arm.

Applied with that accounting, the answer is conditional and the condition is the point. A floored structure is worth 2.1 percentage points a year to an investor who abandons an unprotected position at a ten per cent loss, 1.3 points at twenty per cent, and nothing at thirty; a modest manufacturing spread reverses the ranking for the same investor. Conditional benefits are therefore economically large and sign-changing across investors.

The population economics cannot yet be established, because the trigger distribution relevant to product assignment is not identified: the incidence of panic selling reported in the brokerage literature is not an estimate of the probability that a floor reverses a ranking, and the paper does not treat it as one. What follows is that assignment value exists only where implemented rankings disagree across investors, and that a behavioural instrument earns its place only if it adds decision value over information already collected. Demographic predictors of abandonment are a credible baseline for that comparison, not a substitute for the instrument.

The paper takes no position on whether triggers are mistakes, and the refusal is load-bearing: naming them mistakes requires a normative benchmark for correct action, which is the judgment that keeping utility conventional was meant to exclude. One consequence is stated plainly. Where a default can be imposed, the evidence says impose it, because defaulted arrangements show the lowest trading of any group observed and cost nothing. The design problem here is second-best by construction, for the population where defaults are unavailable.

Two further results survive independently of effect size. Products agreeing on every order-invariant statistic, including average maximum drawdown, differ in implemented value by up to 2.6 percentage points a year, because they differ in the horizon remaining when a trigger binds. And a contractual floor supplies a uniform lower bound on implemented terminal wealth, holding at every trigger value, with the bound itself independent of the population distribution of triggers. That is a different argument for protection from the one conventionally made, and it is a bound rather than a claim that the attained worst case is invariant.

The aggregate investor return gap is not evidence for any of this. A model containing no path dependence, in which investors exit at a constant hazard unrelated to the path, reproduces its relationship to product volatility an order of magnitude better than any trigger specification. That is reported, and it is one reason the case here rests on individual-level evidence rather than on fund aggregates.

Keywords: implementation risk, path dependence, investor return gap, structured products, suitability, optimal stopping, impulse control

JEL: G11, G41, D14, D81, G51

1. Introduction

A protective payoff feature costs money. The case for buying one on behavioural grounds is that the investor will implement the protected product better than the unprotected one, and that the difference exceeds the cost. That is an arithmetic claim, and it has an answer.

Getting the answer requires an accounting that is harder than it looks, for a reason that is the paper's first contribution. Implementation failure is usually modelled as abandonment after losses, and a protective feature is then credited with preventing it. But abandonment is not confined to losses. In a controlled experiment reported in Section 9.1, subjects assigned to sustained-loss paths abandoned them at rates rising with severity, which a loss threshold explains, and subjects assigned to the best-performing path abandoned it at 78 per cent, which a loss threshold cannot explain at all. In brokerage records the probability of selling is a V-shaped function of unrealised profit, rising with the magnitude of gains and of losses alike. A protective feature that reduces exposure to the first arm typically increases exposure to the second, because the participation it surrenders is what causes the product to lag its reference in strong markets. An accounting with one coordinate credits the feature with the first effect and misses the second, and it therefore overstates what protection can recover.

The distinction is not cosmetic, because the two coordinates carry different actions and therefore different mathematics. A loss trigger liquidates: the wealth path terminates in the riskless asset. That is optimal stopping, and Section 5 shows it reduces to classical optimisation over a path-constrained feasible set. A shortfall trigger reallocates: the process continues under new dynamics. The specification used here is one-way, which makes it a state-dependent regime switch rather than repeated intervention, and Section 6 shows that no pathwise constrained-set representation exists for it, not as a matter of difficulty but as a matter of measurability. Recurrent comparative switching generalises to impulse control and inherits the obstruction.

With the accounting in place, the second contribution is the answer. It has three parts and they pull in different directions, which is why the question has resisted a clean resolution.

Conditional on the investor, the effect is large. A floored structure is worth 2.1 percentage points a year to an investor who abandons an unprotected position at a ten per cent loss, and 1.3 points at twenty per cent. The comparison is against what that investor actually does, not against a disciplined buy-and-hold benchmark they do not resemble.

In the population, the effect cannot yet be sized. Published work covering 653,455 brokerage accounts finds panic selling in fewer than one household in ten over thirteen years, at a rate near 0.1 per cent of investors in any month. That is the incidence of a particular liquidation event, and it is not an estimate of $\mathbb{P}(\tau^L \leq 10)$ or of the probability that a floor reverses a ranking. Since the population benefit is $\mathbb{E}_\kappa[B(\tau)]$ and Section 11.1 shows κ is not identified by the available data, the paper does not multiply a conditional benefit by an incidence rate and report the product as a population average.

The aggregate return gap is not the phenomenon. A model with no path dependence at all, in which investors exit at a constant hazard and remain out for a fixed period, reproduces the observed relationship between implementation cost and product volatility an order of magnitude better than any trigger specification tried. Churn accounts for the aggregate; triggers account for dispersion.

Section 11.3 notes a cheaper intervention where the objective is simply to suppress discretionary trading: default architecture, which the evidence says works and which costs nothing. Whether suppressing trading is itself welfare-improving is a judgment this paper declines to make. The design problem below addresses the population for whom defaults are unavailable.

Together these give the paper's economic conclusion, which is about where value can come from rather than how much of it there is. Because conditional benefits change sign across investors, assignment has value exactly where implemented rankings disagree, and a signal is worth acquiring only to the extent that it changes which product is chosen. **The value is in the identification and**

not in the geometry. Demographic predictors of abandonment supply a baseline against which any behavioural instrument must show incremental decision value, and Section 7.2 states the test.

Two results hold regardless of how large the behavioural effect turns out to be. Products agreeing on every order-invariant statistic can differ materially in implemented value, because such statistics record how much risk a product carries and not when it is experienced. And a contractual floor supplies a uniform lower bound on implemented terminal wealth at every trigger value, with the bound independent of a population distribution nobody observes.

Section 2 sets up products and investors. Sections 3 to 6 develop the field, separation, reduction and the impulse-control obstruction. Section 7 turns to design when the trigger is not observable. Section 8 is a hand-solvable worked example. Section 9 reports external evidence, including the estimated trigger function. Section 10 reports what the model implies. Section 11 treats identification, dynamic inconsistency and welfare. Section 12 concludes.

2. Setup

2.1 Products

Time is discrete, indexed $t = 0, 1, \dots, T$. A **product** is a probability measure μ on the space Ω of positive wealth paths $\omega = (\omega_0, \dots, \omega_T)$ with $\omega_0 = 1$. Wealth is expressed in units discounted by the riskless asset where convenient; where it is not, the riskless rate is written r .

A product is characterised for a holding investor by its terminal law μ_T . It is characterised for an implementing investor by considerably more, as Section 3 makes precise.

2.2 Path statistics

A **path statistic** is a measurable functional $g : \Omega \times \{0, \dots, T\} \rightarrow \mathbb{R}$ adapted to the natural filtration, so that $g_t(\omega)$ depends only on $\omega_0, \dots, \omega_t$. The statistics used here:

Symbol	Statistic	Order-invariant	Exit-stable
g_t^L	loss from inception, $1 - \omega_t$	no	yes
g_t^D	drawdown from running peak, $1 - \omega_t / \max_{s \leq t} \omega_s$	no	yes
g_t^G	unrealised gain from entry, $\omega_t - 1$	no	yes
g_t^F	shortfall against a reference asset, $1 - \omega_t / v_t$	no	no
g_t^I	periods since last action	no	no
g_t^R	irregularity of realised outcomes	partly	no

A statistic is **exit-stable** if freezing wealth freezes the statistic: $\omega_s = \omega_t$ for all $s > t$ implies $g_s = g_t$.

Lemma 1 (a sufficient class). Any statistic of the form $g_t = h(\omega_t, \min_{s \leq t} \omega_s, \max_{s \leq t} \omega_s)$ is exit-stable. Dependence on t directly, or on a process v not measurable with respect to $\mathcal{F}^\omega$, can destroy exit-stability, and does so whenever the statistic actually varies with that argument once wealth is frozen. It need not. Compare $g_t = \max_{s \leq t} (1 - \omega_s)$, in which t appears only as the upper limit of a running operation over the path: freezing ω adds no new terms and the statistic stops. Contrast $g_t = t - \sigma_t$, periods elapsed since the last action, which advances whether or not wealth moves. The first is exit-stable and the second is not, though both carry t explicitly.

Proof. Freezing ω freezes its level and its running extrema, so any b of those three is frozen. For the second claim, t advances and v evolves whatever the investor does. $\square$

The displayed class is sufficient, not exhaustive. Cumulative realised variation $\sum_{s \leq t} |\omega_s - \omega_{s-1}|$ is exit-stable, since freezing ω adds no further increments, yet it is not recoverable from the current level and the running extrema. The lemma is used only to certify particular coordinates, which it does. Depth statistics qualify: an investor who liquidates at t has a loss and a drawdown that do not subsequently move. Reference-relative and duration statistics do not: the shortfall against a rising benchmark continues to grow after liquidation, and so does elapsed time. Exit-stability is what permits the reduction of Section 5 and its absence is what blocks it in Section 6.

2.3 The trigger vector

An investor is characterised by a **trigger vector** $\tau = (\tau^L, \tau^D, \tau^F, \tau^I, \tau^R) \in (0, \infty]^5$, one threshold per statistic, with ∞ denoting a coordinate that never fires. Write $\mathcal{T}$ for the space of trigger vectors and κ_x for the law of τ in a population of investors indexed by x .

The **implementation time** for investor τ in product μ is

$$\rho(\tau, \mu) = \inf\{t : g_t^k \geq \tau^k \text{ for some coordinate } k\},$$

with $\rho = \infty$ if no coordinate fires. Where two coordinates cross in the same period, the depth coordinate takes precedence. In discrete time such ties carry positive probability, so the convention is a modelling choice rather than a negligible one; nothing below depends on which way it is resolved, since the results concern one binding coordinate at a time.

Assumption D (coordinate-action mapping). Each coordinate carries a fixed action. The depth coordinates g^L, g^D and the gain coordinate g^G liquidate to the riskless asset. The shortfall coordinate g^F reallocates to the reference asset. The inactivity coordinate reallocates to whatever is salient. The action is a property of the coordinate rather than of the investor.

Assumption D is a modelling choice and not a consequence of anything above. It does most of the work in what follows, because the action determines exit-stability and therefore which coordinates admit Theorem 1.

Two coordinates fire on gains and they are not the same coordinate. g^G records profit against the investor's own entry price and its action is a sale. g^F records performance against a reference asset and its action is a rotation. The evidence in Section 9.6 is for the first: the V-shaped selling function is estimated on unrealised profit per position and what it documents is liquidation. Because g^G is exit-stable, the empirically supported gain arm falls inside the class Theorem 1 covers. It is g^F , the coordinate carried on structural grounds rather than estimated, that falls outside it. The evidential and the mathematical boundaries do not coincide, and it is worth being explicit that the untested coordinate is the awkward one.

The coordinates above are not on an equal evidential footing and the list is not asserted to be exhaustive. The depth coordinates and the gain-realisation behaviour that motivates a second arm have been estimated on transaction data (Section 9.6). The shortfall coordinate g^F has not. It is included because it is the statistic a participation-limited payoff is structurally exposed to, so any accounting of what such a payoff costs its holder must contain something like it. Section 9.6 establishes that a second, non-loss selling mechanism exists and is large, and the coordinate that mechanism corresponds to is g^G , not g^F . Benchmark shortfall is a distinct hypothesis carried because it is what a participation-limited payoff is exposed to, and other mechanisms would serve. It is a hypothesis about the mechanism, carried through the paper as one, and Section 8 demonstrates its consequences rather than its existence.

2.4 Implemented value

Investor τ holds increasing, strictly concave U over terminal wealth. Let $\Phi(\mu, \tau)$ denote the terminal wealth law generated when product μ is implemented by investor τ . Implemented value is

$$J(\mu, \tau) = \mathbb{E}[U(W)], \quad W \sim \Phi(\mu, \tau),$$

and the **implemented certainty equivalent**, annualised, is

$$\text{CE}(\mu, \tau) = [U^{-1}(J(\mu, \tau))]^{1/T} - 1.$$

Intended value $V(\mu) = \mathbb{E}[U(\omega_T)]$ is the special case $\tau = \infty$. The **implementation drag** is $\text{CE}(\mu, \infty) - \text{CE}(\mu, \tau)$, which is non-negative for exit-to-riskless actions under the continuation condition of Section 5 and of ambiguous sign for reallocation actions.

Two remarks on the choice of certainty equivalent over expected return. First, protected products are not bought for their mean, and a mean-based criterion cannot rank them against unprotected ones in a way that reflects why anyone holds them. Second, the certainty equivalent inherits U , which keeps conventional risk aversion in the model and prevents the behavioural layer from silently substituting for it.

3. The stopped-wealth field under a trigger vector

Definition 1 (stopped-wealth field). For product μ , the stopped-wealth field is the family

$$\Psi_\mu = \{\Phi(\mu, \tau) : \tau \in \mathcal{T}\}$$

indexed by the trigger vector. The terminal law μ_T is the single point $\Phi(\mu, \infty)$.

Remark 1 (sufficiency). Implemented value depends on μ only through Ψ_μ . Two products with identical stopped-wealth fields deliver identical implemented value to every investor and every increasing U .

Proof. $J(\mu, \tau)$ is a functional of $\Phi(\mu, \tau)$ alone, and Ψ_μ contains $\Phi(\mu, \tau)$ for every τ . $\square$

The remark is definitional. It states what product information suffices. It is not a claim that the field is estimable; it is a claim that anything less than the field is insufficient, and the point of indexing by a vector is that fields agreeing on one coordinate can differ on another.

Proposition 1 (coordinate-wise non-equivalence). There exist products $\mu \neq \tilde{\mu}$ and coordinates $j \neq k$ such that implemented value agrees across all investors whose only finite coordinate is j , and differs for some investor whose finite coordinate is k .

Proof. Two periods, two equally likely states. Let

$$\mu : 1 \rightarrow 1.20 \rightarrow 1.20 \text{ (good)}, \quad 1 \rightarrow 0.80 \rightarrow 1.00 \text{ (bad)},$$

$$\tilde{\mu} : 1 \rightarrow 1.00 \rightarrow 1.20 \text{ (good)}, \quad 1 \rightarrow 0.80 \rightarrow 1.00 \text{ (bad)},$$

and let the reference follow $\nu : 1 \rightarrow 1.10 \rightarrow 1.15$ in the good state.

The terminal laws are identical, and so is the entire path in the bad state. Loss from inception is zero at every date in the good state for both products and (0.20) at $t = 1$ in the bad state for both. Every loss coordinate therefore fires at the same dates on the same states and produces the same implemented terminal wealth, at any threshold.

In the good state at $t = 1$, μ stands at (1.20) and is ahead of the reference, while $\tilde{\mu}$ stands at (1.00) and is $1 - 1.00/1.10 = 9.1$ behind it. Any shortfall coordinate with $\tau^F < 9.1$ fires on $\tilde{\mu}$ and not on μ . The investor then holds the reference from (1.10) to (1.15) and ends at $1.00 \times 1.15/1.10 = 1.045$, against (1.20) had the trigger not fired. Implemented terminal wealth differs, not merely the trigger event. $\square$

Agreement on one coordinate therefore carries no information about another, at any threshold on either. Section 8 supplies the economically interesting version, in which two products differ on both coordinates in opposite directions.

The practical content: a suitability process that elicits one trigger and matches on it will assign products that are equivalent on that coordinate and materially different on another. This is the

formal version of the observation that a client who says they can tolerate a thirty per cent fall may nonetheless abandon a position that has merely lagged.

4. Order-invariance and separation

Suitability systems in practice summarise products with statistics invariant to the temporal ordering of returns: volatility, Sharpe and Sortino ratios computed from pooled returns, horizon value at risk and expected shortfall, and the terminal wealth distribution.

Definition 2. A criterion C is **order-invariant** if $C(\mu) = C(\mu \circ \pi)$ for every finite permutation π of the period returns.

Proposition 2 (separation). With one risky asset, one riskless asset and an exchangeable return process, permutations of a deterministic exposure schedule agree under every order-invariant criterion. Under the continuation condition of Section 5, two such products have identical intended value and strictly different implemented value for any investor with a finite depth coordinate that binds on one schedule and not the other.

Proof sketch. Exchangeability makes the pooled return law and the terminal law invariant to the schedule permutation. The running-maximum drawdown is not invariant: a front-loaded high-exposure schedule generates its deep excursions early, when the remaining horizon over which the surrendered risk premium compounds is long. Under strict continuation value, terminating early with more horizon remaining is strictly worse. $\square$

The relevance to Proposition 2 is direct. The order-invariant statistics are precisely those that carry no information about *when* an excursion occurs, and timing rather than depth is what determines the cost of a trigger firing. Two products can match on volatility, downside deviation, skewness and upside capture, and even on average maximum drawdown, and still differ in implemented value, because they differ in the expected remaining horizon at the moment the trigger binds.

5. Constraint equivalence, scoped

The following restates the reduction result of the scalar framework, with its scope now made explicit.

Assumption A (exit-stable coordinate). The binding coordinate is exit-stable, so that freezing wealth freezes the statistic.

Assumption B (behavioural exit). The action at ρ is a behavioural response to the experienced path rather than a rational re-optimisation on new information about the return distribution or on a private liquidity shock.

Assumption C (strict continuation). On the trigger event, continuing under the intended strategy is strictly preferred in utility terms to liquidating and holding the riskless asset:

$$\mathbb{E}[U(W_T^\pi) \mid \mathcal{F}_\rho] > U(W_\rho e^{r(T-\rho)}).$$

Assumption C is stated directly on utility rather than on wealth. A submartingale condition on discounted wealth is not sufficient: concavity of U can reverse the comparison, since the continuing strategy carries risk that the liquidated one does not.

Throughout this section wealth and the statistic g are expressed relative to the riskless numeraire, so that holding the riskless asset freezes both. Let $\mathcal{A}$ denote the admissible strategy set, the predictable processes generating wealth paths in Ω with the usual integrability, and let $\mathcal{J}_\tau : \mathcal{A} \rightarrow \mathcal{A}$ be the **implementation operator**, which follows π until $\rho(\tau, \pi)$ and holds the riskless asset thereafter.

The relevant feasible set is defined by the **stopping rule**, not by a bound on the level of g :

$$\mathcal{E}^{\mathcal{A}}(\tau) = \{\pi \in \mathcal{A} : \pi_t \text{ is riskless for all } t \geq \rho(\tau, \pi)\}, \quad \rho(\tau, \pi) = \inf\{t : g_t \geq \tau\}.$$

Two things follow from this definition and both matter. Before ρ the statistic lies strictly below τ automatically, by definition of ρ , so no level constraint has to be imposed. At and after ρ the definition places no restriction on g at all, which is what makes the set well behaved in discrete time.

The level formulation would not survive discrete time. Writing the set as $\{\pi : g_t \leq \tau \forall t\}$ fails on two counts. It admits strategies that touch τ and recover, which an investor who liquidates on first contact does not implement. And in discrete time the statistic generically **overshoots**: the first observation at or beyond τ is strictly beyond it. The worked example in Section 8 is an instance. Investor A carries a twelve per cent threshold and the first loss observed is fifteen per cent, so the implemented strategy violates $g \leq \tau$ at the moment it absorbs. A level-bounded set would exclude the very strategy the investor implements. The absorbing formulation is indifferent to the size of the overshoot and therefore needs no continuity assumption.

Behaviour does not say “never exceed a twenty per cent drawdown”. It says “if twenty per cent is reached, this investor cannot continue”, and the second is what $\mathcal{C}^A(\tau)$ encodes.

Theorem 1 (constraint equivalence, depth coordinates). Under Assumption A, with $\mathcal{J}_\tau(\mathcal{A}) \subseteq \mathcal{A}$,

$$\sup_{\pi \in \mathcal{A}} \mathbb{E}[U(W_T^{\mathcal{J}_\tau(\pi)})] = \sup_{\pi \in \mathcal{C}^A(\tau)} \mathbb{E}[U(W_T^\pi)].$$

Implemented optimisation is therefore ordinary expected-utility maximisation over an investor-specific **absorbing** feasible set.

Proof. Three steps, none requiring continuity. First, $\mathcal{J}_\tau(\pi) \in \mathcal{C}^A(\tau)$ for every $\pi \in \mathcal{A}$: the operator holds the riskless asset from ρ onward by construction, which is the defining requirement, and the value of g_ρ is immaterial to membership. Second, $\mathcal{J}_\tau(\pi) = \pi$ for every $\pi \in \mathcal{C}^A(\tau)$: either $\rho = \infty$, or π is already riskless from ρ onward, so the operator changes nothing. Note that ρ is the same stopping time under π and $\mathcal{J}_\tau(\pi)$, since the two agree strictly before ρ and ρ depends only on the path up to and including that date. Third, terminal wealth under $\mathcal{J}_\tau(\pi)$ equals implemented terminal wealth under π by construction. The first two give $\mathcal{J}_\tau(\mathcal{A}) = \mathcal{C}^A(\tau)$, and the third transfers the objective. $\square$

Exit-stability is what makes the absorbed statistic well defined after ρ : wealth is frozen relative to the numeraire, so g is frozen at g_ρ , whatever the overshoot happened to be. The result holds in discrete and continuous time alike, and the discrete case is the one the worked example uses.

Overshoot does not threaten the equivalence but it does affect magnitudes. The investor exits at g_ρ rather than at τ , and the drag reflects the realised depth rather than the nominal threshold: in the worked example a twelve per cent threshold produces exit at a fifteen per cent loss, so three percentage points of the surrendered wealth are attributable to the observation grid rather than to the investor. The distribution of $g_\rho - \tau$ depends on the product’s return process and coarsens with the observation interval. Everything the paper reports is computed at the realised exit, so D is measured inclusive of overshoot; a finer grid would reduce it, and the elicited object τ is a nominal threshold rather than a realised exit depth.

Assumption C plays no part in this. It is required only to sign the implementation drag, and it is stated here because the drag’s sign is used in Section 10.

Corollary 1 (signed drag). Under Assumption C, $\text{CE}(\mu, \tau) \leq \text{CE}(\mu, \infty)$ for every μ and every τ whose binding coordinate liquidates.

Corollary 2. Behaviour changes the feasible set rather than the welfare criterion. The behavioural contribution is an elicited personal feasibility constraint, not a new utility function.

Assumption B is load-bearing and is not identifiable from observational data. An exit observed after a large fall is consistent with a behavioural threshold and equally consistent with a liquidity shock correlated with market conditions. The insurance literature discussed in Section 9 documents exactly this ambiguity: surrender activity rose in the worst months of the 2008 crisis even among policyholders whose guarantees were deeply valuable, which is attributed to financial distress rather than to the guarantee losing its appeal.

6. Switching triggers, and what does not extend

Theorem 1 does not extend to the shortfall coordinate.

The obstruction is Assumption A. A shortfall statistic $g_t^F = 1 - \omega_t/\nu_t$ compares the position to a reference asset that keeps moving after the investor has acted. Freezing wealth does not freeze the statistic. The stopped counterpart construction fails, and with it the equivalence.

There is a second, deeper obstruction. The action at a shortfall trigger is not liquidation but reallocation, so the wealth process does not terminate; it continues under the dynamics of the new position. The specification used in this paper is one-way, as Section 6.1 makes explicit, and is therefore a state-dependent regime switch rather than repeated intervention. Allowing the investor to switch back generates a sequence of intervention times and post-intervention allocations, which is impulse control, and the value function then satisfies a quasi-variational inequality rather than a free-boundary problem. In neither case does the constrained-set representation of Theorem 1 have an analogue.

6.1 The problem, stated

Two distinct problems are involved and separating them clarifies what is and is not open.

The forward problem is well posed. Given the impulse rule, implemented value is a computation, not an optimisation. Let ν be the reference process, let $\theta_0 = 0$ and

$$\theta_{m+1} = \inf \left\{ t > \theta_m : 1 - \frac{\omega_t / \omega_{\theta_m}}{\nu_t / \nu_{\theta_m}} \geq \tau^F \right\},$$

with holding indicator $i_t \in P, R$ flipping at each θ_m and wealth evolving under the dynamics of the current holding between intervention times. **The comparison rebases at each intervention.** Without rebasing the rule chatters: an investor who switches while behind the reference is still behind it by the same margin an instant later, so the trigger re-fires immediately and the process is not well defined. **The rule as specified is one-way.** Once the investor holds the reference, the rebased ratio is identically one and the shortfall coordinate cannot fire again. What Section 8 uses is therefore a regime-switching problem with post-intervention dynamics, closer to optimal stopping with a changed continuation than to repeated impulse control. Recurrent comparative triggers, in which the investor can also abandon the reference for what it left behind, generalise naturally to impulse control and carry the obstruction in Proposition 3 with them; nothing below relies on the recurrent case.

Rebasing is also the natural reading rather than a technical convenience. An investor who switches into the reference at a moment when they have been lagging it is, from that moment, invested in it; they are no longer behind it by anything. What they can subsequently fall behind is the alternative they left. Rebasing makes the comparison run from the switch rather than from inception, and it is the specification used in every simulation below. Then $\Phi(\mu, \tau)$ is the law of ω_T under that construction and $J(\mu, \tau) = \mathbb{E}[U(\omega_T)]$ is well defined. Every simulation in Section 10 evaluates exactly this.

The optimising benchmark is a quasi-variational inequality. This paragraph is written in continuous time, where the object is standard; the rest of the paper is discrete. If instead the investor chose intervention times to maximise value, with proportional cost c per switch, the state must include the reference, since the trigger reads $1 - \omega_t/\nu_t$. Writing $V(t, w, \nu, i)$, or $V(t, w, w/\nu, i)$ in the scale-invariant case,

$$\min \left\{ -\partial_t V - \mathcal{L}^i V, V - \mathcal{M}V \right\} = 0, \quad \mathcal{M}V(t, w, \nu, i) = \sup_{j \neq i} V(t, w(1-c), \nu, j),$$

with $\mathcal{L}^i$ the joint generator of wealth and reference under the holding in force, and terminal condition $V(T, w, \nu, i) = U(w)$. The enlargement of the state is exactly what Proposition 3 shows

cannot be avoided. The continuation region is where the first term binds and the intervention region where the second does.

The behavioural investor does not solve this. The QVI is stated because it identifies what the behavioural rule departs from, and because the intervention region it defines is a comparison object for the trigger boundary τ^F .

Proposition 3 (no constrained representation). Let the shortfall coordinate be the binding one and let ν not be measurable with respect to $\mathcal{F}^\omega$. Then no set $C \subseteq \Omega$ of wealth paths represents shortfall-triggered implementation **pathwise and uniformly over reference paths**: there is no C for which membership in C coincides with being implemented as intended, simultaneously for every realisation of ν .

Proof. Fix a candidate C and two reference paths ν, ν' that agree at T but differ at an intermediate date. Choose a wealth path ω that crosses the shortfall threshold against ν and not against ν' . Under ν the investor intervenes and ω is not implemented as intended, so $\omega \notin C$; under ν' no intervention occurs and $\omega \in C$. Since C depends on ω alone, both cannot hold. $\square$

The claim is deliberately about representation rather than about the optimum value: a set chosen artificially might happen to attain the same supremum without representing the rule, and nothing here excludes that. What is excluded is the pathwise correspondence that Theorem 1 supplies for depth coordinates.

Proposition 3 is the precise reason Theorem 1 stops where it does. The obstruction is not technical difficulty; the pathwise representation does not exist. Any reduction for switching triggers must enlarge the state to include the reference, which is what the QVI above does and what Theorem 1 avoids.

6.2 Three properties

Three properties of the switching problem bear on what follows.

Remark 4 (sign). The implementation drag from a switching trigger has ambiguous sign. Switching out of a lagging position into a leading one is destructive when the lag is transient and constructive when it is informative. Depth triggers, by contrast, have a signed drag under Assumption C.

Remark 5 (product exposure). A product's exposure to a shortfall trigger is governed by how far and how persistently it can lag its natural reference. A capped or participation-limited structure has a bounded and predictable lag against its underlying, and a large one in strong markets. This is a design variable that the depth-oriented literature does not consider and that Section 8 shows can dominate the depth channel entirely.

Remark 6. Section 9.2 reports that neither a depth coordinate nor a depth-plus-shortfall pair reproduces the observed relationship between aggregate implementation cost and product volatility, and that a trigger-free churn model does. The shortfall coordinate is therefore motivated by the experimental evidence of Section 9.1 and by the construction of Section 8, not by the aggregate cross-section.

7. Ranking reversal, assignment and payoff design

7.1 Ranking reversal

The central result is an accounting identity, and its force comes from what it makes comparable rather than from any difficulty in deriving it.

For product μ and investor τ , define the **implementation drag**

$$D(\mu, \tau) = \text{CE}(\mu, \infty) - \text{CE}(\mu, \tau),$$

and for a pair of products the **technical advantage** of μ over ν ,

$$H = \text{CE}(\mu, \infty) - \text{CE}(\nu, \infty).$$

Theorem 2 (ranking reversal). For any μ, ν, τ ,

$$\text{CE}(\nu, \tau) - \text{CE}(\mu, \tau) = [D(\mu, \tau) - D(\nu, \tau)] - H.$$

Hence ν is preferred to μ on implemented grounds if and only if

$$D(\mu, \tau) - D(\nu, \tau) > H.$$

Proof. Substitute $\text{CE}(\mu, \tau) = \text{CE}(\mu, \infty) - D(\mu, \tau)$ and likewise for ν , then use the definition of H . $\square$

Corollary 3. A technically superior product ceases to be the suitable product for an investor exactly when its excess implementation drag exceeds its technical advantage.

That is the proposition this framework exists to state precisely. It does not assume its conclusion in the welfare function: U is conventional, identical across investors in the worked example, and both terms on the right are measured in the same units of annualised certainty equivalent.

The region is non-empty. Section 8 supplies an instance. The unprotected position beats the note by $H = 7.12 - 5.72 = 1.40$ percentage points a year under holding. For Investor A the unprotected position carries drag of $7.12 - 5.39 = 1.73$ points while the note carries none, since A never triggers in it. Since $(1.73 > 1.40)$, the ranking reverses and the note leads by 33 basis points.

Theorem 2 also organises what each part of the paper contributes. Sections 3 and 4 concern what product information is needed to compute D . Section 5 shows that for depth coordinates D arises from a feasibility constraint rather than a modified objective. Section 6 shows why D is harder to characterise for switching coordinates. Section 7.2 concerns design when τ , and therefore D , is not directly observed.

Proposition 4 (assignment value). For a two-product menu μ, ν and a population law κ of trigger vectors, define $\Delta(\tau) = \text{CE}(\mu, \tau) - \text{CE}(\nu, \tau)$. The value of assigning products to investors rather than offering the population-best single product is

$$M = \frac{1}{2}(\mathbb{E}|\Delta| - |\mathbb{E}\Delta|) \geq 0,$$

with equality if and only if Δ has constant sign κ -almost surely.

Proof. $\mathbb{E} \max(\Delta, 0) = \frac{1}{2}(\mathbb{E}|\Delta| + \mathbb{E}\Delta)$, while the population-best single product yields $\max(\mathbb{E}\Delta, 0)$. $\square$

Assignment value comes from **disagreement**, not from average superiority. A product better for almost everyone creates no assignment value however good it is; the correct response is to sell it to everyone. By Theorem 2 the sign of Δ is governed by the difference in drags relative to H , so disagreement requires the drag difference to straddle H across the population. Payoff features that reduce drawdown tend to move drag in the same direction for most investors and therefore generate little disagreement. Features that trade one coordinate against another generate more.

Definition 3 (implementation-attributable assignment). Assignment value is implementation-attributable if the sign of Δ varies across investors holding U fixed. Otherwise an observed reversal is conventional preference heterogeneity relabelled. Trigger coordinates plausibly correlate with risk aversion, and an assignment computed over the marginal law of τ would attribute to implementation whatever disagreement risk aversion produced on its own. The worked example holds U fixed across both investors for this reason.

The confound is real and it has a location. Repeating the Section 8 comparison under constant relative risk aversion γ , holding the twelve per cent threshold fixed:

γ	H	excess drag $D(\mu, \tau) - D(\nu, \tau)$	reversal
0.5	1.99%	1.94%	no
1 (log)	1.40%	1.74%	yes
2	0.11%	0.96%	yes

γ	H	excess drag $D(\mu, \tau) - D(\nu, \tau)$	reversal
3	-1.24%	-0.13%	not applicable
5	-3.71%	-2.52%	not applicable

Three regions. Below logarithmic utility the excess drag does not cover the technical advantage and no reversal occurs. At $\gamma \geq 3$ the technical advantage is itself negative: the floored structure beats the unprotected one under holding, so there is nothing for behaviour to overturn and a risk-averse investor preferring a floor is what conventional theory already predicts. **The behavioural mechanism does economic work only in the intermediate region**, roughly $\gamma \in [1, 2]$ here, where the structure is technically inferior and implementation reverses it. Outside that band an observed preference for protection is uninformative about implementation, which is Definition 3 with a boundary attached rather than merely a caution.

7.2 Design when the trigger is not observable

For a known trigger vector the implemented-optimal payoff is

$$G^*(\tau) \in \arg \max_{G \in \mathcal{G}} \text{CE}(G, \tau),$$

over a set $\mathcal{G}$ of fairly priced structures at a common investment amount and cost of manufacture. This sits downstream of pricing: arbitrage-free valuation determines which payoffs can be supplied and at what cost, and the rule determines which feasible payoff delivers the highest implemented certainty equivalent.

$G^*(\tau)$ is not an operational rule. The trigger vector is not observable at the point of sale, and the obstruction is not that the investor withholds it. An investor does not know their own break point in advance; it is discovered during the event that provokes it, which is why the problem exists at all.

Unobservability is not ignorance, and the distinction carries the section. Let X denote conventional information already collected, including the demographic fields Section 9.6 shows to be predictive, and let S denote the output of any behavioural instrument. The relevant object is the posterior $\kappa(\tau | X, S)$, and the design problem separates by horizon. **Manufacturing** happens before the holder is known and is conducted against the population law κ . **Assignment** happens after and is conducted against the posterior. Nothing requires τ to be recovered exactly; it requires the posterior to move a decision.

The estimand for a behavioural instrument. With menu $(\{j\})$ and conditional implemented certainty equivalents CE_j ,

$$V(S | X) = \mathbb{E} \left[\max_j \mathbb{E}(\text{CE}_j | X, S) \right] - \mathbb{E} \left[\max_j \mathbb{E}(\text{CE}_j | X) \right] \geq 0,$$

non-negative by Jensen's inequality applied to the maximum. This is what an instrument must earn. It is zero exactly when the signal never changes which product is chosen, which is not the same as the signal carrying no information: a signal can predict τ accurately and be worthless if the conditional ranking never flips, and can be crude and valuable if it flips the ranking near the boundary in Theorem 2. The instrument does not need to recover a break point. It needs to move the posterior across H .

Concretely, in the parameterisation of Section 10.2 the sign of Δ changes between a break point near twenty per cent and one near forty. A signal that separates the shallow tail of κ from the rest clears the bar; a signal that distinguishes a twelve per cent threshold from a fifteen per cent one does not, because both lie on the same side of the crossing. The requirement is therefore a classification into a small number of cells, not a measurement, and an instrument with poor resolution can satisfy it while a precise one that never separates the tail cannot.

The remainder of this section treats manufacturing, where the holder is unknown by construction and the population law is the only object available.

Three criteria are available and they do not agree.

$$\begin{aligned} G^E &\in \arg \max_G \mathbb{E}_\kappa[\text{CE}(G, \tau)], \\ G^M &\in \arg \max_G \min_{\tau \in \text{supp } \kappa} \text{CE}(G, \tau), \\ G^R &\in \arg \min_G \max_\tau [\text{CE}(G^*(\tau), \tau) - \text{CE}(G, \tau)]. \end{aligned}$$

The first maximises average implemented value, the second the worst case over the support, the third the largest shortfall against the payoff that would have been chosen with knowledge of τ . Which criterion is adopted is a decision about how much is known, and Proposition 5 bears on the answer.

Proposition 5 (floor invariance). Let G carry a contractual floor F payable at maturity, let the binding coordinate be exit-stable with liquidation into the riskless asset as its action, let the issuer be default-free, and let liquidation execute at the arbitrage-free value. Then for every τ , implemented terminal wealth satisfies $W \geq F$ almost surely.

Proof. At any date t the arbitrage-free value of G is at least that of its floor component, $Fe^{-r(T-t)}$, since the residual claim is non-negative. Liquidation at t therefore realises at least $Fe^{-r(T-t)}$, and investing the proceeds at r through to T yields at least F . If no coordinate fires, $W \geq F$ by the contract. $\square$

Three scope conditions matter and none is innocuous. The bound fails for switching coordinates, whose action reinvests in a risky asset. It fails on issuer default. And it requires liquidation at fair value: a retail note sold into a dealer bid below the arbitrage-free price realises less than $Fe^{-r(T-t)}$, short by the spread. The proposition holds exactly for a self-manufactured or exchange-cleared structure and approximately, degraded by the bid, for an intermediated one.

The consequence for design under ignorance. A contractual floor supplies a **uniform lower bound** on implemented terminal wealth, and that bound is independent of κ : of its location, its dispersion and its shape. An unprotected position has no such bound, and its worst case deteriorates as the support of κ extends toward shallow thresholds. The bound is a bound; it does not follow that the attained worst-case certainty equivalent is itself invariant to κ .

The proposition bounds the worst case; it does not identify the maximin optimum. A deeper floor raises the bound, but whether it raises $\min_\tau \text{CE}(G, \tau)$ depends on whether the bound is attained and on how much participation the deeper floor surrenders at the trigger values where the minimum is actually taken. In the worked example it does not: the 95 and 100 per cent floors carry the higher bound and the lower implemented certainty equivalent across the range examined.

This is the sense in which protection is worth buying when τ is unknown, and it is a different claim from the one usually made for it. The floor is not primarily insurance against market outcomes, which conventional analysis already prices. It is insurance against the designer's ignorance of who is holding the product, and against the holder's ignorance of themselves.

The criteria conflict. Under G^E deeper floors lose. Under G^M the ranking depends on where the minimum is taken and has to be computed rather than read off the bound. What is invariant across criteria is the insensitivity of the floored payoff's downside to κ , and that is the property being bought.

Proposition 6 (marginal design rule under κ). Let G_0 and G_1 differ by a design modification, with $\Delta \text{CE}^{\text{hold}}$ the change in certainty equivalent under the hold assumption and $\Delta \text{drag}(\tau)$ the change in implementation drag at τ . Under the expected-value criterion the modification improves implemented value if and only if

$$\Delta \text{CE}^{\text{hold}} > \mathbb{E}_\kappa[\Delta \text{drag}(\tau)],$$

and under the maximin criterion if and only if it raises $\min_{\tau} \text{CE}(G, \tau)$.

For the protective structures considered here the left side is negative, since the participation surrendered exceeds what the floor returns in certainty-equivalent terms. That is not automatic: for a sufficiently risk-averse investor a fairly priced protective payoff can raise the hold certainty equivalent, in which case the rule is satisfied trivially. The rule is therefore a test of whether the drag reduction pays for the protection, and it is a test that protective features can fail. Three components of $\Delta \text{CE}^{\text{hold}}$ must be separated in application: the fair value of the embedded option, the issuer's credit spread benefit, and the distribution margin. Only the first is a cost of protection, and the worked example shows the third determining the answer.

What this implies for assignment. Realised assignment value is bounded by what the posterior supports, so it is $V(S | X)$ and not M that the instrument has to deliver. Since Section 9.6 shows X alone already predicts trigger firing, the baseline against which S is judged is a demographic model rather than nothing.

8. Worked example

The example is a three-period binomial, chosen so that every quantity can be computed by hand.

8.1 Primitives

The risky asset moves up by $u = 1.35$ or down by $d = 0.85$, each with real probability $(1/2)$, over $T = 3$ periods. The riskless rate is $r = 0.02$ per period. Expected gross return per period is (1.10) , so the risk premium is eight percentage points.

The risk-neutral probability is

$$q = \frac{1 + r - d}{u - d} = \frac{1.02 - 0.85}{0.50} = 0.34.$$

Both investors have logarithmic utility, so $\text{CE} = \exp(\mathbb{E}[\log W])$ and the annualised figure is $\text{CE}^{1/3} - 1$. Utility is held identical across the two investors, per Definition 3.

8.2 The unprotected position

The eight equally likely terminal values are

path	wealth	$\log W$
(uuu)	2.460375	0.900447
(uud, udu, duu)	1.549125	0.437695
(udd, dud, ddu)	0.975375	-0.024930
(ddd)	0.614125	-0.487611

$\mathbb{E}[W] = 1.331 = 1.10^3$. $\mathbb{E}[\log W] = 0.206391$, so $\text{CE} = 1.229218$ and the annualised hold certainty equivalent is **7.12 per cent**.

8.3 The note

A three-period note with a ninety per cent floor and participation π in the upside of the same underlying. Zero-coupon component: $0.90/1.02^3 = 0.848090$. The at-the-money call has terminal payoffs (1.460375) on (uuu) and (0.549125) on the three paths with two up moves, and zero otherwise, giving

$$C_0 = \frac{q^3(1.460375) + 3q^2(1 - q)(0.549125)}{1.02^3} = 0.172527.$$

With zero manufacturing margin the option budget is $1 - 0.848090 = 0.151910$, so $\pi = 0.151910/0.172527 = 0.881$. The floor is an unsecured claim on the issuer and is treated here as default-free; issuer credit is priced separately and enters as a deduction from the participation rate.

Terminal payoffs: (2.185873) on (uuu), (1.383508) on the two-up paths, (0.90) otherwise. This gives $E[\log W] = 0.166815$, $CE = 1.181534$, annualised **5.72 per cent**.

The note is inferior to the unprotected position under the hold assumption by 1.40 percentage points a year. Under classical criteria the choice is settled.

8.4 The two investors

Investor A has a finite loss coordinate, $\tau^L = 0.12$, and no other finite coordinate. On a twelve per cent loss from inception, A liquidates to the riskless asset for the remaining periods.

Investor B has a finite shortfall coordinate against the underlying, $\tau^F = 0.08$, and a loss coordinate of (0.35). On an eight per cent lag against the underlying, B switches into the underlying.

8.5 Mark-to-market of the note at $t = 1$

After a down move, the note is worth $0.90/1.02^2 + \pi \cdot C(0.85, 2) = 0.9188$, a loss from inception of **8.12 per cent**.

After an up move, the note is worth $0.90/1.02^2 + \pi \cdot C(1.35, 2) = 1.2165$, against the underlying at (1.35), a shortfall of **9.89 per cent**.

The two numbers are the whole example. The floor holds A's loss to 8.12 per cent, below A's 12 per cent threshold, so **A never triggers in the note**. The participation shortfall opens a 9.89 per cent lag, above B's 8 per cent threshold, so **B switches out of the note at $t = 1$** and rides the underlying thereafter.

8.6 Implemented outcomes

Investor / case	unprotected	note	preferred
hold, either investor	7.12%	5.72%	—
Investor A ($\tau^L = 0.12$)	5.39%	5.72%	note
Investor B ($\tau^F = 0.08$)	7.12%	5.77%	unprotected

Investor A liquidates the unprotected position on the four paths beginning with a down move, where the loss from inception reaches 15 per cent at $t = 1$, and on no others: after an up move the position never falls below its starting value within the horizon. That surrenders 1.73 percentage points a year to implementation. In the note, A implements as intended and captures the full 5.72 per cent. The note wins for A by 33 basis points a year despite being 140 basis points worse under the hold assumption.

Investor B implements the unprotected position perfectly, since a 35 per cent loss threshold never binds within three periods, and captures 7.12 per cent. In the note, B switches at $t = 1$ after an up move, which is neither disastrous nor helpful, ending at 5.77 per cent.

The implemented ranking reverses across two investors with identical utility functions, differing only in which trigger coordinate is finite. This establishes Theorem 2 and the non-degeneracy in Proposition 4, and it exhibits the instance promised in Proposition 1.

8.7 The margin

Repeating the calculation with a manufacturing margin deducted from the option budget:

margin	participation	note, hold	note for A	unprotected for A	A's choice
0.0%	88.1%	5.72%	5.72%	5.39%	note
0.5%	85.2%	5.48%	5.48%	5.39%	note
1.0%	82.3%	5.24%	5.24%	5.39%	unprotected
1.5%	79.4%	4.99%	4.99%	5.39%	unprotected

The breakeven margin is approximately seventy basis points. Below it the note is the implemented-preferred choice for A; above it, A is better served by the unprotected position despite abandoning it on four paths out of eight.

This is the practical content of Proposition 6's decomposition. The behavioural benefit available to A is 33 basis points a year at zero margin, and a seventy basis point spread consumes it.

That threshold sits below observed market spreads rather than above them. The structured products literature estimates retail issue prices materially above contemporaneous fair value, with the wedge concentrated in the distribution chain (Henderson and Pearson, 2011; C  lerier and Vall  e, 2017; Vokata, 2021). A three per cent upfront wedge amortises to roughly thirty basis points a year on a ten-year note and a hundred on a three-year one. Exchange-traded defined-outcome funds, which carry no upfront wedge, charge expense ratios near seventy basis points and reset annually, so the spread recurs where the note's does not.

Repeating the breakeven calculation across break points and floor levels shows how narrow the viable region is. Entries are the one-off margin, in basis points of notional, at which the note's implemented certainty equivalent equals the unprotected position's for that investor; a dash means the note cannot match it at any margin.

break point	floor 85%	floor 90%	floor 95%	floor 100%	unprotected, implemented
8%	—	—	—	—	5.39%
12%	203	69	—	—	5.39%
18%	—	—	—	—	6.68%
25%	—	—	—	—	6.68%
35%	—	—	—	—	7.12%

The region has a floor and a ceiling on both axes. Below a 12 per cent break point the note fails because its own mark-to-market loss of 8.12 per cent trips the investor's threshold: protection that is not deep enough to keep the holder inside their own tolerance buys nothing. Above 18 per cent the investor implements the unprotected position well enough that there is no drag to recover. On the floor axis, deeper protection is not better: the 95 and 100 per cent floors consume too much of the option budget, and the participation they leave cannot match the unprotected position even free of charge.

The manufacturing spread, not the payoff geometry, determines whether the product is worth holding, and the geometry that clears is partial protection at long tenor rather than full protection at short.

The model therefore predicts that at prevailing retail spreads these structures should not be bought by the investors they are sold to. That prediction is uncomfortable and it is not novel: it coincides with the finding in the structured products literature that retail issue prices sit materially

above fair value and that product complexity tracks the scope for catering rather than the scope for improvement. The framework here does not explain why the products are nonetheless sold. It supplies a criterion under which most of them should not be, and locates the small region in which they should.

9. External evidence

What follows are facts the model must be consistent with. Two supply the mechanism directly, two constrain it, one concerns the observable, and one rejects a reading of the aggregate evidence that the paper does not make.

9.1 Trigger heterogeneity is real and two-sided

A controlled experiment (Oduwale, 2006) assigned 94 postgraduate finance students to one of six investment paths over ten rounds, with realised sequences deliberately at odds with the probabilities disclosed at the outset, and recorded abandonment. Team-level results:

path	chose	abandoned
flat	28	7% (21% counting reallocations)
two early losses	18	33%
sustained loss	14	57%
deeper sustained loss	12	100%
sustained gain	18	78%

The round-by-round return schedules behind these paths are not recoverable, so the drawdown depths corresponding to each abandonment rate are unknown, and the distribution of thresholds is therefore not identified. Fitting a threshold distribution under alternative depth assumptions moves the implied dispersion parameter by a factor of four, and the two best-fitting assumptions disagree by a factor of three. The data does not select among them.

Two things are nonetheless identified. First, **thresholds are heterogeneous**. A single common threshold predicts a step function, no abandonment then universal abandonment, and cannot generate a graded 7, 33, 57, 100 response at any assignment of depths; a dispersed distribution fits 2.7 times better on root-mean-square error, and this conclusion is invariant to the depth assumption. Second, **abandonment is not confined to loss paths**. The best-performing path was abandoned at 78 per cent. No loss or drawdown coordinate can produce that, since no loss occurred.

The mechanism behind it is not identified, and the ambiguity matters for Assumption B. Subjects who left reported an expectation that the winning path would subsequently reverse, an expectation the design did not fulfil. If that expectation was a belief about the return distribution, the reallocation is rational updating on the subject's own information set and Assumption B fails for those observations; if it was a response to the run itself, it is a trigger. The experiment cannot separate the two. What survives either reading is that the action was a reallocation rather than a liquidation, which places it in the impulse-control class of Section 6 regardless of what drove it.

Two further limitations bear on how much weight this can carry. Subjects chose their team before paths were assigned, so the composition of each group is self-selected and the abandonment rates are not comparable across groups without that caveat. And the assignment of paths was at team level, so a common shock to a group cannot be distinguished from independent individual decisions.

9.2 The aggregate cross-section is explained without any trigger

Industry research covering approximately \$13.6 trillion of fund and exchange-traded fund assets over the decade to December 2025 (Ptak, 2026) estimates that the average dollar earned 8.7 per

cent a year against an aggregate fund return of 9.9 per cent. Sorted by volatility, least-volatile funds show a gap of 0.4 percentage points and most-volatile funds 2.1, with sector equity near 2.6 and target-date vehicles under 0.2. The ratio of cost to volatility across the three volatility-sorted buckets is 7.3, 9.2 and 8.4: close to proportional.

Three models were fitted to those three points, each by grid search on its threshold parameters against squared error, on a common ten-year simulation with drift set to a constant Sharpe ratio.

model	parameters	rms error	fitted
			cost-to-volatility
depth coordinate only	1	0.509 pp	1.1 / 8.1 / 15.1
depth and shortfall coordinates	2	0.441 pp	1.1 / 5.5 / 10.3
constant hazard, no trigger	2	0.068 pp	8.4 / 8.6 / 8.6

The depth coordinate produces a convex response: a fund at 5.5 per cent volatility rarely falls far enough to trip anyone, so the model assigns it six basis points against an observed forty. Adding a shortfall coordinate improves the fit marginally and does not remove the convexity. The third model contains no path dependence whatsoever: investors exit at a constant monthly hazard of 0.8 per cent, unrelated to anything the product has done, and remain in cash for twenty-four months. It reproduces the cross-section to within seven basis points, because the cost of being out of the market is the risk premium and the premium scales with volatility.

The aggregate return gap is therefore not evidence for the mechanism in this paper. It is consistent with ordinary churn at a rate that does not respond to the path at all. This is close to the position of a 2026 Financial Analysts Journal study reworking the same sample, which attributes roughly 0.10 percentage points a year to timing and the remainder to dollar-weighting arithmetic.

What follows for the framework is a sharpening rather than a refutation. A path-dependent trigger and constant churn are indistinguishable in the aggregate mean, because both scale with the premium. They are distinguishable in the cross-section of investors: churn predicts implementation cost uncorrelated with any investor characteristic and uncorrelated across products for the same person, while a trigger predicts both correlations. The observable that separates them is dispersion rather than level, and Section 9.6 reports individual-level evidence on precisely that, from transaction records rather than from fund aggregates.

9.3 A contractual guarantee moves the trigger, not only the path

Analysis of approximately seven million policy years across eight variable annuity issuers (Ruark Consulting, as reported in industry coverage; see also Milliman, 2011) finds that the presence of a living benefit reduces the baseline lapse rate by 40 to 70 per cent, and that in-the-moneyness of the guarantee reduces it by up to a further 80 per cent, taking a ten per cent baseline to under one per cent for deeply in-the-money contracts. Related actuarial practice models lapse as an explicit decreasing function of guarantee moneyness.

Two implications. First, the effect operates on the trigger rather than on the path: policyholders retained fully liquid contracts through a severe drawdown because surrender would have forfeited the guarantee. Second, the **hazard inverts**. In an unprotected position a deeper fall makes exit more likely; in a contract whose protection pays only on survival, a deeper fall makes exit less likely. A model in which protection acts solely by reducing drawdown depth cannot represent this, and understates the value of protection accordingly.

The same source records that surrender activity rose in the worst months of the 2008 crisis even for deeply in-the-money contracts, attributed to acute financial distress. That is Assumption B failing in the data, and it marks the boundary of the mechanism rather than a defect in it.

9.4 Structure suppresses action almost entirely

Retirement plan data covering roughly five million participants records that 5.3 per cent traded between January and April 2020, a four-month window containing a peak-to-trough equity decline of approximately one third. For calendar 2024, a calm year, 5 per cent initiated an exchange, and among pure target-date investors the figure was 1 per cent (Vanguard, 2020; 2025; 2026).

Those figures cover periods of different length and are not a like-for-like comparison; the four-month 2020 rate annualises above the calm-year rate. What they do support is weaker and still useful: through an extraordinary shock, absolute trading remained in single digits, and it was lowest where the arrangement was defaulted. The trigger is not a property of the person alone; it is a property of the person inside a structure, and default architecture, illiquidity and infrequent attention can suppress it to near zero. This bounds the population to which the framework applies: it is a model of investors who look and can act, not of defaulted participants.

9.5 What the available instrument resolves

A five-axis behavioural instrument administered to 63 respondents produces axis distributions that are highly uneven, and the unevenness bears on the scoping of this paper rather than being merely a defect of the sample.

One axis, loss sensitivity, has genuine spread. The other four do not, in three distinct ways. The regularity axis equals exactly 50, the scorer's default, in 39 of 63 records; no result from it should be reported, null or otherwise. The opportunism axis places 43 of 63 records at its ceiling (23 at 93 or above), its floor (11 at 21 or below), or the default of 50 (9 records), leaving 20 anywhere in between; it is behaving as a verdict with a fallback rather than as a graded score. Across roughly fifteen records the planning and engagement axes move in exact lockstep on a fixed ladder, the first falling by three as the second rises by four, so within that subgroup they are one quantity rendered as two coordinates.

The alignment with the theory is closer than it first appears. The coordinate the instrument resolves is the depth coordinate, which is the coordinate for which Theorem 1 holds. The coordinate it fails to resolve is the one that would populate the shortfall trigger, which is the coordinate Section 6 leaves open. Instrument and theory currently reach the same distance and stop in the same place.

Three limitations bound what the sample supports. Forty-three of 63 records sit in the youngest age band and four are above 45, with age recorded in bands rather than years; published work on panic selling locates the behaviour among older investors with dependents, so the sample is drawn from the cohort least likely to exhibit the target event. Three distinct families of stored values are present, consistent with more than one scoring path writing to the same columns, which would make stored distances non-comparable. And no respondent has both a profile and an observed implementation outcome.

Because the instrument resolves the depth coordinate and not the shortfall coordinate, the switching mechanism cannot be tested against it, and it remains untested elsewhere. Section 9.6 establishes that selling behaviour is multi-dimensional, with an arm that no loss threshold generates; it does not establish that the second arm is benchmark shortfall. Two things are therefore open rather than one: the shortfall coordinate itself, and the mapping from any instrument to it.

One quantity survives all three caveats. Sixteen of 63 records, roughly a quarter, are high on both the loss coordinate and the opportunism coordinate. Those are the investors for whom two triggers pull in opposite directions, and by Proposition 4 they are where assignment value, if it exists at all, is located.

9.6 The trigger function has been estimated on transaction data

The two claims this paper needs from outside are that a trigger function exists and that it has more than one arm. Both have been estimated at individual level on real money.

The function is V-shaped in both directions. Ben-David and Hirshleifer (2012), using US discount brokerage records, estimate the probability of selling as a function of unrealised profit and find it rises with the magnitude of gains and with the magnitude of losses alike, with little discontinuity at zero profit. The probability of buying more is V-shaped as well. Kaustia (2010), using the complete record of common stock trades by Finnish household investors from 1995 to 2000, finds selling propensity increasing in the magnitude of gains.

The loss arm is the depth coordinate. The gain arm is a coordinate no depth threshold can produce, and it corroborates the abandonment of the winning path in Section 9.1 on data whose return schedules are identified, which the experiment's are not.

The gain arm is not the shortfall coordinate. Ben-David and Hirshleifer estimate realisation of unrealised profit on a position, measured against its own purchase price. The coordinate g^F used in Section 8 measures a position against a reference asset. These are different statistics carrying different actions: the first sells, the second rotates. The evidence establishes that a second, non-loss selling mechanism exists and is large. It does not establish that the mechanism is benchmark shortfall. The shortfall coordinate is used here because it is what a participation-limited payoff is structurally exposed to, and it is a hypothesised design mechanism demonstrated theoretically rather than an estimated one.

Two qualifications carry through. Ben-David and Hirshleifer read the V-shape as belief revision rather than a preference for realisation, which is Assumption B contested in the best data that exists rather than only in an experiment. And the estimated function is a selling propensity per position, not an implementation threshold per person; recovering τ from it requires an aggregation across a household's positions that the published work does not perform.

The loss arm is predictable from characteristics already collected. Elkind, Kaminski, Lo, Siah and Wong (2022) study 653,455 brokerage accounts held by 298,556 households from 2003 to 2015, defining a panic sale as a fall of 90 per cent in a household's equity assets within a month with at least half attributable to trades. Three findings bear on this paper.

First, the event is predictable from investor attributes: households that are male, older than 45, married, with more dependents, or self-reporting strong investment knowledge, panic sell more often. Cross-investor dispersion in trigger firing is therefore documented and linked to observables, without any psychometric instrument.

Second, the authors document panic selling as distinct from the disposition effect and from overtrading. The literature already separates the loss-liquidation coordinate from the gain-realisation coordinate, on overlapping populations. That separation is the empirical content of the vector.

Third, the base rate is low: on the order of 0.1 per cent of investors at any moment, and fewer than one household in ten over thirteen years.

The event studied there is also more extreme than the coordinates modelled here. A ninety per cent reduction in a household's equity assets within a month is not the same event as crossing a twelve or twenty per cent loss threshold on a single holding, and the incidence of the first bounds neither the incidence nor the demographic profile of the second. What transfers is that trigger firing is predictable from observable characteristics at all; what does not transfer is any number.

The third finding constrains Section 9.2 from the opposite side. An event this rare cannot generate an aggregate gap of 1.2 percentage points a year, whatever its cost when it occurs. Churn accounts for the aggregate and triggers account for dispersion, and the two literatures are consistent rather than competing.

What this implies for the instrument. A prior over κ can be conditioned on age, sex, marital status, dependents and self-reported knowledge, all of which a distribution platform already holds.

This establishes a baseline, not a solution: predicting a liquidation event is not the same as estimating $\kappa(\tau | X)$ well enough to choose between two products, and Section 11.1 shows the second is not available from these data. The relevant question for a behavioural instrument is therefore not whether it predicts trigger firing in the abstract, but whether it changes product decisions relative to that baseline. In the notation of Section 7.2, the conditioning set X already contains those fields, and the instrument must deliver $V(S | X)$ against that baseline rather than against nothing.

10. What the model implies

The results in this section are consequences of the model, not evidence for it. Two of them are the numerical content of Propositions 2 and 4.

10.1 Order-invariant statistics do not determine implemented value

Two products constructed by permuting a deterministic exposure schedule against exchangeable shocks agree to three decimal places on annualised volatility, downside deviation, skewness, upside capture and average maximum drawdown, and on expected hold return. Their implemented values differ by up to 2.64 percentage points a year.

The source of the difference is timing, not depth. At a twelve per cent threshold the front-loaded schedule triggers with 9.01 years of the ten-year horizon remaining; the back-loaded schedule triggers with 5.17 years remaining. The surrendered risk premium compounds over what is left, and no statistic in the conventional set records when the excursion occurs. This is Proposition 2 with a magnitude.

10.2 Where the ranking crosses

Section 9.2 shows the aggregate cross-section cannot identify the threshold distribution, so nothing here is calibrated to it. What follows is a sensitivity exercise: the exit threshold is swept across a range and each row reports the comparison at a fixed value, with the guarantee effect of 9.3 applied at the crossing.

For this exercise only, the deterministic action is relaxed. In the core model a coordinate crossing produces its action with certainty. Here, in the floored product, crossing generates an exit hazard whose probability declines as the guarantee moves into the money, following the lapse pattern documented in Section 9.3. The payoff is unchanged; what changes is the chance that reaching the threshold ends the holding. This is an extension motivated by the annuity evidence and it forms no part of Theorems 1 and 2, both of which treat τ as a deterministic threshold with a deterministic action. The comparison is against what that investor actually does with the unprotected position, not against a disciplined buy-and-hold benchmark they do not resemble.

break point	unprotected, this investor	note, 100% floor
10%	4.69%	6.80%
20%	6.17%	7.49%
30%	7.97%	8.08%
40%	8.92%	8.28%

Negative implementation drag, meaning cells in which triggering appears to improve the certainty equivalent, occurs in well under one per cent of cells and is consistent with simulation noise. That is what Assumption C implies: exiting early is costly in expectation, so ($D > 0$) wherever the trigger binds, and no depth mechanism makes triggering profitable. Positive gaps observed in defined-outcome structures arise from flows clustering at the start of outcome windows, an entry-timing effect outside this model.

The full-floor structure shows zero probability of terminal loss at every break point and a tenth-percentile terminal wealth of exactly par, while partial floors and constant-proportion strategies lose money on 9 to 20 per cent of paths because they can breach their own floor. The crossover sits near a thirty per cent break point under this parameterisation. The number is illustrative and should not be read as an estimate: it moves with the drift, volatility and horizon assumptions, none of which are fitted to anything. What is not illustrative is the sign change, which is Theorem 2 realised numerically, and its presence does not depend on where exactly the crossing falls.

11. Identification, dynamic inconsistency and welfare

11.1 What separates the model from its rivals

Three identification questions arise and they have different answers. None is a question about sample size in the first instance; each is a question about which variation in the data distinguishes the model from an alternative that fits equally well.

Threshold heterogeneity against smooth propensity. A pooled hazard rising with loss depth, of the kind Section 9.6 reports, is equally consistent with two structures: every investor shares one smooth selling propensity in depth, or each investor carries a hard threshold drawn from κ and the pooled curve is the mixture. In cross-section these are observationally equivalent, which is why the V-shaped propensity function does not identify κ , and why Section 9.2's failure to fit the aggregate cross-section was not evidence either way.

The variation that separates them is **within-person repetition across episodes**. Hard thresholds imply that an investor who exits at a given depth in one drawdown exits at a similar depth in the next, so exit depths cluster within person and disperse across. Smooth propensity implies independent draws, so within-person and across-person dispersion coincide. The test is a variance decomposition of exit depth, and it requires accounts observed through at least two separate drawdowns rather than one.

Assumption B against liquidity. A behavioural trigger fires on the position whose statistic crossed. A private liquidity shock, or a rational re-optimisation on macro information, hits the portfolio. The separating variation is therefore **within-person across simultaneously held positions**: under the trigger reading, an investor liquidates the position that has fallen while retaining one that has not; under the liquidity reading, both are reduced together, roughly in proportion. This variation exists in ordinary brokerage records without any experiment, and it is the observation Section 9.6 notes the published work does not perform, because the estimand there is a per-position propensity rather than a per-person rule.

The coordinate structure. Distinguishing a gain-realisation coordinate from a benchmark-shortfall coordinate needs the same within-person, across-position comparison, applied to two positions with the same unrealised profit and different performance relative to a common reference. Where the two statistics move together the coordinates are not separately identified, which in practice is much of the time.

What no design identifies. The map from a behavioural signal S to $\kappa(\tau | X, S)$ requires joint observation of the signal and the implementation outcome for the same person. Transaction records contain no S , and instrument samples contain no outcome. No sample size on either side substitutes for the join. This is the one requirement in the paper that is a data-collection problem rather than an identification argument.

What each test needs. The three requirements differ and only one is demanding. Separating threshold heterogeneity from smooth propensity requires a panel spanning at least two distinct drawdown episodes with exit depth recorded per episode, and since only a minority of accounts fire even once, the doubly-firing subsample governs the sample size. Testing Assumption B requires only accounts holding two or more positions whose performance has diverged, observed at transaction level in a single episode; this is the least demanding of the three and the closest to data that already

exists. Separating the coordinates requires positions on which unrealised profit and performance relative to a reference move in opposite directions, which is a smaller subsample again, and where the two statistics move together the coordinates are not separately identified at all.

Order of magnitude. With a monthly firing rate near 0.1 per cent, detecting a doubling in that rate between two groups at conventional power requires on the order of 2×10^4 account-months per group, and detecting a 25 per cent difference requires roughly 3×10^5 . Because firing concentrates in drawdowns, accounts must be live through one, so the binding constraint is coverage of episodes rather than headcount. The published panel of 653,455 accounts spanning two crises clears this comfortably; a platform sample assembled in a calm decade does not, at any size.

11.2 Relation to dynamic inconsistency

An implementation trigger is a form of dynamic inconsistency: the investor would choose *ex ante* to remain invested and does not do so *ex post*. The relation to the existing literature is close enough to require a statement and different enough that the models are not substitutes.

The Strotz and Laibson formulations concern inconsistency over **dates**, generated by non-exponential discounting of consumption. A trigger is inconsistency over **states**, generated by a rule that reads a path statistic and acts on it. The two can coexist in one agent and neither implies the other: an exponential discounter can carry a loss threshold, and a hyperbolic discounter with no threshold implements a buy-and-hold allocation faithfully.

Gul and Pesendorfer, and Amador, Werning and Angeletos, sit closer, since both concern an agent who anticipates deviating and values restriction. They place the behaviour in preferences, through a temptation term or a cost of self-control. This paper deliberately does not. Its position throughout is that the trigger constrains the feasible set rather than modifying the objective, which is what Theorem 1 formalises and what keeps the welfare comparison in Theorem 2 free of any behavioural component.

The closest antecedent is Heimer, Iliewa, Imas and Weber, who elicit investors' intended stopping rules and then observe realised behaviour, finding systematic departure. That is a direct measurement of the gap this paper prices. The difference in emphasis is the product side: given that the departure occurs, which payoff geometries make it costly and which make it cheap.

The literature on commitment devices asks what the agent should bind themselves to. Section 7.2 asks a different question, because in the retail setting binding is not available: daily liquidity is a regulatory and commercial requirement, and an investor cannot be prevented from acting after a fall. What can be varied is the payoff, and the paper's design results are about payoffs rather than about commitment.

11.3 Are triggers mistakes?

The framework does not answer this, and the refusal is deliberate rather than evasive.

Calling a trigger a mistake requires a normative benchmark for correct action at the moment it fires. Supplying one imports precisely the judgment that a behavioural welfare function would have imported, and it would restore by the back door the circularity that keeping U conventional was meant to avoid. Under Assumption B the trigger is not a preference and not an error; it is a constraint on what this investor can implement, and Theorem 1 says implemented optimisation is ordinary optimisation over the set that constraint defines.

The refusal has a cost and it should be named. Two policy responses follow from the two readings, and the framework does not adjudicate between them. If triggers are mistakes, education and default architecture dominate, and they are cheaper than any payoff feature. If triggers are constraints, product design is the appropriate response. The evidence in Section 9.4 is a strong prior for the first wherever defaults are available: trading in defaulted retirement arrangements is low in absolute terms and lowest among pure target-date investors, at a fraction of the rate elsewhere.

The practical relevance of the framework therefore depends on which reading is correct, and the paper does not settle it. If triggers are mistakes, and the default evidence in Section 9.4 is a reasonable prior for that view, the correct response is better default architecture rather than payoff engineering. If they are constraints, the trigger vector is an input to the design problem and Section 7 applies.

That yields the scope condition the paper operates under. Where the policy objective is to suppress discretionary trading, default architecture is the cheaper first-line intervention, and the evidence in Section 9.4 says it is effective. The framework does not establish that suppressing trading is itself welfare-improving, and it is not entitled to: that is the judgment it declines to make. The design problem in Section 7 addresses the population for whom defaults are unavailable, which is self-directed retail. A framework that recommended manufacturing protection for defaulted participants would be recommending an expensive solution to a problem their plan design has already solved.

12. Discussion and conclusion

An investment product is not a distribution over terminal wealth. For an investor who will act on what they experience, it is a family of distributions indexed by the triggers that might fire, and different triggers fire on different features of the path.

Making the trigger a vector rather than a scalar changes three things. The product information required for suitability is larger than any order-invariant summary, and larger than any single-coordinate one. The reduction of implemented optimisation to classical optimisation over an absorbing feasible set survives for depth triggers and fails for switching triggers, which are a different mathematical object. And the value of matching investors to products comes from disagreement in implemented rankings, which requires the drag difference to straddle the technical advantage across the population.

The normative implication is conditional and should be read as one. If the investor who would otherwise abandon can be identified in advance, a protective feature can pay for itself; if not, it does not, and where the objective is simply to suppress discretionary trading, default architecture is cheaper.

The chain that carries the paper is short. Technical ranking, then implementation drag, then the reversal condition $D(\mu, \tau) - D(\nu, \tau) > H$, then heterogeneous assignment, then the decision value of a signal that moves investors across that boundary. Each link is a result rather than an assertion, and the last one is where a behavioural instrument would have to earn its place.

What the model implies, once its primitives are fixed, is that two products matching on every conventional risk statistic can differ by 2.64 percentage points a year in implemented value, because they differ by nearly four years in the horizon remaining when the trigger binds; and that the ranking of a floored structure against an unprotected one changes sign as the break point deepens, with the full floor delivering a zero probability of terminal loss on either side of the crossing. Where the crossing falls is parameter-dependent and is not estimated here; that it occurs is Theorem 2 realised numerically.

What the external evidence contributes is different in kind, and it cuts both ways. Three sources are consistent with the mechanism: abandonment that rises with loss severity and also reaches 78 per cent on a winning path where no loss occurred, guarantees that cut surrender by an order of magnitude and cut it further as they go into the money, and trading that remained in single digits through an extraordinary drawdown, lowest of all in defaulted arrangements. The fourth constrains it. The aggregate relationship between implementation cost and product volatility is close to proportional, and a model with no trigger at all, in which investors exit at a constant hazard and stay out for a fixed period, reproduces it to within seven basis points, better by an order of magnitude than any trigger specification tried. A path-dependent trigger and ordinary churn are

indistinguishable in the aggregate mean because both scale with the risk premium. They separate in the cross-section of investors, and there the individual-level trading literature is decisive: the selling-probability function is V-shaped in unrealised profit, rising with the magnitude of gains and losses alike, and liquidation on losses is predictable from investor characteristics and documented as a distinct phenomenon from gain realisation.

The case for the vector therefore rests where it should: on a selling-probability function with two arms, estimated on brokerage records rather than asserted; on abandonment observed in the experiment on a path where no loss occurred; and on the construction in Section 8, in which two investors with identical utility rank the same two products in opposite order.

One result changes how the design problem should be posed. A contractual floor supplies a uniform lower bound on implemented terminal wealth at every trigger value, because mark-to-market value never falls below the discounted floor and riskless growth on liquidation proceeds restores it. That bound is independent of the population distribution of triggers, while an unprotected position has no bound at all. What the bound does not settle is which structure wins: under an expected-value criterion deep floors lose, and under a worst-case criterion the answer depends on where the minimum is attained and has to be computed. Choosing the criterion is a decision about how much is known rather than a result.

The worked example makes the economics concrete: a floored note that is 140 basis points a year worse under the hold assumption is 33 basis points better for an investor whose loss coordinate binds, worse for an investor whose shortfall coordinate binds, and worse for both once a distribution margin of seventy basis points is deducted. The payoff geometry sets the sign. The manufacturing spread sets whether the sign survives.

Appendix A. Calibration procedure and estimates

Three models were fitted to the three volatility-sorted points of Section 9.2, namely cost of 0.40, 1.20 and 2.10 percentage points at annualised volatilities of 5.5, 13.0 and 25.0 per cent.

Common set-up. Ten-year horizon, monthly steps, geometric Brownian dynamics with drift $r + 0.55\sigma$ so that the Sharpe ratio is constant across buckets, riskless rate 2 per cent, 24,000 paths per bucket on common random numbers. The reference process used by the shortfall coordinate has the same drift and volatility as the product and correlation 0.85 with it. Implementation cost is reported here with the sign convention of the targets, that is as a negative number, and equals the negative of the drag D defined in Section 7.1.

Fitting. Grid search on the threshold parameters, minimising squared error against the three targets. No weighting. Reported error is the root mean square over the three points.

model	parameters	fitted values	rms	fitted cost	cost/volatility
depth coordinate	τ^L	0.12	0.509	-0.06 / -1.06 / -3.77	1.1 / 8.1 / 15.1
depth and shortfall	τ^L, τ^F	0.80, 0.12	0.441	-0.06 / -0.72 / -2.58	1.1 / 5.5 / 10.3
constant hazard	λ , months out	0.008/month, 24	0.068	-0.46 / -1.11 / -2.14	8.4 / 8.6 / 8.6

Targets are -0.40 / -1.20 / -2.10 and 7.3 / 9.2 / 8.4.

How the depth model fails. Across the search range the fitted threshold trades the low-volatility bucket against the high-volatility one and cannot satisfy both. Shallow thresholds fit the 5.5 per cent bucket and overshoot the 25 per cent bucket by more than a percentage point; deep thresholds fit the 25 per cent bucket and assign the 5.5 per cent bucket essentially zero. A fund at 5.5 per cent

volatility over ten years has a mean maximum drawdown well inside any threshold that leaves the high-volatility bucket in range, so the barrier is rarely crossed and the model has no mechanism to generate cost.

Why the shortfall coordinate does not repair it. With the action set to rotation between two statistically identical assets, implementation cost is essentially zero at every volatility: switching between assets with the same drift is a wash under independent increments. With the action set to exit into cash, cost is generated but the firing rate still rises faster than linearly in volatility, and the low-volatility bucket remains near zero.

Why the constant-hazard model fits. Cost is the risk premium forgone over the time spent out of the market, which is $\lambda \cdot m \cdot (\mu - r)$ to first order, with m the months out per firing. Under a constant Sharpe ratio $\mu - r$ is proportional to σ , so cost is proportional to σ by construction. The fitted values imply roughly one exit per decade with two years spent in cash. The fit is a property of the premium specification and does not depend on any behavioural assumption.

References

- Amador, M., Werning, I. and Angeletos, G.-M. (2006). Commitment vs. flexibility. *Econometrica*, 74(2), 365-396.
- Barber, B. M. and Odean, T. (2000). Trading is hazardous to your wealth. *Journal of Finance*, 55(2), 773-806.
- Barberis, N. and Xiong, W. (2009). What drives the disposition effect? An analysis of a long-standing preference-based explanation. *Journal of Finance*, 64(2), 751-784.
- Basak, S. and Shapiro, A. (2001). Value-at-risk-based risk management. *Review of Financial Studies*, 14(2), 371-405.
- Ben-David, I. and Hirshleifer, D. (2012). Are investors really reluctant to realize their losses? Trading responses to past returns and the disposition effect. *Review of Financial Studies*, 25(8), 2485-2532.
- Black, F. and Scholes, M. (1973). The pricing of options and corporate liabilities. *Journal of Political Economy*, 81(3), 637-654.
- Calvet, L. E., Campbell, J. Y. and Sodini, P. (2009). Fight or flight? Portfolio rebalancing by individual investors. *Quarterly Journal of Economics*, 124(1), 301-348.
- C  lerier, C. and Vall  e, B. (2017). Catering to investors through security design. *Quarterly Journal of Economics*, 132(3), 1469-1508.
- De Giorgi, E., Hens, T. and Post, T. (2005). Prospect theory and the size and value premium puzzles. Working paper.
- Dichev, I. D. (2007). What are stock investors' actual historical returns? *American Economic Review*, 97(1), 386-401.
- Elkind, D., Kaminski, K., Lo, A. W., Siah, K. W. and Wong, C. H. (2022). When do investors freak out? Machine learning predictions of panic selling. *Journal of Financial Data Science*, 4(1), 11-39.
- Friesen, G. C. and Sapp, T. R. A. (2007). Mutual fund flows and investor returns. *Journal of Banking & Finance*, 31(9), 2796-2816.
- Grinblatt, M. and Han, B. (2005). Prospect theory, mental accounting, and momentum. *Journal of Financial Economics*, 78(2), 311-339.
- Grossman, S. J. and Zhou, Z. (1996). Equilibrium analysis of portfolio insurance. *Journal of Finance*, 51(4), 1379-1403.
- Gul, F. and Pesendorfer, W. (2001). Temptation and self-control. *Econometrica*, 69(6), 1403-1435.
- Heimer, R., Iliewa, Z., Imas, A. and Weber, M. (2025). Dynamic inconsistency in risky choice. *American Economic Review*, 115(1), 330-363.
- Henderson, B. J. and Pearson, N. D. (2011). The dark side of financial innovation. *Journal of Financial Economics*, 100(2), 227-247.
- Kahneman, D. and Tversky, A. (1979). Prospect theory. *Econometrica*, 47(2), 263-291.
- Kaustia, M. (2010). Prospect theory and the disposition effect. *Journal of Financial and Quantitative Analysis*, 45(3), 791-812.
- Laibson, D. (1997). Golden eggs and hyperbolic discounting. *Quarterly Journal of Economics*, 112(2), 443-477.
- Markowitz, H. (1952). Portfolio selection. *Journal of Finance*, 7(1), 77-91.
- Merton, R. C. (1971). Optimum consumption and portfolio rules in a continuous-time model. *Journal of Economic Theory*, 3(4), 373-413.
- Millman (2011). *Variable annuity dynamic lapse study: a data mining approach*.
- Odean, T. (1998). Are investors reluctant to realize their losses? *Journal of Finance*, 53(5), 1775-1798.

- Oduwale, O. (2006). Observed violations of asset pricing models and the implications of prospect theory for optimal portfolio selection. MSc dissertation, Imperial College London.
- Prak, J. (2026). *Mind the Gap 2026*. Morningstar Portfolio and Planning Research.
- Strotz, R. H. (1955). Myopia and inconsistency in dynamic utility maximization. *Review of Economic Studies*, 23(3), 165-180.
- Thaler, R. H. and Benartzi, S. (1995). Myopic loss aversion and the equity premium puzzle. *Quarterly Journal of Economics*, 110(1), 73-92.
- Vanguard (2020, 2025, 2026). *How America Saves*.
- Vokata, P. (2021). Engineering lemons. *Journal of Financial Economics*, 142(2), 737-755.
- Bad timing does not cost investors 15% of their funds' returns: an examination of Morningstar's "Mind the Gap" study.* (2026). *Financial Analysts Journal*.

OLD SHOREHAM INVESTMENT MANAGEMENT

Traditional | Quantitative | Global

www.oldshoreham.com