

Pre-Analysis Plan

Using Equivalent Offsets to Test Reference Dependence

Ao Wang

National University of Singapore

1 The experiment

1.1 Structure

Each participant makes 60 binary decisions, one per round. The 60 rounds form two stocks of 30 rounds each, and each stock forms two situations of 15 consecutive rounds. I index a decision by the participant n , the stock $s \in \{1, 2\}$ and the round $t = 1, \dots, 30$ within that stock, and I write $j(t)$ for the situation containing round t : $j(t) = 1$ for $t \leq 15$ and $j(t) = 2$ for $t \geq 16$. A 60-second break separates the two stocks. Every round is a decision, and the sequence of values continues whether or not the participant sold in an earlier round.

1.2 Values

Values are in points. For each participant and stock the design draws a baseline b_{ns}^j and a benchmark shock η_{ns}^j for each situation $j = 1, 2$, which together form the situation's benchmark m_{ns}^j :

$b_{ns}^j \in \{80, 100, 120\}$ the baseline. The first situation of each stock draws its baseline uniformly from the three values; the second situation draws uniformly from the two values the first did not take, so $b_{ns}^1 \neq b_{ns}^2$.

$\eta_{ns}^j \sim N(0, 6^2)$ the benchmark shock, one draw per situation.

$m_{ns}^j := b_{ns}^j + \text{round}(\eta_{ns}^j)$ the benchmark of situation j of stock s , an integer constant over the situation's 15 rounds.

In each round the current value is the benchmark of the round’s situation plus a round shock,

$$\Delta_{nst} := m_{ns}^{j(t)} + \varepsilon_{nst}, \quad \varepsilon_{nst} \sim \text{Uniform}\{-15, -14, \dots, 15\},$$

with ε_{nst} drawn independently in every round. The baseline pairs are independent across participants and stocks, all shocks are mutually independent and independent of the baselines, and the value sequence does not depend on the participant’s choices. Where the indices are not needed, I write Δ for the current value, m for the benchmark of its situation and ε for the round shock.

1.3 Decision and earnings

In each round the participant sells the stock at the current value or keeps it. Each participant has a die whose six faces carry the market movements $\{-15, -10, -5, 5, 10, 15\}$, one movement per face, in an order drawn once per participant; each face is equally likely. The earnings of the round are

$$\text{sell: } \Delta, \quad \text{keep: } \Delta + x, \quad x \text{ the movement of the face rolled.}$$

The movement affects the earnings of that round only; the next round’s value is the fresh draw of Section 1.2. One of the 60 rounds is selected at random for payment, at ten points per dollar, with a floor of USD 3.50 and a cap of USD 17.00.

1.4 Arms

The two arms use the same decision task but display the current value differently. The **Framed** screen shows the current value as its difference from the situation benchmark, $\Delta - m_{ns}^{j(t)}$, with $m_{ns}^{j(t)}$ itself available behind a question-mark button. The **Plain** screen shows the raw current value Δ , and the benchmark never appears.

1.5 Sample

The plan collects 300 participants per arm, 36,000 decisions in total; the counts below use 300 and scale with the realised number. I analyze only participants who pass the comprehension questions within the allowed attempts and finish the experiment. Every analysis below runs separately in each arm, and every constant computed from the data (standard deviations, bandwidths, knots, centring constants, fold assignments) is computed within the arm.

2 Equivalent Offsets and the estimand

A reference point, or referent, is a quantity against which the participant judges the current value. Take one candidate referent R , for example the value in the previous round. This plan defines reference dependence with respect to R as the property that the probability of selling depends on the current value Δ and on R only through their difference $\Delta - R$. Raising Δ and R by the same amount then leaves the choice unchanged. This property is Equivalent Offsets (EO).

Let $Y = 1$ when the participant sells and $Y = 0$ when the participant keeps the stock, and define the probability of selling given the current value and the referent,

$$P(\Delta, R) := \mathbb{E}[Y \mid \Delta, R].$$

The estimand, the quantity the analysis targets, is the average derivative of P in the direction of an equal increase in Δ and R , evaluated at each decision's own (Δ, R) and averaged over the decisions in the sample, subject to any trimming specified in the relevant procedure,

$$S := \mathbb{E} \left[\frac{\partial P}{\partial \Delta} + \frac{\partial P}{\partial R} \right].$$

Two further quantities enter the analysis: the average marginal effect of the current value alone,

$$\text{ME}_\Delta := \mathbb{E} \left[\frac{\partial P}{\partial \Delta} \right],$$

and, when $\text{ME}_\Delta \neq 0$, the offset ratio

$$Q := 1 - \frac{S}{\text{ME}_\Delta}.$$

EO gives $S = 0$, equivalently $Q = 1$. A choice that ignores the referent has $\partial P / \partial R = 0$, so $S = \text{ME}_\Delta$ and $Q = 0$. The hypothesis tested for every candidate is

$$H_0 : S = 0.$$

The plan tests 23 candidates per arm, 46 in total, in four groups, one per procedure: 12 single referents (Procedure A, Section 3); 6 functional rules, formulas that map the values seen so far in the current stock to one referent (Procedure B, Section 4); 1 mixture allowing different participants to use different referents from a set of six candidates (Procedure C, Section 5); and 4 lagged distributions, in which the K most recent values each enter as a referent, for $K = 2, 3, 4, 5$ (Procedure D, Section 6). For a candidate with

K referents $R_1, \dots, R_K$, write $P(\Delta, R_1, \dots, R_K) := \mathbb{E}[Y \mid \Delta, R_1, \dots, R_K]$; the common shift raises Δ and every R_j by the same amount, and

$$S := \mathbb{E} \left[\frac{\partial P}{\partial \Delta} + \sum_{j=1}^K \frac{\partial P}{\partial R_j} \right], \quad \text{ME}_\Delta := \mathbb{E} \left[\frac{\partial P}{\partial \Delta} \right], \quad Q := 1 - \frac{S}{\text{ME}_\Delta}.$$

Under EO, the sum of partial derivatives is zero at each payoff-referent tuple, so it is also zero on average under any weighting: for a weight $\omega(\Delta, R_1, \dots, R_K)$ fixed by the design and for which the expectations exist,

$$S_\omega := \mathbb{E} \left[\omega \left\{ \frac{\partial P}{\partial \Delta} + \sum_{j=1}^K \frac{\partial P}{\partial R_j} \right\} \right] = 0, \quad Q_\omega := 1 - \frac{S_\omega}{\text{ME}_{\Delta, \omega}}, \quad \text{ME}_{\Delta, \omega} := \mathbb{E} \left[\omega \frac{\partial P}{\partial \Delta} \right],$$

and $Q_\omega = 1$ under EO whenever the denominator is nonzero. Procedures A, B and C use the unweighted averages; Procedure D (Section 6) uses a weight supplied by the value process, which lets the test lean on the comparisons for which the design provides more information.

2.1 Predictions

H1. In the Framed arm, where the screen shows the current value as its difference from the benchmark $m_{ns}^{j(t)}$ of Section 1.2 and makes $m_{ns}^{j(t)}$ available through a question-mark button, H_0 is not rejected for the benchmark: behaviour is reference dependent with respect to the quantity the screen makes salient.

H2. In the Plain arm, where the benchmark exists in the design and the screen never shows it, H_0 is rejected for the benchmark.

The Plain screen gives no earlier value prominence, so which candidates built from the participant's own value history satisfy H_0 in that arm is an open question.

3 Procedure A: single referents

3.1 Candidates (12)

Write $v_1, v_2, \dots, v_{30}$ for the values of the current stock in round order, and consider the decision in round t of the stock, where the current value is $\Delta = v_t$. Write R for the candidate referent under test. Every history quantity uses the values of the current stock in rounds before t ; the history restarts at the first value of each stock.

Benchmark $R := m_{ns}^{j(t)}$, the benchmark of the situation that contains round t (Sec-

tion 1.2): a constant set by the design, the same for the 15 rounds of the situation, and the one candidate that is not a function of the value history. The Framed screen displays $\Delta - m_{ns}^{j(t)}$; the Plain screen never shows it. This candidate carries the predictions H1 and H2 of Section 2.1.

Lag 1 $R := v_{t-1}$, the value in the previous round.

Lag 2 $R := v_{t-2}$, the value two rounds earlier.

Lag 3 $R := v_{t-3}$, the value three rounds earlier.

Lag 4 $R := v_{t-4}$, the value four rounds earlier.

Lag 5 $R := v_{t-5}$, the value five rounds earlier.

Initial 1 $R := v_1$, the first value of the stock.

Initial 2 $R := v_2$, the second value of the stock.

Initial 3 $R := v_3$, the third value of the stock.

Prior minimum $R := \min\{v_1, \dots, v_{t-1}\}$, the lowest value so far in the stock.

Prior maximum $R := \max\{v_1, \dots, v_{t-1}\}$, the highest value so far in the stock.

Prior mean $R := \frac{1}{t-1} \sum_{r=1}^{t-1} v_r$, the mean of the values so far in the stock.

3.2 Sample

The 60 rounds of each of the 300 participants per arm give 18,000 decisions per arm. Each candidate uses the decisions on which it is defined: the benchmark uses all 18,000; lag k , $k = 1, \dots, 5$, uses rounds $t \geq k+1$ of each stock, giving 600 $(30-k)$ decisions across the 600 stocks in an arm; the three initial values and the prior minimum, maximum and mean use rounds $t \geq 4$ of each stock, 16,200 decisions, a common sample for these six candidates.

3.3 Estimator

The estimator is a local linear regression. Pick one decision. Take the decisions whose (Δ, R) lie close to it, and fit to them, by weighted least squares, a plane

$$Y \approx a + b_{\Delta} (\Delta - \Delta_0) + b_R (R - R_0),$$

where (Δ_0, R_0) is the picked decision's own pair. The two slopes b_Δ and b_R are the estimates of $\partial P/\partial \Delta$ and $\partial P/\partial R$ at that decision, in probability per point. Repeating this for the decisions retained by the trimming rule below and averaging the slopes gives the estimates of ME_Δ and of

$$\text{ME}_R := \mathbb{E} \left[\frac{\partial P}{\partial R} \right].$$

Let N be the number of decisions in the candidate's sample and index them by $i = 1, \dots, N$. Δ and R enter in raw points throughout; their standard deviations serve to define which decisions count as near and to scale the local solve, as stated in steps 1 and 4.

1. **Distance.** With $\text{sd}(\Delta)$ and $\text{sd}(R)$ computed on the observed sample, the distance between any two decisions $i \neq i'$ in the candidate's sample is

$$d(i, i') := \sqrt{\left(\frac{\Delta_i - \Delta_{i'}}{\text{sd}(\Delta)} \right)^2 + \left(\frac{R_i - R_{i'}}{\text{sd}(R)} \right)^2},$$

so that one unit of distance means one standard deviation in either coordinate.

2. **Bandwidth.** Let

$$q := \max(10, \lceil k_{\text{frac}} N \rceil), \quad k_{\text{frac}} = 0.045,$$

and let $d_q(i)$ be the distance from decision i to its q th nearest decision, counting i itself as the first. The single global bandwidth is the median of $d_q(i)$ over the N decisions,

$$h := \text{median}\{d_q(1), \dots, d_q(N)\},$$

floored at 10^{-8} in the degenerate case where that median is smaller.

3. **Trim.** Let τ be the 97.5th percentile of $\{d_q(i)\}$. Decision i is an evaluation point when $d_q(i) \leq \tau$; write $\mathcal{E}$ for the set of evaluation points and $|\mathcal{E}|$ for its size. This drops roughly 2.5% of decisions whose q th neighbour is farthest away.

4. **Fit.** At each evaluation point $i_0 \in \mathcal{E}$, give every other decision i the weight

$$w_i := \begin{cases} \exp(-\frac{1}{2} d(i, i_0)^2 / h^2) & \text{if } d(i, i_0) \leq 4h, \\ 0 & \text{if } d(i, i_0) > 4h, \end{cases}$$

and solve

$$\min_{a, b_\Delta, b_R} \sum_{i=1, i \neq i_0}^N w_i [Y_i - a - b_\Delta (\Delta_i - \Delta_{i_0}) - b_R (R_i - R_{i_0})]^2.$$

The sum omits i_0 : the evaluation point's own observation is left out of its fit, as in Härdle and Stoker (1989). Write the minimisers as $b_\Delta(i_0)$ and $b_R(i_0)$: the estimated partial derivatives of P with respect to Δ and R at decision i_0 , in probability per point. The solve forms the 3×3 weighted normal matrix with the two regressors divided by $\text{sd}(\Delta)$ and $\text{sd}(R)$, adds 10^{-8} to its diagonal, and expresses the resulting slopes per point.

5. **Average.**

$$\widehat{\text{ME}}_\Delta := \frac{1}{|\mathcal{E}|} \sum_{i_0 \in \mathcal{E}} b_\Delta(i_0), \quad \widehat{\text{ME}}_R := \frac{1}{|\mathcal{E}|} \sum_{i_0 \in \mathcal{E}} b_R(i_0) :$$

the estimated derivative at each evaluation decision, averaged with equal weight over the evaluation points of step 3. This is the sample version of the expectation in Section 2, taken over the trimmed set $\mathcal{E}$.

6. **Statistics.**

$$\hat{S} := \widehat{\text{ME}}_\Delta + \widehat{\text{ME}}_R, \quad \hat{Q} := 1 - \frac{\hat{S}}{\widehat{\text{ME}}_\Delta} = -\frac{\widehat{\text{ME}}_R}{\widehat{\text{ME}}_\Delta},$$

the estimates of S and Q of Section 2.

The standard deviations, the bandwidth h and the set $\mathcal{E}$ are computed once on the observed sample; every bootstrap draw (Section 7) reuses them.

4 Procedure B: functional rules

4.1 Candidates (6)

A functional rule builds one referent from the value history through a formula with a parameter vector θ . As in Section 3.1, $v_1, \dots, v_{30}$ are the values of the current stock in round order, the decision is in round t , the current value is $\Delta = v_t$, and the history available at that decision is $v_1, \dots, v_{t-1}$, indexed by r . The sample of Section 4.2 gives $t \geq 4$. Each rule defines the referent R_t as follows.

Adaptive expectations, $\theta = \rho \in (0, 1)$. Nerlove (1958), the adaptive reference point of Thakral and Tô (2021): an exponentially weighted average of the history,

$$A_1 := v_1, \quad A_r := \rho A_{r-1} + (1 - \rho) v_r \quad (r = 2, \dots, t-1), \quad R_t := A_{t-1}.$$

Moving average, $\theta = M \in \{1, \dots, 29\}$. The reference point of DellaVigna, Lindner, Reizer and Schmieder (2017): the mean of the last M values in the history, or of

all of them when fewer than M exist,

$$R_t := \frac{1}{t - r_0} \sum_{r=r_0}^{t-1} v_r, \quad r_0 := \max(1, t - M).$$

Peak-end, $\theta = \iota \in (0, 1)$. Fredrickson and Kahneman (1993): a weighted average of the highest value in the history and the most recent one,

$$R_t := \iota \max_{1 \leq r \leq t-1} v_r + (1 - \iota) v_{t-1}.$$

Baucells–Weber–Welfens weighting, $\theta = (\varsigma, \nu)$, $\varsigma, \nu \in [e^{-5}, e^5]$. Baucells, Weber and Welfens (2011), who apply a Prelec-form weighting function to chronological position r :

$$R_t := \sum_{r=1}^{t-1} \pi_r v_r, \quad \pi_r := \varpi\left(\frac{r}{t-1}\right) - \varpi\left(\frac{r-1}{t-1}\right),$$

$$\varpi(z) := \exp[-(-\ln z)^\varsigma / \nu] \text{ for } 0 < z < 1, \quad \varpi(0) := 0, \quad \varpi(1) := 1.$$

The π_r sum to one.

Extrapolative expectations, $\theta = \varrho \in \mathbb{R}$. Metzler (1941): the most recent value plus ϱ times the most recent change,

$$R_t := v_{t-1} + \varrho(v_{t-1} - v_{t-2});$$

$\varrho > 0$ extrapolates the last change and $\varrho < 0$ reverses it.

Experience weights, $\theta = \zeta \in \mathbb{R}$. Malmendier and Nagel (2011): a weighted average of the history whose weights rise with recency when $\zeta < 0$ and fall with recency when $\zeta > 0$,

$$R_t := \sum_{r=1}^{t-1} \psi_r v_r, \quad \psi_r := \frac{(t-r)^\zeta}{\sum_{r'=1}^{t-1} (t-r')^\zeta}.$$

4.2 Sample

Rounds $t \geq 4$ of each stock: 16,200 decisions per arm, the same decisions for all six rules.

4.3 Estimator

Procedure B fits θ once, builds the column of referent values $R_t(\hat{\theta})$, and applies the estimator of Procedure A to that column.

1. **Fit θ by an auxiliary logit**, a first-stage fit used only to choose θ . With $\Lambda(z) := 1/(1 + \exp(-z))$, the model is

$$\Pr(Y = 1 \mid \Delta, R_t) = \Lambda(a_0 + a_1 (\Delta - R_t(\theta))).$$

I fit a_0 , a_1 , and θ jointly on the decisions of Section 4.2, pooled across participants in the arm. I use BFGS from one starting value, restrict θ to the domain stated in Section 4.1, and take the estimate the optimiser returns. For the moving average, M is discrete: the fit is run at every $M = 1, \dots, 29$ and the M with the highest returned likelihood is taken, the smallest M on a tie. The fit is done separately in each arm.

2. **Build the column.** Evaluate R_t at the returned estimate $\hat{\theta}$ for every decision in the sample of Section 4.2.
3. **Test.** Apply the estimator of Section 3.3, unchanged, to the pair $(\Delta, R_t(\hat{\theta}))$, giving $\hat{S}$ and $\hat{Q}$ as there. The estimand is S of Section 2 for the referent column $R_t(\hat{\theta})$ held fixed.

Every bootstrap draw (Section 7) reuses $\hat{\theta}$ and the column built from it.

5 Procedure C: the six-referent mixture

5.1 Candidate (1)

With $v_1, \dots, v_{30}$ the values of the current stock in round order and t the round of the decision, six referents of Section 3.1 enter jointly:

R_1 Lag 1, v_{t-1} .

R_2 Initial 1, v_1 .

R_3 Prior minimum, $\min\{v_1, \dots, v_{t-1}\}$.

R_4 Prior maximum, $\max\{v_1, \dots, v_{t-1}\}$.

R_5 Prior mean, $\frac{1}{t-1} \sum_{r=1}^{t-1} v_r$.

R_6 Benchmark, $m_{ns}^{j(t)}$ for the situation containing round t (Section 1.2).

The estimand follows Section 2 with $K = 6$ and $P = P(\Delta, R_1, \dots, R_6)$.

5.2 Sample

Rounds $t \geq 4$ of each stock: 16,200 decisions per arm.

5.3 Estimator

The estimator is a penalised spline sieve: a least-squares fit of Y on a fixed set of spline basis functions. Its inputs are the current value and its gap from each referent. Define the gap to referent j ,

$$g_j := \Delta - R_j, \quad j = 1, \dots, 6,$$

and write $f(\Delta, g_1, \dots, g_6)$ for the fitted function, so that the fitted probability is

$$\hat{P}(\Delta, R_1, \dots, R_6) := f(\Delta, \Delta - R_1, \dots, \Delta - R_6).$$

The fitted function is

$$f = \alpha_0 + \sigma_\Delta(\Delta) + \sum_{j=1}^6 [\sigma_j(g_j) + u_j(\Delta, g_j) + \kappa_j(\Delta, g_j)] :$$

an intercept α_0 , a curve σ_Δ in the current value, and for each referent j a gap curve σ_j , an interaction surface u_j between the current value and the gap, and a kink term κ_j at zero gap, defined in step 3. Each referent's gap interacts with Δ only.

1. **Curves.** A cubic spline with a given set of knots is a function of one variable that is a cubic polynomial on each interval between consecutive knots and has a continuous value, first derivative and second derivative at every knot. With 9 interior knots and the boundary knots below, the cubic splines form a 13-dimensional space; the cubic B-splines are the basis of that space in which each basis function is positive on at most four adjacent knot intervals and zero elsewhere, built by the recursion of de Boor (2001, ch. IX). Each curve is a linear combination of them:

$$\sigma_\Delta(\Delta) := \sum_{a=1}^{13} \alpha_a^\Delta \beta_a^\Delta(\Delta), \quad \sigma_j(g_j) := \sum_{a=1}^{13} \alpha_a^j \beta_a^j(g_j), \quad j = 1, \dots, 6,$$

with β_a^Δ built on the knots for Δ and β_a^j on the knots for g_j . The interior knots sit at the $1/10, \dots, 9/10$ quantiles of the variable on the observed sample, Δ for σ_Δ and g_j for σ_j ; the boundary knots sit at its minimum and maximum, each repeated four times and each padded outward by $10^{-6} \max(\text{range}, 1)$, with range the maximum minus the minimum.

2. **Surfaces.** Let $\tilde{\beta}_1^\Delta, \dots, \tilde{\beta}_8^\Delta$ be the cubic B-spline basis in Δ with 4 interior knots (the reduced Δ basis, common to all j) and $\tilde{\beta}_1^j, \dots, \tilde{\beta}_8^j$ the cubic B-spline basis in g_j with

4 interior knots, built as in step 1. The 4 interior knots of each margin sit at the $1/5, 2/5, 3/5, 4/5$ quantiles of that variable on the observed sample, with boundary knots as in step 1. Each surface is the tensor product of the two bases: a linear combination of the 64 pairwise products,

$$u_j(\Delta, g_j) := \sum_{a=1}^8 \sum_{b=1}^8 \alpha_{ab}^{u,j} \tilde{\beta}_a^\Delta(\Delta) \tilde{\beta}_b^j(g_j), \quad j = 1, \dots, 6.$$

3. **Kinks.** $\chi_j(g_j) := \max(g_j, 0)$, which has a corner at $g_j = 0$; its derivative is taken as 0 at $g_j = 0$. Each enters once on its own and once multiplied by each of the 8 reduced Δ basis functions, 9 regressors per referent:

$$\kappa_j(\Delta, g_j) := \alpha_0^{\chi,j} \chi_j(g_j) + \sum_{a=1}^8 \alpha_a^{\chi,j} \tilde{\beta}_a^\Delta(\Delta) \chi_j(g_j), \quad j = 1, \dots, 6.$$

4. **Coefficient count.** $1 + 7 \times 13 + 6 \times 64 + 6 \times 9 = 530$: seven curves (σ_Δ and six σ_j), six surfaces and six kinks.
5. **Penalty.** Each surface u_j of step 2 has the 8×8 coefficient matrix $\alpha^{u,j} = (\alpha_{ab}^{u,j})$, with a indexing the Δ margin and b the gap margin. Its penalty is λ_{int} times the sum of squared second differences of that matrix along each margin,

$$\begin{aligned} \text{Pen}_u(\alpha^{u,j}) := \lambda_{\text{int}} & \left[\sum_{b=1}^8 \sum_{a=1}^6 (\alpha_{a+2,b}^{u,j} - 2\alpha_{a+1,b}^{u,j} + \alpha_{a,b}^{u,j})^2 \right. \\ & \left. + \sum_{a=1}^8 \sum_{b=1}^6 (\alpha_{a,b+2}^{u,j} - 2\alpha_{a,b+1}^{u,j} + \alpha_{a,b}^{u,j})^2 \right], \quad \lambda_{\text{int}} = 1000. \end{aligned}$$

The curves and the kinks carry only the ridge penalty below. Write $\alpha = (\alpha_0, \dots, \alpha_p)$ for the 530 coefficients counted in step 4, with α_0 the intercept and $p = 529$. The total penalty is

$$\text{Pen}(\alpha) := \sum_{j=1}^6 \text{Pen}_u(\alpha^{u,j}) + 10^{-3} \sum_{\ell=1}^p \alpha_\ell^2.$$

Write Ψ for the symmetric matrix satisfying $\text{Pen}(\alpha) = \alpha^\top \Psi \alpha$.

6. **Fit.** Write $\varphi_1, \dots, \varphi_p$ for the basis functions of steps 1 to 3 stacked in one list, $p = 13 + 6 \times 13 + 6 \times 64 + 6 \times 9 = 529$: the 13 of σ_Δ , then the six 13-column gap curves, the six 64-column interaction surfaces, and the six 9-column kink blocks. The coefficient vector α stacks the block coefficients in the same order, α_a^Δ , α_a^j , $\alpha_{ab}^{u,j}$ and

$\alpha_a^{x,j}$ of steps 1 to 3; α_ℓ is the coefficient on φ_ℓ whichever block it belongs to. The fitted function is then

$$f(\Delta, g_1, \dots, g_6; \alpha) := \alpha_0 + \sum_{\ell=1}^p \alpha_\ell \varphi_\ell(\Delta, g_1, \dots, g_6),$$

$$\frac{\partial f}{\partial \Delta}(\cdot; \alpha) = \sum_{\ell=1}^p \alpha_\ell \frac{\partial \varphi_\ell}{\partial \Delta}, \quad \frac{\partial f}{\partial g_j}(\cdot; \alpha) = \sum_{\ell=1}^p \alpha_\ell \frac{\partial \varphi_\ell}{\partial g_j}.$$

Let X be the matrix whose rows are the decisions and whose columns are the constant 1 and $\varphi_1, \dots, \varphi_p$ evaluated at each decision, with the curve and surface columns centred on the observed sample, and $\mathbf{y}$ the vector of Y over the same decisions. On any set of decisions $\mathcal{S}$, the coefficients minimise the penalised sum of squares,

$$\hat{\alpha} := \arg \min_{\alpha} \sum_{i \in \mathcal{S}} (Y_i - f(\Delta_i, g_{i1}, \dots; \alpha))^2 + \text{Pen}(\alpha),$$

which in matrix form is $\hat{\alpha} = (X^\top X + \Psi)^{-1} X^\top \mathbf{y}$. After centring, the columns of each curve block sum to zero, so each block is rank deficient by one; the ridge floor in Ψ makes the solve well posed. Step 7 states the sets of decisions on which this is done.

7. **Cross-fitting.** A random permutation assigns participants to 5 folds $\mathcal{F}_1, \dots, \mathcal{F}_5$, with sizes differing by at most one. For each fold $\xi = 1, \dots, 5$, let $X_{-\xi}$ and $\mathbf{y}_{-\xi}$ collect the decisions of participants outside $\mathcal{F}_\xi$, and estimate

$$\hat{\alpha}_{-\xi} := (X_{-\xi}^\top X_{-\xi} + \Psi)^{-1} X_{-\xi}^\top \mathbf{y}_{-\xi}.$$

Define the fitted function that excludes fold ξ ,

$$f_{-\xi}(\cdot) := f(\cdot; \hat{\alpha}_{-\xi}), \quad \frac{\partial f_{-\xi}}{\partial \Delta} = \sum_{\ell=1}^p \hat{\alpha}_{-\xi, \ell} \frac{\partial \varphi_\ell}{\partial \Delta}, \quad \frac{\partial f_{-\xi}}{\partial g_j} = \sum_{\ell=1}^p \hat{\alpha}_{-\xi, \ell} \frac{\partial \varphi_\ell}{\partial g_j}.$$

Every decision's derivative in step 8 comes from the fit that excludes its own participant's fold. For the point estimate, I randomly assign participants to five groups once. I make a new assignment for each bootstrap sample (Section 7).

8. **Average derivatives.** Index the decisions in the sample by $i = 1, \dots, N$, with values Δ_i and $g_{i1}, \dots, g_{i6}$, and define

$$\xi(i) := \text{the fold that contains decision } i\text{'s participant,}$$

so that $f_{-\xi(i)}$ is the fitted function of step 7 estimated without that fold. Take the partial derivative of $f_{-\xi(i)}$ with respect to Δ , with the gaps held fixed, at decision i 's

own values, and average over decisions:

$$\hat{D}_\Delta := \frac{1}{N} \sum_{i=1}^N \frac{\partial f_{-\xi(i)}}{\partial \Delta}(\Delta_i, g_{i1}, \dots, g_{i6}).$$

Do the same for the partial derivative with respect to each gap g_j :

$$\hat{D}_{g,j} := \frac{1}{N} \sum_{i=1}^N \frac{\partial f_{-\xi(i)}}{\partial g_j}(\Delta_i, g_{i1}, \dots, g_{i6}), \quad j = 1, \dots, 6.$$

9. **Statistics.** By the chain rule applied to $\hat{P} = f(\Delta, \Delta - R_1, \dots, \Delta - R_6)$, a unit increase in Δ moves every gap by one unit as well:

$$\frac{\partial \hat{P}}{\partial \Delta} = \frac{\partial f}{\partial \Delta} + \sum_{j=1}^6 \frac{\partial f}{\partial g_j}, \quad \frac{\partial \hat{P}}{\partial R_j} = -\frac{\partial f}{\partial g_j}.$$

Averaging over decisions,

$$\widehat{\text{ME}}_\Delta := \hat{D}_\Delta + \sum_{j=1}^6 \hat{D}_{g,j}, \quad \hat{S} := \widehat{\text{ME}}_\Delta + \sum_{j=1}^6 (-\hat{D}_{g,j}) = \hat{D}_\Delta, \quad \hat{Q} := 1 - \frac{\hat{S}}{\widehat{\text{ME}}_\Delta}.$$

The bootstrap of Section 7 repeats steps 7 to 9 on each draw, with the knots of steps 1 and 2 and the centring constants of step 6 fixed on the observed sample.

6 Procedure D: lagged distribution

6.1 Candidates (4)

With $v_1, \dots, v_{30}$ the values of the current stock in round order and t the round of the decision, each candidate takes the last K values,

$$L_d := v_{t-d}, \quad d = 1, \dots, K,$$

jointly as the support of a single stochastic reference distribution in the sense of Kőszegi and Rabin (2007):

$K = 2$ refers to $L_1 = v_{t-1}, L_2 = v_{t-2}$.

$K = 3$ refers to $L_1 = v_{t-1}, L_2 = v_{t-2}, L_3 = v_{t-3}$.

$K = 4$ refers to $L_1 = v_{t-1}, L_2 = v_{t-2}, L_3 = v_{t-3}, L_4 = v_{t-4}$.

$K = 5$ referents $L_1 = v_{t-1}$, $L_2 = v_{t-2}$, $L_3 = v_{t-3}$, $L_4 = v_{t-4}$, $L_5 = v_{t-5}$.

The candidate has K referents $R_d = L_d$, and the K values form one reference distribution: the participant's choice may depend on all K gaps $\Delta - L_d$ jointly, through a selling probability $P(\Delta, L_1, \dots, L_K)$ that is a general, possibly non-additive function of the gaps. The relative probabilities the participant attaches to the K values are left unrestricted and are not identified by the test; the only condition is that they do not change under a common shift of the value and the K referents. They are distinct from the estimation weight ω below, which the design supplies.

Values are integers, so Procedure D works with one-unit changes rather than derivatives. Write $L + \mathbf{1}$ for the referents all raised by one point, with $\mathbf{1}$ a vector of K ones, and define the payoff-only change and the common-shift change,

$$A(\Delta, L) := P(\Delta + 1, L) - P(\Delta, L), \quad D(\Delta, L) := P(\Delta + 1, L + \mathbf{1}) - P(\Delta, L),$$

and the referent change after the payoff increase, $\mathcal{R}(\Delta, L) := P(\Delta + 1, L + \mathbf{1}) - P(\Delta + 1, L)$, so that $D = A + \mathcal{R}$ exactly. EO says $D = 0$ at every (Δ, L) . With a weight ω fixed by the design and defined in Section 6.3, the estimand is the finite-difference counterpart of the weighted estimand of Section 2,

$$S_\omega := \mathbb{E}[\omega D], \quad \text{ME}_{\Delta, \omega} := \mathbb{E}[\omega A], \quad Q_\omega := 1 - \frac{S_\omega}{\text{ME}_{\Delta, \omega}} = -\frac{\mathbb{E}[\omega \mathcal{R}]}{\mathbb{E}[\omega A]},$$

and EO gives $S_\omega = 0$ for every weight and $Q_\omega = 1$ whenever $\text{ME}_{\Delta, \omega} \neq 0$. I report $\hat{Q}_\omega$ with its interval and the sign test of $\hat{S}_\omega$ (Sections 8 and 9).

6.2 Sample

Rounds $t \geq K + 1$ of each stock: 600 $(30 - K)$ decisions per arm.

6.3 Estimator

Procedure D tests $S_\omega = 0$ and reports Q_ω without fitting the selling probability. The randomisation of Section 1.2 gives the distribution of the current value given the gaps in closed form, and that distribution turns two averages of unit steps in P into averages of observed choices multiplied by known numbers.

1. The objects.

- Δ : the current value; $L = (L_1, \dots, L_K)$: the K lagged values, the referents; $g := (\Delta - L_1, \dots, \Delta - L_K)$: the gap vector.

- Y : the sell indicator; $P(\Delta, L) := \mathbb{E}[Y \mid \Delta, L]$: the probability of selling at value Δ and referents L . Every conditional quantity in this section, P and the two value laws below, also conditions on the round position t within the stock, which fixes how many of the K lagged values belong to the current situation; the notation suppresses t .
- $A(\Delta, L) = P(\Delta + 1, L) - P(\Delta, L)$, the payoff-only change of Section 6.1: the value rises by one point and every referent stays.
- $D(\Delta, L) = P(\Delta + 1, L + 1) - P(\Delta, L)$, the common-shift change of Section 6.1: the value and every referent rise by one point, so the gaps stay. EO says $D = 0$ at every (Δ, L) .
- $p(k \mid g)$: the probability, under the value process of Section 1.2, that the current value equals k given the gap vector g and the round position, and $\mu_g := \sum_k k p(k \mid g)$. Within one situation the gaps are differences of round shocks, so $\Delta - m$, with m the situation's benchmark of Section 1.2, is uniform on the integers $[-15 + \max(0, g_1, \dots, g_K), 15 + \min(0, g_1, \dots, g_K)]$ and m is independent of g ; in rounds 16 to $15 + K$ of a stock at least one lag belongs to the preceding situation, and $p(k \mid g)$ is the finite sum over the two benchmarks' joint distribution. Both are computed from Section 1.2.
- $p(k \mid L)$: the probability that the current value equals k given the K lagged values and the round position; the benchmark is confined to the integers $[\max L - 15, \min L + 15]$ under its law of Section 1.2 and Δ adds an independent round shock; in rounds 16 to $15 + K$, the two-benchmark analogue. Computed from the same design law.

The target has the form $Q_\omega = 1 - (\text{weighted average of } D) / (\text{weighted average of } A)$ with the same weight in both averages.

2. Numerator: the weight the design supplies. Fix a gap vector g and define the running sum over the value support

$$c_g(k) := \sum_{j \leq k} (\mu_g - j) p(j \mid g) \geq 0,$$

which converges to zero in both tails. Applying summation by parts to finite sums and taking limits turns the conditional covariance of value and selling into a weighted sum

of common-shift changes: with $f_g(k) := P(k, k\mathbf{1} - g)$,

$$\begin{aligned}\mathbb{E}[(\Delta - \mu_g) Y \mid g] &= \sum_k (k - \mu_g) p(k \mid g) f_g(k) \\ &= \sum_k c_g(k) \{f_g(k+1) - f_g(k)\} = \sum_k c_g(k) D(k, k\mathbf{1} - g).\end{aligned}$$

Define the weight $\omega_S(k, g) := c_g(k)/p(k \mid g)$ on supported steps. Then $\mathbb{E}[(\Delta - \mu_g) Y] = \mathbb{E}[\omega_S D]$: the average of $(\Delta - \mu_g)Y$ over decisions is the common-shift effect under the weight ω_S , and it is zero under EO whatever the weights. D is never estimated.

3. Denominator: the same weight on the payoff-only change. Fix the lag vector L . At value k the gaps are $k\mathbf{1} - L$, so the weight carried over is $\tilde{\omega}(k, L) := \omega_S(k, k\mathbf{1} - L)$. Define the coefficient $w_L(k) := p(k \mid L) \tilde{\omega}(k, L)$ and the multiplier

$$h_L(k) := \frac{w_L(k-1) - w_L(k)}{p(k \mid L)},$$

with w_L zero outside the support. Summation by parts in the other direction gives $\mathbb{E}[h_L(\Delta) Y \mid L] = \sum_k w_L(k) A(k, L)$, hence $\mathbb{E}[h_L(\Delta) Y] = \mathbb{E}[\tilde{\omega} A]$, the payoff-only effect under the same weight, from observed choices multiplied by a known number. A is never estimated.

4. The trimming rule. The multiplier h_L divides by $p(k \mid L)$ and can be large where that probability is small. For each L let $p_{\max}(L) := \max_k p(k \mid L)$ and

$$\underline{k}(L) := \min\{k : p(k \mid L) \geq 0.10 p_{\max}(L)\}, \quad \bar{k}(L) := \max\{k : p(k \mid L) \geq 0.10 p_{\max}(L)\},$$

and retain the steps from $\underline{k}(L)$ to $\bar{k}(L) - 1$: $I(k, L) := \mathbf{1}\{\underline{k}(L) \leq k < \bar{k}(L)\}$. The trimmed weight is

$$\omega_{\text{tr}}(k, L) := \tilde{\omega}(k, L) I(k, L),$$

and it is applied to both sides. Denominator: replace w_L by $w_{L,\text{tr}}(k) := p(k \mid L) \omega_{\text{tr}}(k, L)$ and h_L by

$$h_{L,\text{tr}}(k) := \frac{w_{L,\text{tr}}(k-1) - w_{L,\text{tr}}(k)}{p(k \mid L)}, \quad \text{so that} \quad \mathbb{E}[h_{L,\text{tr}}(\Delta) Y] = \mathbb{E}[\omega_{\text{tr}} A].$$

Numerator: fix g ; at value k the lag vector is $k\mathbf{1} - g$, so define $w_{g,\text{tr}}(k) := c_g(k) I(k, k\mathbf{1} - g)$ and

$$h_{g,\text{tr}}(k) := \frac{w_{g,\text{tr}}(k-1) - w_{g,\text{tr}}(k)}{p(k \mid g)}, \quad \text{so that} \quad \mathbb{E}[h_{g,\text{tr}}(\Delta) Y] = \mathbb{E}[\omega_{\text{tr}} D].$$

The constant 0.10 is fixed before the confirmatory sample is collected. The estimand of Section 6.1 is the one with $\omega = \omega_{\text{tr}}$:

$$S_\omega = \mathbb{E}[\omega_{\text{tr}} D], \quad \text{ME}_{\Delta, \omega} = \mathbb{E}[\omega_{\text{tr}} A], \quad Q_\omega = 1 - \frac{\mathbb{E}[\omega_{\text{tr}} D]}{\mathbb{E}[\omega_{\text{tr}} A]} = -\frac{\mathbb{E}[\omega_{\text{tr}} \mathcal{R}]}{\mathbb{E}[\omega_{\text{tr}} A]},$$

with $\mathcal{R}$ the referent change of Section 6.1, $D = A + \mathcal{R}$.

5. Centring. Both multipliers average to zero given their own conditioning set: $\sum_k \{w_{g, \text{tr}}(k-1) - w_{g, \text{tr}}(k)\} = 0$, and likewise for $w_{L, \text{tr}}$. So replacing Y by $Y - \hat{H}(g)$ in the numerator and by $Y - \hat{H}_L(L)$ in the denominator, for any fitted functions of the gaps and of the lags, leaves both expectations unchanged; the centring serves to improve precision. $\hat{H}$ is a penalised least-squares fit of Y on an intercept and one cubic B-spline curve in the gap. The same curve is evaluated at each of the K gaps and summed, so that $\hat{H}(g) = \hat{\alpha}_0 + \sum_{d=1}^K \hat{\sigma}(g_d)$. Its 9 interior knots are the $1/10, \dots, 9/10$ quantiles of a subsample of N values, where N is the number of decisions. This subsample is drawn once at random without replacement from the $N \times K$ pooled gaps. The boundary knots are placed at the subsample's minimum and maximum, each repeated four times and padded outward by $10^{-6} \max(\text{range}, 1)$, where range is the subsample's maximum minus its minimum. The fit uses a second-difference penalty on the curve's coefficients with weight 10^4 and a ridge of 10^{-3} . It includes no kink or current-value term, and the fitted probability is clipped to $[0, 1]$. $\hat{H}_L$ is the same construction with the K lagged values $L_1, \dots, L_K$ in place of the gaps, its knot subsample drawn the same way from the pooled lagged values. Both are cross-fitted on 5 participant folds, assigned once by a random permutation of participants, so that each decision is centred by a fit that excludes its own participant's fold. The curve is additive across its K arguments and need not fit the selling probability well; the precision gain depends on the fit. Neither fit supplies D or A ; the identities of steps 2 to 4 recover their weighted averages, and the fits are used to improve precision.

6. Computation.

1. For every decision i compute from Section 1.2 the two multipliers $h_{g, \text{tr}}(\Delta_i)$ at its gap vector and $h_{L, \text{tr}}(\Delta_i)$ at its lag vector.
2. Form $Z_{S, i} := h_{g, \text{tr}}(\Delta_i) \{Y_i - \hat{H}(g_i)\}$ and $Z_{A, i} := h_{L, \text{tr}}(\Delta_i) \{Y_i - \hat{H}_L(L_i)\}$.
3. Average within each participant, then across participants: $\hat{S}_\omega := \overline{Z}_S$ and $\widehat{\text{ME}}_{\Delta, \omega} := \overline{Z}_A$. Dividing both averages by one positive constant leaves $\hat{Q}_\omega$ and the sign test of Section 8 unchanged.

$$4. \hat{Q}_\omega := 1 - \hat{S}_\omega / \widehat{\text{ME}}_{\Delta, \omega}.$$

The centring functions $\hat{H}$ and $\hat{H}_L$, with their fold assignment, are computed once on the observed sample, as are the multipliers, which depend on the design and the observed values only. Each bootstrap draw of Section 7 resamples participants and repeats steps 3 and 4 on the observed per-decision contributions $Z_{S,i}$ and $Z_{A,i}$, holding the centring functions and the multipliers fixed at their observed-sample values; the draw refits and recomputes nothing else. Draws are retained by the rule of Section 7, which requires $\hat{S}_\omega^{(b)}$ and $\hat{Q}_\omega^{(b)}$ both finite. The centring cannot move either expectation (step 5). The test is on $\hat{S}_\omega$ (Section 8) and the interval for Q_ω follows Section 9.

7 Bootstrap

Inference uses a block bootstrap with the participant as the block, resampled within the arm. One draw samples as many participants as the arm has (300 by design) with replacement and keeps all of each drawn participant's decisions in the candidate's sample, so the dependence among one participant's decisions is preserved inside the draw. Every procedure uses $B = 999$ bootstrap draws.

One retention rule applies to every procedure. A draw is retained when both its estimated offset sum $\hat{S}^{(b)}$ and its offset ratio $\hat{Q}^{(b)}$ are finite numbers, and is dropped and counted otherwise. Write B_{fin} for the number of retained draws, so $B_{\text{fin}} = 999$ when no draw is dropped; the same retained draws serve the test of Section 8 and the interval of Section 9.

8 Test statistic and p -value

The test statistic is $\hat{S}$ for Procedures A–C and $\hat{S}_\omega$ for Procedure D. I write both as $\hat{S}$ below and compute the p -value from its bootstrap distribution. Write $\hat{S}^{(1)}, \dots, \hat{S}^{(B_{\text{fin}})}$ for the estimates from the retained bootstrap draws of Section 7, and $\#\{\cdot\}$ for the number of draws satisfying a condition. Every procedure uses

$$p := \min \left\{ 1, \frac{2 \left(\min \left(\#\{\hat{S}^{(b)} \leq 0\}, \#\{\hat{S}^{(b)} \geq 0\} \right) + 1 \right)}{B_{\text{fin}} + 1} \right\}.$$

Each candidate is tested on its own; no multiplicity adjustment is applied across candidates.

9 The interval for Q

Each retained draw b carries its own $\hat{S}^{(b)}$ and $\widehat{\text{ME}}_{\Delta}^{(b)}$, the average marginal effect of the current value computed inside the draw (in Procedure D, $\hat{S}_{\omega}^{(b)}$ and $\widehat{\text{ME}}_{\Delta, \omega}^{(b)}$), and

$$\hat{Q}^{(b)} := 1 - \frac{\hat{S}^{(b)}}{\widehat{\text{ME}}_{\Delta}^{(b)}}$$

is formed inside the draw. The 95% interval for Q , the offset ratio of Section 2, is the 2.5th and 97.5th percentiles of $\{\hat{Q}^{(b)}\}$ over the retained draws of Section 7, the same draws that give the p -value of Section 8; the observed $\hat{Q}$ is reported as the point estimate.

10 Reduced-form evidence on path dependence

The regression answers two questions. The first is whether earlier values move the current sell decision once the current value is known, which is path dependence within a stock and the prerequisite for any backward-looking reference point. The second is whether the benchmark of the previous stock still acts as a reference after the break that separates the two stocks, the 60 seconds the participant waits after round 30 of the first stock and before round 1 of the second. The second question needs a yardstick, and the regression supplies two: the current benchmark's own effect within a stock, and the effect of the benchmark the participant leaves at the change of situation inside a stock, where no break intervenes. I fit the sell indicator by least squares, separately in the Framed and Plain arms, on the current value, the current benchmark, the five earlier shocks and the two departed benchmarks. One regression answers both questions and supplies both yardsticks.

10.1 Variables

Index a decision by participant n , stock $s \in \{1, 2\}$ (the two blocks of 30 rounds of Section 1.1) and within-stock round $t = 1, \dots, 30$. Let $j(t)$ denote the situation containing round t , so $j(t) = 1$ for $t \leq 15$ and $j(t) = 2$ for $t \geq 16$, as in Section 1.1.

Y_{nst} 1 when the participant sells, 0 when the participant keeps the stock.

Δ_{nst} the current value, in points.

m_{ns}^1, m_{ns}^2 the benchmarks of the two situations of stock s : m_{ns}^1 is the benchmark of situation 1, which governs rounds 1 to 15, and m_{ns}^2 is the benchmark of situation 2, which governs rounds 16 to 30. The randomisation of Section 1.2 draws both

for the participant, and each stays constant over its 15 rounds.

$m_{ns}^{j(t)}$ the benchmark of the round's own situation: $m_{ns}^{j(t)} := m_{ns}^1$ for $t \leq 15$ and $m_{ns}^{j(t)} := m_{ns}^2$ for $t \geq 16$.

ε_{nst} the round shock of Section 1.2, recovered as $\varepsilon_{nst} := \Delta_{nst} - m_{ns}^{j(t)}$.

$\varepsilon_{n,s,t-d}$, $d = 1, \dots, 5$ the shock d rounds earlier, under one convention: for the second stock a round index $t - d$ at or below zero continues into the first stock, so $\varepsilon_{n,2,t-d}$ with $t \leq d$ means the first stock's shock $\varepsilon_{n,1,30+t-d}$.

Two of the benchmarks above also enter the regression after the participant has left them: m_{ns}^1 , the first-situation benchmark of the stock the participant is in, which the participant leaves at round 16 of that stock; and m_{n1}^2 , the first stock's second-situation benchmark, which the participant leaves at the 60-second break. All values and benchmarks enter in points.

10.2 Sample

Rounds 6 to 30 of the first stock and all 30 rounds of the second, 55 decisions per participant, 16,500 decisions per arm at 300 participants. Rounds 1 to 5 of the first stock have no five earlier shocks, so I drop them. Every round of the second stock has five, because the lags continue into the first stock. The regression runs separately in each arm.

10.3 Model

$$Y_{nst} = a_{ns} + \lambda_{st} + \beta \Delta_{nst} + \gamma m_{ns}^{j(t)} + \sum_{d=1}^5 \tau_d \varepsilon_{n,s,t-d} \\ + \gamma^{\text{within}} m_{ns}^1 \mathbb{I}\{16 \leq t \leq 20\} + \gamma^{\text{across}} m_{n1}^2 \mathbb{I}\{s = 2, t \leq 5\} + e_{nst}.$$

Each coefficient answers one question.

β the response to the value the participant sees.

γ the response to the current benchmark with the value held fixed, which is how strongly that benchmark acts as a reference. A participant who judges the value on its own leaves γ at zero; a participant who judges the value against its benchmark makes γ negative; $\gamma = -\beta$ means the choice depends only on the distance $\Delta_{nst} - m_{ns}^{j(t)}$.

τ_d the response to the shock d rounds earlier, which is path dependence through past deviations from the benchmark.

γ^{within} the response, in the five rounds after the benchmark changes at round 16 of a stock, to the first-situation benchmark the participant has just left, with the current value and the current benchmark held fixed. It equals γ if the departed benchmark still acts as the reference exactly as the current one does, and zero if it has no effect.

γ^{across} the same quantity for the first stock's second-situation benchmark in the first five rounds of the second stock, the rounds that follow the break. It equals γ if that benchmark still acts as the reference exactly as the current one does, and zero if it has no effect.

a_{ns} participant-by-stock fixed effects, one for each participant and stock. They absorb any permanent influence of a departed benchmark on the participant's selling in that stock, so γ^{within} and γ^{across} each describe the five rounds after the change relative to the rest of that stock.

λ_{st} round-by-stock fixed effects, one for each of the 55 positions (s, t) in the sample of Section 10.2.

e_{nst} the regression residual.

The design of Section 1.2 randomises the benchmarks and the shocks, so each coefficient is a causal effect.

10.4 Estimator

I absorb the two sets of fixed effects jointly by alternating projections and estimate the nine remaining coefficients by least squares. The covariance is clustered by participant and carries the small-sample factor $\frac{N_P}{N_P-1} \cdot \frac{N-1}{N-9}$, where N_P is the number of participants, N the number of decisions in the arm, and nine the number of estimated coefficients.

This clustered covariance supplies the standard errors of the nine coefficients and of the sum $\tau_1 + \dots + \tau_5$, the joint Wald test that the five lagged-shock coefficients are zero, referred to the $F(5, N_P - 1)$ distribution, and delta-method standard errors for the two ratios $\gamma^{\text{within}}/\gamma$ and $\gamma^{\text{across}}/\gamma$. The block bootstrap of Section 7, the sign test of Section 8 and the percentile interval of Section 9 govern Procedures A to D alone.

11 Predictive gain from EO-constrained referents

11.1 Purpose

Every candidate referent has an offset ratio Q , defined in Section 2 as the fraction of a change in the current value that an equal change in the referent offsets, and Procedures A and B estimate it. This section asks whether candidates with higher estimated offset ratios $\hat{Q}$ also predict selling decisions better. Equivalent Offsets says that the probability of selling depends on the current value Δ and the candidate referent R only through the single index $\Delta - R$. I therefore fit, for each candidate, a flexible binary model of the sell decision on that index and compare its held-out fit with the fit of the same model built on the current value alone. The comparison delivers one number per candidate: the held-out predictive gain of the EO-constrained index over the current-value-only model. I record whether candidates with a higher $\hat{Q}$ earn a larger gain.

11.2 Candidates

The exercise runs separately in each arm and covers the 18 candidates of Procedures A and B. Twelve are the single referents of Section 3.1: the benchmark, lags 1 to 5, initial values 1 to 3, and the prior minimum, maximum and mean. Six are the functional rules of Section 4.1: adaptive expectations, the moving average, peak-end, the Baucells–Weber–Welfens weighting, extrapolative expectations and experience weights. The plan specifies 23 candidates per arm in all; the mixture of Procedure C and the four lagged distributions of Procedure D carry more than one referent each, so no single scalar index $\Delta - R$ exists for them and this section leaves them out.

Each functional rule carries a parameter vector θ , and I hold θ at $\hat{\theta}$, the estimate that step 1 of Section 4.3 returns in that arm, so every rule enters as a fixed column of referent values. Write R_{nst} for the value the candidate takes at the decision of participant n , stock s and within-stock round t , indexed as in Section 10.1, and write $\hat{Q}$ for that candidate’s estimated offset ratio from Procedure A or B in that arm.

11.3 Sample

All 18 candidates share one set of decisions: within-stock rounds $t = 6, \dots, 30$ of each stock, which is 50 decisions per participant and 15,000 per arm at 300 participants. Round 6 is the first round at which lag 5 exists, so the common sample is the largest on which every candidate is defined, and differences in gain across candidates come from the candidates rather than from the decisions they are scored on. The outcome is Y_{nst} , equal to 1 when the participant sells and 0 when the participant keeps the stock, as in

Section 10.1.

This set is the sample of the gain alone. Each $\hat{Q}$ keeps the sample its own procedure gives it in Section 3.2 or Section 4.2, where the benchmark uses all 18,000 decisions of the arm, lag k uses 600 $(30 - k)$, and the three initial values, the prior minimum, the prior maximum, the prior mean and the six functional rules use 16,200. The two members of each pair are therefore estimated on different sets of decisions.

11.4 The two models

Both models apply a smooth function to one scalar regressor and pass its value through the logistic function to obtain the probability of selling. They differ only in which regressor they use.

Value-only model the regressor is the current value Δ_{nst} . This model lets the sell probability bend with the payoff in any smooth way and knows nothing of the candidate. It is the same model for all 18 candidates within an arm.

Candidate model the regressor is the index $G_{nst} := \Delta_{nst} - R_{nst}$, the gap between the current value and the candidate referent. Writing the probability as a function of G alone imposes Equivalent Offsets for that candidate: a rise of one point in the referent moves the probability exactly as a fall of one point in the current value does.

I represent the smooth function with a spline basis. Let x denote the regressor of the model at hand, Δ in the value-only model and G in the candidate model. The basis is a cubic B-spline in x with ten knots placed at the 0, 1/9, 2/9, ..., 1 quantiles of x , so the outermost two knots sit at the smallest and the largest observed value and eight knots fall inside that range. The basis functions are cubic between consecutive knots and have a continuous value, first derivative and second derivative at every knot; beyond the two boundary knots each basis function continues linearly, so a held-out value outside the fitted range receives a finite score. Write $\varphi_1(x), \dots, \varphi_p(x)$ for the basis functions, $p = 12$ in number, which is the ten knots plus the degree three less one, and write $\Lambda(z) := 1/(1 + \exp(-z))$ for the logistic function. The model is

$$\Pr(Y = 1 \mid x) = \Lambda\left(b_0 + \sum_{\ell=1}^p b_{\ell} \varphi_{\ell}(x)\right),$$

fitted by penalized maximum likelihood on the decisions of the arm, pooled across participants.

The basis functions sum to one and are therefore collinear with the intercept, which leaves

the split between b_0 and the b_ℓ undetermined. A ridge penalty on the basis coefficients alone fixes that split: the fit minimises one half of the sum of squared basis coefficients, $\frac{1}{2} \sum_{\ell=1}^p b_\ell^2$, weighted 10^{-6} against the summed log-likelihood of the decisions, with the intercept b_0 unpenalised. The penalty selects a stable coefficient vector for prediction.

11.5 Cross-fitting

I score each model on decisions it did not see.

1. **Folds.** I randomly assign participants to 5 folds $\mathcal{F}_1, \dots, \mathcal{F}_5$. Each participant's decisions stay together. I draw the assignment once and use it for both models and all 18 candidates in the arm; the analysis code records the assignment.
2. **Fits.** For each fold $\xi = 1, \dots, 5$, I fit both models on the decisions of participants outside $\mathcal{F}_\xi$. The knots of the basis come from those training decisions alone, placed at the quantiles of x among them, so the held-out decisions enter neither the coefficients nor the knots.
3. **Predicted probabilities.** Each decision is predicted by the fit that excludes its own participant's fold. Write $\hat{p}_{nst}^{\text{cand}}$ and $\hat{p}_{nst}^{\text{val}}$ for the sell probabilities the candidate model and the value-only model assign to decision (n, s, t) , each clipped to the interval $[10^{-9}, 1 - 10^{-9}]$ so that its logarithm is finite. Every predicted probability this section scores is clipped this way, held-out and in-sample alike.

11.6 The held-out gain

The log score of a predicted probability $\hat{p}$ at outcome Y is $Y \log \hat{p} + (1 - Y) \log(1 - \hat{p})$, the log-likelihood the model awards that single decision. The predictive gain at a decision is the difference of the two log scores,

$$g_{nst} := \left[Y_{nst} \log \hat{p}_{nst}^{\text{cand}} + (1 - Y_{nst}) \log(1 - \hat{p}_{nst}^{\text{cand}}) \right] - \left[Y_{nst} \log \hat{p}_{nst}^{\text{val}} + (1 - Y_{nst}) \log(1 - \hat{p}_{nst}^{\text{val}}) \right],$$

positive when the index predicts the held-out decision better than the current value alone. Let $\bar{g}_n$ be the mean of g_{nst} over the 50 decisions of participant n , and let N_P be the number of participants in the arm. The reported gain and its standard error are

$$\widehat{\text{Gain}} := \frac{1}{N_P} \sum_{n=1}^{N_P} \bar{g}_n, \quad \text{SE}(\widehat{\text{Gain}}) := \frac{1}{\sqrt{N_P}} \sqrt{\frac{1}{N_P - 1} \sum_{n=1}^{N_P} (\bar{g}_n - \widehat{\text{Gain}})^2}.$$

The gain averages the participant means, which equals the mean over held-out decisions because every participant contributes the same 50 decisions. The standard error treats

the participant as the independent unit, as the bootstrap of Section 7 does, and carries the dependence among one participant’s decisions into the spread of the $\bar{g}_n$.

I also compute the in-sample gain, the same difference of mean log scores with both models fitted on all the arm’s decisions and scored on those same decisions, their predicted probabilities clipped as in Section 11.5. The in-sample gain measures how much the index improves the description of the data the models saw; the held-out gain is the quantity the comparison rests on.

11.7 Ranks and their correlation

Within an arm I rank the 18 candidates twice: by $\hat{Q}$, from 1 for the highest estimate to 18 for the lowest, and by $\widehat{\text{Gain}}$, from 1 for the largest gain to 18 for the smallest. A tie in either ordering is broken by the order of Section 3.1 followed by Section 4.1, so the 18 ranks are always the integers 1 to 18. The association between the two orderings is the Spearman rank correlation between the 18 pairs $(\hat{Q}, \widehat{\text{Gain}})$, reported with the Kendall rank correlation of the same pairs. Both run over candidates within an arm, and the two arms give two pairs of correlations.

I compute the same rank, the same two rank correlations and the same count of candidates with a negative gain on the in-sample gain as on the held-out gain.

11.8 Computed quantities

For each arm and each of the 18 candidates, the computation returns $\hat{Q}$ from Procedure A or B with its 95% interval from Section 9, the held-out gain $\widehat{\text{Gain}}$ with its standard error, the candidate’s rank by $\hat{Q}$ and by held-out gain, and, for the arm, the Spearman and Kendall rank correlations of Section 11.7. The in-sample gain, its ranks and its rank correlations are computed and archived with the output.

Each candidate’s gain stands on its own. No multiplicity adjustment is applied across the 18 candidates, and no hypothesis test is run on the gains or on the rank correlations.

References

- Baucells, M., M. Weber and F. Welfens (2011). Reference-point formation and updating. *Management Science* 57(3), 506–519.
- Cleveland, W. S. (1979). Robust locally weighted regression and smoothing scatterplots. *Journal of the American Statistical Association* 74(368), 829–836.

- de Boor, C. (2001). *A Practical Guide to Splines*, revised edition. Springer.
- DellaVigna, S., A. Lindner, B. Reizer and J. F. Schmieder (2017). Reference-dependent job search: evidence from Hungary. *Quarterly Journal of Economics* 132(4), 1969–2018.
- Fredrickson, B. L. and D. Kahneman (1993). Duration neglect in retrospective evaluations of affective episodes. *Journal of Personality and Social Psychology* 65(1), 45–55.
- Härdle, W. and T. M. Stoker (1989). Investigating smooth multiple regression by the method of average derivatives. *Journal of the American Statistical Association* 84(408), 986–995.
- Kőszegi, B. and M. Rabin (2007). Reference-dependent risk attitudes. *American Economic Review* 97(4), 1047–1073.
- Loader, C. (1999). *Local Regression and Likelihood*. Springer.
- Malmendier, U. and S. Nagel (2011). Depression babies: do macroeconomic experiences affect risk taking? *Quarterly Journal of Economics* 126(1), 373–416.
- Metzler, L. A. (1941). The nature and stability of inventory cycles. *Review of Economics and Statistics* 23(3), 113–129.
- Nerlove, M. (1958). Adaptive expectations and cobweb phenomena. *Quarterly Journal of Economics* 72(2), 227–240.
- Thakral, N. and L. T. Tô (2021). Daily labor supply and adaptive reference points. *American Economic Review* 111(8), 2417–2443.