

Compute Market

Research report

Curves, claims, operating assets and financed control

1 Research design

One coupled curve law generates compute benchmarks and basis, power, rates and funding. Physical service, hardware appraisal, financial claims and operating decisions use the resulting state. Execution liquidity adds idiosyncratic and shared funding shocks. The study evaluates the full chain rather than only price forecasts.

All inputs are synthetic. The portfolio experiment trains on **4,096 paths**, evaluates **1,024 independent paths**, and uses **256 conditional innovations per decision**. Thirteen named stresses use **512 paths each**. The horizon is one year with eight operating intervals. Policies share every held-out market, physical, liquidity and counterparty event.

The financial universe has eleven dated futures, two paid puts, a transferable fixed-strike power forward and a periodic compute cap. The adaptive controller jointly chooses current dispatch, financing and financial inventory against its fitted continuation. Separate experiments examine delivery-bucket risk, operating options, parameter identification, latent observations and finite-control reservation prices.

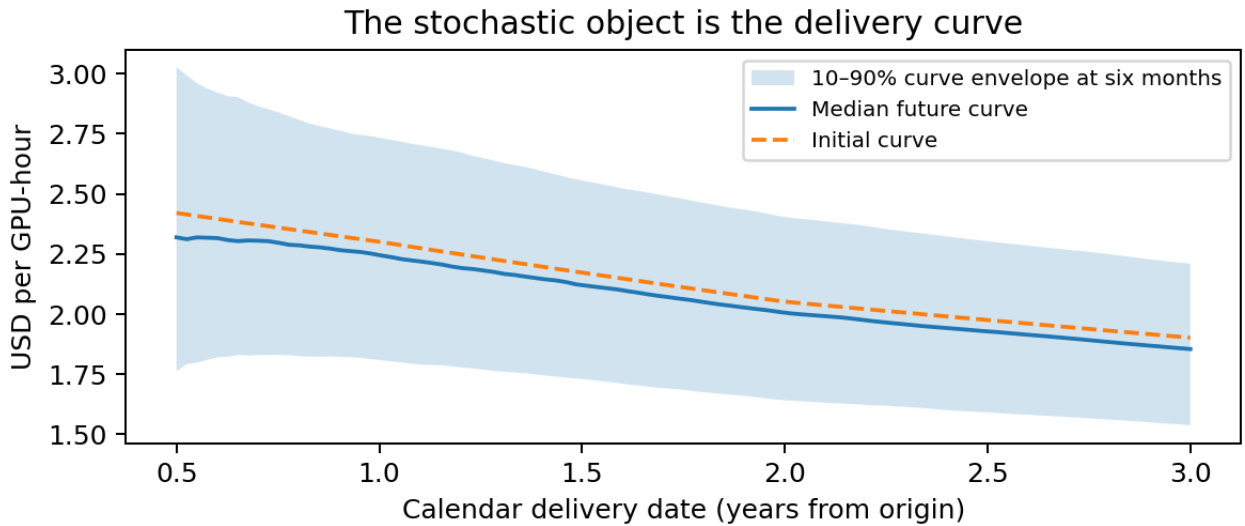


Figure 1: Distribution of the entire G1 delivery curve at six months. The band shows pointwise future market variation under the physical law.

2 Curve and claim pricing identities

The positive futures curves include the complete pricing-law cumulant, while electricity uses an additive reference. Across **65,536 pricing-law paths**, the largest absolute discrepancy from the futures martingale identity is **2.42 Monte Carlo standard errors** across twelve market/date checks. The two-year analytical discount factor is **0.92890389**, versus **0.92886497** in simulation with standard error **0.00003815**.

The two-year G1 futures reference is **2.050000**; the corresponding terminal-forward strike is **2.050820**. This difference is generated by rate–compute covariance. Setting the rate loading to zero removes it. Splitting and recombining an arithmetic delivery strip changes its mark by **4.4e-16**.

The option registry uses the discounted affine transform for deterministic batch pricing. Doubling Fourier quadrature from 128 to 256 nodes changes the three tested put prices by at most **1.2e-14**. Independent pricing-law Monte Carlo, Gaussian-limit formulas, put–call parity and expiry settlement provide separate checks. Deterministic marks let the operating controller compare funded protection with futures without introducing option-pricing simulation noise into every optimization.

3 Identifying model parameters

The initial-curve fit recovers eight knots from dated and strip observations. Its data Jacobian has rank **8**, with weighted quote RMSE **0.387** in measurement-standard-error units.

The structural recovery experiment uses a two-factor affine model, **4,096 independent physical histories**, 24 transitions per history and nine noisy option prices. Four declared coordinates are estimated while nuisance jump parameters and loading anchors remain fixed. Six option prices at excluded expiries form the pricing holdout. This controlled identification test is separate from the nine-factor operating study, whose physical drift is fitted from its own training sample.

Free coordinate	Generating value	Estimate	Bootstrap 2.5–97.5%
level mean reversion	0.80000	0.79301	0.77156–0.80616
level diffusion	0.25000	0.24972	0.24722–0.25141
physical level drift	0.03000	0.03265	0.02784–0.03613
pricing jump intensity	0.30000	0.27375	0.22896–0.31822

All **24 of 24** cluster-bootstrap refits converged. The data-only Jacobian has full rank **4**. Held-out option-price RMSE is **0.000569 per quoted unit**. A separate example attempts to infer physical drift from a time-zero futures reference: its data rank is **0**, even with a prior penalty.

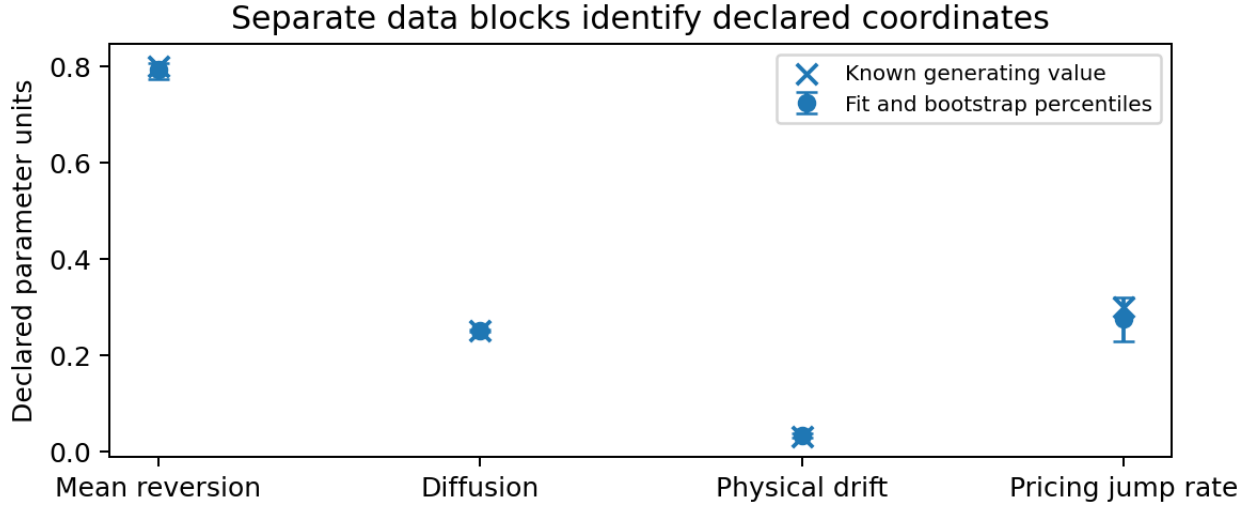


Figure 2: The recovery experiment fixes factor gauge and nuisance coordinates. Bars are percentiles across 24 refits.

The predictive variance calculation separates **0.278340** of within-model process variance from **0.000048** of between-model variance. These values describe this well-observed synthetic design; a short market history need not have the same balance.

4 Hidden factors and a short history

A separate latent experiment sees three noisy clean prices at each of 145 dates over three years. It estimates one physical drift coordinate with the remaining law fixed. The fitted drift is **-0.1203**, against the generating value 0.0300. Its two factor-state RMSEs are **0.0092** and **0.0229**.

The retained quasi-likelihood profile explains the difference. Over the entire tested range from -0.20 to $+0.20$, twice the objective change stays below **3.21**. The short history tracks states more sharply than it identifies long-run drift. This profile measures sensitivity with nuisance coordinates fixed; finite-sample coverage would require a design-specific bootstrap.

The nine-factor observation experiment uses fifteen prices per date and predicts two withheld local prices. Its held-out quote RMSE is **0.0094 per GPU-hour** and the minimum observation rank is **8 of nine**. The operating policy study observes its simulated factors; filtering performance is measured independently.

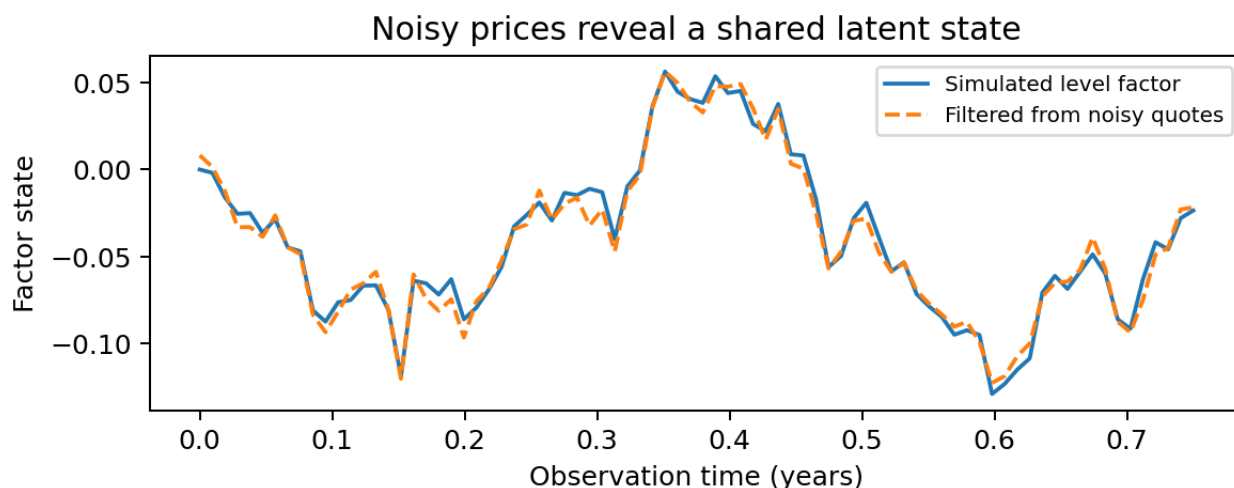


Figure 3: One component of the nine-factor noisy-quote experiment. The filter and the process history are retained separately.

5 Delivery buckets, factors and unspanned shapes

A clean physical-margin strip is perturbed in **21 calendar-delivery buckets** across two compute references and one power reference. Compute buckets use log coordinates; power buckets use price coordinates. Their Jacobian spans **8 of nine** factor directions. For these node-supported claims, the largest absolute bridge error versus analytic factor Greeks is **0.0042 currency per factor unit**, on million-dollar sensitivities. Half-size bumps change a bucket delta by at most **0.00062**.

A constructed bucket shock has factor-projection norm **2.2e-16**, yet produces **4,648.82** of fully revalued P&L. The local bucket estimate is **4,634.73**. This is explicit shape risk outside the affine factor span. The experiment assigns no occurrence probability to that off-span shock.

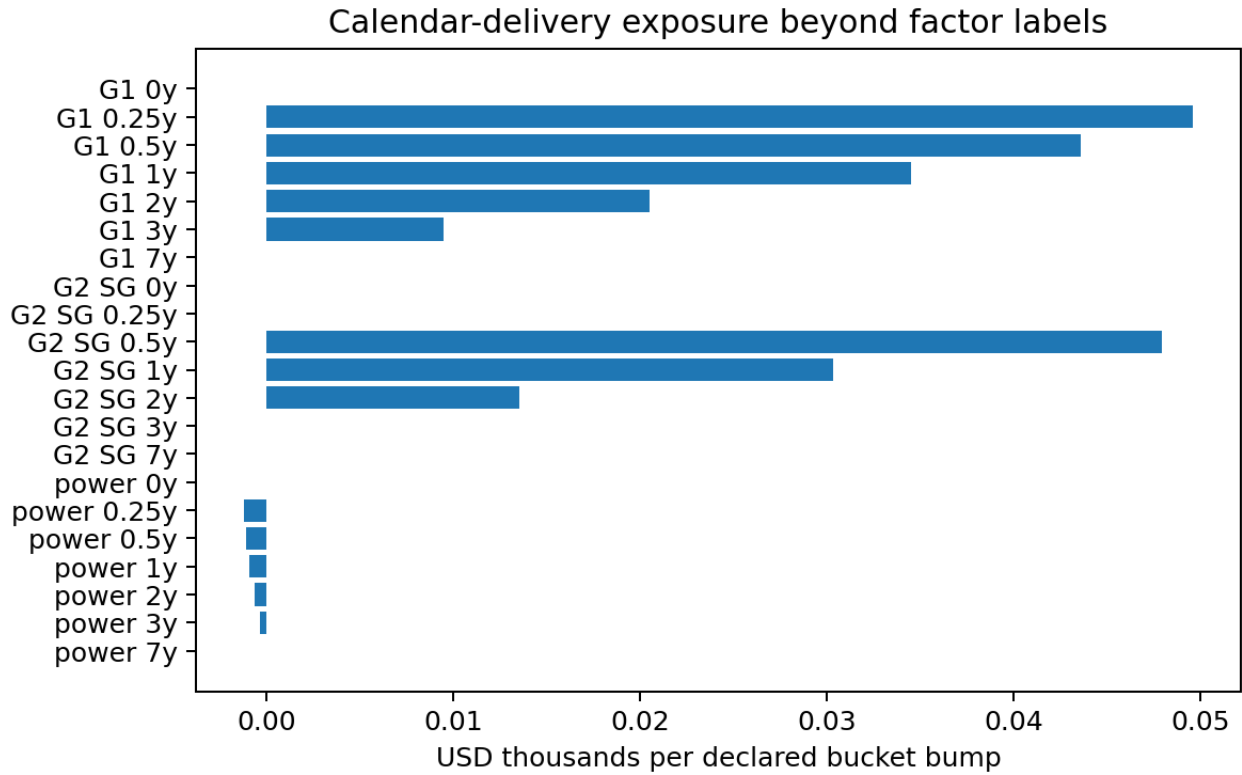


Figure 4: Currency P&L per declared bucket bump. Compute bumps are 0.0001 in log price; power bumps are 0.01 currency/MWh.

A second callback reoptimizes a financed book after changing the G1 US initial curve. All bumps reuse the same conditional innovations. Its small numerical fixture has three branches and two intervals; it measures policy response on that declared support.

Delivery center	P&L for a 0.0001 log bump	Delta change after half bump
0.5 years	34.08	-0.0006
1 years	41.62	-0.0004
3 years	10.92	0.0000

The covariance-weighted hedge projection is different from the bucket-shape projector. The full financial universe spans 7 local covariance directions. Its constrained version imposes signed positions, turnover and a **430,000** cash budget. Nonlinear jump and funding risks remain in the full conditional revaluation.

6 Economic hardware options

The asset experiment values a 64-GPU cohort over three years, using 24 operating intervals. It includes startup and shutdown costs, standby expenses, two-interval minimum runs and an independent retirement-sale function. The interval rental rate is fixed when operating capacity is committed. The policy trains on **8,192 pricing-law paths** and is evaluated on **8,192 independent paths**.

Operating rule	Discounted value (USD)	Standard error (USD)
Switching and retirement	817,178	5,201
Always operate	706,838	6,477
Sell immediately	500,000	—
Perfect-information upper comparison	863,864	4,871

The fitted policy operates for an average **16.26 intervals** and retires on **55.5%** of paths. Its mean gap to the information upper bound is **46,686**, with paired standard error **924**. Every policy path satisfies its operating constraints and remains below the corresponding upper payoff up to rounding.

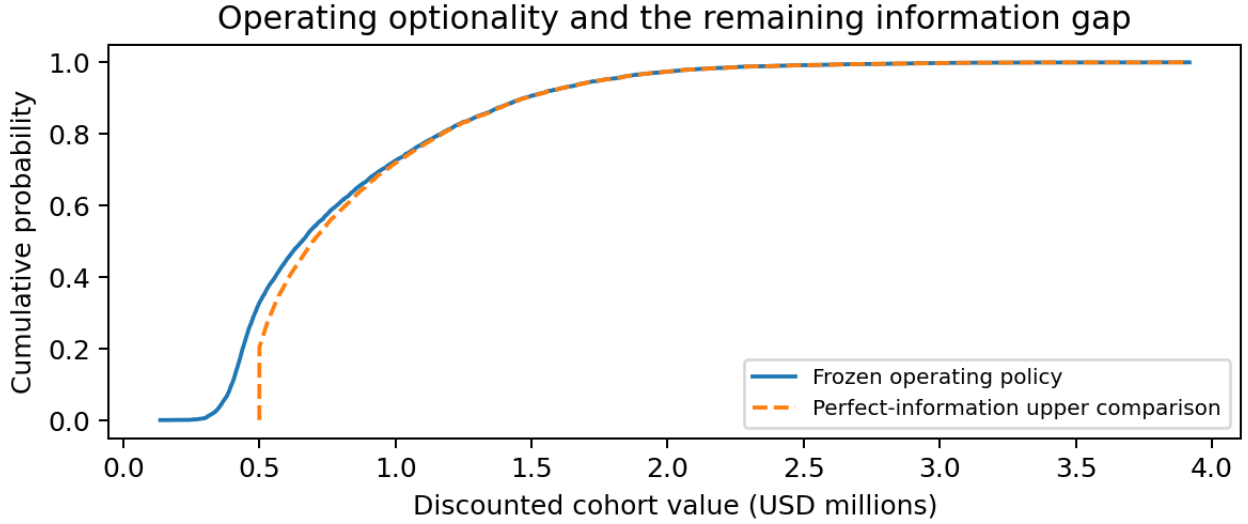


Figure 5: Independent operating-policy outcomes and the pathwise information upper comparison.

The software separately records lender marks and executable liquidation proceeds for the financed fleet. Those totals have a different asset scope from the 64-GPU experiment. Borrowing uses the lender convention; disposal uses executable proceeds; investment analysis uses the operating policy value.

7 Portfolio policy outcomes

All policies start with the same capacity, commitments, debt and cash. The worst-5% column averages terminal wealth over exactly 5% of probability mass, with fractional weight at the boundary observation. Higher mean and tail wealth are preferable; lower standard deviation indicates less dispersion.

Policy	Mean wealth (USD m)	Worst-5% wealth (USD m)	SD (USD m)
Unhedged	3.566	1.578	1.088
Static financial inventory	3.653	2.708	0.517
Adaptive joint policy	3.704	3.013	0.362

Adaptive control raises worst-5% mean wealth by **1,434,742** versus unhedged, with a paired bootstrap interval of **1,309,347–1,553,503**. Mean wealth rises by **137,476** and standard deviation falls by **66.8%**. Its worst-tail improvement over static inventory is **304,865**. These intervals measure sampling uncertainty under the specified market law.

Each path carries **350,400 GPU-hours** of firm service over the year. Delivery shortfall includes service left unfulfilled after financing failure.

Policy	Mean shortfall (GPU-hours)	Share of promised service
Unhedged	38.88	0.0111%
Static financial inventory	38.88	0.0111%
Adaptive joint policy	38.88	0.0111%

Initial margin is recorded after every trade and returned at settlement. Average posted margin weights each balance by the time it remains posted; the peak statistic is the 95th percentile of each path’s maximum requirement. Borrowing interest, debt fees, premiums and trading costs enter terminal wealth.

Policy	Average posted IM (USD m)	Peak IM, p95 (USD m)	Failed paths
Unhedged	0.000	0.000	0 / 1,024
Static financial inventory	0.352	0.744	0 / 1,024
Adaptive joint policy	0.492	0.952	0 / 1,024

A matched ablation restricts the financial universe to eleven futures. The full universe adds two paid puts, a transferable power forward and a periodic cap, with the same training, tape and conditional innovations. Their joint contribution is **11,039** to mean wealth (**0.30%**) and **41,451** to worst-5% mean wealth (**1.40%**). The paired interval for the tail improvement is **23,408–61,108**. Peak margin at p95 changes by **-69,972**, and average posted margin by **-39,916**. The incremental benefit is modest in this configuration; practical instrument selection also depends on access, funding terms and implementation costs.

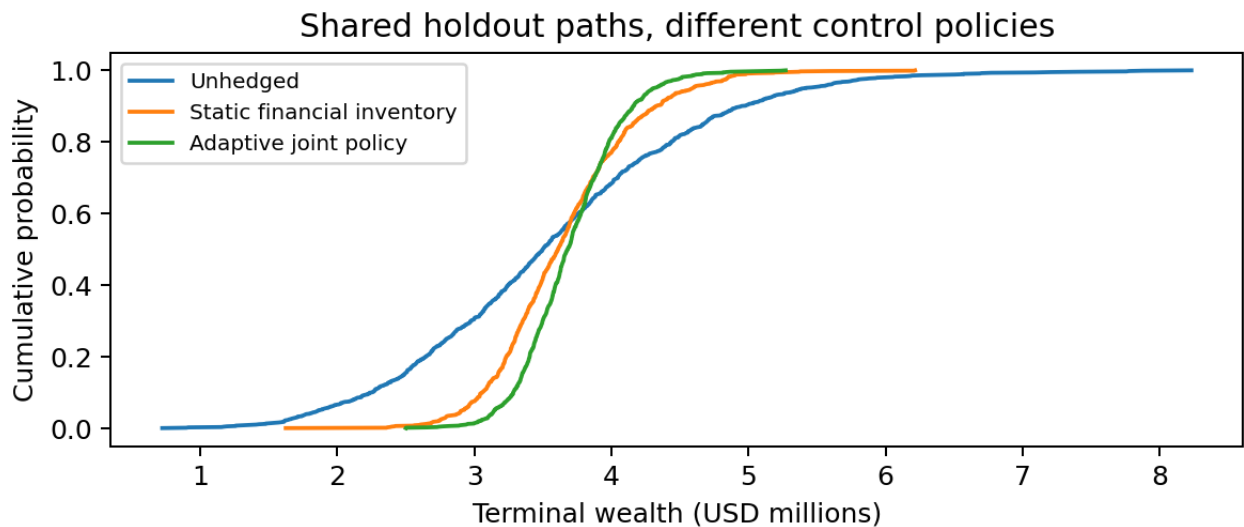


Figure 6: Held-out terminal wealth. Saved policies include financial positions, cash, debt and accounting components.

8 Liquidity, credit and stress

The held-out tape includes stochastic spreads, executable depth and collateral multipliers coupled to funding and market downside. Supplier survival is observed causally. Funded inventory, segregated margin, payable liabilities and default recovery follow their actual ledger entries.

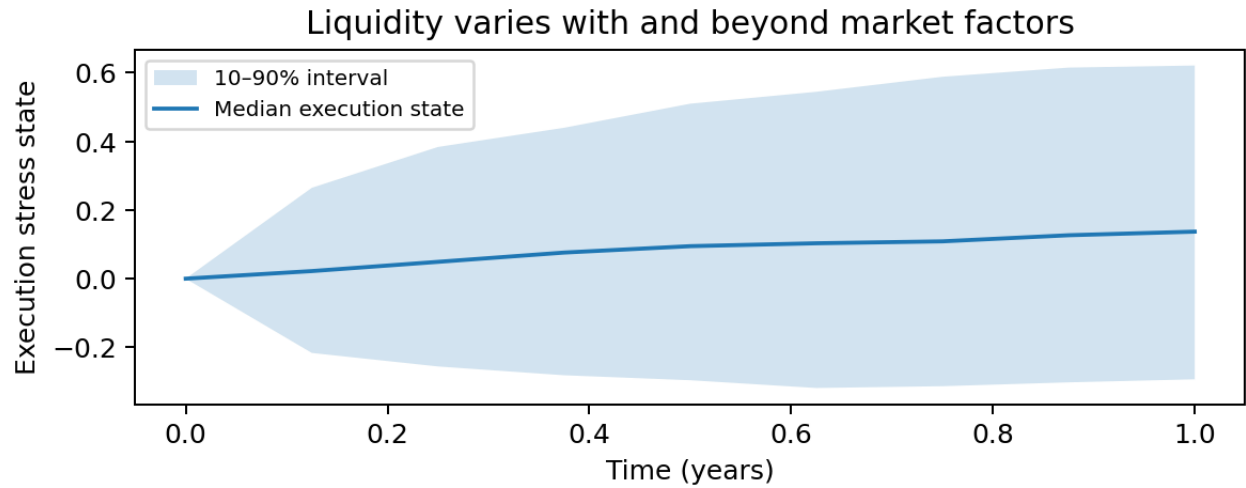


Figure 7: The stochastic execution state. Named stress interventions act on top of ordinary market-access variation.

Named intervention	Adaptive gain in worst-5% wealth (USD m)	Adaptive failures / 512
Level collapse	+2.313	0
Long end repricing	+2.375	0
Generation transition	+1.631	0
Regional scarcity	+1.440	0
Power and rates	+1.668	0
Basis break	+1.314	0
Benchmark dislocation	-0.257	1
Margin call	+1.177	9
Execution freeze	+1.326	0
Lender haircut	+1.501	0
Outage	+1.208	0
Liquidity spiral	+1.442	19
Counterparty default	+1.503	0

Benchmark dislocation tests whether the hedge remains aligned with the physical book. Margin, execution and liquidity interventions test whether the operator can finance its positions through adverse cash demands. Failure counts belong to each named intervention and remain in its wealth and service outcomes.

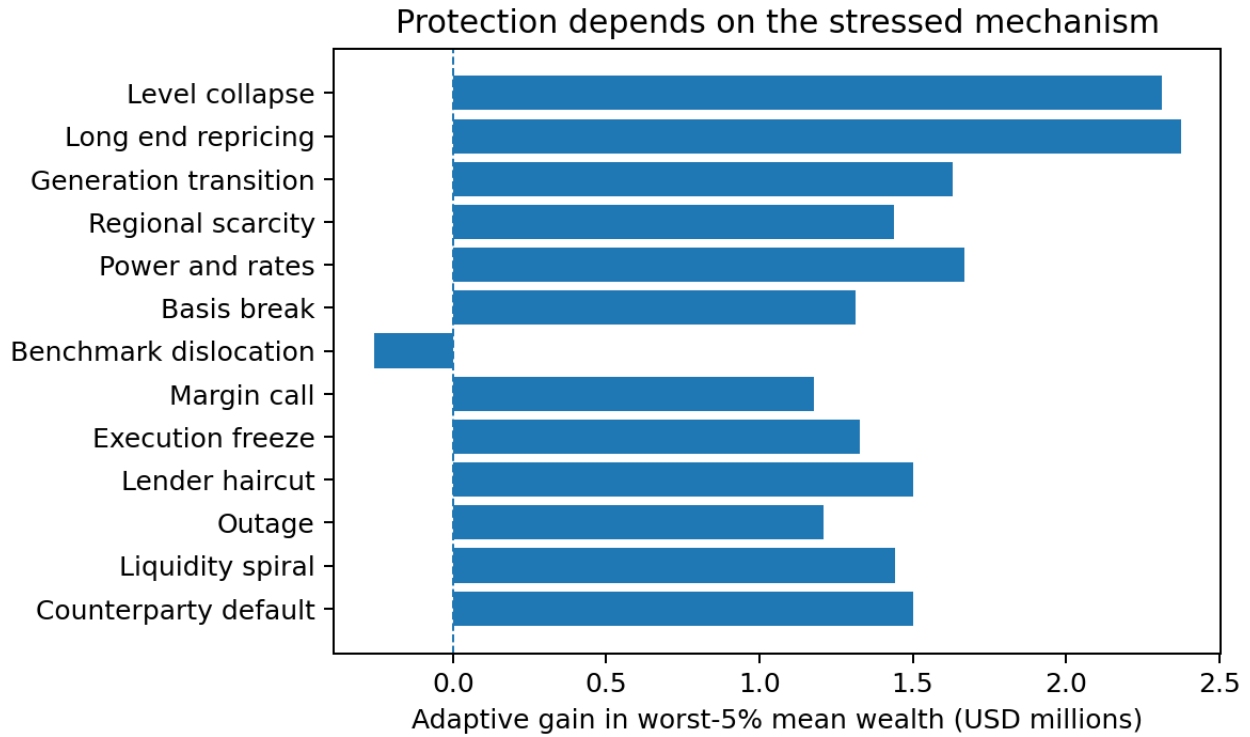


Figure 8: A negative bar means the hedge underperforms the unhedged physical book under that intervention.

9 Reservation pricing and numerical resolution

Complete financed cases are sampled from the same curve law. The study fixes two decision intervals over one year and changes only branching and independent seed. The adapter retains aligned futures, European options and the fixed-strike forward. It records the quarterly cap's exclusion from this coarser grid rather than moving its cash events.

Branches	History nodes	Payment range, three seeds	
		(USD)	
8	73	102,998–109,244	0.8
16	273	91,914–113,508	1.6
32	1,057	100,414–119,549	3.2

Numerical payment certificates are tighter than the spread across scenario constructions. The latter measures support sensitivity. Increasing branching does not enforce monotone price convergence across independent draws. Conditional tail-weight and worst-state saturation diagnostics accompany every result.

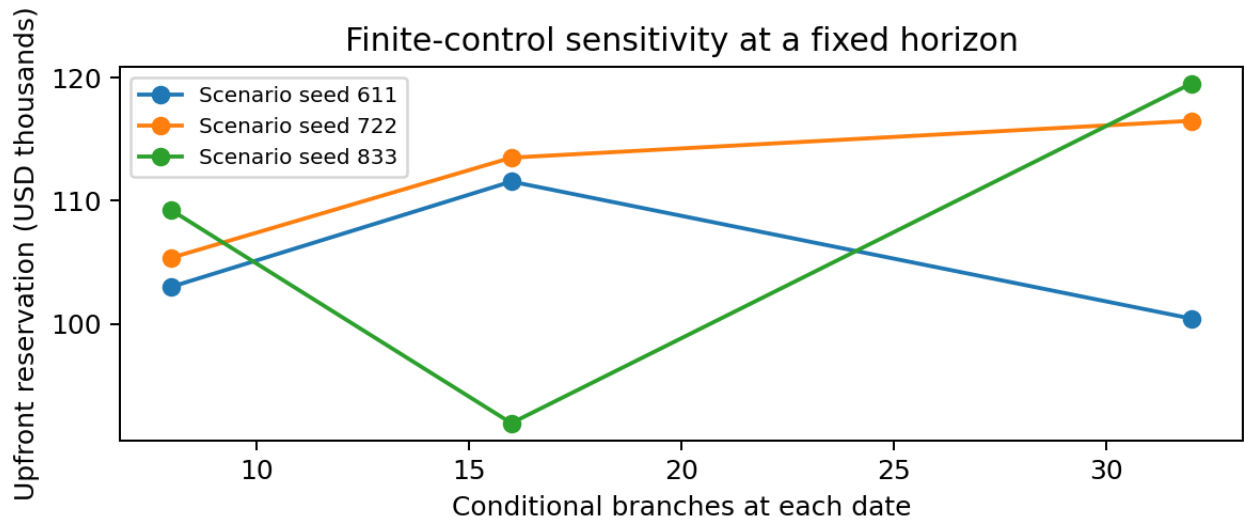


Figure 9: Reservation payments from independent finite scenario constructions at a fixed horizon.

10 Evidence and reproducibility

The largest cash-equation residual across all portfolio cases is **9.3e-10 currency units**; the largest independent terminal-wealth reconstruction error is **2.8e-09**. Maximum numerical prefunding residual is **1e-06**. Distributions, premium exchanges, debt principal, interest, transaction costs and hardware sales remain separate accounting components.

`results/study/` contains complete valuation cases and all training, holdout and stress tapes. `assets/` stores the operating rule, independent outcomes and upper comparison. `risk/` stores bucket shocks, factor bridges, financed-book bumps and transform checks. `calibration/` stores raw synthetic histories, noisy quotes, anchors, parameter draws, held-out repricing and the latent profile. `pricing/` retains complete finite programs and solver certificates.

Software checks, coverage, extracted-package tests, installation and replay outcomes are in `docs/VALIDATION.md`.