
INDEX PROSPECTUS

GM! Index

Index of futures conviction
A full-distribution lattice and machine-native data suite

Primary: core-6 / 1h / multi-horizon Fine tier: crypto-3 / 5m / 90M

Index Sponsor, Calculation Agent and Administrator
ConfigSecretAI, Inc. — MUTANTDEFI Research Lab

Document version: 2.4.4

Publication date: Saturday, 4 July 2026 (2026-07-04); revised Sunday, 19 July 2026 (2026-07-19)

Index distribution: live and continuous — gm.mutantdefi.com (HTTP + x402 + MCP)

Historical replay floor: 6 January 2026 (universe data floor)

GM! Index is a calculated market-information index: a distribution of directional conviction across instruments and horizons, not a portfolio. This document describes the Index, its current primary and fine-tier composition, its additive hurdle-net Omega representation, its agent-native data suite, the evidence program that validates its information content, and its risk factors. Construction internals remain proprietary; the published contract, governance gates, validation record, and evidence boundary are disclosed. Capitalized terms are defined in Appendix [D](#).

Important Notices

Nature of the Index. GM! Index (the “Index” or “GM!”) is a real-time *information index*: a lattice of bounded directional-conviction and confidence values calculated continuously from HYPERLIQUID futures market data and published as data. Like the VIX, the Index itself is not investable; it is a measurement. References in this document to trading results concern the Index’s *utilization validation benchmark* — a deliberately minimal rule used to prove that the published values carry extractable forward information — and not the Index.

No offer or solicitation. This prospectus is an index methodology and disclosure document. It does not constitute an offer to sell, or the solicitation of an offer to buy, any security, futures contract, swap, or other financial instrument in any jurisdiction. Access to Index values is sold as machine-readable data over the x402 micropayment protocol; no investment advice, discretionary management, or trading service is offered.

Regulatory posture. The Sponsor’s validation traders operate exclusively in dry-run (paper) mode as a deliberate United States compliance posture, while the research program maintains an aggressive optimization cadence against predeclared falsification gates. This design is described in Section 14 and should be read before the benchmark record in Section 9.

Source of record. Every quantitative claim in this document is drawn from the published MUTANTDEFI research archive (dfai/publications/journal/) and the DFY platform repository (dfy/). Where two published runs report different values for the same window, both are shown and reconciled (Section 13, Appendix A).

Restricted information. The construction recipe of the Index (component weights, scales, and source decomposition) is proprietary, protected by a codified projection boundary (Section 7). This prospectus discloses only the public surface of the Index contract.

Contents

Important Notices	1
1 The Index at a Glance	6
1.1 Headline record	6
2 Introduction and Index Objective	7
2.1 Objective	7
2.2 Conviction distribution, not current price	7
2.3 The three-layer separation	7
2.4 Plural sources of authority	8
2.5 Research philosophy: an evidence-bound iterative process	8
2.6 What distinguishes this Index	9
2.7 The Index as a standalone information product	9
3 Index Sponsor, Administration and Governance	9
3.1 Sponsor and calculation agent	9
3.2 Promotion gates	10
3.3 Versioning and restatement	10
4 Index Definition and Published Values	10
4.1 The lattice	10
4.2 Published channels	11
4.3 Symbol contract	11
4.4 Calculation profile	11
5 The GM! Family: Breadth of Conviction Cells	12
5.1 Horizon ladder	12
5.2 Individual cells	12
5.3 Basket and universe composites	12
5.4 Production and research families	13
5.5 Breadth is consumed, not just published	13
6 Index Universe and Eligibility	13
6.1 Constituent markets	13
6.2 Eligibility and changes	13
6.3 Per-instrument signal quality	14
7 Calculation, Distribution and Data Suite	14
7.1 Topology	14
7.2 Data freshness and liveness	15
7.3 Distribution surfaces	15
7.4 Comprehensive GM! toolset	16
7.5 Harvest rungs: selectable admission ladders (FC-27)	16
7.6 The projection boundary (construction protection)	17
8 Index Performance: Signal Efficacy	18
8.1 Registered objective-profile evidence	18

8.2	Efficacy metrics defined	18
8.3	Interpretation	19
9	The Utilization Validation Benchmark	19
9.1	Why a benchmark exists	19
9.2	Development and holdout record (legacy profile study)	19
9.3	Rolling four-window record	19
9.4	Powered extension window (current lineage)	20
9.5	Summary statistics	21
9.6	Statistical significance and sample power	22
9.7	Live capture	23
10	The Validation Program: A Ten-Arm Utilization Battery	23
10.1	Arms under test	24
10.2	Primary holdout ladder	24
10.3	Rolling ladder across arms	25
10.4	Interpretation: the naive readout is the demonstrated optimum of the admission layer	25
10.5	Alpha: separating the benchmark's return from market return	26
10.6	Subsequent registered sizing and exit results	27
11	Governance Case Study: Rejection of the UTIL250 Candidate	27
12	Strategy Lineage and Research Provenance	28
13	Index Maintenance, Data Integrity and Restatement Policy	28
13.1	Data sources	28
13.2	Documented incidents and corrections	28
13.3	Restatement policy	29
14	Regulatory Posture and Validation Design	29
14.1	The Index as a data product	29
14.2	Why validation runs in paper mode	30
14.3	Benchmark performance disclosure	30
15	Research and Product Roadmap	30
16	Risk Factors	31
16.1	Measurement and statistical risks	31
16.2	Market and instrument risks	32
16.3	Venue and operational risks	32
16.4	Model and governance risks	32
16.5	Legal and regulatory risks	32
17	Conflicts of Interest and Disclosure Boundary	33
17.1	Conflicts	33
17.2	Disclosure boundary	33
A	Full Battery and Reconciliation Tables	33
A.1	Interim (superseded) rolling run	33

A.2	Train-window stability across runs	34
B	Statistical Significance Tables	34
B.1	Benchmark, primary holdout	34
B.2	Overlays vs. the benchmark, primary holdout	34
B.3	April 2026 stress window	35
C	Reproducibility Manifest	35
D	Glossary	35
	References	37

List of Tables

1	Summary of key Index terms.	6
2	The three-layer separation (standing invariant of the research program).	8
3	Standing promotion gates governing any change to the Index or its benchmark.	10
4	Lattice cell addressing.	10
5	Published horizon ladder. Extended horizons bridge the legacy four-rung ladder where analysis requires finer resolution.	12
6	Published and researched profile families across the GM! breadth.	13
7	Index universe: core-6 HYPERLIQUID 1h perpetual futures.	13
8	Per-instrument 1h/1D forward IC, macro legs, 6 Jan–18 Jun 2026.	14
9	Production topology of the GM! data stack.	14
10	Canonical public GM! routes and discovery resources.	15
11	Current GM! Index and companion data capabilities.	16
12	Harvest ladder rungs on the frozen 1h core-6 surface, window 2026-06-19 to 2026-07-18 (benchmark-style utilization read; paper mode).	17
13	Registered objective-profile gate results (index level, trader-agnostic; not the current production identifier).	18
14	Public efficacy metrics computed per cell and per profile.	18
15	Legacy profile-selection study (battery profile; superseded — retained for audit).	19
16	Benchmark rolling four-window out-of-sample record (unified run, final).	20
17	Powered extension window, 1 May–8 Jul 2026 (identical benchmark admissions; pre-registered reference seams).	20
18	Benchmark summary statistics (four windows; June partial).	21
19	Utilization battery arms. The validation benchmark is control arm A0.	24
20	Full ladder, train and primary holdout (unified run, 4 Jul 2026).	24
21	Rolling ladder by arm (unified run; June overlay fold pending at archive close).	25
22	Benchmark versus equal-weight core-6 buy-and-hold, per window. Separation is benchmark minus market.	26
23	UTIL250 vs. incumbent P00, trade-level battery (G4).	27
24	Lineage champions (best-in-class per series). Do not rank by profit % alone; windows and venues differ.	28
25	Interim rolling run for the benchmark (superseded by Table 16).	33
26	Train-window drift between June waves and July unified run.	34
27	Overlay arms vs. benchmark on primary holdout. No arm is favorable.	34

28 Key artifacts underlying this prospectus (paths relative to the df ai research root and the dfy platform root). 35

MUTANTDEFI

1 The Index at a Glance

Table 1. Summary of key Index terms.

Item	Description
Index name	GM! Index — Index of Futures Conviction
Index type	Calculated market-information index (non-investable measurement, VIX-analogous)
Sponsor / calculation agent	ConfigSecretAI, Inc. — MUTANTDEFI Research Lab
What it measures	The cross-instrument distribution of bounded directional conviction and confidence per (instrument, granularity, horizon) cell
Published values	value $\in [0, 1]$; signed_conviction $\in [-1, +1]$; confidence $\in [0, 1]$; direction; additive hurdle-net Omega fields and hurdle stamps
Primary composition	Core-6 at 1h: BTC, ETH, HYPE, XYZ-GOLD, XYZ-SILVER, XYZ-CL; six published horizons, with 1D heavily used
Fine composition	Crypto-3 at 5m: BTC, ETH, HYPE; primary horizon 90M; routed through the same GM! origin
Production profiles	GM!-FEED-8.01-1H-CORE6-P01 (primary) and GM!-FEED-8.03-5M-CRYPT03-P02 (fine)
Distribution	gm.mutantdefi.com — HTTP, x402, MCP, free discovery, service export, and machine-readable metadata
Data suite	Conviction snapshot, lattice, allocation weights, Omega book, exit hazard, harvest rungs, ticker/OHLCV adapters, MCP execution, and provisioned access
Harvest rungs	Selectable admission ladders (purist / balanced / active) on paid reads via ?rung=; default balanced (Section 7.5)
Agent surface	DFY GM! Oracle on Virtuals, with fixed-fee USDC services and a machine-readable service catalog
Historical replay floor	6 January 2026 (XYZ-CL listing)
Validation benchmark	“A0” minimal utilization rule (paper-mode; Section 9)
Governance gates	G1 holdout, G2 rolling, G3 live reconcile, G4 profile battery (Section 3)

1.1 Headline record

The Index’s record has two legs, in order of authority:

Signal efficacy (the Index itself). On the registered objective profile gate—the trader-agnostic measurement of whether conviction covaries with subsequent market movement—the research record reports a **positive holdout forward information coefficient of +0.032** with **harvest of +14.5 basis points** per active observation (Section 8). On the longer six-month benchmark window (6 Jan–18 Jun 2026, $n = 23,320$ cell-periods) the registered objective-profile study reports **harvest of +15.26 basis points per active period at $t = +5.13$** . A t -statistic is the number of standard errors by which a measured average exceeds zero: at $t = +5.13$, the harvest estimate sits more than five standard errors above the no-information value, an excursion that pure chance produces roughly one time in several million. At the Index level, where more than twenty-three thousand observations average the noise away, the forward information clears conventional significance thresholds decisively. These are Index-level statistics computed across cell observations, not return claims.

Utilization validation (the benchmark). The canonical benchmark lineage is the unified cfi-val A0 record. A deliberately minimal long–short rule reading only the public 1D cells converts the Index’s information into aligned out-of-sample results: **+8.28% cumulative** across the four-window rolling ladder (Mar–Jun 2026, 117 trades, three of four windows positive) with **+1.82%** on the 18-day June holdout slice ($n = 23$), and — on the subsequent *powered* extension window (1 May–8 Jul 2026, which overlaps the May–June rolling months and is therefore not additive with the ladder) — **+11.64% at $n = 27$** with an **81.5% win rate** (22 of

27; one-sided binomial $p \approx 8 \times 10^{-4}$) and 0.17% maximum drawdown (Section 9). *Powered* is a specific, pre-registered property, not an adjective: the evaluation calendar was extended, before results were seen, until the expected trade count crossed the program's declared decision floor ($n \geq 23$) — the sample size at which the registered verdict rules are allowed to apply. Results on powered windows are decidable evidence; results on sub-power slices are labeled as such and never promoted.

The reference-tool program then validated, at pre-registered power on that same extension window, improvements *below* the admission layer: lattice-marginal **sizing** at +14.78% versus +11.64% equal-stake (SUPERIORITY, $n = 27$), and a feed-conditioned **hazard exit** at +11.47% versus +9.73% hand tail control (SUPERIORITY, $n = 30$, three of three windows) — while learned entry overlays remained negative or inconclusive (Section 10). A gate re-derivation on an independent venue and timeframe (Kraken Futures, 5-minute bars) produced +1.93% at trade-Omega 1.71 ($n = 81$) on its own holdout — evidence the Index's information is not one venue's artifact. Reference findings are distributed as tools without redefining the Index itself.

Two disclosures are integral to these figures. The powered extension results share a single, favorable market regime and count as one regime's worth of evidence, not several independent confirmations. And while the June-18 slice remains under-powered at trade level (Section 9.6), the powered extension window now supplies registered $n \geq 23$ verdicts, so the sample-power caution is confined to the short slice rather than to the benchmark record as a whole.

2 Introduction and Index Objective

2.1 Objective

GM! Index measures, in real time, the direction, magnitude, and cross-sectional distribution of conviction embedded in futures market state, together with an explicit reliability channel. Where the VIX distills an option surface into an expected-variance number, GM! publishes a *lattice*: signed conviction, confidence, full-distribution hurdle-net information, and change across instruments and horizons.

The Index is a measurement, not a strategy. Its value proposition is that the published numbers carry demonstrable forward information — a claim the Sponsor treats as falsifiable and tests continuously through the validation program of Sections 9–11.

2.2 Conviction distribution, not current price

Current market prices are inputs to the measurement, not the object GM! claims to replace. The Index asks how conviction is distributed across direction, instrument, and horizon, and how that distribution is changing. A horizon-bound future-state valuation is one possible downstream use; it is not a requirement of the theory or of the Index.

2.3 The three-layer separation

The research program separates three objects that must not be conflated. **The Index is the first layer**; the other two exist to validate it:

Table 2. The three-layer separation (standing invariant of the research program).

Layer	Symbol	Role
Index lattice	F	The Index itself: published cells, composites, Omega fields, freshness, and provenance. Gated by forward IC and harvest.
Minimal utilization	M_0	The validation benchmark (“A0”): a minimal, fully transparent rule proving lattice information converts into aligned out-of-sample trades.
Reference utilization	M	Entry, sizing, and exit seams tested against M_0 to locate where extractable information lives.

The standing invariant behind this separation: trader profit-and-loss alone is an insufficient statistic for signal quality — a good index can be consumed badly, and a bad one can look profitable under overfit gates. By publishing the index layer separately and validating it with the most transparent possible rule, the program keeps every claim attributable.

2.4 Plural sources of authority

Current product facts are reconciled from four authority classes:

1. the deployed GM! runtime contract and live discovery/health surfaces;
2. the DFY GM! Oracle and its Virtuals agent registry;
3. the versioned DFY service export and operational runbooks; and
4. frozen registrations and published statistical records.

The Virtuals agent is one source of authority and is authoritative over deprecated name-associated material for current public scope. It does not override a registered result or a contradictory live response. Historical documents remain evidence for their stated windows, but deprecated names, hosts, and route aliases are not carried into this prospectus.

2.5 Research philosophy: an evidence-bound iterative process

The Index is not the output of a single modeling campaign. It is the current fixed point of an iterative, evidence-bound research process that operates in the spirit of Bayesian updating: each publication in the v1–v8 lineage introduces predeclared hypotheses with explicit verdict rules, subjects them to out-of-sample proof or refutation, and carries the surviving posterior — and only the surviving posterior — into the next cycle. Concretely:

1. **Philosophical foundation.** The program begins from the standing invariant that trader PnL is an insufficient statistic for signal quality, which forces the three-layer separation of Table 2 and makes every subsequent claim attributable to a specific layer.
2. **Hypothesis introduction.** Each candidate mechanism enters as a named, predeclared hypothesis (H0–H9, CP) with its verdict rule fixed *before* the evaluation windows are run (Section 10).
3. **Proof or refutation.** Verdicts are published in both directions. The archive retains confirmed effects (v5 bounded exits, $n=481$, $p<0.001$), refuted claims (early RL superiority, $p=0.24$, withdrawn), and rejected candidates (UTIL250, gates G2/G4), together with the non-significant results of the current record.
4. **Posterior revision.** Beliefs are updated when evidence demands it: the June-only holdout was demoted to co-equal status with the rolling ladder once its variance was quantified, and the interim rolling run was

superseded and publicly reconciled rather than silently replaced (Section 13).

The Index published today is therefore the surviving posterior of that process: the lattice definition and public contract supported by the evidence available at each promotion decision.

2.6 What distinguishes this Index

1. **It is live.** The Index is calculated and distributed continuously today over production market data (Section 7); it is not a research proposal.
2. **Falsification-first governance.** Changes to the Index profile require passing predeclared statistical gates; a documented optimization candidate (UTIL250) was rejected by exactly this machinery (Section 11).
3. **Two production cadences on one origin.** The primary 1h core-6 surface and 5m crypto-3 fine tier share a public contract while retaining per-cell feed provenance.
4. **Multi-asset breadth.** The lattice spans crypto-native contracts (BTC, ETH, HYPE) and synthetic macro perpetuals (gold, silver, WTI crude) on a single margin framework.
5. **Public auditability.** Every claim cited here is reproducible from named artifacts, down to the parameter hash of the profile and the timestamped backtest summaries (Appendix C), with a live parity gate proving that the distributed values equal the reference calculation.

2.7 The Index as a standalone information product

Although this prospectus is written as an index methodology and record, the GM! INDEX is also a commercial subject in its own right, and can be evaluated as one without reference to anything else the Sponsor operates:

- **A new market-data category.** A normalized, cross-instrument reading of directional *conviction and confidence* is a distinct data class — positioned beside price, volatility, funding, and open interest — for which there is no incumbent standard. The Index defines and operates that category at production quality today.
- **A machine-native data suite.** Atomic reads, lattice composites, sizing, exit hazard, and Omega books are available through HTTP, x402, MCP, Virtuals, and provisioned credentials. The commercial surface is not a human-dashboard retrofit.
- **Trust that is verifiable, not asserted.** A prospective customer or investor can confirm the Index's quality independently — measured efficacy (Section 8), paid-versus-reference parity (Section 7), and per-reading data freshness (Section 7.2) are all observable on the live surface at gm.mutantdefi.com.

GM! Index is the flagship of the Sponsor's MUTANTDEFI information-product line at gm.mutantdefi.com; the company-level investment case is set out in the separate *ConfigSecretAI Investor Prospectus*. This document keeps its focus on the Index itself, and nothing in the Index's evaluation depends on anything beyond it.

3 Index Sponsor, Administration and Governance

3.1 Sponsor and calculation agent

The Index is sponsored, calculated, and administered by ConfigSecretAI, Inc. through its MUTANTDEFI Research Lab (together, the "Sponsor"). The Sponsor operates the calculation service, the distribution service, the research archive, and the promotion machinery described below. There is no external index committee;

governance is codified as executable gates rather than discretionary review (see Section 17 for the associated conflicts discussion).

3.2 Promotion gates

Any change to the Index profile, its parameters, or the validation benchmark must clear four predeclared gates, published in the research archive and applied mechanically:

Table 3. Standing promotion gates governing any change to the Index or its benchmark.

Gate	Name	Criterion
G1	Primary holdout	Benchmark profit % > 0; trades ≥ 15 ; win rate $\geq 50\%$; maximum drawdown $\leq 2\%$ on the primary holdout window.
G2	Rolling ladder	Challenger must beat the incumbent benchmark's profit % on at least 2 of 4 rolling monthly windows.
G3	Live reconciliation	Direction of live paper-capture agrees with the offline backtest before any feed deployment.
G4	Profile battery	Any offline profile-optimization winner must be confirmed by the trade-level battery before the production profile hash is swapped.

As of this revision, the production primary profile is GM!-FEED-8.01-1H-CORE6-P01; the production fine profile is GM!-FEED-8.03-5M-CRYPT03-P02. A separate objective build led on offline information coefficient but was not silently substituted for the production profile. A0 remains the admission benchmark. A14b exit hazard and A15b lattice sizing cleared their registered superiority tests as reference utilizations; their combination did not establish an additional interaction effect. UTIL250 remains rejected.

3.3 Versioning and restatement

The Index inherits the research archive's versioning discipline: every published result is tied to a document version, a build date, a profile parameter hash, and a timestamped machine-generated summary file. Data corrections follow the restatement policy in Section 13.

4 Index Definition and Published Values

4.1 The lattice

The Index is a lattice of discrete long–short conviction cells. Each cell is addressed by a triple key:

(instrument, granularity, horizon)

Table 4. Lattice cell addressing.

Key	Meaning	Examples
instrument	One traded underlying or a named basket composite	BTC, XYZ-GOLD, DFY-BLUE
granularity	OHLCV sampling timeframe of the calculation surface	5m, 1h
horizon	Forward evaluation window (wall-clock)	15M, 90M, 4H, 1D

Granularity and horizon are independent: a 1h surface publishes 90M, 4H, and 1D conviction simultaneously, so the Index is genuinely multi-horizon rather than a single number.

4.2 Published channels

Every cell publishes a primary conviction frame:

$\text{value} \in [0, 1]$	long-axis bounded likelihood
$\text{signed_conviction} \in [-1, +1]$	$(\text{value} - 0.5) \times 2$
$\text{confidence} \in [0, 1]$	data-support / reliability channel

The coarse direction label is long, short, or neutral. The production primary mode is:

$$\text{value_mode} = \text{composite}$$

Alongside that primary frame, the cell publishes the full-distribution hurdle-net representation:

$$\begin{aligned} \omega &= \Omega(\theta), & \omega_{\text{mass}} &= G + L, \\ \omega_{\text{value}} &= \frac{G}{G + L}, & \omega_{\text{direction}} &= \frac{G - L}{G + L}, \\ \omega_{\text{opportunity}} &= \tanh(\log(1 + G + L)), & \omega_{\text{delta}} &= \Delta\Omega. \end{aligned}$$

The fields `value_mode`, `omega_hurdle`, and `omega_hurdle_mode` state the active representation and hurdle. These additive fields extend the observable information set without silently changing the primary conviction series.

A reading of `signed_conviction = +0.60` with `confidence = 0.70` on BTC1D/DFY states: strong upward one-day conviction on BTC, on well-supported data. Consumers are free to threshold, size, or combine cells; the discrete reading used by the validation benchmark (long above +0.32, short below −0.32, confidence floor 0.45) is one published, tested interpretation — not part of the Index definition.

Alongside the three Index channels, every cell carries timeliness metadata: `staleness_ms` (wall-clock age of the underlying market data behind the reading, i.e. now minus the last underlying bar timestamp) together with the raw `underlying_last_ts` and `feed_last_ts` epochs, so a consumer can always verify how current a reading is before acting on it. *Restatement note (6 Jul 2026, v2.1.4)*: prior to this revision `staleness_ms` reported internal generation lag (near zero by construction) rather than data age; the field now carries the wall-clock semantics stated here, and the change is recorded as a published-field restatement in Section 13.

4.3 Symbol contract

Individual cells follow a CCXT-shaped symbol contract, making the Index directly consumable by standard trading tooling:

$$\{\text{BASE}\}\{\text{HORIZON}\}/\text{DFY} \quad \text{e.g. BTC1D/DFY, HYPE90M/DFY, XYZ-GOLD4H/DFY}$$

Basket composites aggregate constituent convictions under non-negative weights and enter the lattice as first-class instruments:

$$\text{DFY}-\{\text{NAME}\}\{\text{HORIZON}\} \quad \text{e.g. DFY-BLUE1D} = \text{weighted } \{\text{BTC, ETH, HYPE}\} \text{ at 1D}$$

4.4 Calculation profile

The production Index is dual-profile. The primary core-6 surface is GM!-FEED-8.01-1H-CORE6-P01; the crypto-3 fine tier is GM!-FEED-8.03-5M-CRYPT03-P02. Each response identifies its own `feed_version`,

parameter hashes, underlying timeframe, and freshness. Profiles are selected by trader-agnostic information gates and then held fixed until promotion criteria are met.

5 The GM! Family: Breadth of Conviction Cells

Each lattice cell is one instrument–horizon conviction index publishing its own bounded conviction, confidence, Omega, freshness, and provenance under the contract of Section 4. GM! Index is the family and its cross-sectional distribution, not merely one headline composite.

5.1 Horizon ladder

Table 5. Published horizon ladder. Extended horizons bridge the legacy four-rung ladder where analysis requires finer resolution.

Horizon	Ladder role	Notes
15M	Legacy rung	Shortest forward window; anti-divergence input in multi-horizon utilization (arm A4)
30M	Extended bridge	Fine-resolution bridge rung
90M	Legacy rung	Primary horizon of the 5m crypto surface; alignment input in A4
3H	Extended bridge	Mid-horizon bridge rung; dedicated mid-horizon efficacy reports
4H	Legacy rung	Alignment input in multi-horizon utilization
1D	Legacy rung	Primary horizon of the production core-6 surface; benchmark horizon

Granularity and horizon are independent: a single 1h surface publishes 90M, 3H, 4H, and 1D conviction simultaneously for every instrument it covers. Mid-horizon efficacy reports (90M, 3H) run the identical scoring pipeline with shorter forward windows for granularity-sensitivity analysis.

5.2 Individual cells

Every (base, horizon) pair with profile parameters and adequate history is an individual cell under the CCXT-shaped contract {BASE}-{HORIZON}/DFY. On core-6, the six-rung ladder spans up to $6 \times 6 = 36$ addressable cells. BTC, ETH, and HYPE use the 5m fine profile on the public origin; the remaining core names use the 1h primary profile. Response provenance makes the routing explicit.

5.3 Basket and universe composites

Basket indices compose constituent signed convictions under non-negative weights and enter the lattice as first-class instruments under the contract DFY-{NAME}-{HORIZON} (e.g. DFY-BLUE1D = weighted {BTC, ETH, HYPE} at 1D). Basket cells inherit the same three published channels. Instrument sets may be any subset of the traded universe for which OHLCV and profile parameters exist at the chosen granularity — the basket space is combinatorial, bounded only by the data-coverage validity rule below.

5.4 Production and research families

Table 6. Published and researched profile families across the GM! breadth.

Profile	Gran./hor.	Universe	Status
GM! primary P01	1h / multi-horizon	Core-6	Production ; 1D-heavy primary surface
GM! fine P02	5m / 90M primary	BTC, ETH, HYPE	Production fine tier ; same public origin
Objective successor	1h / 1D	Core-6	Offline IC leader; not substituted into primary
Conviction-confidence	1h / 1D	Core-6	Staged representation; hardened confidence gate passed
Omega primary	1h / 1D	Core-6	Staged primary-mode candidate; additive Omega fields already published
Macro7 preset	1h / 1D	CL, GOLD, SILVER, SP500	Available basket grammar; default remains core-6 and evidence is reported separately

Validity rule. A (universe, granularity, horizon) combination is publishable whenever the granularity supports the horizon's candle count and instrument history covers the evaluation window. The lattice is therefore extensible by construction; what is *production* at any moment is whatever subset has cleared the G1–G4 gates.

5.5 Breadth is consumed, not just published

The validation program exercises the family's cross-horizon structure directly: battery arm A4 consumes 1D + 90M/4H alignment with 15M anti-divergence; arm A5 fuses horizons under a linear weighted score. That neither yet beats the minimal 1D benchmark (Section 10) is a finding about fusion design, not about breadth: the multi-horizon cells are live, published, and independently scored, and their joint utilization is a declared research priority (A4 ablation, A5 weight sweeps). Consumers are not limited to the headline slice — the full ladder is readable today through the same distribution surfaces (Section 7).

6 Index Universe and Eligibility

6.1 Constituent markets

Table 7. Index universe: core-6 HYPERLIQUID 1h perpetual futures.

Contract	Type	Notes
BTC	Crypto-native perpetual	Deepest liquidity in universe
ETH	Crypto-native perpetual	
HYPE	Crypto-native perpetual	Venue-native token
XYZ-GOLD	Synthetic macro perpetual	Tracks spot gold
XYZ-SILVER	Synthetic macro perpetual	Tracks spot silver
XYZ-CL	Synthetic macro perpetual	Tracks WTI crude; listed 6 Jan 2026 14:00 UTC (data floor)

6.2 Eligibility and changes

The universe is fixed by the calculation profile: a market is eligible only if (i) OHLCV history exists at 1h granularity covering the evaluation window, (ii) profile parameters exist for the instrument within P00,

and (iii) per-instrument signal diagnostics do not disqualify it. Universe changes are profile changes and therefore subject to the full G1–G4 promotion machinery. A macro-4 extension profile (adding XYZ-SP500) was evaluated and *failed* its holdout gates at both the index level (harvest $t = -0.69$) and the validation level (-0.63% holdout, 36.8% win) and remains excluded — the eligibility machinery declining to expand the universe.

6.3 Per-instrument signal quality

Per-instrument forward IC on the long 1h/1D window (6 Jan–18 Jun 2026) explains the universe boundary:

Table 8. Per-instrument 1h/1D forward IC, macro legs, 6 Jan–18 Jun 2026.

Instrument	<i>n</i>	forward_ic	ic_high_conf	harvest_bps
XYZ-GOLD	3889	+0.054	+0.200	+13.7
XYZ-SILVER	3889	+0.029	+0.186	+24.2
XYZ-CL	3875	−0.031	−0.148	−9.0
XYZ-SP500 (excluded)	2172	−0.108	+0.095	−9.1

Precious-metal legs carry pooled positive IC; crude is weak on long windows but remains in the pooled core-6 profile, which passes its gate in aggregate; SP500 is excluded.

7 Calculation, Distribution and Data Suite

The Index runs on independently deployed public and restricted services. The GM! service never embeds feed feathers in production; it reads the primary and fine profiles from separate Tier-C upstreams. Desk posture is a sister surface, not part of the GM! conviction path.

7.1 Topology

Table 9. Production topology of the GM! data stack.

Service	Domain	Port	Role
dfy-feeder	<i>internal</i>	8677	Collects and pushes validated market-data bundles
dfy-feed	feed.mutantdefi.com	8655	Tier-C 1h calculation feed (bearer-gated)
dfy-feed-5m	<i>internal</i>	8656	Tier-C 5m crypto-3 fine feed
gm	gm.mutantdefi.com	8666	GM! Index, x402, MCP, discovery, and agent export
desk	desk.mutantdefi.com	8645	Posture, outcomes, journal, and ingest; separate from GM! conviction
dfy-web	apps.mutantdefi.com	managed	Public live showcase and developer reference
virtuals	app.virtuals.io/virtuals/14925	managed	DFY GM! Oracle — ACP checkout, chat, and provisioned-access front

The public service routes BTC, ETH, and HYPE to the 5m fine upstream and other primary names to the 1h upstream. Both profiles appear on the same origin. Embedded calculation exists only for tests and local development.

7.2 Data freshness and liveness

A real-time index is only as good as the currency of its inputs, so freshness is measured and published independently at every stage of the chain rather than assumed:

- **Producer.** dfy-feeder runs a continuous hourly cycle — collect the core-6 underlying futures data, validate it, push it to the calculation agent — and reports on its own health surface whether the last push succeeded recently.
- **Calculation agent.** dfy-feed publishes a freshness block: the wall-clock age of the newest underlying bar, judged against a three-hour liveness threshold keyed to a 24/7 reference instrument (BTC), so weekend gaps in market-hours instruments do not read as pipeline failure. When the producer stalls, the service reports stale instead of silently serving old values.
- **GM! distribution.** gm relays the upstream freshness verdict on its own health surface (feed_upstream), and every published cell carries staleness_ms (Section 4) so the data age is visible per reading, not just per service.

Each layer is independently observable; a break anywhere in producer → calculation → distribution is visible at the exact stage it occurs, and pushed data is validated in staging and atomically swapped, so a bad push cannot take down the live feed.

7.3 Distribution surfaces

Table 10. Canonical public GM! routes and discovery resources.

Route	Access	Description
/, /.well-known/x402, /llm.txt	free	Service discovery for human and agent consumers
/services.json, /mcp/tools	free	Virtuals/Hermes service export and Model Context Protocol registry
/mcp/execute	x402	Model Context Protocol interface for AI-agent consumers
/v1/gm/conviction, /v1/gm/lattice	x402	Atomic and cross-sectional conviction reads
/v1/gm/weights, /v1/gm/omega_book, /v1/gm/hazard	x402	Validated sizing, mean–Omega book, and exit-hazard reference tools
/v1/gm/harvest_rung	x402	Rung coverage and distance-to-band under a selected harvest rung (FC-27); ?rung= also annotates conviction, lattice, weights, and hazard reads
/v1/gm/ticker, /v1/gm/ohlcv	x402	CCXT-shaped integration adapters
/v1/feed/metadata	free	Profile registry + validation telemetry

Index values are sold per read via the x402 micropayment protocol — data distribution in the pattern of a market-data vendor, not a managed product. The MCP surface makes the Index natively consumable by autonomous agents.

7.4 Comprehensive GM! toolset

Table 11. Current GM! Index and companion data capabilities.

Capability	Surface	Purpose
Conviction snapshot	HTTP / x402 / MCP / Discord	One cell with primary conviction, confidence, direction, Omega, hurdle, provenance, and freshness
Lattice composite	HTTP / x402 / MCP / Discord	Universe–horizon distribution, cell list, composite, and lattice marginals
Allocation weights	HTTP / x402 / MCP / Discord	A15b relative stake multipliers and self-financing diagnostics
Omega book	HTTP / x402 / MCP / Discord	Mean–Omega long–short or long-only research book; distinct from stake multipliers
Exit hazard	HTTP / x402 / MCP / Discord	A14b “when to stop” read for a caller-supplied open position
Harvest rung	HTTP / x402 / MCP	FC-27 rung selection (purist/balanced/active) with coverage and distance-to-band; MCP tool <code>get_gm_harvest_rung</code>
CCXT adapters	HTTP / MCP	Ticker and OHLCV-shaped compatibility for quantitative systems
Agent discovery	free HTTP	x402 manifest, <code>llm.txt</code> , MCP tools, and <code>services.json</code>
Provisioned access	Virtuals agent	Day pass, full suite, and live-pro credentials under fixed-fee terms
Desk suite	separate origin	Regime, posture, realized outcomes, attribution, guidance, and journal; no GM! cell values

The live `services.json` export and the DFY GM! Oracle on Virtuals are current machine-readable authorities for feature availability. The ACP v2 planning catalog is supporting documentation, not by itself the completeness boundary. The public contract refuses percentage fees, fund-transfer jobs, custody, and managed-capital framing.

7.5 Harvest rungs: selectable admission ladders (FC-27)

The Index publishes one frozen conviction lattice. Consumers differ, however, in how much *activity* they want from it: a research desk may prefer only the most concentrated admissions, while an agent paying per read may prefer more frequent actionable states. The **harvest ladder** (registration PR-2026-07-19-HARVEST-LADDER-01) resolves this without touching the Index: three registered admission rungs — *purist*, *balanced*, and *active* — each a frozen (band, `min_confidence`) pair applied *on top of* the same published values. Rung selection is a consumption choice, in the same sense that a chart consumer chooses a timeframe; it never alters calculation, profile parameters, or the published value channels.

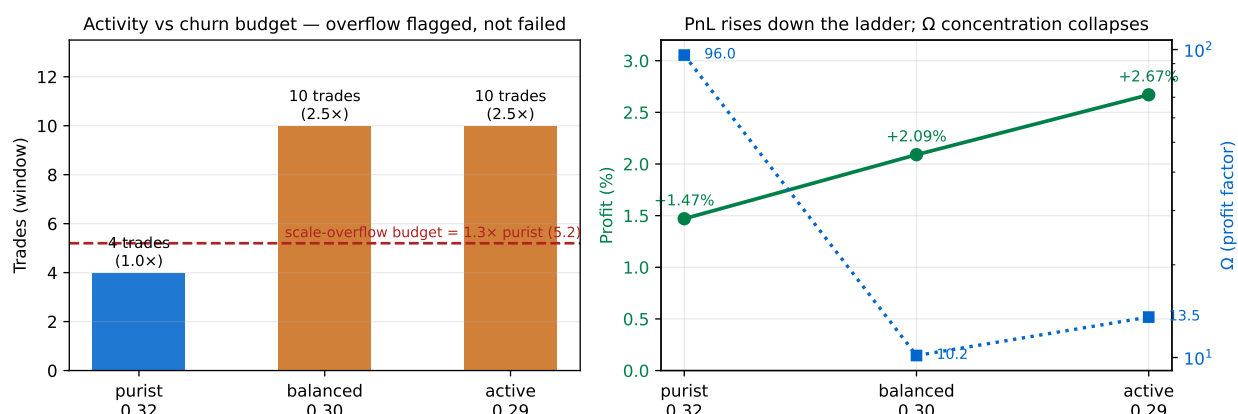
Product surface. Paid reads accept `?rung=purist|balanced|active`; responses are annotated with the selected rung so purity is always labeled. A dedicated endpoint, `/v1/gm/harvest_rung` (MCP: `get_gm_harvest_rung`), returns rung coverage and distance-to-band for a cell — “how far is this instrument from an actionable state under my chosen rung.” The service default when a choice is required is *balanced*.

Registered rung evidence. The first full quantitative read of the ladder (frozen 1h core-6 F-layer, out-of-sample window 19 June–18 July 2026) is shown in Table 12 and Figure 1. The shape is a *trade-off frontier*, not a ranking: stepping down the ladder raises activity and window profit while concentration (per-trade profit factor) falls by roughly an order of magnitude.

Table 12. Harvest ladder rungs on the frozen 1h core-6 surface, window 2026-06-19 to 2026-07-18 (benchmark-style utilization read; paper mode).

Rung	Band	Trades	Profit %	Omega (PF)	Churn vs. purist
purist	0.32	4	+1.47	95.95	1.00×
balanced	0.30	10	+2.09	10.15	2.50× (overflow flagged)
active	0.29	10	+2.67	13.53	2.50×

FC-27 harvest ladder, frozen 1h F-layer, window 20260619–20260718

**Figure 1.** The harvest ladder as a trade-off frontier (same window as Table 12). Left: activity against the registered 1.3× churn budget — lower rungs exceed it at 2.5× purist, disclosed as *scale overflow* (a budget flag, not a performance failure). Right: window profit rises down the ladder while Omega concentration collapses. Rung choice prices activity against concentration; it does not change the Index.

Governance boundary. Two disclosures keep the ladder honest. First, *scale overflow*: when an operating rung's trade count exceeds 1.3× the next-higher rung (as balanced does above), that is disclosed as an activity-budget breach — a lesson institutionalized from the UTIL250 rejection (Section 11) — rather than laundered into a performance claim. Second, rungs are *not* promotion candidates: the A0 benchmark on the purist rung remains the sole promotion clock, and choosing balanced or active for activity is an operating-point selection on the same feed, never a governance event. Rung parameters change only by registered amendment to FC-27.

A standing parity harness (`services/fci/scripts/validate_fci_parity.py`) verifies that the x402-served values equal the free reference value-frame for the same cell — exit code 0 is the Phase-1 release gate. Consumers therefore never depend on trust that the paid surface matches the calculation: the equality is continuously testable.

7.6 The projection boundary (construction protection)

The platform enforces a three-tier disclosure boundary in code. Tier A (paid public posture), Tier B (journal), and Tier C (raw feed and strategy source) are separated by an allowlist projection with a leak auditor: any projected output containing strategy source, thresholds, or raw feature values raises `LeakError`, and the boundary test suite is a release gate. The Index's published channels are exactly the tier-boundary-safe surface;

its construction internals are structurally unable to leak through the public routes. This is the enforcement mechanism behind the disclosure boundary of Section 17.2.

8 Index Performance: Signal Efficacy

For an information index, “performance” means measurement quality: does the published value anticipate what it claims to anticipate? For the VIX that question is correspondence with realized volatility; for the GM! it is covariance of signed conviction with realized forward returns over the cell horizon. This section reports those statistics — they are the Index’s primary record.

8.1 Registered objective-profile evidence

The objective profile used for the registered efficacy study passed its trader-agnostic gate on the build holdout (30 Apr–18 Jun 2026):

Table 13. Registered objective-profile gate results (index level, trader-agnostic; not the current production identifier).

Metric	Train	Holdout
Forward information coefficient	+0.021	+0.032
IC, high-confidence quartile	+0.092	+0.004
Harvest (bps)	+16.7	+14.5
Gate verdict	pass	pass

Holdout IC *exceeds* train IC—the profile was not tuned into its evaluation window. The high-confidence lift degrades on holdout, a known calibration-drift pattern on short windows that is disclosed, monitored, and one reason the validation benchmark applies confidence as a soft threshold rather than a hard gate. These statistics support the measurement family; they do not imply that the historical profile label is the current production profile.

8.2 Efficacy metrics defined

Table 14. Public efficacy metrics computed per cell and per profile.

Metric	Definition
forward_ic	Spearman rank correlation of signed conviction with realized forward return over the horizon
ic_high_confidence	IC restricted to top-quartile confidence observations
conf_lift	$\text{ic_high_confidence} - \text{forward_ic}$
harvest_bps	Mean signed conviction $\times$ forward return (bps) over bars with $ \text{signed} \geq 0.05$
coverage	Fraction of bars with an active signed opinion

These metrics are recomputed on every profile build and published with bin diagnostics (IC within confidence quantiles; confidence within price quantiles), so consumers can evaluate cell-level measurement quality directly rather than trusting pooled numbers.

8.3 Interpretation

A pooled holdout IC of +0.032 on hourly cells is a persistent, positive, out-of-sample rank correlation between the Index and subsequent market direction, sustained across six markets and thousands of observations, with positive harvest confirming that the correlation lives where the Index actually voices an opinion. Per-instrument heterogeneity (Table 8) is disclosed rather than averaged away, and drives universe governance.

9 The Utilization Validation Benchmark

9.1 Why a benchmark exists

Correlation statistics alone leave a gap: a consumer’s practical question is whether Index readings, taken at face value under transparent rules, align with tradable outcomes. The Sponsor therefore maintains the **GM! Minimal Utilization Benchmark** (“A0”): the simplest defensible consumption of the Index — long when 1D signed conviction $\geq +0.32$ with confidence ≥ 0.45 , short on the mirror condition, one-horizon hold, equal notional, no other logic.

The benchmark is a validation instrument, not the Index and not a product. It runs in paper mode by design (Section 14); its parameters were selected once by an 80-epoch search on the development window — which notably rejected a higher-raw-profit configuration (+18.3%, 505 trades) in favor of selectivity (+14.22%, 158 trades), a choice subsequently vindicated out of sample — and are frozen pending the G1–G4 gates.

9.2 Development and holdout record (legacy profile study)

The parameter-selection study below was run on the legacy battery profile and is retained for audit of the 80-epoch selection choice; the canonical out-of-sample record is the `cf i -val` lineage of Table 16 and the powered extension of Table 17, whose June figure (+1.82%) supersedes the legacy study’s holdout row.

Table 15. Legacy profile-selection study (battery profile; superseded — retained for audit).

Statistic	Train (6 Jan–31 May)	Holdout (1–18 Jun)
Total return	+14.22%	+0.85%
Trades	158	23
Win rate	51.3%	65.2% (15 of 23)
Maximum drawdown	4.03%	0.83%

The program itself flags the 18-day holdout as informative but insufficient — each trade moves the window return by roughly 3.7 percentage points — which motivated elevating the rolling ladder to co-equal promotion authority.

9.3 Rolling four-window record

The canonical rolling record is the unified re-run consolidated in the *Strategy Lineage Standing Reference* (final, built 4 July 2026, source `summary_20260704T182745Z.md`):

Table 16. Benchmark rolling four-window out-of-sample record (unified run, final).

Window	Return	Trades	Benchmark level
Base (1 Mar 2026)	—	—	1,000.00
1–31 Mar 2026	+3.88%	41	1,038.80
1–30 Apr 2026	−2.18%	31	1,016.15
1–31 May 2026	+4.65%	22	1,063.41
1–18 Jun 2026	+1.82%	23	1,082.76
Cumulative (110 days)	+8.28%	117	

9.4 Powered extension window (current lineage)

The registered exit- and sizing-seam programs (ladder stamps 20260709T223801Z, 20260709T225641Z) evaluated the same `cfi-val` benchmark on a longer, powered holdout, 1 May–8 Jul 2026. **What “powered” means here.** The program’s power ledger, fixed at registration, sets $n \geq 23$ trades as the minimum sample at which a verdict may be issued at all — below it, outcomes are labeled INCONCLUSIVE regardless of how favorable they look. The original 18-day June slice repeatedly under-shot that floor; the extension window was declared in a dated amendment, *before execution*, precisely to cross it. That ordering matters for the reader: because the calendar and the rules were frozen first, the statistics below are *confirmatory* tests of a pre-stated claim, not favorable windows selected after the fact. This window *overlaps* the May–June months of Table 16 and extends 20 days past it; it is reported as a separate record, not added to the ladder. All rows share identical admissions — the benchmark rule unchanged — and differ only in the reference seam under test:

Table 17. Powered extension window, 1 May–8 Jul 2026 (identical benchmark admissions; pre-registered reference seams).

Arm	Return	Trades	Win rate	Max DD	Registered verdict
Benchmark (A0, horizon exit)	+11.64%	27	81.5%	0.17%	powered reference
+ lattice-marginal sizing	+14.78%	27	81.5%	0.25%	SUPERIORITY
+ hazard exit (vs. tail control +9.73%)	+11.47%	30	70.0%	0.28%	SUPERIORITY

At this sample the benchmark’s win record is individually significant: 22 of 27 (one-sided binomial against $p = 0.5$: $p \approx 8 \times 10^{-4}$). Two cautions accompany the table. First, all three rows share one favorable regime — the program’s own trajectory record treats them as a single regime’s evidence. Second, the sizing arm passed on the profit branch of its endpoint, not the drawdown branch; concentration was paid in this regime, and the composition study (registered NO-INTERACTION) found the hazard layer adds nothing on top of sizing in calm conditions — protection is bought for adverse regimes the window does not contain.

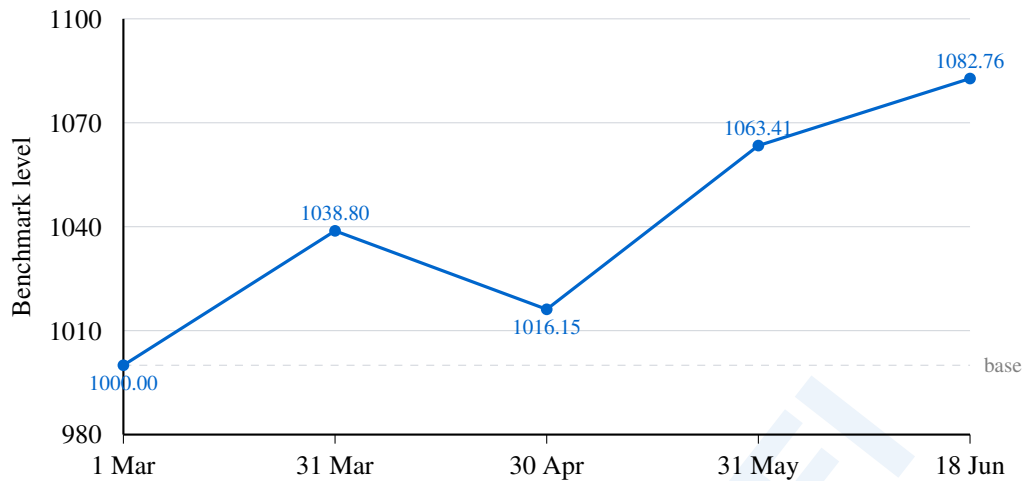


Figure 2. GM! Minimal Utilization Benchmark level, window-close observations, base 1,000.00 at 1 March 2026. Validation instrument in paper mode; see Section 14.

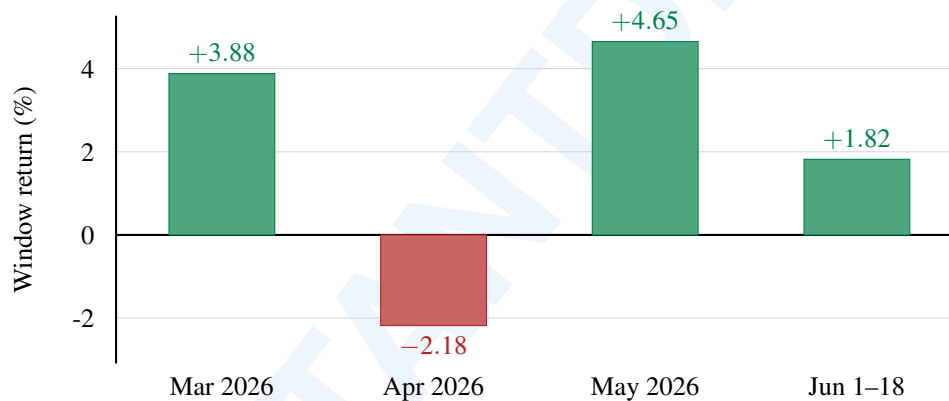


Figure 3. Benchmark window returns across the rolling ladder (unified run). April 2026 is the documented stress window.

9.5 Summary statistics

Table 18. Benchmark summary statistics (four windows; June partial).

Statistic	Value
Cumulative return (1 Mar–18 Jun 2026)	+8.28%
Annualized return (geometric, 110 days)	+30.2%
Mean window return	+2.04%
Window return standard deviation (sample)	3.06%
Annualized volatility of window returns ($\times \sqrt{12}$)	10.6%
Positive windows	3 of 4
Worst window	April 2026, -2.18%
Best window	May 2026, +4.65%

Four windows cannot characterize a return distribution; the annualized figures are reported for convention and carry limited predictive standing on their own. Their function here is validation — three of four independent out-of-sample windows align positively with the Index, consistent with the IC evidence of Section 8.

9.6 Statistical significance and sample power

The Sponsor publishes significance tests alongside point estimates, including the unfavorable ones:

- **Holdout win rate** 15/23 (65.2%): one-sided binomial test against $p = 0.5$ gives $p = 0.1050$ — *not significant* at the conventional $\alpha = 0.05$ threshold.
- **Holdout per-trade mean return** +0.367%: one-sample t -test gives $p = 0.0963$ — likewise not significant at $\alpha = 0.05$.
- Several *historical* claims in the lineage are strongly significant (Section 12), including the v5 bounded-exit treatment effect ($n = 481$, $p < 0.001$).
- The **powered extension window** (1 May–8 Jul 2026, Table 17) is significant at trade level: 22 of 27 wins, one-sided binomial $p \approx 8 \times 10^{-4}$; the sizing and exit reference seams on the same window carry pre-registered SUPERIORITY verdicts at $n \geq 23$. The under-power caution of this section is therefore confined to the 18-day June slice, not the benchmark record as a whole — while the single-regime caveat of Section 9.4 still applies.

What the powered-window values mean for the reader. The same yardstick applied to the June slice above applies here, and now cuts the other way. A p -value of 8×10^{-4} states that if the benchmark had *no* directional edge (a true 50% win probability), a record at least as favorable as 22 wins in 27 trades would arise by chance roughly *one time in thirteen hundred* — against one time in ten for the June slice. Chance ceases to be an admissible explanation at conventional thresholds. Two properties license taking this number at face value rather than discounting it as selection: the window and decision rules were frozen in a dated amendment before the run (a confirmatory test, not a window found by searching), and the sample exceeds the pre-registered power floor, so the verdict rules applied mechanically. What the p -value does *not* establish is generality across market regimes: it measures the improbability of luck *within* this window, and the single-regime disclosure of Section 9.4 governs how far beyond it the reader should extrapolate. Statistical significance and economic breadth are separate claims; this record now has the first and, by design, says so plainly about the second.

What these values mean for the reader. A p -value of 0.1050 states that if the benchmark had *no* directional edge (a true 50% win probability), a record at least as favorable as 15 wins in 23 trades would still arise by chance roughly one time in ten. The null hypothesis of no edge therefore *cannot be rejected* on the trade-level record alone: chance remains an admissible explanation of the benchmark's holdout profitability, and readers must not treat the benchmark figures in this prospectus as statistically established.

Why the tests are underpowered. At $n = 23$, the one-sided test requires 16 or more wins to reach significance ($\alpha_{\text{actual}} = 0.047$); the observed 15 falls one win short. More importantly, even if the benchmark's true win probability were exactly the observed 65.2%, a sample of 23 trades would clear that bar only about 42% of the time (statistical power ≈ 0.42). Approximately 69 trades at the same win rate would be required for 80% power. The test's failure to reach significance is therefore expected behavior at this sample size whether or not the edge is real — it is weak evidence in *both* directions, not evidence of absence.

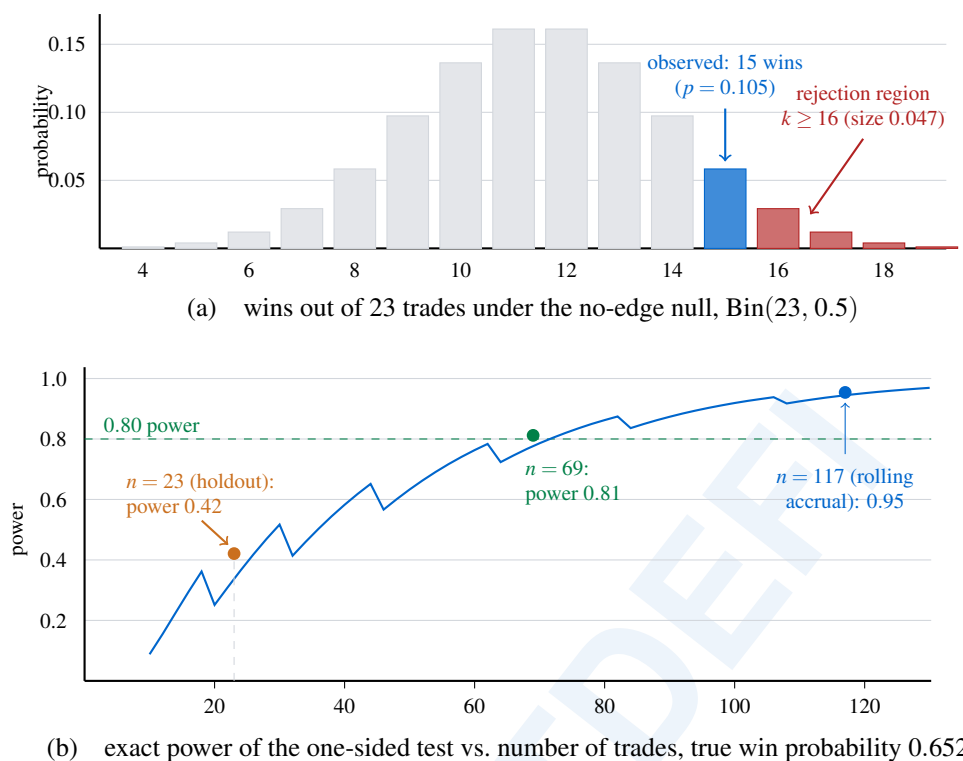


Figure 4. Why the trade-level tests are underpowered. (a) Under the no-edge null, the observed 15 wins falls one short of the 16-win rejection threshold; the record is favorable but chance is not excluded. (b) Exact power of the same test if the true win probability equals the observed 65.2%: at $n = 23$ the test detects the edge only 42% of the time; roughly 69 trades reach the conventional 0.80 standard, and the 117 trades accrued across the rolling ladder correspond to 0.95 power once evaluated jointly. The sawtooth reflects the discreteness of the binomial critical value.

The correct weighting of the evidence. The statistically load-bearing leg of the Index’s record is the index-level measurement: positive holdout forward IC and harvest computed across thousands of cell observations per instrument (Section 8). The trade-level benchmark, which by design trades rarely, is consistent with that evidence and accrues sample — and hence power — with every rolling window (117 trades across the four-window ladder to date). Until the accumulated trade-level record reaches conventional significance, the benchmark should be read as corroborating measurement, not as independent proof.

9.7 Live capture

A live capture trio runs the benchmark and two overlay arms against the production Index in paper mode. As of 4 July 2026 the live sample is six trades (−1.36%) against +0.85% offline on the overlapping window — far below any inferential threshold, and exactly why gate G3 exists: the production profile hash does not move until live capture direction reconciles with offline results. The capture trio and its automated watcher run continuously; reconciliation status is telemetry, published at `/v1/feed/metadata`.

10 The Validation Program: A Ten-Arm Utilization Battery

The benchmark’s authority is earned, not assumed. The initial nine richer utilization designs (A1–A9) were tested on identical universes, windows, and data under predeclared verdict rules; none cleared promotion.

Later seam-isolated studies changed the conclusion more precisely: learned entry overlays remain unproven, while registered sizing and exit mechanisms can add value without redefining the Index.

10.1 Arms under test

Table 19. Utilization battery arms. The validation benchmark is control arm A0.

Arm	Class	Mechanism
A0	Control (benchmark)	1D signed conviction + confidence threshold
A1	v8.04 overlay	Execution overlay, hard gate off
A2	v8.04 overlay	Hard confidence gate on remote profile
A3	v8.04 overlay	Hard confidence gate on objective profile
A4	Hand multi-horizon	1D + 90M/4H alignment + 15M anti-divergence
A5	Hand multi-horizon	Linear weighted score $S = \sum_h w_h s_h$
A6	Learned (FREQAI)	Meta-label admission over A1 surface
A7	Learned (FREQAI)	Confidence-based position sizing
A8	Learned (FREQAI, v8.07)	Store-aware sizing over A0 surface
A9	Learned (FREQAI, v8.07)	Store-aware meta admission over A0 surface

10.2 Primary holdout ladder

Table 20. Full ladder, train and primary holdout (unified run, 4 Jul 2026).

Arm	Train %	Train trades	Holdout %	Holdout trades	Holdout win %
A0	+14.22	158	+0.85	23	65.2
A5	+5.98	190	+0.13	26	46.2
A1	+1.28	55	+0.35	8	37.5
A2	+1.96	44	+0.02	3	33.3
A3	+2.41	54	+0.01	4	25.0
A9	−1.06	77	+0.00	16	37.5
A8	−0.07	138	−0.22	27	40.7
A7	+0.22	133	−0.30	21	22.7
A6	+2.29	71	−0.90	7	14.3
A4	+6.46	112	−1.60	17	35.3

10.3 Rolling ladder across arms

Table 21. Rolling ladder by arm (unified run; June overlay fold pending at archive close).

Arm	Mar %	Apr %	May %	Jun %	Wins vs. A0
A0	+3.88	−2.18	+4.65	+1.82	—
A1	+2.06	+1.12	+0.13	—	1/3
A2	+2.04	+0.94	+0.33	—	1/3
A3	+2.13	+0.45	+0.31	—	1/3
A4	+2.05	−1.59	+3.18	—	1/3
A5	+1.61	−0.52	+2.57	—	1/3
A6	+0.57	−0.09	−1.83	—	1/3
A7	+0.65	−0.44	−0.45	—	1/3
A8	+1.20	−0.91	−0.23	—	1/3
A9	+0.90	−1.09	−0.42	—	1/3

No arm reaches the G2 requirement of 2-of-4 wins. The April result is disclosed prominently: in the one negative window, overlays A1 (+1.12%) and A2 (+0.94%) beat the benchmark with Mann–Whitney $p < 0.001$ on per-trade returns. The standing decision is to weight April equally with June in all future promotion decisions — the program treats stress-window behavior as an active optimization frontier.

10.4 Interpretation: the naive readout is the demonstrated optimum of the admission layer

The battery’s finding is easy to under-read as “the baseline held up.” The evidence supports a materially stronger statement, and this subsection states it with its exact scope.

What was thrown at the benchmark. Across the battery and the subsequent pre-registered program, the two-threshold readout faced: hard confidence gates at neighboring thresholds; multi-horizon temporal consensus and soft weighted fusion; learned admission (meta-labeling) and learned sizing over the same surface, with and without an upgraded feature store; reinforcement-learned policies over raw multi-asset state, over Index state, and over Omega-transformed Index state; a distributionally motivated replacement of the learning objective at two hurdle levels; a hybrid objective; a pooled cross-pair learner; and the same learner supplied with maximal training history. Every arm ran under frozen windows, seeds, endpoints, and prohibited-move rules. None beat the naive readout at admission.

Why “the challengers were weak” is not an available objection. The program closed that escape with a controlled experiment. Scaling training data proved the pooled learner class genuinely capable on raw market state (holdout $-0.29\% \rightarrow +3.12\%$, a registered SCALE verdict). Handing that same, demonstrably capable learner the Index’s features then *degraded* it to -3.54% with $1.8\times$ the churn. The failure to out-admit the naive readout is therefore a property of the information, not of the optimizers: what the Index knows about *which trades to admit*, two thresholds already extract.

The mechanism is measured, not conjectured. The confidence channel triples subset information coefficient, yet applying it as an additional admission filter degrades returns monotonically — at the admission margin the auxiliary channel is redundant with the value channel. Residual ordering information demonstrably exists *inside* the admitted set, but it is a sizing resource, not an admission resource (Section 10, sizing row).

The proper interpretation. On current evidence the naive utilization is not merely acceptable or sufficient; it is *optimal in the demonstrated sense*: resilient against every registered attempt at optimization, across overlay, learned, and objective-engineered extractor classes, under pre-registration discipline that prevented

favorable selection. This is an empirical optimality claim — bounded by the classes tested and falsifiable by the standing hooks (any powered admission extractor beating the gate; a materially positive residual conditional IC) — not a mathematical proof. It has survived every scheduled opportunity to fail.

Why this matters to a reader of an index prospectus. The finding inverts the usual burden of sophistication. Ordinarily a naive consumption of a signal is the floor, and realized value depends on proprietary machinery layered above it — machinery the consumer must build, tune, and trust. Here the naive consumption is the demonstrated ceiling of the admission layer: the Index’s published channels deliver their full entry-decision value at the simplest possible reading, so the sophistication burden on the consumer is two thresholds. That is a statement about the *Index* — its information is complete at its interface — and it is why the Sponsor’s reference tools improve *where to size* and *when to stop* (Section 10.6) rather than promising better entries: the battery did not establish that every downstream transformation must fail, and the registered record shows the ones that succeed operate below the admission layer, on trades the naive readout selects.

10.5 Alpha: separating the benchmark’s return from market return

The returns of Sections 9–10.4 invite an obvious objection: 2026 crypto markets moved a great deal — how much of the benchmark’s record is skill, and how much is simply exposure to markets that went up? The CAPM decomposition answers this by splitting any return into the part explained by market exposure ($\beta \times$ market return) and the residual, α — the part market exposure *cannot* explain. A long–short rule that admits longs and shorts by the same mirrored condition should carry $\beta \approx 0$ by construction; the table below tests whether it does so in fact.

The market leg is the equal-weight buy-and-hold of the same six instruments, computed close-to-close from the registered OHLCV data (risk-free rate taken as zero over these horizons; funding excluded on both legs):

Table 22. Benchmark versus equal-weight core-6 buy-and-hold, per window. Separation is benchmark minus market.

Window	Market (EW core-6)	Benchmark	Separation
1–31 Mar 2026	+6.53%	+3.88%	–2.65 pp
1–30 Apr 2026	+4.66%	–2.18%	–6.84 pp
1–31 May 2026	+8.39%	+4.65%	–3.74 pp
1–18 Jun 2026	–12.30%	+1.82%	+14.12 pp
1 May–8 Jul 2026 (powered)	–5.79%	+11.64%	+17.43 pp

Regressing the four rolling windows on the market leg gives $\hat{\beta} = +0.06$ with correlation $+0.19$ and estimated alpha of $+1.9\%$ per window — the benchmark is market-neutral in *measurement*, not merely by design. At four observations the regression is descriptive, not inferential, and is offered as such; the individual windows carry the interpretable evidence:

- **In up-months the benchmark trails the market** (Mar, May) — as a two-sided book must: it holds shorts against the rally and declines to harvest beta. May is the sharpest illustration: the equal-weight market leg gained 8.39% largely on a single instrument’s $+80\%$ melt-up, exposure the benchmark deliberately does not take. April’s loss in a rising market is likewise idiosyncratic, not directional — disclosed in Table 16 and weighted equally in promotion decisions.
- **In down-markets the separation is the product.** Over 1–18 June the six markets lost 12.30% held passively; the benchmark made $+1.82\%$. Over the powered window the market leg lost 5.79% while the

benchmark returned +11.64% — a +17.4 pp separation, unlevered, at equal notional. With $\hat{\beta} = 0.06$, market exposure “explains” roughly −0.4 points of that window’s return; effectively the entire +11.64% is alpha in the CAPM sense.

- **This is what the Index’s information coefficient looks like when converted.** A signed conviction signal harvested symmetrically produces return with little market footprint — which is why the Index-level statistics of Section 8 (IC, harvest) and not market direction are the record’s primary explanatory variables.

The claim, precisely scoped: over the evaluated windows the benchmark’s return is statistically and economically separable from market return — it is not repackaged beta. The regression sample is four windows; the per-window separations above, particularly the two adverse-market windows, are the durable exhibits, and all inputs (registered OHLCV, ladder rows) are reproducible from the artifact manifest. For reference: since April 2026 is the sole negative-alpha window ($r_s - \hat{\beta} r_m = -2.47$ points), the at-rest yield on the book’s neutral collateral (USDC) would have to run about +2.5% per 30-day window — roughly 30% annualized, versus the $\approx 0.36\%$ per window that a representative 4.5% stablecoin rate provides — before every window in Table 22 cleared its CAPM-required return, so the record is presented as what it is: positive alpha in four of five windows, unaided by cash-yield assumptions.

10.6 Subsequent registered sizing and exit results

The later program isolated two seams and evaluated them under frozen holdouts. A14b (exit hazard) increased the primary holdout from +1.76563% to +3.25024% (+1.48461 percentage points) and passed its registered superiority test. A15b (relative lattice sizing) increased that holdout from +1.76563% to +1.834319% and reported a one-sided paired-bootstrap $p < 0.001$ with positive confidence interval for incremental profit. The combined A16 study did not establish an additional interaction effect ($p = 0.2137$).

These results define the current public toolset: A15b is exposed as relative stake multipliers; A14b as exit-hazard guidance; and their combination remains informational rather than an asserted synergistic strategy. The distinction between stake multipliers and mean–Omega portfolio weights is maintained in every surface.

11 Governance Case Study: Rejection of the UTIL250 Candidate

The optimization stance is aggressive by policy: profile-generation campaigns run continuously against the incumbent. In July 2026 one such campaign produced UTIL250, a 250-iteration profile candidate (build `utilization_a0`, parameter hash `a91e77d5725fadbc`) that outscored the incumbent on the offline harvest objective in 2 of 4 rolling windows and expanded actionable coverage six-fold. Under gate G4 it was then run through the trade-level battery:

Table 23. UTIL250 vs. incumbent P00, trade-level battery (G4).

Window	P00 %	UTIL250 %	UTIL250 trades	Winner
Train (Jan–May)	+14.22	−2.44	510	P00
Mar 2026	+3.34	+2.35	119	P00
Apr 2026	−2.86	+5.59	98	UTIL250
May 2026	+3.87	+0.28	103	P00
Jun 1–18	+0.85	−3.03	60	P00

UTIL250 won 1 of 4 rolling windows and was **rejected**. The methodological lesson is recorded as a standing invariant: offline harvest objectives without trade-count guardrails over-trigger (510 vs. 158 development

trades) and can invert realized outcomes. The successor objective (`utilization_a0_v2`) already incorporates the correction — rolling trade-profit as primary term, over-trigger and drawdown penalties, harvest demoted to auxiliary. For consumers, the episode demonstrates both halves of the program’s character: the optimization pressure is real and continuous, and the gates that protect the production Index from unproven “improvements” bind in practice.

12 Strategy Lineage and Research Provenance

The Index is the current endpoint of a documented v1–v8 research lineage. Earlier champions are listed for provenance; windows and venues differ across rows and returns are *not* comparable:

Table 24. Lineage champions (best-in-class per series). Do not rank by profit % alone; windows and venues differ.

Series	Champion	Headline result	Status
v1.x	v1.29/v1.30 Surf SVM-RBF	ETH 100% win rate (11 trades); +4.13% / 96% WR (3d)	Superseded
v2.x	v2.08 PPO-RL	ETH +\$251, 35% WR (20 trades)	Deprecated
v3.x	v3.03 SVM-RBF futures	+1.54% live, 66.7% WR (12 trades)	Superseded
v4.x	v4.02 spot RL	Validation fork (v4.01: 0% WR)	Archived
v5.x	v5.14/v5.15 bounded HyperSurf	Treatment +\$888 vs. control −\$30k ($n = 481$)	Thesis lives in exits
v6.x	Natural ensemble	+\$2,268 composite (512 ETH trades)	Orthogonal
v7.x	v7.02 HyperSurf composition	+0.27%, PF 1.18, 113d Kraken 5m	Not on this stack
v8.x	v8.01/v8.02 tail-control	+2.63%, 48 trades, 1.26% DD	Beaten by benchmark
Present	GM! lattice + A0 benchmark	IC +0.032; harvest $t = +5.1$; +8.28% cum.; powered window +11.64% ($n=27$)	Authority

Significance tests survive from the historical record where sample sizes permit:

- v5 bounded-exit treatment effect: $n = 481$ A/B, $p < 0.001$.
- v3.03 vs. v1.34 SVM-RBF superiority: $p < 0.01$.
- v7.04 omega tail compression: Mann–Whitney $U = 702$, $p = 0.029$.
- Early reinforcement-learning “superiority” (v2.01): 6 trades, $p = 0.24$ — claim falsified and withdrawn.

The falsified RL claim is retained in the archive deliberately: the practice of publishing negative and non-significant results is itself part of the diligence record supporting this prospectus.

13 Index Maintenance, Data Integrity and Restatement Policy

13.1 Data sources

Index calculation uses HYPERLIQUID OHLCV maintained by automated collectors with self-healing backfill: 1h for the core-6 primary profile and 5m for the crypto-3 fine profile. Each routed response carries its feed version and freshness. Validation summaries are machine-generated, timestamped, and retained.

13.2 Documented incidents and corrections

The archive records its own data incidents, which prospective users should review:

- **OHLCV truncation (Jun 2026).** Between 23–26 Jun 2026, automated backfill rounds truncated 1h history to ~72 bars due to a `keep_days` merge defect. The defect was fixed (`keep_days=None` on backfill merge) and all published battery results were re-run on post-fix data. Pre-fix and post-fix holdout numbers for the benchmark are identical to the basis point.
- **XYZ-CL listing floor.** No crude candles exist before 6 Jan 2026 14:00 UTC; window definitions respect this floor and learned arms auto-shift their training windows accordingly.
- **Interim vs. final rolling run.** The first rolling batch (Mar +3.34/Apr −2.86/May +3.87/Jun +0.85) was superseded by the unified re-run (Mar +3.88/Apr −2.18/May +4.65/Jun +1.82) after a runner defect was fixed and all arms were re-executed on one configuration. Both are published; the interim publication carries an explicit “do not cite” banner and a corrective ledger. This prospectus uses the final run as canonical and reproduces the interim numbers in Appendix A. Run-to-run differences (up to 0.97 pp on a monthly window) are themselves a disclosed measure of evaluation sensitivity.
- **Calculation-input staleness (Jul 2026).** In early July 2026 the calculation agent’s underlying data stopped refreshing (the external producer pipeline had no continuous operator), and the feed continued to serve values computed from days-old bars while its health surface reported healthy. Compounding the incident, the `staleness_ms` field then measured internal generation lag rather than data age, so the served cells self-reported as fresh. Corrections, deployed 6 Jul 2026: (i) wall-clock freshness gates on the calculation and distribution health surfaces (three-hour threshold, 24/7 reference instrument); (ii) restated `staleness_ms` semantics — the field now reports true data age, and cells additionally publish `underlying_last_ts` / `feed_last_ts` (Section 4); (iii) a dedicated producer service (`dfy-feeder`) making data refresh a supervised, health-checked production function rather than an operator habit (Section 7.2).
- **Distribution misrouting and access-control gap (Jul 2026).** A configuration-precedence defect in the deployment platform’s config-as-code resolution caused the public Index domain to route to the desk service (Index routes answered 404) for a period in early July 2026, and the calculation agent’s bearer-gated data routes were briefly reachable without a token. Corrections: per-service deployment configuration made explicit and de-aliased (the “gateway” naming collision that enabled the misroute was retired), read-token enforcement verified (data routes now answer 401 without a bearer), and domain-to-service routing added to the deploy verification checklist. Distribution incidents of this class cannot corrupt Index values — calculation and distribution are separate services (Section 7) — but they interrupt availability, and are disclosed as such.

13.3 Restatement policy

If a data or code defect is found to have affected published values, the Sponsor re-runs the affected windows on corrected data, publishes a delta table (past vs. current, per window, per arm), and version-bumps the affected documents rather than silently editing them. The July 2026 corrective ledger (`battery_rerun_compare_20260704.md`) is the template.

14 Regulatory Posture and Validation Design

14.1 The Index as a data product

The Index is calculated live and sold as machine-readable data: paid reads are fixed-fee x402 micropayments for index values or reference analytics, in the pattern of a market-data vendor. The Sponsor does not execute trades for customers, transfer customer funds, manage assets, or accept percentage-of-proceeds compensation. Directional and utilization fields are impersonal informational outputs, not individualized advice. This is the

same structural position occupied by any published index: the publisher of the VIX measures the market; it does not manage capital against its own measurement.

14.2 Why validation runs in paper mode

All utilization validation — the A0 benchmark, the capture trio, and every battery arm — executes in dry-run (paper) mode against live market data. This is a deliberate, two-sided design choice:

1. **United States compliance.** By keeping every Sponsor-operated trader non-executing, the program stays cleanly on the information-provider side of U.S. regulatory lines: no discretionary trading of customer or proprietary funds is conducted or offered in connection with the Index, and no claim of live trading performance is made. The validation record is presented as what it is — systematic, timestamped measurement of how Index-aligned rules would have performed — without the regulatory surface of an operating trading business.
2. **Aggressive optimization without production risk.** Paper mode removes the cost of being wrong, which is precisely what allows the optimization stance to be aggressive: continuous profile-build campaigns (UTIL250 and its successors), a standing ten-arm battery, rolling re-runs on every data correction, and predeclared gates that the incumbent must keep winning. The program can challenge its own production configuration at full intensity because failed challengers cost analysis time, not capital.

The two motives reinforce each other: the compliance posture is not a limitation the program tolerates but the enabling condition of its falsification cadence. Should the Sponsor or any third party ever seek to operate live capital against the Index, that step would carry its own regulatory analysis, disclosures, and the G3 live-reconciliation gate — which is precisely why G3 exists and is reported as open (Section 9).

14.3 Benchmark performance disclosure

Consistent with the above, all benchmark results herein derive from simulated execution over recorded market data. Simulated results are prepared with the benefit of a controlled environment and do not reflect the impact of execution risk, market impact, or the discipline required to trade through losses; realized results of any live implementation would differ. The Sponsor's mitigation is structural rather than rhetorical: predeclared gates, published negative results, run-to-run sensitivity disclosure, and a live paper-capture reconciliation requirement that must close before any production promotion.

15 Research and Product Roadmap

The roadmap is evidence-gated rather than date-promised:

1. **Accumulate powered evidence.** For confirmatory promotion claims, predeclare the smallest effect of interest and target at least 80% power at a two-sided 5% significance level. Freeze holdouts, account for serial dependence with block or clustered resampling, and control multiplicity across horizons and instruments.
2. **Promote full-distribution reporting.** Extend validation beyond mean and variance to skewness, kurtosis, tail mass, Omega surfaces, and changes in those objects. This is especially important in cryptocurrency and other markets with discontinuities, leverage, reflexive liquidations, heavy tails, and regime mixtures.
3. **Complete the Omega-primary gate.** Keep hurdle-net Omega fields additive while the primary series remains composite. Promotion to `value_mode=omega` requires frozen invariants, horizon matching, holdout superiority, and non-degradation across registered stress slices.

4. **Advance the fine-tier evidence ladder.** Maintain GM!-FEED-8.03-5M-CRYPT03-P02 as the live 90M fine surface while accumulating independent trade-level and distribution-level power for any subsequent parameter or utilization promotion.
5. **Expand composition only through gates.** Macro7/SP500 and additional basket grammars remain candidates until data coverage, efficacy, freshness, and utilization evidence clear the same governance boundary as core-6.
6. **Unify agent distribution.** Keep HTTP, x402, MCP, Virtuals, Discord, and provisioned access on one canonical service catalog, with identical semantics, feed provenance, prices, and refusal policy.
7. **Publish calibration diagnostics.** Add reliability curves, confidence calibration, per-regime efficacy, and rolling decay monitors so users can distinguish conviction magnitude from confidence quality.

Roadmap goals are not current-performance claims. A result enters the public production contract only after its preregistered gate and reproducibility record are complete.

16 Risk Factors

Use of the Index, and any future product referencing it, involves substantial risk. The following list is extensive but not exhaustive.

16.1 Measurement and statistical risks

Signal decay. The Index's efficacy record (holdout IC +0.032, harvest +14.5 bps) is measured over a specific period and regime set. Market structure evolves; conviction signals that covary with forward returns today may weaken, invert, or migrate across instruments. The program's response — continuous re-measurement and gated profile evolution — mitigates but cannot eliminate this risk.

Calibration drift. High-confidence IC lift degrades on the holdout window (a documented short-window calibration pattern). The confidence channel is informative in aggregate but should not be consumed as a hard probability.

Limited history. The universe data floor is 6 January 2026; all efficacy and validation statistics rest on under six months of market data. The benchmark's trade-level tests are individually underpowered ($n = 23$ holdout trades; binomial $p = 0.105$; power ≈ 0.42 against the observed win rate) and accumulate power only as rolling windows accrue. Standing alone, the trade-level record does not exclude chance as an explanation of the benchmark's profitability (Section 9.6).

Regime concentration. A single month (May 2026, +4.65%) contributes the majority of the benchmark's cumulative record; the single negative month (April 2026) is one where simpler overlays beat the benchmark with high per-trade significance. Stress-window behavior is an open optimization frontier, disclosed as such.

Evaluation sensitivity. Two published runs of identical configurations differ by up to 0.97 percentage points on a monthly window (Section 13). Consumers should treat published window returns as carrying at least this much reproducibility noise.

Data timeliness. The Index is calculated from underlying data that must be continuously refreshed by a producer pipeline. If that pipeline stalls, published values describe an older market state until the freshness gates surface the condition (three-hour liveness threshold; Section 7.2) — a failure mode that materialized in July 2026 and is documented, with its corrections, in Section 13. Consumers with latency-sensitive uses should check `staleness_ms` on each reading rather than relying on service health alone. Market-hours

instruments (XYZ-*) also carry legitimately older bars over weekends and holidays even when the pipeline is fully live.

16.2 Market and instrument risks

Digital-asset volatility. The underlying crypto-native perpetuals (BTC, ETH, HYPE) are among the most volatile liquid instruments traded; Index readings over them inherit that regime volatility.

Synthetic macro perpetuals. XYZ-GOLD, XYZ-SILVER, and XYZ-CL track reference assets through venue-specific oracle and funding mechanisms. They may deviate from their reference markets, may trade thinly, and have short listing histories (XYZ-CL: under six months at publication).

Funding and carry. Perpetual futures embed periodic funding payments. Any consumer converting Index readings into positions bears funding carry not reflected in the Index values themselves.

16.3 Venue and operational risks

Single-venue concentration. All six constituent markets trade on HYPERLIQUID. A venue outage, oracle failure, delisting, or protocol exploit would affect the entire Index simultaneously; there is no fallback venue in the methodology.

Calculation and distribution dependency. The Index depends on the Sponsor's calculation and distribution services. The separated topology (Section 7) isolates failure domains, and the parity gate detects divergence between served and reference values, but an extended calculation outage would interrupt Index publication.

Data pipeline defects. A documented collector defect truncated price history for three days in June 2026; it was caught, fixed, and restated (Section 13). Similar incidents could cause gaps or erroneous readings before detection.

16.4 Model and governance risks

Restricted construction. The Index's construction recipe is proprietary and structurally protected (LeakError projection boundary). External parties can verify the published surface, the parity of distribution, and the efficacy record — but not the internal recipe. This asymmetry is inherent to the product and disclosed plainly.

Self-administration. The Sponsor calculates, distributes, validates, and governs the Index (Section 17). Mechanical gates and published falsifications mitigate but do not eliminate the conflict inherent in self-graded research.

Methodology evolution. Governance explicitly contemplates profile promotion when a challenger clears the gates. Consumers of Index values bear transition risk at promotion events, mitigated by parameter-hash disclosure on every published value.

16.5 Legal and regulatory risks

Regulatory evolution. Decentralized derivatives venues, synthetic asset contracts, and crypto market-data products face evolving, jurisdiction-dependent regulation. Adverse action against the venue or the product category could impair Index calculation, distribution, or any future referencing product. The Sponsor's current posture (Section 14) is designed for the present U.S. landscape and will be revisited as that landscape moves.

No investor protections. HYPERLIQUID positions are not covered by deposit insurance, SIPC, or comparable protection schemes; consumers acting on Index readings do so at their own risk.

17 Conflicts of Interest and Disclosure Boundary

17.1 Conflicts

The Sponsor performs every role in the Index's lifecycle: construction, calculation, distribution, validation, gate adjudication, publication, and administration. Mitigants, all verifiable in the archive: (i) predeclared mechanical gates with published criteria; (ii) publication of negative, non-significant, and self-correcting results (falsified RL claim, UTIL250 rejection, live-capture gap, interim-run supersession); (iii) machine-generated, timestamped result artifacts; (iv) a versioned restatement policy; (v) a continuously testable parity gate between served and reference values.

17.2 Disclosure boundary

This prospectus, like the underlying publications, may cite: lattice cell keys; the public value, `signed_conviction`, and confidence semantics; feed-version IDs and parameter hashes; profile-level and per-instrument forward IC, harvest, and coverage; and validation trade statistics. It must not and does not cite: component recipes, weights, scales, internal profile construction steps, or non-public source decomposition. The boundary is enforced in code (Section 7), not merely by editorial policy.

A Full Battery and Reconciliation Tables

A.1 Interim (superseded) rolling run

The first rolling batch, published in the interim battery report before the unified re-run completed. Superseded but retained per the restatement policy:

Table 25. Interim rolling run for the benchmark (superseded by Table 16).

Window	Trades	Return %	Win %	Max DD %
1–31 Mar 2026	38	+3.34	55.3	1.01
1–30 Apr 2026	33	−2.86	30.3	3.25
1–31 May 2026	23	+3.87	52.2	0.45
1–18 Jun 2026	23	+0.85	65.2	0.83

Cumulative interim record: +5.16%; the unified re-run improved every window (+0.54, +0.68, +0.78, +0.97 pp respectively). Intramonth maximum drawdown figures are published only for the interim run; the worst observed is 3.25% (April).

A.2 Train-window stability across runs

Table 26. Train-window drift between June waves and July unified run.

Arm	Jun-wave train %	Jul train %	Δ
A0 (benchmark)	+14.22	+14.22	0.00
A1	+1.65	+1.28	−0.37
A2	+1.96	+1.96	0.00
A3	+2.41	+2.41	0.00
A4	+6.46	+6.46	0.00
A5	+5.94	+5.98	+0.04
A6	+0.37	+2.29	+1.92
A7	+0.14	+0.22	+0.08
A8	−0.04	−0.07	−0.03
A9	−0.68	−1.06	−0.38

Deterministic arms (A0–A5) are stable to ≤ 0.37 pp; learned arms (A6–A9) drift with FREQAI artifact regeneration — a further reason the deterministic minimal rule serves as the validation benchmark.

B Statistical Significance Tables

B.1 Benchmark, primary holdout

Observed 15/23 wins (65.2%), profit +0.85%. One-sided binomial against $p = 0.5$: $p = 0.1050$. One-sample t -test on per-trade returns (exit-bucket synthesis, $n = 23$): mean +0.367%, $p = 0.0963$. Neither is significant at $\alpha = 0.05$; both are directionally favorable.

Power note. The binomial test at $n = 23$ requires ≥ 16 wins for significance (attained size 0.047); observed wins were 15. Power against a true win probability equal to the observed 0.652 is ≈ 0.42 ; roughly $n = 69$ trades are required for 0.80 power. Chance is therefore not excluded by this record, and non-significance at this sample size is uninformative about the presence or absence of edge (Section 9.6).

B.2 Overlays vs. the benchmark, primary holdout

Table 27. Overlay arms vs. benchmark on primary holdout. No arm is favorable.

Arm	Δ profit (pp)	Win-rate p	Return p	Favorable?
A1	−0.50	0.1712	0.3227	No
A2	−0.83	0.2857	0.2500	No
A3	−0.84	0.1308	0.0493	No
A4	−2.45	0.0610	< 0.001	No
A5	−0.72	0.1806	< 0.001	No
A6	−1.75	0.0180	0.0010	No
A7	−1.15	0.0059	0.0242	No
A8	−1.07	0.0842	0.3765	No
A9	−0.85	0.0368	0.0629	No

B.3 April 2026 stress window

When the benchmark was negative (-2.18%), several overlays showed significant per-trade advantage (Mann-Whitney on exit-bucket returns): A1 $p < 0.001$; A2 $p < 0.001$; A4 $p < 0.001$; A5 $p < 0.001$. Disclosed as the principal open question in stress-window utilization.

C Reproducibility Manifest

Table 28. Key artifacts underlying this prospectus (paths relative to the dfai research root and the dfy platform root).

Artifact	Path
Standing reference (final)	dfai/publications/journal/2026-strategy-lineage-standing-reference/
Rolling battery report (interim + ledger)	dfai/publications/journal/2026-fci-utilization-battery-a0-a9-rolling/
Seven-arm battery report	dfai/publications/journal/2026-cfi-utilization-battery-a0-a7/
Lattice definition and validation	dfai/publications/journal/2026-dfy-conviction-forward-index/
Unified train/holdout summary	summary_20260704T031357Z.md
Final rolling summary	summary_20260704T182745Z.md
Corrective ledger	battery_rerun_compare_20260704.md
Significance ledger	significance_tables.md (standing reference)
Battery matrix (51 rows)	battery_matrix.tsv (rolling battery report)
Efficacy tooling	user_data/feed_versions/confidence_forward_index.py, user_data/feed_versions/profile_benchmark.py
Battery runner	user_data/versions/scripts/run_cfi_utilization_battery.sh
Feed version registry	user_data/feed_versions/dfy_derivative_feed_versions.json
GM! distribution service	dfy/services/fci/README.md; live /services.json
Parity harness	dfy/services/fci/scripts/validate_fci_parity.py
Data producer + freshness gates	dfy/services/feeder/; dfy/packages/dfy_iq/dfy_iq/synthetic_feed/freshness.py
Projection boundary + leak auditor	dfy/packages/dfy_iq/dfy_iq/guidance_projection.py; boundary test in dfy/packages/dfy_iq/tests/
Deployment identity	dfy/packages/dfy_iq/dfy_iq/stack_identity.py

D Glossary

A0 / benchmark

The GM! Minimal Utilization Benchmark: 1D signed conviction with entry band 0.32 and confidence floor 0.45, one-horizon hold, paper mode. A validation instrument — not the Index.

Arm One candidate utilization design in the battery (A0–A9).

Cell One addressable Index reading: (instrument, granularity, horizon).

Confidence

Published channel in $[0, 1]$ expressing data-support / reliability of a cell's conviction.

Coverage

Fraction of bars on which a cell publishes an active signed opinion.

GM! lattice

The full set of conviction cells — **the Index itself**.

Forward IC

Spearman rank correlation between signed conviction and realized forward return over the cell horizon.

G1–G4

The four standing promotion gates (Table 3).

Harvest (bps)

Mean of signed conviction $\times$ forward return, in basis points, over bars with $|\text{signed}| \geq 0.05$.

Harvest rung

A registered (band, min_confidence) admission pair (purist / balanced / active) selectable on paid reads via ?rung=; a consumption choice over the frozen Index, not a calculation change (Section 7.5).

Scale overflow

Disclosure flag raised when an operating rung's trade count exceeds $1.3 \times$ the next-higher rung; an activity-budget breach, not a performance failure.

Holdout

A time window excluded from all parameter fitting, used only for evaluation.

Horizon ladder

The published forward-window rungs (15M, 30M, 90M, 3H, 4H, 1D); legacy four-rung core plus extended bridge horizons.

GM! cell

One instrument–horizon member of the conviction distribution, with primary conviction, confidence, Omega fields, freshness, and provenance.

MCP

Model Context Protocol — the agent-native interface through which the Index is consumable by AI systems.

Production profiles

The primary GM!-FEED-8.01-1H-CORE6-P01 and fine GM!-FEED-8.03-5M-CRYPT03-P02 profiles.

Parity gate

Standing harness proving x402-served values equal the free reference value-frame.

Perpetual future

A futures contract with no expiry, kept near spot by periodic funding payments.

Rolling ladder

The four monthly out-of-sample windows (Mar–Jun 2026) used as co-equal promotion authority.

Signed conviction

Published channel in $[-1, +1]$: $(\text{value} - 0.5) \times 2$.

Tier A / B / C

The platform's codified disclosure tiers: paid public posture / journal / private construction, separated by the LeakError projection boundary.

UTIL250

A 250-iteration profile-optimization candidate rejected under gate G4 in July 2026.

Utilization

Any rule converting published Index values into positions; distinct from the Index itself.

x402 Per-request micropayment protocol (Base network) through which paid Index reads are settled.

References

1. *Strategy Lineage Standing Reference*, v1.0.0 (final), ConfigSecretAI, Inc. / MUTANTDEFI Research Lab, 4 Jul 2026.
dfai/publications/journal/2026-strategy-lineage-standing-reference/
2. *The Utilization Ladder: A Nine-Arm Rolling Battery for GM! Index Extraction*, v1.0.0 (interim, with corrective ledger), Jul 2026.
dfai/publications/journal/2026-fci-utilization-battery-a0-a9-rolling/
3. *From Vanilla to Optimal: A Seven-Arm Battery for GM! Index Utilization*, v0.4.0, Jun 2026.
dfai/publications/journal/2026-cfi-utilization-battery-a0-a7/
4. *Conviction Forward Index over the DFY CCXT Feed*, research preprint, 21 Jun 2026.
dfai/publications/journal/2026-dfy-conviction-forward-index/
5. *DFY Derivative CCXT Feed Construction* (restricted).
dfai/publications/journal/2026-dfy-derivative-ccxt-feed-construction/
6. *DFY GM! platform* — topology, tier boundary, service export, and deployment identity.
dfy/README.md; dfy/SOUL.md; dfy/ops/notes/2026-07-14-gm-origin-cutover.md
7. *FC-27 Harvest Ladder Pre-Registration* (PR-2026-07-19-HARVEST-LADDER-01), 19 Jul 2026.
dfai/publications/pre-registrations/2026-07-19-fc27-harvest-ladder-preregistration/
8. *Dual Feed, Empty Funnel: Ops Coherence After CRYPTO2 Cutover and Engines 1–2*, v1.2.0 (rung quantification and confirmation adjudication), 19 Jul 2026.
dfai/publications/journal/2026-07-19-dual-feed-empty-funnel-ops-coherence/
9. William Lloyd Gleim, *GM Futures*, first edition, ConfigSecretAI, Inc. / MUTANTDEFI Research Lab, Jul 2026.
10. Vu Ngoc Nguyen, *Omega Function: A Theoretical Introduction*, 2009.

Document state: v2.4.4 | Index prospectus | Build date: 2026-07-19

v2.4.0 change note: headline and benchmark records unified on the canonical cfi-val lineage (legacy battery-profile study relabeled and retained for audit); added the powered extension window (1 May–8 Jul 2026: benchmark +11.64%, $n=27$, $p \approx 8 \times 10^{-4}$; sizing and hazard reference seams at pre-registered SUPERIORITY), six-month harvest significance (+15.26 bps, $t = +5.13$), cross-venue 5m generalization, and the single-regime and window-overlap disclosures. **v2.4.1:** reader-facing explanations of the powered-window concept and of the t/p statistics added at the point of use (headline, Section 9.4, and Section 9.6); no figures or claims changed. **v2.4.2:** Section 10.4 expanded from a two-sentence note to the full interpretation of the battery: the naive readout as the *demonstrated optimum* of the admission layer (empirical-optimality scope, falsification hooks, and the sophistication-burden inversion stated for the reader); no figures changed. **v2.4.3:** new Section 10.5 — CAPM separation of benchmark return from market return (equal-weight core-6 buy-and-hold leg computed from registered OHLCV; four-window $\hat{\beta} = +0.06$, per-window separations including +17.4 pp on the powered window; descriptive scope disclosed). **v2.4.4:** one reference sentence

appended to Section 10.5 — the USDC at-rest rate required for positive alpha in every window ($\approx 30\%$ annualized; April binding).

Disclosure: public research citations only; restricted construction excluded and structurally protected. This document is not an offer of securities. The Index is a calculated information index; benchmark results are simulated validation measurements — see Section 14.

MUTANTDEFI