

From closed bars to governed signals.

A mathematical methodology, exact signal-flow audit, real BTC/ETH/SOL benchmark, market landscape, and prospective evaluation plan for a local crypto research system.

Implementation-grounded

Negative results included

Research AI has no alert authority

No wallet or order path

Published: 28 July 2026 Evidence cutoff: 14:07 UTC 67,608 public hourly bars

Research methodology, not investment advice

IMPLEMENTED POLICIES 8 TSMOM, Donchian, compression-expansion, fragility	PROSPECTIVE RESEARCH EMITS 3 BTC, ETH, SOL; research-only and provisional	PRODUCTION ALERTS 0 120 evaluations stopped before candidate creation	ROBUST POSITIVE TSMOM IC 0 / 9 Every 95% dependence-aware interval crossed zero
---	--	--	--

01 / ABSTRACT

The system moat is visible. The alpha moat is not proved.

Crypto Signal Lab implements deterministic signal definitions, point-in-time evidence, a ten-gate decision path, cost-aware labels, uncertainty controls, replay, and promotion governance. A separate research lane can produce descriptive insights, but it cannot create a production alert.

The central finding

Three prospective research insights were emitted for BTC, ETH, and SOL. The production lane produced zero candidates and zero alerts. A 2024-2026 public-data benchmark found weak and unstable TSMOM correlation, with eight of nine mean net episode returns below zero. No production-alpha claim is supported today.

The value of this result is falsifiability. The system can now distinguish an implemented mechanism, a prospective runtime event, a retrospective experiment, and an unproven product claim. Those categories are not added together to manufacture confidence.

02 / THE MOAT, STATED PRECISELY

The difficult part is not another indicator. It is preserving authority across the whole evidence chain.

Every mathematical ingredient has precedent. The present differentiation is the orchestration: data availability, deterministic hypotheses, independent challenge, calibrated authority, cost-aware utility, uncertainty, attention, immutable outcomes, and delivery remain separate but reproducible.

IMPLEMENTED TODAY

Governed systems moat

One point-in-time identity follows a claim from closed bar to feature, candidate, gate, alert or suppression, outcome label, replay, and promotion record. Research models can add context, but types and routing prevent them from acquiring alert authority.

EXPLICITLY UNPROVEN

Predictive alpha moat

The public benchmark does not establish durable directional edge. No promoted production champion exists. The system is an auditable research process, not evidence that its present thresholds outperform after costs.

CONDITIONAL AND PROSPECTIVE

Compounding evidence moat

A non-emitting shadow tape can accumulate scores, modeled costs, realized outcomes, calibration, regime behavior, and live/replay parity under stable identities. That evidence can become difficult to replicate, but only after it exists.

What is novel, and what is not

Momentum, breakouts, CUSUM, volatility scaling, triple barriers, calibration, block bootstrap, and multiple-testing controls are established methods. The engineering contribution is a local system that composes them while making missing data, suppression, uncertainty, narrative AI, and promotion boundaries durable and inspectable.

Demand is established across four adjacent product categories. Their centers of gravity remain different.

This is a category scan based on official product pages, not a feature-by-feature competitor scorecard. It validates that users pay attention to alerts, derivatives state, on-chain intelligence, and attention analytics. Signal Lab is positioned around the governed process connecting those observations to a local research record.

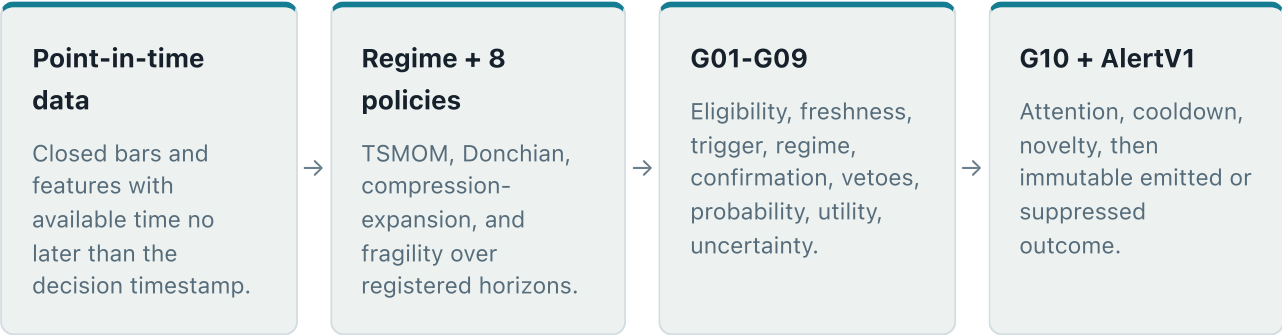
CATEGORY EXAMPLE	OFFICIALLY DESCRIBED CENTER	RELEVANCE TO SIGNAL LAB
TradingView alerts	Price, technical, watchlist, strategy, and chart-pattern conditions with cross-device notifications.	Validates demand for timely customizable monitoring. Signal Lab focuses on what must happen after a condition fires: independent evidence, costs, uncertainty, suppression, and outcome lineage.
CoinGlass liquidation heatmap	Visual estimates of potential liquidation concentrations across prices, pairs, and exchanges.	Validates demand for derivatives-state context. Signal Lab treats fragility as typed evidence or risk context rather than allowing a heatmap or crowding feature to become direction by itself.
Nansen	On-chain wallet intelligence, portfolio monitoring, AI research, and integrated spot/perpetual execution.	Validates demand for labeled on-chain intelligence. Signal Lab deliberately takes the opposite authority boundary: local read-only research with no wallet, custody, or execution path.
Kaito	Attention and mindshare analytics across voices, sectors, regions, and social followings.	Validates demand for information prioritization. Signal Lab allocates a finite attention budget only after deterministic evidence and uncertainty gates, and does not use social attention as ground truth.

Official pages accessed 28 July 2026. Product surfaces change. The table describes stated category emphasis and does not claim that any named product lacks unlisted capabilities.

There are two lanes, and only one can create a production alert.

The separation is the most important architectural fact. Production detectors create governed candidates. Pattern Studio, VIC, Chronos, and Qwen create research artifacts or explanations that cannot cross into production authority.

PRODUCTION DECISION LANE



RESEARCH INTELLIGENCE LANE



Four registered hypotheses, normalized for risk and challenged after detection.

Time-series momentum

The implemented score averages five volatility-normalized momentum horizons: 24, 72, 168, 336, and 720 hours. Each component is capped before averaging; a candidate requires absolute score of at least 0.25.

Donchian breakout

Prior-range breaks use 24, 168, or 360 hourly bars. Efficiency ratio must reach 0.30 and governed cross-venue breadth must reach 0.60. Strength combines path efficiency, breadth, and distance beyond the channel.

Compression-expansion

A 24-hour hypothesis requires a prior squeeze, a 24-bar Donchian break, current volume at least twice its trailing median, and breadth of at least 0.60. The 90-day compression percentile creates a 2,160-hour warm-up requirement.

Derivatives fragility

A weighted score combines funding extremity, positive open-interest change, basis magnitude, and liquidation intensity. It emits typed risk-off context at 0.70, with at most one missing component and weight renormalization.

Implementation is not evidence of edge.

These equations define reproducible hypotheses. Their existence proves what the software computes. Only out-of-sample and prospective outcomes can establish whether the hypotheses predict anything useful after costs.

06 / MANAGED RUNTIME EVIDENCE

Three real research emits. Zero production alerts.

At the 28 July 2026 00:00 UTC closed bar, the research lane emitted down-continuation insights for BTC, ETH, and SOL. They were delivered at information priority and remained

provisional because historical same-side outcome and utility evidence were unavailable.

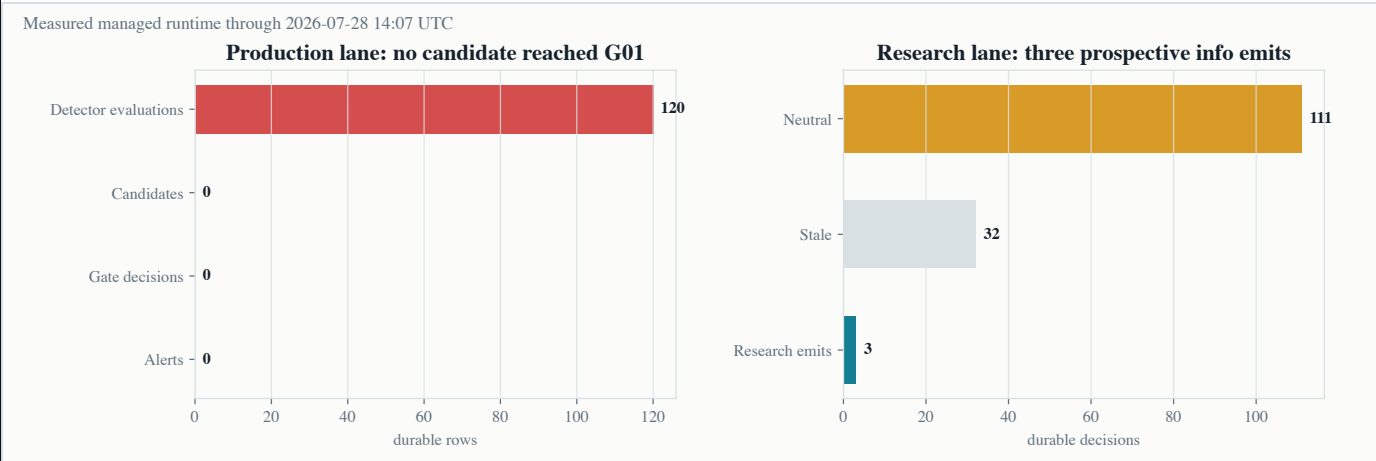


Figure 1. Production and research events are different contracts. They must never be summed into one emit count.

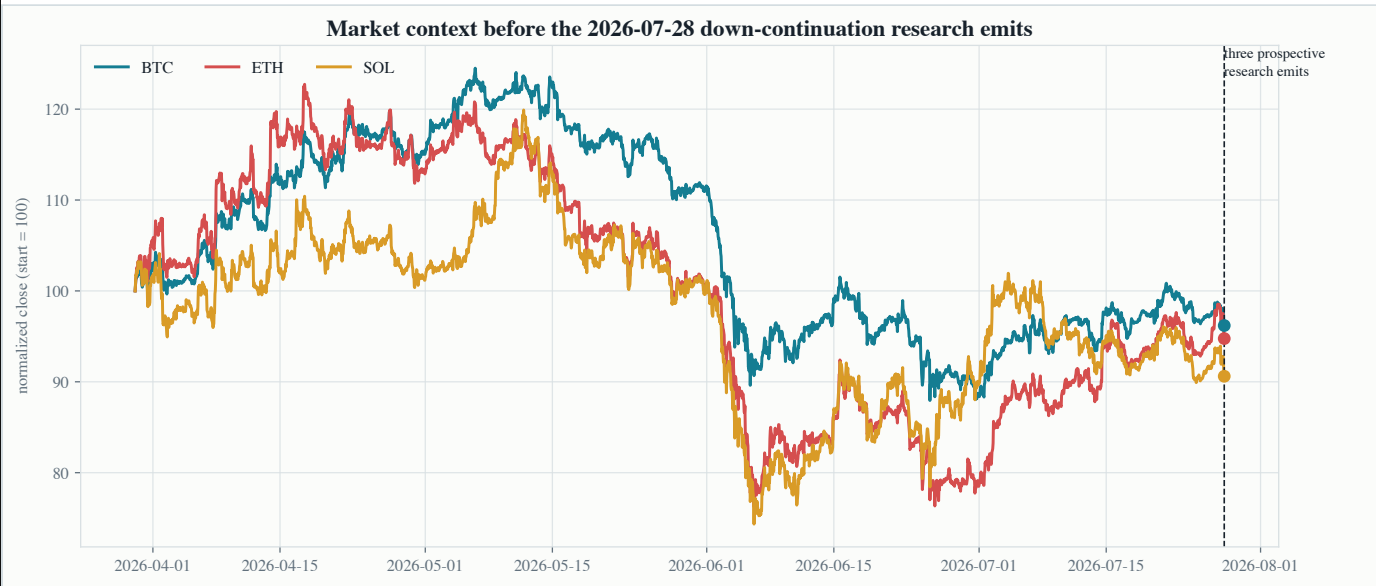


Figure 2. Public Binance spot context around the prospective research emits. The markers are descriptive, not entries or outcome evaluations.

Production telemetry recorded 120 detector evaluations: 33 each for 4-hour TSMOM and Donchian, and nine for each 24-hour or 72-hour policy. Every row ended precondition_failed / detector_row_missing. There were no production candidates, gate decisions, promoted champions, or alerts.

The implemented TSMOM score is weak in this first public-data benchmark.

The experiment used 22,536 contiguous Binance spot hourly bars for each of BTC, ETH, and SOL from 1 January 2024 through the 28 July 2026 boundary. Canonical raw-row hashes are retained in the public manifest.

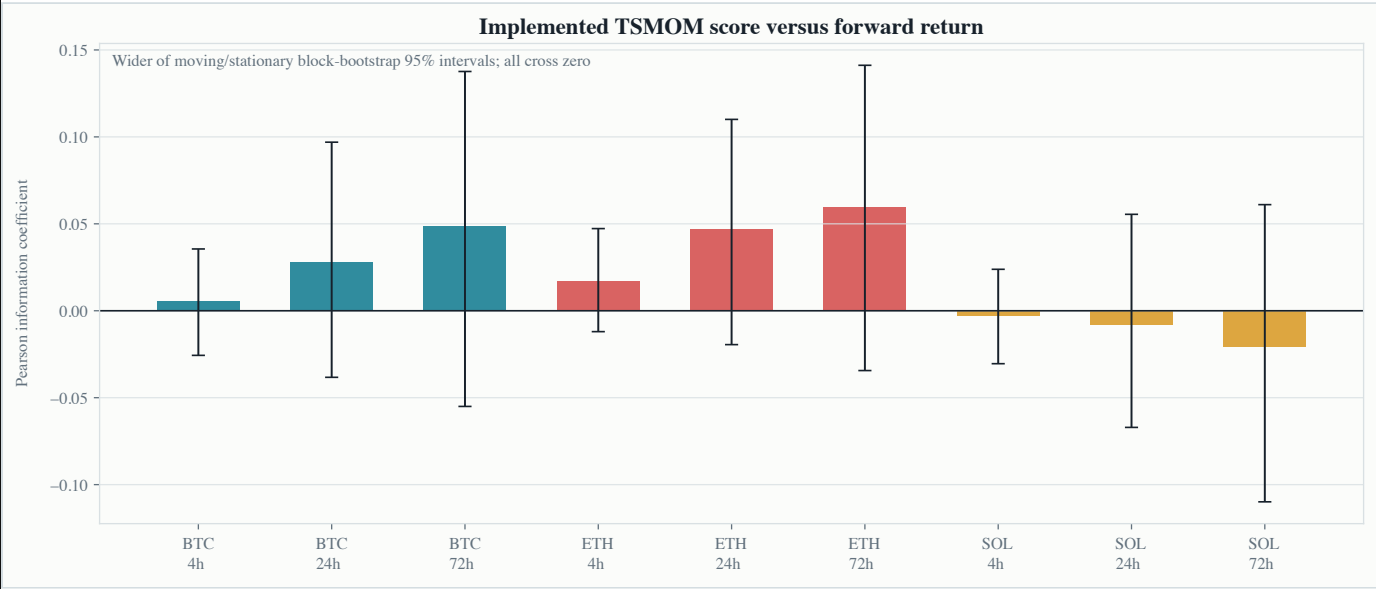


Figure 3. All nine 95% moving-block or stationary-block bootstrap intervals cross zero. SOL is negative at every tested horizon.

ASSET	HORIZON	PEARSON IC	95% BLOCK INTERVAL	EPISODES	MEAN NET BPS
BTC	4h	.0057	[-.0257, .0355]	759	-26.20
BTC	24h	.0280	[-.0383, .0970]	398	-36.13
BTC	72h	.0487	[-.0550, .1376]	395	-10.65
ETH	4h	.0171	[-.0120, .0473]	728	-28.66
ETH	24h	.0469	[-.0195, .1101]	371	-25.29
ETH	72h	.0595	[-.0344, .1412]	370	19.73
SOL	4h	-.0031	[-.0304, .0238]	798	-24.99
SOL	24h	-.0078	[-.0671, .0555]	391	-14.00
SOL	72h	-.0207	[-.1099, .0610]	390	-46.39

An episode is the first candidate bar after a non-candidate state or direction change. Net return uses a transparent research proxy: 20 bps round-trip taker fees, 2 bps round-trip fallback spread, and two square-root-impact legs at USD 10,000. It is less complete than the governed minute-fill label.

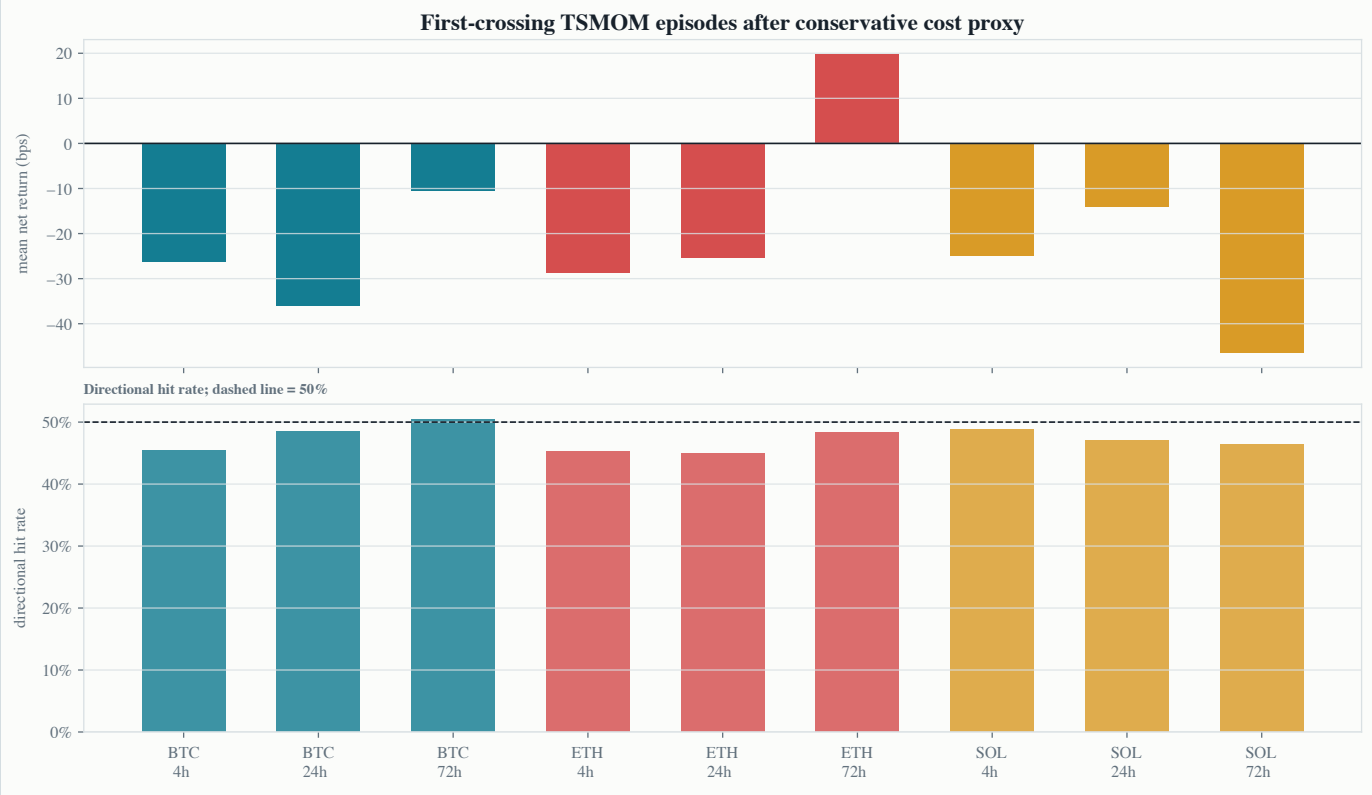


Figure 4. Eight of nine mean net episode returns are negative. ETH 72-hour is positive, but its HAC statistic is only 0.48 and its IC interval crosses zero.

Interpretation

The score may still be useful as one feature in a regime- and confirmation-aware model. This experiment rejects any simple claim that the registered threshold alone creates stable, universal directional alpha across BTC, ETH, and SOL.

08 / PUBLIC CLAIM LEDGER

What the evidence permits us to say.

IMPLEMENTED

The signal definitions are deterministic and versioned.

Strict contracts, frozen configuration, definition hashes, and deterministic identities support this implementation claim.

IMPLEMENTED

Research AI cannot create production authority.

Chronos can corroborate or brake research context; Qwen can explain frozen evidence. Neither can create a candidate, probability, promotion, wallet action, or order.

OBSERVED

The research lane can emit governed descriptive insights.

Three prospective information-priority emits exist for BTC, ETH, and SOL. They are provisional and are not production alerts.

OBSERVED

The measured production lane has not produced a candidate.

All 120 measured evaluations ended before candidate creation. A zero gate count is therefore not evidence that the market lacked setups.

UNPROVEN

The present thresholds predict positive net returns.

The first public-data TSMOM benchmark does not support this claim. Donchian, compression-expansion, and fragility still require full governed-input evaluation.

UNPROVEN

The full system creates durable user utility.

This requires prospective shadow outcomes, live/replay parity, cost parity, calibration, and expert review. Calendar time alone cannot establish it.

Use HFT-style correlation discipline, then demand calibration and economics.

Information coefficient evaluates the continuous score before a threshold creates selection effects. It is the right first test, but not a complete product metric.

M1	M2	M3	M4	M5
Detector truth Classify every policy-asset tick as missing, warm-up, precondition, no-trigger, candidate, or error.	Reachable gates Prove every enabled policy has a satisfiable path and real required data producers.	Historical challenger Walk-forward IC, calibration, costs, stresses, baselines, and multiplicity controls must pass.	Prospective shadow Collect every eligible score without notifying, then mature outcomes at 4, 24, and 72 hours.	Governed production Only a promoted champion surviving G01-G10 may spend attention as a production alert.

Primary statistical tests

- Pearson and Spearman score-to-forward-return IC by asset, horizon, side, regime, volatility, and data-quality state.
- Joint timestamp-block bootstrap across assets so BTC/ETH/SOL dependence is retained.
- Residualization against broad market direction, realized volatility, and BTC beta to distinguish exposure from incremental information.
- Brier skill, reliability slope/intercept, selective-risk curves, and interval coverage for probabilistic outputs.
- Realized net return after fees, spread, impact, funding, gas, and user-specific local costs.
- Romano-Wolf correction for production claims; false-discovery controls only for exploratory discovery.

Minimum prospective proof

Current governance requires at least 30 unchanged days and 50 closed shadow outcomes, feature parity of at least 99.5%, bounded probability deltas, decision-to-completion p95 no more than 60 seconds, p99 no more than 180 seconds, modeled-versus-observed cost deviation no more than 25%, and no corruption incidents.

The methods have lineage. The composition still has to earn evidence.

1. Moskowitz, Ooi, and Pedersen (2012), Time Series Momentum. Continuation hypothesis.
2. Moreira and Muir (2017), Volatility-Managed Portfolios. Volatility scaling and risk normalization.
3. Liu and Tsyvinski (2021), Risks and Returns of Cryptocurrency. Dedicated crypto factor and risk evidence.
4. Makarov and Schoar (2020), Trading and Arbitrage in Cryptocurrency Markets. Cross-venue fragmentation.
5. Cont, Kukanov, and Stoikov (2014), The Price Impact of Order Book Events. Flow and impact motivation.
6. Page (1954), Continuous Inspection Schemes. Two-sided CUSUM lineage.
7. Newey and West (1987), HAC covariance estimation. Overlapping-return inference.
8. Politis and Romano (1994), The Stationary Bootstrap. Dependence-aware resampling.
9. Brier (1950), Verification of Forecasts Expressed in Terms of Probability. Probability scoring.
10. Bailey and Lopez de Prado (2014), The Deflated Sharpe Ratio. Search-aware performance evidence.
11. Romano and Wolf (2005), Stepdown Methods for Multiple Hypothesis Testing. Production-claim multiplicity control.
12. Howard et al. (2021), Time-uniform Confidence Sequences. Prospective monitoring without fixed-horizon peeking.
13. Franc, Prusa, and Voracek (2023), Optimal Strategies for Reject Option Classifiers. Abstention as a first-class decision.
14. Ansari et al. (2024), Chronos. Probabilistic pretrained forecasting lineage; research context only here.

The downloadable PDF preserves this selected research grounding. The public evidence bundle publishes the experiment metrics, source hashes, runtime counts, and derived event datasets used by the argument.

Public evidence is published beside the argument.

The experiment used frozen public Binance hourly snapshots with 22,536 rows per asset and 2,000 moving/stationary block-bootstrap resamples. The evidence bundle contains no keys, tokens, wallet data, private endpoints, or execution records.

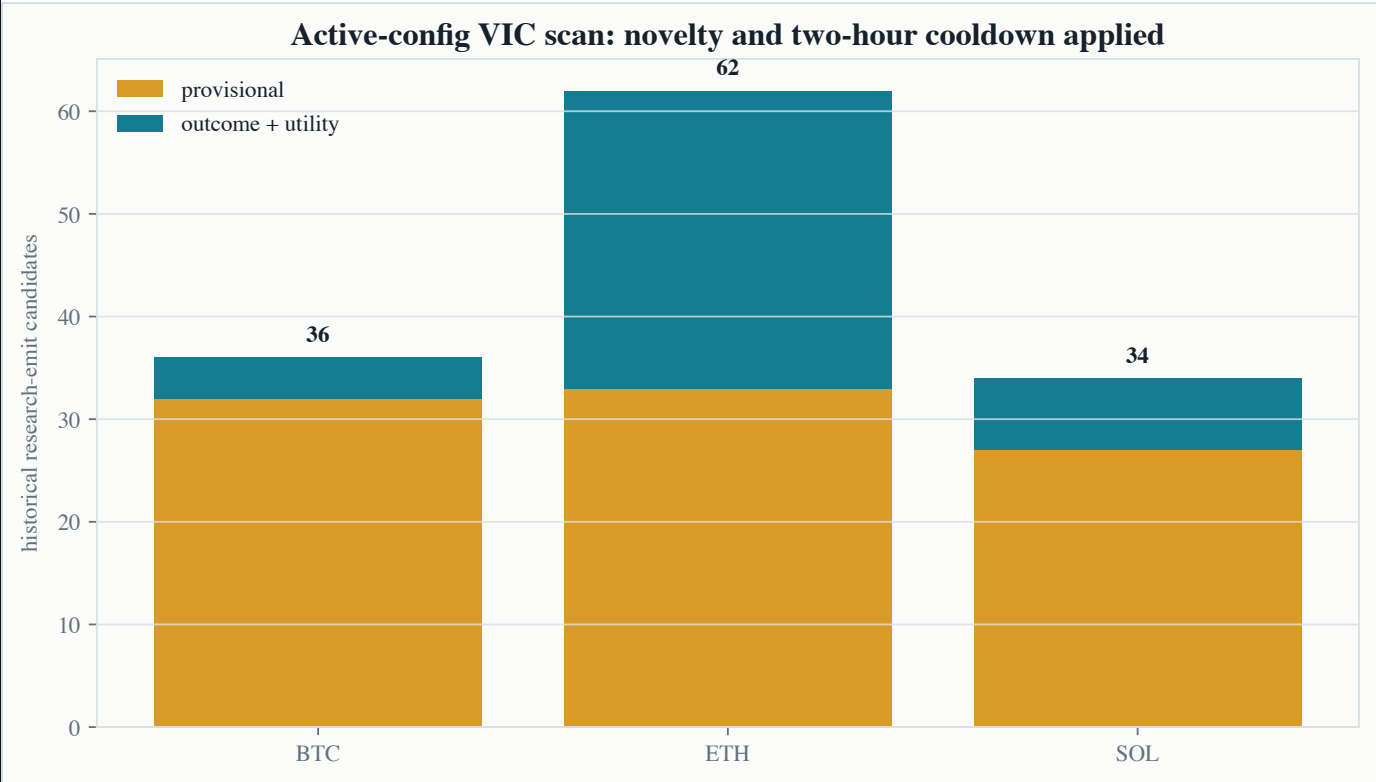


Figure 5. Historical active-configuration VIC candidates after novelty and two-hour cooldown. These are retrospective research events, not prospective emits or production alerts.

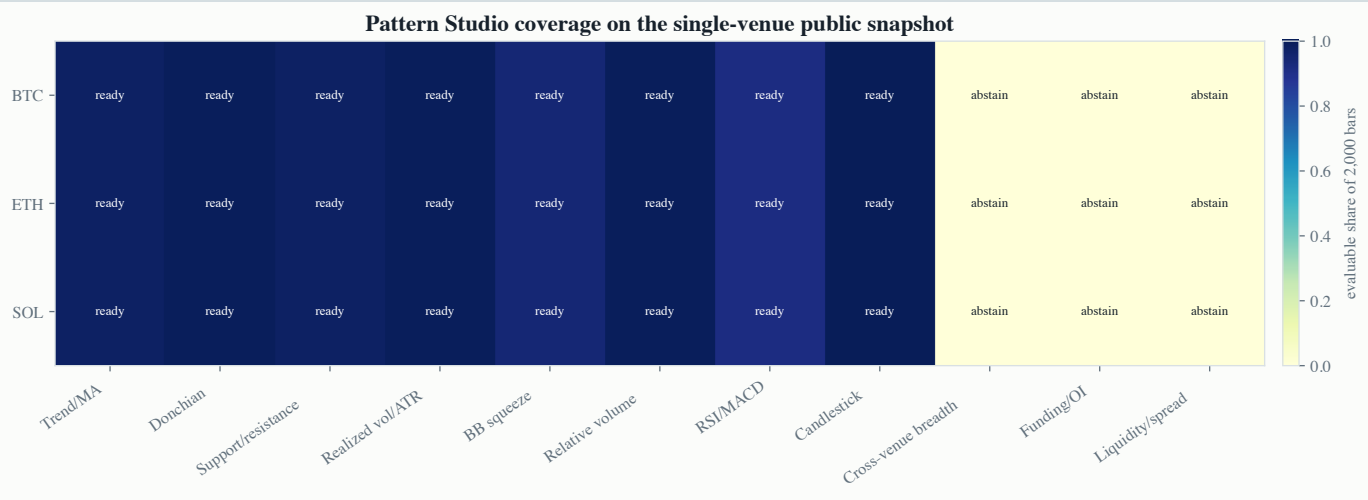


Figure 6. Price and volume patterns evaluate on the single-venue public snapshot. Multi-venue, funding/OI, and liquidity patterns correctly abstain when their required evidence is absent.

Download the evidence package

Public methodology whitepaper

PDF

Sanitized managed-runtime evidence

JSON

Experiment manifest and metrics

JSON

Artifact checksums and byte sizes

JSON

TSMOM first-crossing episodes

PARQUET

Historical VIC candidate events

PARQUET

Privacy and data boundaries

HTML

Review standard

A useful expert review should challenge the equations, signal-family semantics, label geometry, cost assumptions, correlation gate, and claim ledger. Claims that cannot survive independent review should be removed rather than tuned into apparent success.