

A Falsificationist Cycle Analysis of Cryptocurrencies

A Popperian validation protocol for the cycle theory of markets, with empirical verification on cryptocurrencies, stock indices and thirty-seven instruments in six classes

Luigi Garone · Independent Researcher · Cryptoverso — Cyclical Research

September 2026 · Licensed CC BY-NC-ND 4.0 · Correspondence: luigi@cryptoverso.net

Disclosure boundary. Thesis, method, null hypotheses, verdict rules and statistical results are published and verifiable; the dataset and the verdicts in JSON format are provided for independent verification, so that the reader can rerun the tests on the delivered output and challenge the verdicts. The cycle detector and its calibration remain proprietary — thresholds, parameter values, tuning procedure, operational rules and code — which the document **describes** without delivering them.

Abstract

For more than half a century, the cycle analysis of markets has produced rich conceptual frameworks that have rarely been subjected to empirical falsification. The three historical schools — classical American (Hurst), Italian (battleplan and inverse cycle) and quantitative DSP (Ehlers) — are adopted as alternative doctrines, almost never tested jointly on a statistically relevant sample. This work formalizes a Popperian protocol that puts each founding claim of the three schools to the test on a multi-asset dataset — Bitcoin and Ethereum since 2017, Solana since 2020, extended to thirty-seven instruments in six classes and to six stock indices with up to ninety-nine years of daily history — over several thousand phased hierarchical cycles.

The methodology is causal by construction and deterministic; both properties are formalized in automated tests and demonstrated empirically, in particular on the Bayer-tongue detection module, verified as non-repainting in walk-forward mode on five asset/timeframe combinations. To this scaffolding the protocol adds four requirements, which are the part of the work that outlives the individual results: **the null value of a measure must be measured, not assumed**, with surrogate series that preserve the power spectrum and destroy the temporal structure; **a conditional rate must be accompanied by its own marginal**, that is, by the same condition counted on the cases the event does not select; **a negative outcome does not count until it has been shown that the test would recognize the phenomenon if it were there**; and **every level must be measured on the timeframe that resolves it**, because the daily timeframe does not separate the shortest levels of the hierarchy and measuring them there changes the very numbers under discussion.

The results refute the most exposed canonical assertions: Hurst's fixed harmonic nesting and the principle of harmonicity — the latter with a test whose power is verified on a true synthetic scale, where on the daily ruler none of the markets with enough peaks turns out to be closer to the rungs than a set of periods drawn at random from the same grid —, the sequential mean-reversion of durations, Bayer's canonical threshold for connecting cycles, the peak hold, and — by tautology — two claims of the Italian school: the inverse cycle, whose operational definition is not falsifiable *ex ante*, and the unconditional swing, whose success condition is true almost always even without the swing — in all the parent cycles of the daily, and in at least nine out of ten on every ruler. They confirm, with statistical significance and cross-asset convergence, the cyclic commonality of durations — strong at the short levels and, on long history, equally strong at the high levels —, cyclical translation as a predictor of polarity, the violation of the previous peak and, in a rewritten form that eliminates a set-theoretic inclusion, the two bearish signals of the Italian corpus.

The same yardstick applies to the regularities the work could claim as its own, and there it is more severe. The **short-cycle tail**, measured against one hundred phase-randomized surrogates and one hundred AAPT per market with each level on the ruler that resolves it, **leaves the cloud in one asset/level combination out of thirty-three**, below the falsification threshold: it is what the detector produces, not what the market does.

The same holds for the deviation from the nominal period and for the position of the peak within the cycle. And the work's **starting axiom**, according to which a quaternary structure would be a binary of binaries, **cannot be decided** with this instrument: the trough that would separate the two binaries is the lowest of the three internal ones in 1.3% of the 557 measured cycles, within the band that the same detector produces on noise; and on synthetic series in which the binary of binaries is implanted, the detector never recovers it in numbers sufficient to vote.

The contribution is threefold: a falsifiable and reproducible protocol applied uniformly to the whole corpus, regardless of the origin of the assertion; the demonstration of causality and non-repainting of the central analytical component; and the systematic distinction between what the **market** produces and what the **detector** produces, which on several occasions changes the sign of a conclusion — to the point of removing, among the regularities this work could have claimed as its own, those that seemed most solid.

1. Introduction and motivation

The idea that financial prices oscillate around recurring periodicities has accompanied technical analysis since the 1930s (Wheeler, Dewey, Bayer) and reaches a mature formulation with J. M. Hurst (*The Profit Magic of Stock Transaction Timing*, 1970) through the Nominal Model and the seven core principles. The Italian strand unfolds in two phases: a first one (Migliorino, Sartorelli) introduces the battleplan as a context tool and **cyclical structures** as units of measurement of time, with nesting by two and by three, integrated with Gann on equity and crypto markets; a second one (Elico64, later taken up for teaching purposes by Marini) formalizes in operational terms the inverse cycle, the polarity rules, the definitions of conditional and unconditional swing and the violation of cycle tops.

A debt must be acknowledged here, because it concerns the method before the results. Of everything this work has been able to put to the test, the most testable part comes from the Italian school: the idea that the cycle structure is a **unit of measurement** — and not a figure to be recognized by eye — is what makes it possible to count, compare and falsify. Sartorelli, in particular, brings an engineer's approach to that idea: explicit definitions, a repeatable procedure, extension to a new market without changing the rules to make them fit. And the corpus of **Elico64** reaches a degree of formalization — the inverse cycle, polarity, the conditional and unconditional swing, the violation of tops — that here could be translated into tests **without having to interpret it**: they are operational definitions, written by someone who did not have the means to verify them on tens of thousands of cycles. Putting them to the test is, in large part, the content of this work.

A doctrine can be falsified only if someone has formulated it with enough precision, and on this ground the merit does not belong to whoever measures: it belongs to whoever wrote clearly enough to be proven wrong. Six of the claims falsified here would not even have been formulable without that work. In parallel, J. F. Ehlers brings classical DSP into the financial world: the Hilbert transform, MESA, the SuperSmoother, the autocorrelation periodogram.

Three schools, three vocabularies, three doctrines. None of the three — in the founding texts or in market practice — submits its core claims to a falsifiable empirical test on a statistically relevant sample of cycles and on a modern asset. The cryptocurrency market offers the opportunity to close this gap: continuous trading, high volatility, a complete minute-by-minute digital history since 2017, a visually recognizable cycle structure, an abundance of correlated assets for external validation.

The work pursues three objectives: to falsify formally (hypothesis → test → *p*-value → verdict) the core claims of the three schools; to synthesize what survives into a causal, multi-level, regime-adaptive protocol; to document the empirical regularities that emerged in the process, **including negative results and the refutations of its own hypotheses**.

Premise. The cycle structure displayed on price is not a deterministic predictive system: it is a context tool that reduces decision variance by giving the analyst a quantitative reference for the position within the cycle. Any predictivity emerges from lower-level signals conditioned by the structural context, not from the cyclical overlay alone.

This premise is itself a claim, and as such it has been put to the test: § 6.6 reports the outcome, which supports it on seven rulers out of nine — with a reservation that defines its scope, because the measured context is known only after the fact.

2. State of the art: three schools compared

School	Reference	Focus	Approach	Key limitation
Classical American	J. M. Hurst (1970)	Nominal Model, seven principles, harmonic nesting	Deterministic top-down	Fixed nesting never tested outside US equities
Italian, phase 1	Migliorino, Sartorelli	Battleplan, structures by two and by three, integration with Gann	Visual bottom-up, manual-based	Battleplan as a mechanical system hard to falsify
Italian, phase 2	Elico64 (taken up by Marini)	Inverse cycle, polarity, swing, violation of tops	Bottom-up with formalized rules	Tautological definition of “inverse cycle”
Quantitative DSP	J. F. Ehlers (2001, 2013)	Adaptive filters, instantaneous frequency	Signal processing	Assumes local stationarity; blind to the structural context
Synthesis (this work)	—	Fusion and falsification	Popperian data-driven	—

Map of the three historical schools and proposed falsificationist synthesis.

Each carries a limitation of its own, and it is the limitation that makes its claims interesting to test: in Hurst the fixed nesting, never tested outside US equities; in the first Italian phase a battleplan hard to falsify because it is not mechanized; in the second a tautological definition of “inverse cycle”; in Ehlers an assumed local stationarity, blind to the structural context.

From the three corpora the most characteristic operational claims were extracted and reduced to a set of falsifiable hypotheses, each accompanied by a statistical test, a decision threshold fixed a priori and a minimum sample. To these are added the hypotheses specific to the Elico64 corpus on the swing and on the violations of cycle tops, and the hypotheses of this work itself. The protocol is applied uniformly: **the author’s hypotheses are subjected to the same rigor as those of the historical schools**, and in the verdict table (§ 8) they appear with the same detail, including those that fell. The “Source” column of that table reports the separate count by origin, and it is the most uncomfortable line in the document.

2.1. Related work: the statistical literature this protocol answers to

The claims tested here come from practitioner traditions, but the objections they must survive come from a statistical literature that this work takes as its standard.

Cycles that are not there. The oldest objection to any cyclical reading of a price series is that a purely stochastic process can look periodic. Yule (1927) showed that a second-order autoregressive process driven by random shocks produces visible pseudo-periodicity — the effect later named after him and Slutsky (1937) — and Fisher (1929) gave the classical significance test for the largest ordinate of a periodogram against white noise. The surrogate approach used here (Theiler et al. 1992; Schreiber and Schmitz 2000) is the computational descendant of that test, and it is preferred for a reason that is methodological rather than practical: Fisher’s *g*-test assumes a stationary Gaussian white-noise null, an assumption that non-stationary, heavy-tailed financial series violate, whereas phase-randomized and AAFT surrogates fix the null at the empirical spectrum of the series actually under study.

Many tests, one dataset. A protocol that tests dozens of claims on a handful of markets is exposed to exactly the failure the finance literature has spent two decades documenting: White (2000) formalized the bootstrap reality check for data snooping, Sullivan, Timmermann and White (1999) applied it to technical trading rules and found that the best in-sample rule does not survive out of sample, Harvey, Liu and Zhu (2016) argued that the accumulated search over the cross-section of returns demands far stricter thresholds than the conventional ones, and Harvey (2017) made the case institutionally. This work answers that objection with thresholds fixed before execution, with the Holm (1979) correction inside each declared family, and with negative results reported under the same evidentiary standard as positive ones. Two choices deserve to be stated rather than assumed. First, Holm and not the stepwise procedure of Romano and Wolf (2005), which dominates it in power under dependence: the stepwise method requires resampling the joint distribution of the whole family of statistics, while here the surrogate clouds are generated per market and per hypothesis, so the joint distribution is not available. The cost is power, and it is declared: under Holm a **NOT PROVEN** verdict may be a false negative, which is precisely why every negative outcome in this document is accompanied by a power check. Second, neither the deflated Sharpe ratio (Bailey and López de Prado 2014)

nor the probability of backtest overfitting (Bailey, Borwein, López de Prado and Zhu 2017) is used, and their absence is not an oversight: both correct the selection of one strategy out of many by its performance statistic, whereas the object tested here is a detector measured against a null, not a strategy measured by its Sharpe ratio. Should the framework later produce operational signals with a performance statistic attached, those corrections would become necessary and their absence would have to be argued again.

The market under study. The claim that cryptocurrencies are informationally inefficient, and therefore a plausible place to look for exploitable structure, was first made systematically by Urquhart (2016) and immediately contested by Nadarajah and Chu (2017), who reached the opposite conclusion on the same data with a power transformation of the returns. That exchange is cited here deliberately and in full: it is the clearest demonstration, inside the crypto literature itself, that the verdict follows the test, which is the reason this work measures the null value of each statistic instead of assuming it. Brauneis and Mestel (2018) show that efficiency is not uniform across the crypto universe, which is part of why the widest test here spans thirty-seven instruments in six classes rather than one asset. Caporale and Plastun (2019) test a calendar periodicity — a day-of-the-week effect — on cryptocurrencies; the object is different from the one studied here, an exogenous calendar regularity rather than endogenous multi-scale cycles, and the distinction is worth making explicitly because the two are easily conflated.

The closest methodological precedent. Brock, Lakonishok and LeBaron (1992) tested moving average and trading-range rules on ninety years of Dow Jones data against four explicitly stated null models — random walk, AR(1), GARCH-M, EGARCH — estimated by bootstrap. That design, null models declared in advance and measured rather than assumed, is the one this protocol follows; what is added here is the power requirement before a negative outcome counts, the marginal reported next to every conditional rate, and the rule that every level is measured on the timeframe that resolves it.

The epistemological frame is Popper (1959), and it is taken literally rather than decoratively: a claim earns the verdict **CONFIRMED** only by surviving a test that could have refuted it, and the document distinguishes **FALSIFIED** from **NOT PROVEN** throughout.

3. Principle of the method

Disclosure boundary. This section describes the conceptual principle of the method: what the system does, on which theoretical sources it rests and with which causal logic. It does not disclose calibration thresholds, parameter values, the calibration procedure, the set of decision rules or code. Methodology and statistical results are instead fully documented, and the dataset is provided for independent verification.

3.1. Conceptual architecture: three layers

The method splits the problem into three layers, and the separation is itself a conceptual contribution: **L1**, cycle structure — where we are in the hierarchy of cycles; **L2**, DSP timing — when the current cycle opens and closes; **L3**, order-flow confirmation, provided as an extension. Each layer answers a different question and uses different evidence: this is what makes the problem mechanizable.

3.2. Cycle hierarchy and expected cycle window

The first layer answers the question “where are we now?”. Price is decomposed into a hierarchy of nested cycles, from the dominant long-period cycle down to the minor sub-cycles, according to the principle of cyclic commonality and synchronicity of troughs of the Hurstian tradition. Candidate periods are discovered spectrally, with a criterion that does not assume stationarity over the whole window; the tracking of the dominant cycle and of its instantaneous phase follows Ehlers’s causal DSP.

The **expected cycle window** is the notion that makes everything else applicable, and it can be defined symbolically. Let P_L be the nominal period of level L and t_0 the previous trough confirmed at that level. The current cycle phase is

$$\varphi(t) = \frac{t - t_0}{P_L},$$

which is 0 at the trough, ≈ 0.5 at mid-cycle and ≈ 1 at the expected close. A new trough of level L is expected in the window

$$W_L = [t_0 + (1 - \tau) P_L, t_0 + (1 + \tau) P_L],$$

where τ is the relative tolerance that absorbs cyclical elasticity. The structure of the window — anchoring to the previous trough and a width proportional to the nominal period — is sufficient to reproduce the concept; the value of τ and the criterion for admitting the nominal period per asset remain proprietary, as do the criteria for promoting and removing levels and the alignment tolerances between levels.

3.3. Causal detection of tongues: the mathematics of the criterion

A **Bayer tongue** is an anomalously short cycle that acts as a connector: it signals a liquidity grab followed by a restart, or a transition cycle between two larger structures. The system detects it with a rule-based classifier — not an opaque model — that is causal: every evaluation uses only data up to the current instant. The logic is replicable at the level of families of criteria, organized in two steps.

Step 1 — brevity within the window. The cycle is a tongue candidate if its duration lies below a **low percentile** of the duration distribution observed *up to* t for that level and that asset. **Step 2 — causal price criterion.** A short cycle alone is not enough: the criterion, in the style of Dow and Migliorino, requires a rally of sufficient amplitude from the candidate trough and the breaking of recent highs within the window.

In symbols, let D be the duration of the candidate cycle measured between two consecutive troughs: the cycle is a candidate if $D < \delta_{\text{short}}(L)$, with δ_{short} estimated adaptively and causally on level L and on the asset. **The family of the criterion — a threshold given by a low percentile, estimated causally — is what the reader must replicate; the value of the percentile is reserved**, and it is one of the points where the disclosure boundary bites. The difference from Bayer's canonical threshold, which fixes brevity at a constant fraction of the mean, is substantial and is verified in § 6.10.

Both conditions of step 2 are causal by construction: they look at bars already closed, and the window within which they look for the breakout is determined by the level of the candidate cycle, not by its outcome.

Of the two steps, only one carries information. Step 1, isolated from step 2, **carries none** about the restart (§ 6.10): the discriminating power of the criterion lies entirely in the price step. This matters for anyone wishing to replicate the method, and it matters also as a warning — brevity serves to **select** candidates, not to **classify** them, and using it alone would give a criterion that seems to work because it decides nothing.

3.4. Confirmation at window close

A tongue is confirmed when its detection window closes, exactly as a swing is confirmed at the end of the cycle. From that moment its historical label no longer changes. The only moving entity is the last tongue whose window is still open, which by construction can update: this is expected behavior and not a causal repaint. All the results of this work rest exclusively on confirmed tongues.

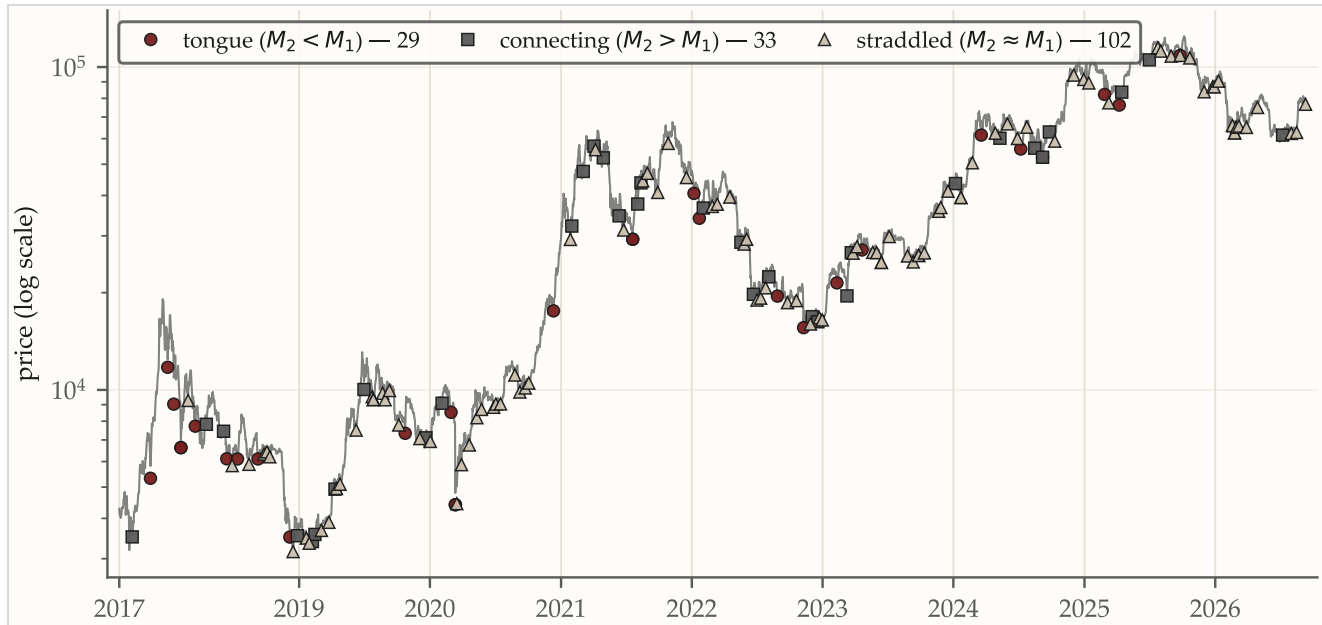


Figure 1: The tongues confirmed by the causal detector over nine years of Bitcoin, on a logarithmic scale. Each marker is a tongue at its own trough, colored by category according to the relation between the trough of the short cycle and that of the previous one. It is the only figure in this document that shows the **output** of the system instead of a statistic about it, and it serves a check that no number provides: the detection is **dense and distributed** over the 2020–2021 bull, the 2022 bear and the subsequent recovery, not concentrated in a regime.

4. Unified cycle nomenclature

The three schools use different names for the same objects. The relative scale, of Italian usage, is adopted here as the vocabulary, with the **measured** durations next to the listed nominal periods. The origin of the scale is not arbitrary and is worth recalling: **T stands for trading cycle**, the level on which one operates, and it is the term with which Migliorino introduces the numbering. The scale is relative in its distances, not in its origin. The scale **does not stop at the ten-day cycle**: it goes down to the daily cycle and below, and the column of the timeframe that resolves each level says why the daily timeframe is not enough to measure them.

Level	Name	Nominal	Resolved on	Measured duration
T+9	—	~15 years	W	outside the available history
T+8	Hyper	~8 years	W	1 duration per market: an interval, not a distribution
T+7	Super	~4 years	W	1–2 durations per market
T+6	—	480 d (<i>duplicate of T+5</i>)	D	not produced by any ruler
T+5	Long	480 d	D	Table 2
T+4	Major	160 d	D	Table 2
T+3	Minor	90 d	H4	Table 2
T+2	Short	45 d	H4	Table 2
T+1	Swift	20 d	H2	Table 2
T	Base	10 d	H1	Table 2
T-1	Micro	5 d	M30	Table 2
T-2	Nano	2.5 d	M15	Table 2
T-3	daily	1 d	M5	Table 2
T-4	half-daily	~12 h	M5	Table 2
T-5	—	~7 h	M1	Table 2
T-6	—	~5 h	M1	Table 2
T-7	—	~3 h	M1	Table 2
T-8	—	~1.5 h	none	no admitted ruler resolves it

Table 1: Unified cycle nomenclature, from the multi-year cycle to the three-hour one, with the **ruler that resolves** each level. The measured durations are in a single table — Table 2 — and are not repeated here: the duration of a cycle read where the detector does not separate it is the duration of another object, so a level without its ruler beside it is not a comparable number. At the two top levels the duration does not appear because **with one or two durations there is no distribution, there is an interval**: over nine years of history a four-year cycle fits twice, and the third time it appears on Ethereum it exists only because the detector accepted two cycles shorter than the nominal. T-5, T-6 and T-7 have their own troughs but **are not distinct from one another** (§ 5.6.1).

The “resolved on” column is not an editorial choice: it is computed by the same production rule that § 5.6 describes, looking for the timeframe in which the level falls within the resolution window and closest to its geometric center. It is the reason why a complete nomenclature **must** have that column: without it, the reader assumes that the numbers of every row are comparable with one another, and they are not.

The scale has **eighteen steps**, from T+9 (~15 years) to T-8 (~1.5 hours), and they are not all measurable from the same timeframe: **T+6 is never produced** on any of the three markets, T+7, T+8 and T+9 do not have a sufficient sample on the crypto sample alone (§ 6.11), T-8 has no admitted ruler that resolves it, and below the ten-day cycle the Daily no longer separates — there the measurement must be made where the level is resolved (§ 5.6). **Thirteen measured levels** remain, and they are the table that follows.

The list describes the measured durations well, and the deviation has a direction. Measured where the level is resolved, the median lies close to the nominal across the whole short scale — at T it is 9.5 days against 10, at T-1 4.7–4.9 against 5, at T+1 18.8–19.8 against 20 — and the deviation is almost always **downward**, by a few percentage points. The gap widens only at T+4, where the median lies between 10% and 24% below the nominal. It is the direction one expects from a detector that closes cycles on troughs and imposes a minimum duration: it cuts the long tails more than it lengthens the short ones. It is not a confirmation of the list as a model of the market — it is the observation that **our** detector, read on the right ruler, returns the catalog durations within a few points.

Level	Ruler	Nominal	Bitcoin	Ethereum	Solana	Deviation
T-7	M1	0.13 d	0.13 (4,870)	0.13 (4,931)	0.13 (4,816)	0% ... -3%
T-6	M1	0.20 d	0.18 (3,438)	0.18 (3,433)	0.18 (3,394)	-8% ... -9%
T-5	M1	0.30 d	0.27 (2,318)	0.27 (2,298)	0.27 (2,281)	-9% ... -10%
T-4	M5	0.50 d	0.48 (6,613)	0.48 (6,617)	0.48 (4,407)	-4% ... -5%
T-3	M5	1.00 d	0.94 (3,297)	0.94 (3,276)	0.95 (2,175)	-5% ... -6%
T-2	M15	2.50 d	2.37 (1,318)	2.38 (1,317)	2.43 (873)	-3% ... -5%
T-1	M30	5 d	4.81 (648)	4.74 (664)	4.94 (433)	-1% ... -5%
T	H1	10 d	9.46 (333)	9.46 (328)	9.58 (218)	-4% ... -5%
T+1	H2	20 d	19.79 (158)	18.83 (162)	19.25 (105)	-1% ... -6%
T+2	H4	45 d	39.67 (73)	44.08 (70)	41.42 (50)	-2% ... -12%
T+3	H4	90 d	88.17 (35)	96.83 (35)	90.17 (23)	-2% ... +8%
T+4	D	160 d	144.0 (21)	121.0 (21)	137.0 (13)	-10% ... -24%
T+5	D	480 d	495.0 (5)	481.0 (5)	— (4)	+0.2% ... +3%

Table 2: The durations **established by the system**, each measured on the ruler that resolves its level. Thirteen levels out of eighteen, from the minute to the day. The medians are in **days**; in parentheses the number of cycles on which the median is computed, because a median without its denominator cannot be verified. The deviation is between the median and the listed nominal, over the three markets. The Solana row at T+5 is empty because four durations do not make a distribution, and the measurement declares it instead of reporting a number. Every value comes from **CR-009**, which carries the same durations **in bars of their own timeframe** with the nominal converted into the same unit: at T, for example, 227 hourly bars against 240.

The six bottom levels enter now, and they are the concrete result of the floor lowered to one minute: up to the previous run T-2 and the five below it had no admitted ruler and were not measured. Their deviation from the nominal lies between zero and -10%, that is, in the same band as the high levels, on samples of thousands of cycles instead of tens.

T+6 does not appear. The procedure produces no level in that range of periods: between T+5, which on the three markets lies between 368 and 495 bars, and T+7, which on Ethereum is 1,153 bars, there is no intermediate step with its own troughs and spectral support.

4.1. Level T+3, and why its measurement was double

For level T+3 alone, the supporting material long contained **two incompatible measurements** produced by the same hierarchy: 28/24/19 cycles with medians of 92.5 / 104 / 89 bars from one module, 38/37/28 cycles with medians of 75 / 72 / 75 from two others. The other levels coincided. A discrepancy of this kind is not resolved by choosing the more convenient series: one isolates the cause.

The measurement isolates it in two steps. **First:** the durations of T+3 are read from the same hierarchy along the paths of the three modules, with the level filters disabled. **Second:** the filter is switched back on and one measures how much it shifts count and median. Both steps run on each of the nine rulers, and the verdict is read on **H4** — the ruler that resolves T+3 (§ 5.6).

On the ruler that resolves the level the readings coincide digit for digit: 35 cycles and a median of 529.0 bars on Bitcoin, 35 and 581.0 on Ethereum, 23 and 541.0 on Solana — in days 88.2 / 96.8 / 90.2, the same durations as Table 2. There is no second cause to look for: no divergence of hierarchy, no different trough, no different count. With a caveat that the measurement declares itself: two of the three paths are **identical by construction** — both are the difference between consecutive troughs — so the distinct witnesses are two, not three.

And the filter, on that ruler, moves nothing: 35 cycles remain 35, 529.0 bars remain 529.0, and likewise on the other two markets. The threshold is **60 bars** and the median of T+3 on H4 is **529**: a floor of sixty bars beneath a distribution that lives around five hundred meets nothing to cut.

The cause exists, and it is deeper than a parameter. The double measurement originated on the **Daily**, where a T+3 cycle lasts about ninety bars and a floor of sixty cuts part of them; on the ruler that resolves the level the same floor is invisible. It was not a disagreement between modules and it was not merely a parameter applied without declaring it: it was **a level measured on a ruler that does not resolve it**, that is, the thesis that § 5.6 measures on its own. The series to use is the unfiltered one, read on H4; the other

must be withdrawn, and the conclusions that depended on T+3 — cross-asset convergence at that level and the nesting ratio $T+3 \rightarrow T+4$ — are no longer provisional.

The verdict rule, applied to the letter, falsifies. To declare *cause established*, a reduction of at least 20% of the cycles **and** an increase of at least 15 bars in the median **on all three markets** had been fixed a priori; the rule declared the hypothesis falsified if no market reproduced the effect. On H4 the reduction is **zero**: the formal verdict is **FALSIFIED**. It must be read for what it says: not that the filter was innocent, but that **on that ruler it cannot act**. The rule, written when the measurement lived on a single ruler, does not distinguish *the effect is not there* from *the effect cannot be there*, and the distinction cannot be recovered by rewriting it after seeing the number. The historical question — did the filter explain the double measurement **on the Daily**? — is no longer computable: there the reading that goes through the set of levels no longer produces any T+3 cycle, and only the raw difference between troughs remains (32 / 34 / 24 cycles, medians of 98.5 / 102.5 / 88.0 days).

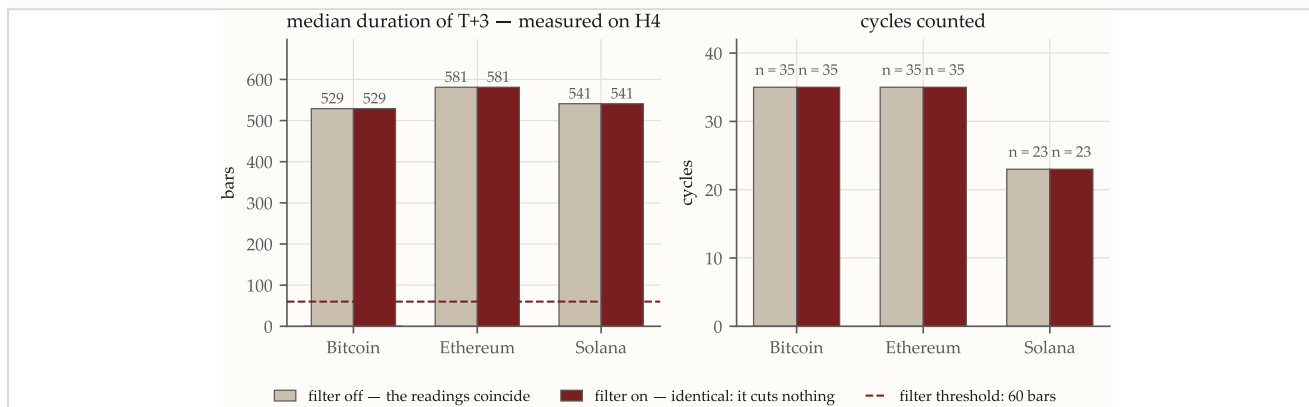


Figure 2: Level T+3 on the ruler that resolves it, read along three paths with the filter off — the readings coincide, same count and same median — and the same quantities with the filter on, which are identical: at a median of five hundred bars a floor of sixty cuts nothing. The double measurement of the earlier material was not a disagreement between modules, it was a level read on a ruler that does not resolve it.

5. Validation methodology

5.1. Dataset and coverage

Item	Value
Primary markets	BTCUSDT, ETHUSDT and SOLUSDT, since 2017 (Solana since 2020), granularity from one minute to weekly
Broad cross-asset test	37 instruments in six classes: cryptocurrencies, stocks, indices, commodities, currencies, bonds (§ 6.14)
Long history	6 equity indices: ^GSPC 24,775 bars (since December 1927), ^N225 15,149, ^IXIC 14,001, ^FTSE 10,769, ^GDAXI 9,770, ^DJI 8,719 (§ 6.11)
Timeframe of the claims	Daily for the historical framework; nine rulers — W, D, H4, H2, H1, M30, M15, M5 and M1 — for the re-measurements on the timeframe that resolves the level (§ 5.6) and for the cross-timeframe check (§ 6.13)
Atomic datum	one-minute OHLCV from an exchange source, partitioned in Parquet; every higher timeframe is derived from the minute with a certified no-repainting resampler
Data cut-off	last trade of the run that redid thirty claim cards in a single read: 5 September 2026 (§ 5.7)

Table 3: Reference dataset. The number of bars of each measurement, with the price fingerprint, is declared in its own evidence file and summarized in § 5.7.

5.2. Popperian protocol

- 1. Causality.** Every filter, feature and signal produces output at time t using exclusively data available up to t . An automated *prefix invariance* test verifies that the output at time t does not change when subsequent bars are added. It is the non-negotiable premise: a result computed with future data produces false confidence and invalidates any validation.

2. **Falsifiability.** Every hypothesis is expressed with a statistical test, a decision threshold fixed a priori and a minimum sample. The admissible verdicts are CONFIRMED, FALSIFIED, PARTIAL and NOT PROVEN. The last verdict is necessary and is not a nuance: “**refuted**” and “**not supported**” are two **different outcomes**, and confusing them passes off as refutation what is only absence of proof — or, worse, as confirmation what is only absence of refutation.
3. **Thresholds before numbers.** Every measurement declares its own verdict thresholds in the module that computes it, **before** being run, and does not retouch them afterwards. In more than one case this produces a verdict more severe than the numbers would suggest: where this happens, the document declares it openly as a fragility, stating plainly the value that would have changed the outcome (§ 4.1, § 6.3). It is the only way to keep the discipline without hiding information.
4. **Determinism and reproducibility.** Fixed seeds and locked dependencies: repeated runs produce identical output. Every published number lives in a versioned measurement file that contains the code commit, the seed, the list of datasets and the SHA-256 fingerprint of the prices. **No number appears in this document unless it exists in a measurement file**, and the files are delivered as supporting material.
5. **External validation.** Cross-asset convergence as a test of generality, and walk-forward analysis for the temporal stability of the classifications.

5.3. The null value is measured: the role of surrogates

It is the requirement that holds up everything else, and it must be stated before the results because none of them would make sense without it: **a reference value is measured, not assumed.**

Many claims of cycle analysis take the form “quantity X and quantity Y are related”. The naive test compares the observed measurement with the value it would have under independence. The comparison, however, is almost always wrong, because **the cycle detector induces a relationship on its own**: any procedure that closes a cycle on a trough produces, even on noise, long cycles that are on average wider than short ones.

The correct null value is obtained with **surrogate series**: synthetic replicas that preserve some properties of the true series and destroy others.

- **Phase-randomized surrogates**: the Fourier transform of the series is computed, the phases are replaced with uniform random values and the result is inverse-transformed. They preserve the power spectrum exactly — that is, the periodicities — and destroy the temporal position of the extremes.
- **AAFT surrogates** (*amplitude adjusted Fourier transform*): they also preserve the distribution of the values of the original series, and are therefore more severe.

The protocol is: **the same detector** is applied to the true data and to one or two hundred replicas per market and per type, the same quantity is measured, and confirmation is declared only if the observed value falls outside the surrogate cloud.

One null, however, is not always the same thing as another, and the choice decides the sign of the result. This holds in both directions, and the document meets both cases:

A null that also destroys the geometry of the detector measures the detector, not the market. A null that preserves precisely what is under discussion measures nothing.

§ 6.12 is the first case — shuffling all polarities together destroys the position within the parent cycle, which is a property of the detector, and produces a p -value of 0.0005 where the null that preserves position gives 0.71. § 6.2 is the second: surrogates that preserve the spectrum also preserve the peaks, that is, exactly the set whose arrangement is under discussion, and on that question they have no power at all.

A note on empirical p -values. Those estimated by resampling — surrogates, permutations, bootstrap — are reported with the estimator $(k + 1)/(n + 1)$ and not as the raw fraction k/n : with n draws the minimum that can be declared is $1/(n + 1)$, and writing 0.000 would declare a precision the sampling does not have. Where instead a threshold had been fixed a priori on the **empirical quantile**, that quantity remains the decision rule and the valid p -value appears next to it in the evidence.

5.4. The marginal next to the conditional

The second requirement concerns success rates, and it is the one this work learned at its own expense on two claims it had considered settled:

Every success rate must be accompanied by the rate of the same condition on the cases the event does NOT select. If the two coincide, the event selects nothing and the rate is not evidence, whatever its value.

It is the general form of the criticism that § 8.1 levels at Elico64's inverse cycle: a definition that applies only a posteriori is not falsifiable ex ante. Applied to the unconditional swing it shows that the event does not select (§ 6.7); applied to the break of the starting trough it shows that the publishable number was inflated by a **set inclusion**, that is, by a set of events that contains by construction all the positive outcomes (§ 6.8). The same safeguard, applied in the same way to two neighboring claims of the same corpus, withdraws one and strengthens the other: it is not a tool for demolition.

5.5. A negative outcome is valid only if the test would have power

The third requirement is symmetrical to the second and closes the door on the other way of being wrong:

A test that never falsifies and one that always falsifies are equally useless. Before reading a negative outcome, one must verify that the test would recognize the phenomenon if it were there, on a synthetic sample built to contain it.

§ 6.2 carries it out explicitly, and without that check its verdict would have no value: the first formulation of the harmonicity test gave a negative outcome that did not distinguish "there is no harmonicity" from "this test would not see it anyway". The same logic governs the identity and inadequacy checks of § 5.6.

5.6. Every level on the timeframe that resolves it

The fourth requirement is the most recent and the most invasive, because it does not concern a single measurement but the **ruler** with which almost all of them were taken.

A cycle level has a nominal period expressed in units of time, not in bars. On a given timeframe that period becomes a number of bars N , and N decides what the detector can see: below a certain resolution two adjacent levels are no longer two different detectors, they are the same detector called by two names. The resulting admissibility grid **factorizes** into two independent axes, and this is a measured fact, not an assumption: the number of cycles C that a level possesses is **invariant to the timeframe** for a given market — maximum relative deviation 0.0015 over 54 market/level pairs, against a declared threshold of 1% — while the resolution N depends only on the timeframe. **Going down in timeframe does not give you more sample: it only increases the resolution with which that sample is seen.**

The control for the criterion is the group of timeframes in which a cap on the bars served bites: there the history is a declared fraction of the available one and C must fall in the same proportion. It does in fact fall by 0.70–0.80 on all eighteen levels of each market, two orders of magnitude above the threshold: the criterion discriminates, and the invariance is not saturation.

Hence an **owner rule**: the timeframe that owns a level is the one in which N falls within the ideal window and closest to its geometric center. The choice is made on the N axis and **only** on that axis, because C is invariant and cannot discriminate.

Timeframe	Levels it owns	Lowest level it truly separates
W	T+7, T+8, T+9	T+7
D	T+4, T+5, T+6	T+3
H4	T+2, T+3	T+1
H2	T+1	T
H1	T	T-1
M30	T-1	T-2
M15	T-2	T-3
M5	T-3, T-4	T-5
M1	T-5, T-6, T-7	T-7

Table 4: Who owns what, with the floor lowered to one minute. Seventeen levels out of eighteen have an owner according to the N-axis rule; **thirteen** are actually measured, and they make up the table of durations. Left out are T+6 — which the procedure does not produce, see below — and the three high levels T+7, T+8 and T+9, for which nine years of history are not enough. The rule is “on the timeframe that **resolves** it”, not “on the most central one”: it changes the set of candidates, not the selection criterion.

The cap has been raised, and with it the grid. With the previous limit of 100,000 bars no ruler below the hour was admitted and six levels out of eighteen were declared not resolvable *by this grid*; with the cap at one million the four fine rulers come in and those levels find their own. **They were not unmeasurable levels: they were levels that grid did not resolve**, and it is the distinction the rest of the document uses.

The consequence is not theoretical, and it is measured. The same four statistics of a level — how many cycles, median duration, coefficient of variation of the durations, share of bearish cycles — computed on the Daily and on the timeframe that resolves that level, with the **same** production chain and on the **same** price snapshot, **disagree in 12 pairs out of 12**. The agreement thresholds were declared beforehand and were generous (10% on the median, 0.05 on the CV, 10% on the share, ratio between 0.80 and 1.25 on the number of cycles). Two controls bound the criterion from both sides: the Daily against itself agrees in 12 pairs out of 12 — the criterion does not fire blanks — and the Daily against the Weekly, where those levels are worth between 2 and 13 bars, disagrees 12 times out of 12 — the criterion is not saturated.

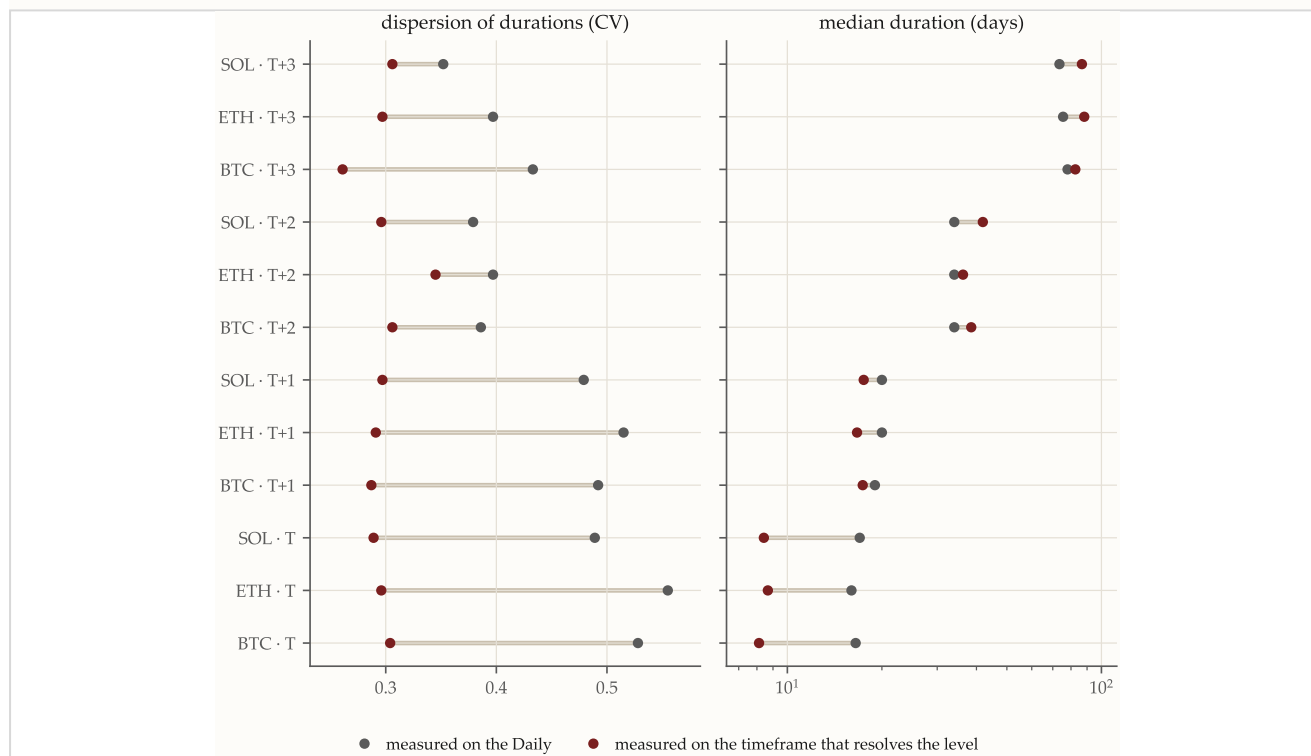


Figure 3: The same four statistics, measured on the Daily and on the timeframe that owns the level. On the left the coefficient of variation of the durations, which is higher on the Daily in **twelve pairs out of twelve**; on the right the median duration in days, where level T changes by a factor of two. The share of bearish cycles, which does not appear because it never disagrees, is the statistic that is robust to the change of ruler.

The coefficient of variation measured on the Daily is inflated, and in one direction only. The ratio between the owner timeframe and the Daily has a median of 0.61: measured where it is resolved, the same level has a

dispersion of durations **about 40% lower**. This bears directly on the reading of cycles as “quasi-fixed with an elasticity of twenty to forty percent”: part of that dispersion is not elasticity of the market, it is the resolution of the ruler. How much, this measurement does not say — it says that the ruler alone produces enough of it to shift the number by a third.

Level T measured on the Daily is not T. Its median on the Daily is 16.5, 16.0 and 17.0 days on the three markets, against 8.1, 8.7 and 8.4 on the timeframe that resolves it: a ratio between **1.85 and 2.03**, that is, a factor of two, and the number of cycles more than doubles (from 176, 175 and 114 to 374, 365 and 245). Sixteen days is not the nominal of T, which is ten: it is that of T+1, which is twenty. On the Daily level T has ten bars at its disposal, which the grid marks as not resolvable, and the detector does the predictable thing — it finds the nearest structure those bars allow it to see. It is the same fact that § 6.1 records as “T and T+1 are not distinct”, seen from the side of the cause instead of that of the symptom.

The four numbers of this comparison **must not be read together with those of Table 2**, where the same T on H1 is worth 9.46 days: that table measures with the production chain and on the current snapshot, this comparison with the level filters disabled and on an earlier snapshot. What the measurement asserts is the **ratio between the two rulers**, which is computed within the same chain; the absolute order of magnitude is to be taken from the table.

What does NOT follow from this. It does not follow that the numbers of the owner timeframe are the “true” ones. It follows that they are **different**, and that the difference is not noise: it is systematic in direction, on three markets and four levels. Which of the two measurements better describes the market is a question about the market and calls for an external criterion, which this work does not have and does not pretend to have. It also follows, very concretely, that **the polarity measurements are robust to the change of timeframe and the duration ones are not** — the share of bearish cycles never disagrees, with a maximum deviation of 0.083 against a threshold of 0.10 — and this says which results of § 6 the ruler affects and which it does not.

The document therefore keeps the numbers of the Daily **and** those of the resolving timeframe: the former as a control arm, the latter as the measurement of the level. Where the level is high — from T+4 upwards — the Daily *is* the owner, and the question does not arise.

5.6.1. THE RULE HOLDS FOR PROPERTIES, NOT YET FOR RELATIONS

The owner rule is well defined for the **properties of a level** — how many cycles it has, how long they last, how dispersed the durations are — because those quantities belong to one level only, and that level has a ruler. It is not defined, in the same way, for the **relations between two levels**: a link between a cycle and its sub-cycle belongs to neither of the two rulers, and when the two members are resolved by different timeframes the rule, taken literally, does not say what to do.

The measurements in this document that concern relations — the link between a level and its sub-level — **do not solve** the problem: they **declare** it. Cell by cell it is recorded whether the sub-level is also resolved on the same timeframe as the parent, and the verdicts are also reported on the complete pairs alone. On the run of 7 September 2026 the cells with a sub-level are seventeen, and the complete pairs are **ten**: in the other seven the sub-level would be resolved elsewhere, and the cell enters the general count but not the restricted one.

It is a restriction, not a bridge, and must be read as such: where the two members happen to be on the same ruler the relation is measured as it should be; where they do not, the relation is either measured on the ruler of the parent — and then the sub-level is seen below its own resolution — or it leaves the restricted count. The construction that measures **each member on its own ruler**, aligning the troughs on absolute time instead of on indices, is implemented and under validation; the relation measurements reported here do not go through it yet, and for this reason the restriction is declared with its number instead of being kept silent.

5.6.2. THE DAILY CYCLE IS NOT THE DAILY TIMEFRAME

One case is worth isolating, because it is the one in which the name misleads the most. In the catalog, the **daily cycle** is the level with a nominal period of one day — T-3 on the relative scale — and the half-daily is the half-day one, T-4. The timeframe called *Daily* is not theirs: it owns T+4, T+5 and T+6, that is, cycles of 160 to 480 days, and on it the daily cycle is worth **one bar**. The only timeframe that brings T-3 within the ideal resolution window is M5, where it is worth 288 bars against a geometric center of 282.8.

But the five counts are not five levels, and it is the most important result of the measurement. Applying to the newly produced troughs the same criterion as § 6.1 — bidirectional overlap above eighty percent, threshold declared and not retouched — the three shortest levels **are not distinct from one another**: on M1 T-7 and T-6 have **identical** sets of troughs and T-6 and T-5 stand at 0.949. The defensible reading is

therefore narrower than the one the count suggests: below the floor of the Daily troughs of their own exist, and the distinct levels are **two** — T-3 and T-4 — plus a third object that the catalog calls by three names.

The limit of the sample remains declared: on M5 and M1 the store covers a **window** and not the entire history, and the counts are those of that window.

5.7. Price snapshots, declared

The data are live: the current bar updates and every so often a new one arrives. A work made up of dozens of measurements run at different times therefore does not rest on **one** price snapshot, and this document does not pretend otherwise.

What is frozen, and what is not. What is frozen are the **prices**: the series has a last trade and it is declared. At the date of this writing the store contains, for the three primary markets, the trades **up to 5 September 2026**. What is not frozen are the **measurements**: no measurement was redone to align its snapshot with that of the others, because redoing a measurement to make it coincide with another means choosing the snapshot by looking at the result.

What is declared instead of a single date. Every evidence file carries, for every market and timeframe it uses, the number of bars and the **SHA-256 fingerprint of the prices**: two measurements rest on the same snapshot if and only if the two fingerprints coincide, and the check requires no trust. At the time of writing the delivered evidence files are spread over seventeen distinct snapshots; the **majority** one — that of the run that redid thirty claim cards in a single read of the prices — is worth 3,308 daily bars on Bitcoin and gathers 29 evidence files out of 53, and the next ones are worth 3,124, 3,299 and 1,894 bars. The descriptive tables of § 4 and § 6 that come from the historical framework rest on the snapshot of 15 June 2026.

Five cited claim cards lie outside that majority by construction, and they must be named: CR-021 (long history on the indices), CR-041, CR-046, CR-049 and LG-004 are redone with a runner of their own, hence with an invocation of their own and a snapshot of their own. It is not a defect that can be corrected by relaunching the collection — no run includes them — and it is the reason why the consistency check excludes them from the comparison instead of asking the impossible of them. But the exclusion must be **stated here**, because their numbers end up on the page next to the others: where this may mislead, the paragraph declares it on the spot (§ 5.6).

5.8. Non-repainting gate and determinism

Since the tongue module is the central analytical component, its causality is verified with a dedicated gate, on the real pipeline and in walk-forward mode.

The method is walk-forward on the real pipeline: bars are progressively added and it is verified that the labels **already confirmed** do not change. **Five combinations out of five pass** — the three markets on Daily plus Bitcoin on H1 and H4 — with identical runs on identical data and zero mature labels missing or in excess.

The limit is declared here and holds for the whole document: the gate covers the **Daily**, where the historical framework of § 6 is measured, and two intraday timeframes on a single market. The property is demonstrated where the results need it, not everywhere; a denser check, with the confirmation latency measured instead of assumed, shows that the hold is not uniform across all intraday combinations, and this remains declared as open work (§ 7.2).

6. Statistical results

The numbers in this section come from the versioned measurement files, each with its own declared price snapshot (§ 5.7), and can be recomputed from the supporting material.

6.1. How many levels really exist

Before asking whether cycles nest by two or by three, one must know how many levels exist. A level exists, in the data, if it has troughs of its own and a period distinguishable from that of its neighbors. The falsification criterion is declared beforehand: **a level whose troughs coincide with those of a neighbor in more than eighty percent of cases in both directions is not a distinct level**.

Bidirectionality is the part that matters. In a synchronized hierarchy every trough of a high level is by construction also a trough of all the levels below: the overlap from parent to child is always one hundred

percent and distinguishes nothing. The information lies in the opposite direction — how many troughs of its own the child has — and if that too exceeds the threshold, the child is not a level but another name for the same object.

A level turns out to be a duplicate of its neighbor only when it is read on a ruler that does not resolve it. Measured with each level on the timeframe that resolves it — and the children lent by theirs, according to the bridge of § 5.6.1 — the scale contains no non-autonomous pair: **thirty-eight adjacent pairs, zero flagged**, with maximum overlap 0.71. The same hierarchy read entirely on the base ruler, which is the earlier form, flags **forty out of one hundred and eighty-four**; and each of those forty has at least one member that ruler does not resolve. The count from both sides:

	flagged as duplicate	not flagged	total
the ruler resolves both levels	0	45	45
at least one level below its resolution	40	99	139

Table 5: The adjacent pairs of three markets on nine rulers, by the declared criterion and by a property that does not look at the troughs — whether or not the ruler resolves each member. Fisher's exact test, direction declared beforehand: $p = 2.6 \cdot 10^{-6}$.

The two conditions are not symmetric. “The ruler does not resolve it” is **necessary and not sufficient** — forty pairs out of one hundred and thirty-nine — while the opposite direction has no exceptions: when the ruler resolves both members, the pair is never flagged. And the two clouds do not touch, 0.27–0.71 against 0.81–1.00: the threshold falls in the gap between them, and raising it to 0.85 does not change the conclusion. **With the floor lowered to one minute the result has strengthened without changing shape**: the pairs with both members resolved go from twenty-one to forty-five and the flagged ones remain zero, while the p -value drops by two orders of magnitude.

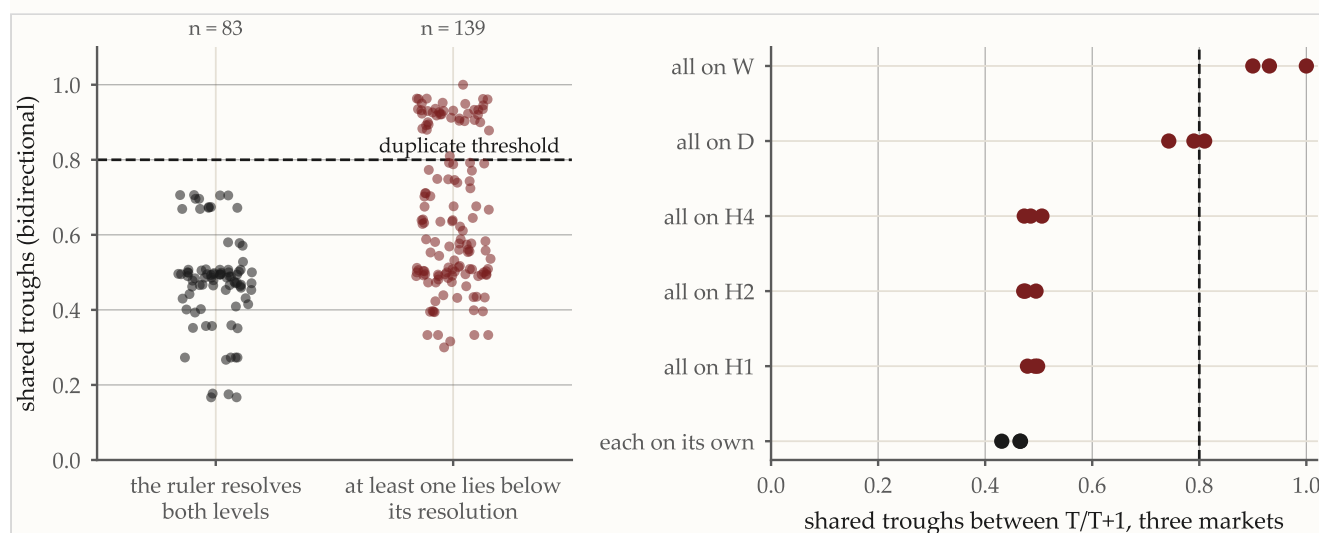


Figure 4: On the left all the measured adjacent pairs, separated by a property that does not look at the troughs: the two clouds do not touch and the threshold falls in the gap between them. On the right the T/T+1 pair, read on each ruler and on the three markets — above the threshold where the ruler resolves neither, at half where it resolves at least one.

The case of T and T+1 is the clearest demonstration. The same pair, on the same three markets and with the same function, gives 0.90–1.00 on the weekly, 0.74–0.81 on the daily — which resolves neither one nor the other — and collapses to 0.47–0.51 on every ruler that resolves at least one of them, down to 0.43–0.47 in the bridge reading. The ratio between median durations moves with it: 1.24–1.35 on the daily, **1.99–2.10** with each level on its own. The two that the hierarchy requires is there, and it appears as soon as one stops reading it on the wrong ruler.

And it is not a crypto fact. On the six long-history equity indices, where at the high levels the sample is ten times larger — up to 251 and 64 cycles on the S&P 500, against 22 and 6 on Bitcoin — none of the fifteen pairs with both members resolved turns out to be a duplicate, in either of the two duration conventions.

But the eighty percent threshold still has to be compared with a null value. It is an inherited threshold, and until it is measured against what the detector produces on noise it remains a choice. The check is the same as all the others: the production function applied to the true hierarchy and to the one the same detector produces on one hundred phase-randomized surrogates and one hundred AAFT per market. The first round measured it on the daily, and it is the table that follows.

The direction of the conclusion, and its scope. On the daily T and T+1 are indistinguishable, and there the detector does not separate them even on noise: the eighty percent threshold is exceeded by 98–100% of the surrogates, and the observed value lies even *below* the median of the cloud. On that ruler, therefore, T+1 should not be offered: not because the market does not contain it, but because the daily does not have the resolution to see it. Where the ruler has it — from H2 down — the same two levels turn out distinct, with ratio 2.0 and overlap halved. And on the chain of the run the thirty-three rows that a ruler owns entirely all have margin, the worst of **0.096** — Ethereum on T-7/T-6 of M1, with Solana at 0.097 and Bitcoin at 0.099. The verdict stays NOT PROVEN, but **not because of coverage**: in none of the thirty-three does the null exceed the threshold, and the blocker is that margin, which falls below the 0.10 declared before the measurement. The distinction matters, because coverage can be widened by measuring more, whereas a margin below the threshold is an outcome.

The validity criterion of a level must therefore be declared **per timeframe**, like the catalog of periods: the same threshold, on the same levels, is empty on the weekly, undecided on the daily and has 29–32 points of margin on H4 and H1.

T+6 does not exist and T+8 is not measurable on crypto. The first is a duplicate in the catalog — it has the nominal period of T+5 — and no ruler produces it (§ 4); the second coincides with T+7 on Bitcoin, where both have only two troughs. On the long-history indices, instead, T+8 counts between four and nine cycles on the weekly: it is a limit of history, not of scale.

Verdict. Of the nine steps of the inherited scale, none turns out to be a duplicate of its neighbor when measured on the ruler that resolves it. Two do not exist as autonomous levels **on the daily timeframe** (T+1 and T+6) and two have no sample on crypto (T+7 and T+8). The scale must be used as a vocabulary, not as an inventory: levels must be declared per market **and per timeframe**, not presupposed — and their autonomy must be measured where they are resolved, or one measures the ruler instead of the market.

6.2. Nesting: measured ratios and the principle of harmonicity

Hurst states that the periods of the waves stand to one another in constant harmonic ratios, typically two; the Italian school also admits three. It is the most exposed claim of the tradition, and it is measured in three independent ways.

Second way: how many sub-cycles each parent cycle actually contains.

The two measurements agree, and the table is carried cell by cell in the evidence files. **The two-part structure is the most frequent at every level up to T+4 on all three markets** (49–75% of parent cycles) and the ratio between durations stays between 1.62 and 2.27 up to the step T+3 → T+4. The ratio moves away from two only at the last measurable step, T+4 → T+5 (2.42–4.11), where the parent cycles are **four or five per market** and Solana still shows the two-part structure.

The form that has power: the scale, not the pairs. If the periods are harmonic, there exists a **fundamental** P_0 such that the peaks fall near the rungs of a scale — binary ($P_0 \cdot 2^k$, Hurst's doubling) or integer ($P_0 \cdot n$, the small integer ratios). The statistic is the mean distance of the peaks from the nearest rung, in octaves, minimized over the fundamental; the null is a set of periods drawn **from the same grid as the detector**, with the same cardinality. And the **power check comes before the result**: the same chain on a synthetic series with components at 20, 40, 80 and 160 bars **recognizes the scale** — distance 0.0000 octaves on the integer scale, beating chance in 100% of comparisons, and 0.0233 on the binary one, in 99.85%.

On the daily ruler, where the measurement was born, the scale is not there. Bitcoin beats chance in **92.95%** of comparisons and Ethereum in **93.6%**, against a declared threshold of 95%; on the binary scale both lie **below** the median of chance, Bitcoin lower than nine draws out of ten. Solana produces only three peaks and does not enter the count. On the six intraday rulers the integer scale seems to give the opposite outcome, and the verdict below says why it is not a confirmation.

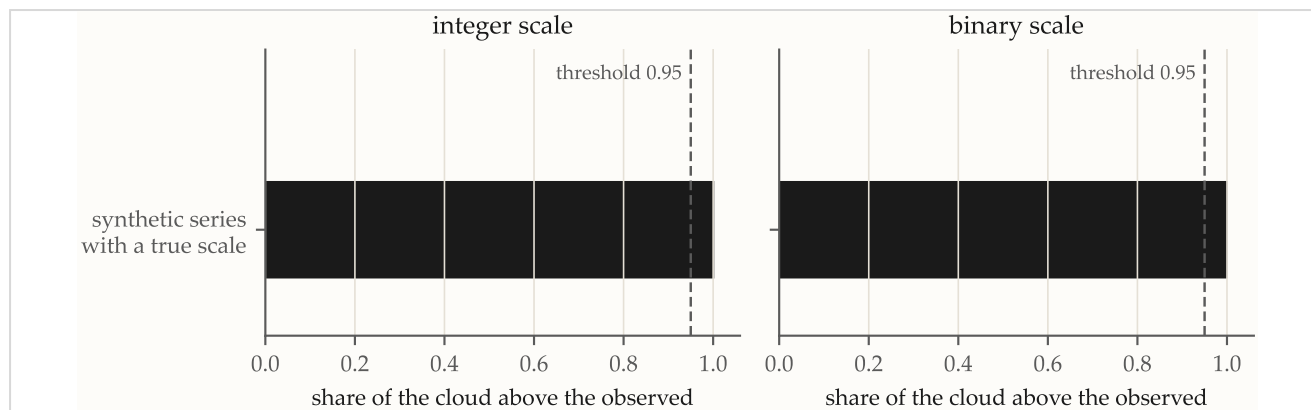


Figure 5: Where the observed distance falls within the null cloud, by market and by scale. The quantity is the share of the cloud that lies **above** the observed value: since the statistic measures a distance from the rungs, the longer the bar, the more harmonic that market. The dashed line is the threshold declared beforehand. In black the **power check**: the same chain on a synthetic series built with components at 20, 40, 80 and 160 bars.

An objection to the null, measured instead of argued. Free drawing from the grid can produce adjacent periods, something a prominence detector never does, and two nearby periods almost always fall near the same rung: the free null could turn out more harmonic than it should. Imposing on the null the same minimum separation as the observed peaks, Bitcoin falls from 92.95% to 92.3% and Ethereum rises from 93.6% to **95.8%**, above the threshold. This null is declared as a robustness check and does not enter the verdict; if it did, the daily ruler would move from falsified to not proven — one market out of two — and not to confirmed. It must be written down, because half a point separates the two verdicts.

On the intraday rulers the integer scale seems to reverse the outcome, and it is the null. The same measurement — same null, same threshold — on the six rulers from H4 to M1 remeasured with this floor finds sixteen to twenty-four peaks per market, and the integer scale beats chance in **eighteen contexts out of eighteen**; the binary one in none. Before reading it as a confirmation there is a competing explanation to rule out: on those rulers no peak lies below 53 bars, while the null draws from the whole grid between 10 and 500 — almost half of the nodes lie below that threshold — and in octaves the rungs of an integer scale grow denser as the period lengthens. A null that occupies a band the peaks do not occupy makes whatever sits high look more harmonic. The measurement that separates the two readings is the same one with the null restricted to the band of the peaks, declared before running it: the integer scale stays above the threshold on **one ruler out of six** (M1, two markets out of three), and the near-pass of the daily ruler disappears too — Bitcoin falls from 92.95% to 49.35%, Ethereum from 93.6% to 40.15%. The eighteen out of eighteen was the geometry of the null.

The four numbers in this paragraph come from the measurement with the restricted null **redone on 12 September** on the remeasured evidence, whose fingerprint the result declares: this time the fingerprint matches the file it was read from. The denominator is **six** and not seven because the floor of the remeasurement does not include M5 — a ruler that in the previous measurement already sat among the unconfirmed ones, and whose departure indeed moves no verdict: ruler by ruler the results are identical. **The outcome, however, falls exactly on the threshold**: the declared rule asks for five unconfirmed rulers, and there are five out of six. A single ruler more on the confirming side would have made the outcome indeterminate. It is said here because a verdict that passes by one unit is not a robust verdict, and declaring it is more honest than leaving it to be inferred.

And there remains a choice that no measurement can arbitrate, and it is the most important one in this section: the null draws from the detector's **grid**, not from a continuum. With continuous log-uniform drawing the integer scale turned out confirmed on both markets. The grid is the correct choice — the observed value can only fall there, and a null that lives elsewhere measures the grid instead of harmonicity — but the outcome of this measurement rests on an argument of principle, not on a numerical margin.

Verdict. Fixed harmonic nesting is falsified as a law: the measured ratio is not constant and the two-part structure, although dominant, does not exceed 75% of parent cycles at any level. And the **principle of harmonicity is FALSIFIED on the daily ruler**, not merely unsupported: with a test whose power is verified on a true scale, **zero markets out of two** turn out more harmonic than a set of periods drawn at random from the same grid. Doubling does not pass on any of the six measured rulers; the integer scale, which passes on all the intraday ones against this null, passes on only one against the null restricted to

the band of the peaks: the advantage belonged to the null, not to the markets. The difference between the two verdicts is not formal — *unsupported* means that one does not know, *falsified* means that one has looked with an instrument capable of seeing. The operational substitute is not a second law, but a measurement: **the ratio must be estimated per pair of levels and per market**, and on the scale examined it stays close to two as far as the sample allows it to be measured.

And there is no pivot. The idea that the ratio is two below a certain level and three above — a “bimodal nesting” with a transition step — is the first thing the numbers suggest to the eye, and the measured ratios and the structure counts do not support it: the transition to three is observed **only at the last step**, where the sample is not enough to distinguish it from noise.

6.3. The distribution of durations: a single mode and a tail of short cycles

The claim under examination. That cycle durations are **bimodal**: at every level a group of short cycles and one of long cycles, two distinct populations instead of a single one. It is an attractive claim, because a two-humped distribution suggests two regimes and therefore an operational rule to tell them apart.

The test, and why this one. The question “how many modes does this distribution have?” has a test of its own, **Hartigan’s dip**, which measures the maximum distance between the empirical distribution function and the nearest unimodal distribution. It must not be confused with a test of **uniformity**: a Kolmogorov–Smirnov statistic against the uniform ramp answers “is this distribution flat?”, and a single-mode distribution — which is anything but flat — maximizes it. The two questions have opposite answers on the same data, and it is an easy confusion to make: Hartigan’s dip lives by construction in $[0, 1/4]$, so **a value outside that interval is the telltale sign that one is looking at another statistic**.

Hartigan’s dip on the duration series: two rejections out of thirty-nine, and they fall in the same place. Measured with each level on the ruler that resolves it, the test covers **39** asset/level combinations on nine rulers. Thirty-seven do not reject unimodality, with p between 0.064 and 0.988. The two that reject lie **on the same level and on the same ruler** — T-4 on M5 — and the third market, on the same cell, is at the edge.

The fragility must be stated here, because it touches the verdict. The metric declared before the measurement was “rejection of unimodality on at least one asset/level combination”: on the Daily alone, with 18 combinations, none rejected and the **FALSIFIED** verdict followed from the letter of the criterion. With 39 combinations two reject, and **the letter of the criterion is met**. Two facts stand on the other side, and they are declared without being used to choose: with 39 tests at 5% two rejections are as many as one expects by chance, and with the Holm correction — threshold 0.00128 for the smallest — **none** rejects. But the two are not scattered: they fall on the same cell (T-4 on M5) of two markets, with the third at 0.052. The threshold has not been retouched and the verdict has not been changed: it remains **FALSIFIED** with this fragility written down, and the question that follows — whether at that level two populations of durations really coexist, or whether it is the delimitation that produces them, as for the short-cycle tail — is settled with the same measurement on the surrogates of the same detector.

What remains, and is measurable. The distribution has a single mode and a **tail of short cycles**, whose size is stable across markets and grows with the level. The threshold — 0.65 of the nominal duration — is declared beforehand, and the two groups are built by a known rule, not found in the data.

Level	Bitcoin	Ethereum	Solana
T	3.0%	6.1%	—
T+1	18.8%	19.0%	16.8%
T+2	32.9%	35.0%	35.3%

Table 6: Share of cycles that close below 0.65 of the nominal duration of their own level.

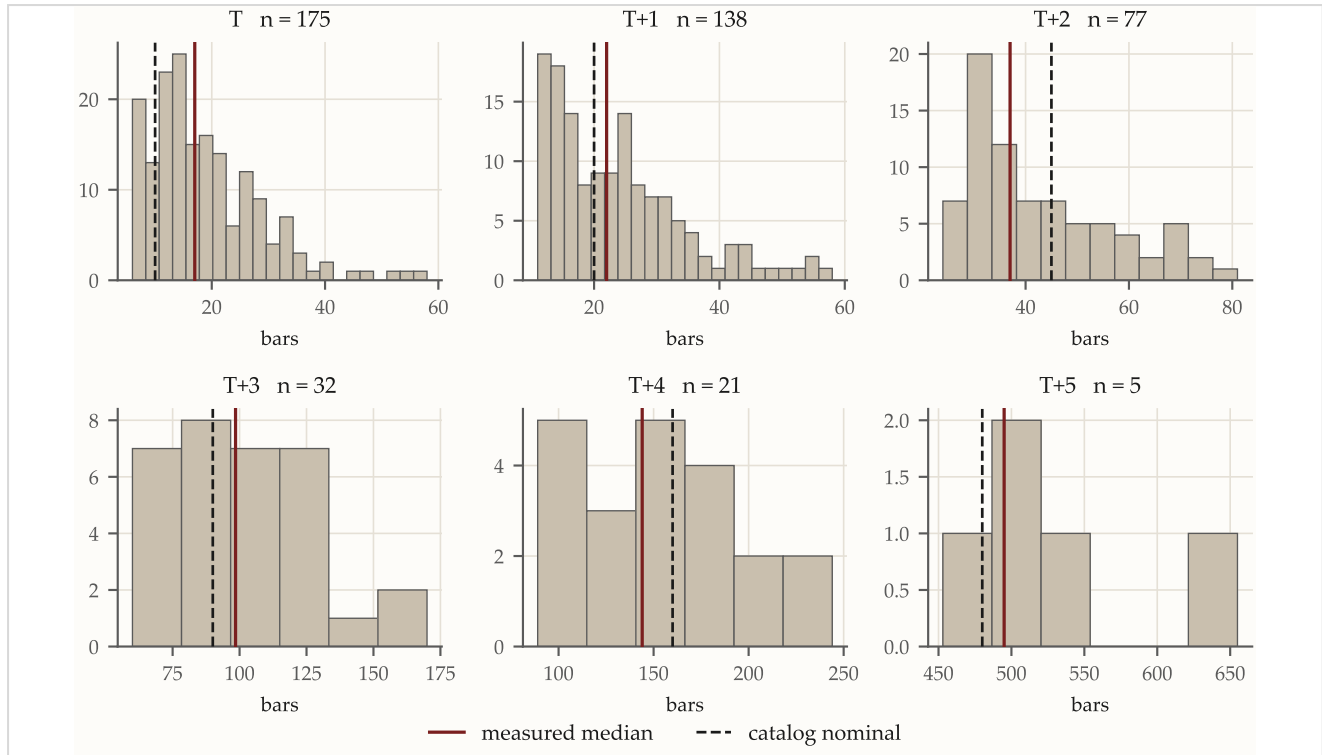


Figure 6: The distribution of the durations of each level, with the measured median and the catalog nominal. A single mode and a left tail: none of the six shows the two separate groups that bimodality required.

Verdict. Bimodality of durations is **FALSIFIED**: it is absent in thirty-seven cells out of thirty-nine, and where the sample is that of the daily it is absent at every level. The distribution has **a single mode with a tail of short cycles**, and the share of that tail grows with the level in an orderly way on all three markets. **The two exceptions are named and not absorbed**: they lie on the same T-4/M5 cell of two markets, the third is at the edge, and the fragility that follows for the letter of the criterion is declared above. There also remains the **systematic deviation** between measured median and nominal period (§ 4), which is however a property of the **measurement system** and not of the market: see § 6.3.2.

6.3.1. IS THE SHORT-CYCLE TAIL THE MARKET'S OR THE DETECTOR'S?

Table 6 describes an orderly fact: the share of short cycles is stable across markets and grows with the level. It is the kind of regularity one is tempted to claim as a discovery — and the question that § 5.3 requires to be asked **before** claiming it is what its null value is.

It is measured by applying the detector, with the production arguments, to one hundred phase-randomized surrogates and one hundred AAFT per market, with each level on the ruler that resolves it, and comparing the observed share with the cloud that comes out. The thresholds are declared before execution: the tail belongs to the market if it exits the cloud in at least 60% of the asset/level combinations, it belongs to the detector if it exits in no more than 10%.

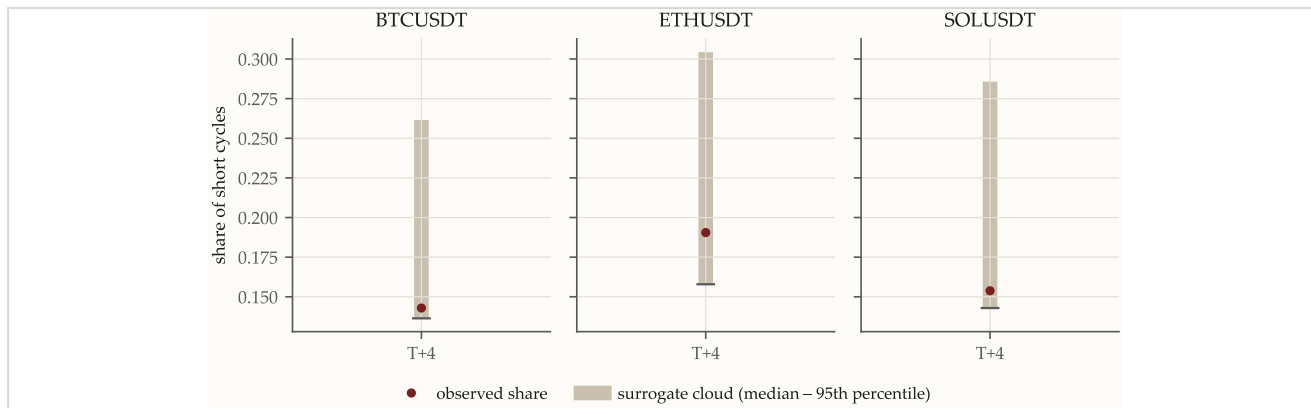


Figure 7: The observed share of short cycles (dot) against the cloud produced by the same detector on series without temporal structure (band: from the median to the 95th percentile), level by level on the three markets, with the ruler that resolves each level under it. Of the two surrogate families — phase-randomized and AAFT — the **more severe** one is drawn, that is, the one with the higher band. One dot out of thirty-three exits above its own band — Ethereum at T-1, on the M30 ruler —; the others lie inside or below, with no direction.

The outcome. On the chain — eleven levels for three markets, thirty-three combinations — **one** exits the cloud against both surrogate families: Ethereum at T-1 on M30, with a share of 0.1551 against a 95th percentile of 0.1367 (phase) and 0.1328 (AAFT). The share is $1/33 = 0.030$, below the declared falsification threshold. On six of the seven evaluable rulers none exits; only M30 has one, and read on its own it would give NOT PROVEN. Bitcoin at T-5 on M1 exits against a **single** family, and the verdict for M1 is FALSIFIED; Bitcoin at T-4 on M5 **is not measured**, because M5 is outside the floor of this run: it is a missing datum, not an exception. The observed value lies below the median of the more severe family — the more severe is the one with the higher band, as in Figure 7 — in 18 combinations out of 33: the market does not systematically produce more or fewer short cycles than the synthetic. The definition must be stated because the count rests on it: taking instead the family with the higher median, the combinations would be 19. The first round, on the daily alone, gave 0 out of 15 and the observed value below the median at T+2 and T+3 on all three markets; but the daily does not resolve those two levels, and on the ruler that resolves them the direction disappears.

Verdict: FALSIFIED. The short-cycle tail **is not a property of the market**. It is what the detector produces anyway: it exits the cloud once out of thirty-three, and otherwise lies above or below the surrogates' median with no direction.

It is the result this work would have had most interest in keeping — it was listed among the regularities that earlier schools do not document — and it is the reason why the null value is measured instead of assumed: **the wrong null hypothesis does not get the number wrong, it gets the conclusion wrong**. It also explains why the other three hypotheses failed: if the tail is not a market phenomenon, no market cause can explain it.

The tail remains a **correct description of the output**, useful to whoever calibrates, and the shares in Table 6 are the right ones. It leaves the list of market regularities and enters that of the properties of the phasing procedure.

6.3.2. DEVIATION FROM THE NOMINAL: SAME QUESTION, SAME ANSWER

That median durations lie **below** the nominal period of the catalog is, in this work, the most frequently repeated fact: the nomenclature of § 4 shows it at all the low levels, and it comes naturally to read it as a flaw of the inherited catalog — durations conceived for the equity market of the 1960s, applied to a market that did not exist.

The question, however, is the same as for the short-cycle tail, and it must be posed with the same instrument: is the ratio $r = \text{median}(\text{duration})/P_{\text{nominal}}$ what the detector would produce **on a series that has no cycles at all**? The primary null is a cloud of paths with zero-mean Gaussian increments and dispersion estimated from the asset's log-returns — no cycle, no drift — with the close alone on both sides, because a synthetic series has no bar highs and lows, and comparing it with the real ones would measure the presence of bar extremes.

On six long-history equity indices and three crypto markets, the observed value falls **inside** the cloud without cycles in **28 cells out of 37** (share 0.757; the threshold declared beforehand was 0.667). And where it does not — T+2 and T+3, six cells out of nine each — the direction is the **opposite** of the one that would

strengthen the market thesis: the observed value is **higher** than the null (0.733–0.800 against medians of 0.544–0.600), that is, the market produces **longer** cycles than noise, not shorter ones.

Verdict: the deviation from the nominal belongs to the detector. The lattice exists — the durations really do fall below the catalog — but what imposes it are the floors and proportions of the detector, not a property of the markets. The nominal catalog still needs revising; **the reason to revise it is no longer “markets are faster”**, and it is a distinction that changes who has to do the work.

6.3.3. EHLERS'S LOCAL STATIONARITY, WITH A TEST OF ITS OWN

The absence of bimodality would have been tempting to use against Ehlers, and it cannot be: bimodality, had it been there, would have contradicted local stationarity, but its absence does not confirm it. A test of its own is needed, and it has been built.

It must be built on the form Ehlers actually uses. Not “cycle duration is constant”, which nobody holds, but: **within the working window of an adaptive filter the dominant period is stable enough to be treated as constant.** The operational form is the drift of the dominant period across sliding windows — windows of 750 bars with step 100, roofing filter on log-prices, Goertzel spectrum on a log-spaced grid of 60 periods between 10 and 250 bars, MAD/median statistic of the dominant periods — compared with 100 phase-randomized surrogates and 100 AAFT, which are spectrally stationary by construction.

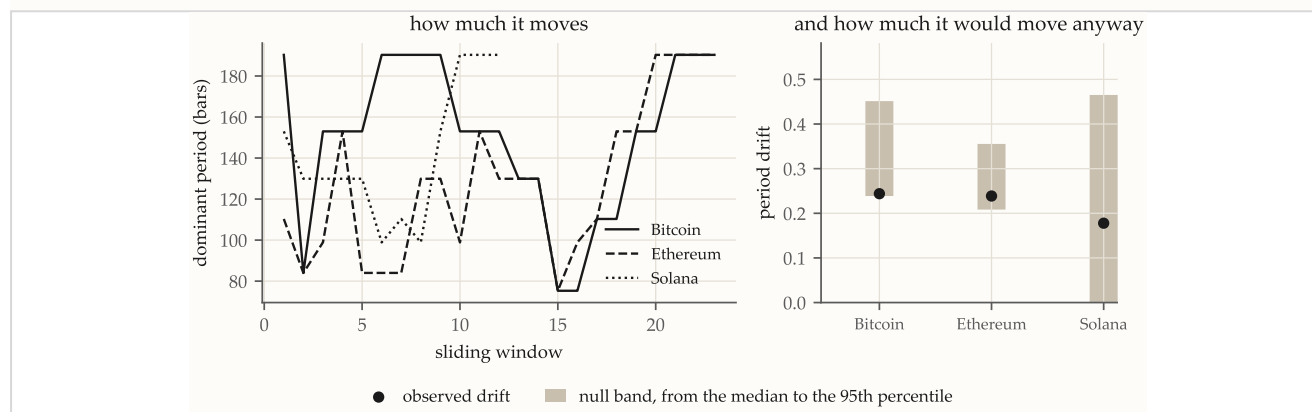


Figure 8: At the top the dominant period measured window by window on the daily: it moves, and a lot. At the bottom the drift set beside the band that the same detector produces on series without structure, ruler by ruler, and for each ruler Bitcoin, Ethereum and Solana from the left. The question is not whether the period moves — it moves — but whether it moves more than it would anyway: on the daily it does not; on the intraday rulers the drift lies below the cloud in 14 cases out of 21, inside in 6 and above only once, Bitcoin on M5. The color is the claim card's verdict for that market; the drift is a statistic with discrete values, and a marked point can sit on the edge of the band.

Verdict NOT PROVEN on the daily, CONFIRMED on five rulers out of nine. On the daily none of the three markets exits the cloud, neither above nor below: the measurement does not distinguish the market from a spectrally stationary process, and the verdict is NOT PROVEN. But the same measurement was redone on nine rulers, and that is where the picture forms: on H2, H1, M30, M5 and M1 the observed drift lies **below** the cloud in two or three markets out of three — the dominant period is more stable than what the same detector produces on series without structure — and the rule declared beforehand (two markets out of three below the cloud) gives **CONFIRMED**. On H4 and M15 only one market, hence NOT PROVEN; on the weekly the series is not long enough.

The direction matters more than the count, and it favors Ehlers: on no ruler and in no market does the drift lie **above** the cloud. There is not a single context in which the dominant period drifts more than a stationary process — which is what would falsify the hypothesis on which adaptive filters are built. Where the measurement has a sample (the intraday rulers, where the windows number hundreds instead of twenty) local stationarity not only holds: it turns out **tighter** than chance.

The limit, declared: the scale on which the measurement has the most power is also the one on which the filter works, so the result says that *within the working window* the assumption holds — not that the dominant period of a market is stable in the long run.

Declared fragility. The period grid is discrete, so the statistic lands on a lattice of possible values, and on two of the six clouds the surrogate median is exactly zero — over half of the surrogates have a constant dominant period across all windows. With a finer grid the null distribution would be smoother and the *p*-values could

shift. It is the main fragility of this measurement. A second one is added, of a different nature: a roofing filter has a long transient, and a spectrum computed on too short a window carries its settling instead of the signal. The windows used here are long enough to make it negligible, but the same measurement on short windows would not be readable.

6.4. Does duration depend on the market regime?

The hypothesis is the author's: short cycles would belong to market phases different from those of long cycles, with the direction **inverted** relative to intuition — compression in sideways markets, lengthening in trends. It is the hypothesis that, among the author's own, looked most likely to hold.

The first formulation compared the volatility of short cycles with that of long cycles, that is, it **built the regime starting from the cycles it wanted to explain**. Here the regime exists before the cycles. Two axes, both functions of prices alone: realized volatility over 30 bars and trend strength over 90 bars, $|\log(P_t/P_{t-90})|$. Both are transformed into tertiles **causally** — at time t the two cut quantiles are computed on the history up to t only, with a warm-up period and no percentile computed on the whole sample and applied backwards. Each cycle receives the tertile of its own opening trough and is classified as short if it lasts less than 0.65 of its own nominal — the threshold already declared in § 6.3, not retouched.

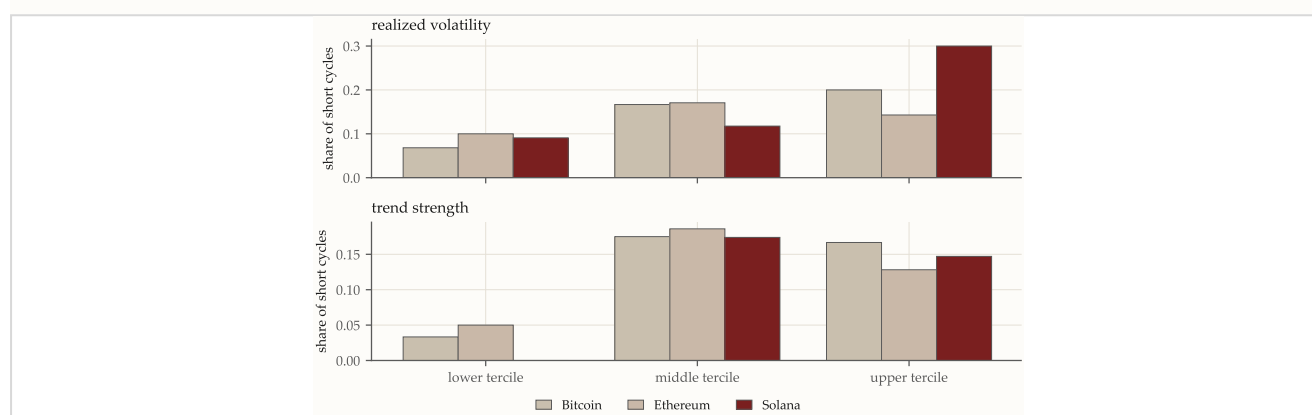


Figure 9: The share of short cycles by regime tertile, on the two exogenous axes and on the three markets, measured on H4. The columns rise together — the direction is the same in six cells out of six — but the difference between the first and the third column never exceeds the noise: none of the six is significant. It is the image of a hypothesis that has the right shape and lacks the strength.

Verdict NOT PROVEN. On H4 none of the six combinations is significant: the p -values lie between 0.17 and 0.65. The measurement was repeated on all nine rulers — 48 cells in all — and the picture does not change: **three cells out of 48** fall below 0.05, fewer than chance produces at the same threshold, and no ruler reaches the declared rule. Two rulers (H2 and the minute) give **FALSIFIED**, five **NOT PROVEN**, and on the weekly and on the daily the sample is not enough.

One thing, however, must be said, and it is not a result: on H4 the direction is the same in **six cells out of six** — more regime, more short cycles — and the same happens on the five-minute ruler. On the other five rulers the direction splits (two or four cells out of six). A concordance of sign without a single significant cell does not support the hypothesis: it is recorded because it would be dishonest to keep quiet about it, and because it indicates where to look if one day it is to be remeasured with a statistic that uses the sign instead of ignoring it.

6.5. Amplitude and duration: the slope, not the correlation

The hypothesis, original in its quantitative formulation, is that cycles behave like an accordion: greater amplitude, greater duration. It is consistent with the literature on *duration dependence* (Lunde and Timmermann, 2004; Claessens, Kose and Terrones, 2011) and with Hurst's Principle of Proportionality.

The rank correlation between the two quantities, measured on the system's cycles, is strong and replicated: Spearman's ρ +0.724 on Bitcoin (452 cycles), +0.765 on Ethereum (445) and +0.725 on Solana (300). **The number is solid; the inference one draws from it by instinct is not.** Comparing that ρ with the null hypothesis of independence — the natural reflex — is here the wrong hypothesis (§ 5.3): **the detector induces the correlation even on noise**, so the expected value under chance is not zero, and must be measured. Compared with the surrogate cloud, the picture changes.

The comparison, over ten asset/level combinations: the surrogate median is $+0.554 - +0.660$, that is, **as much as the observed value or more**, and in **eight combinations out of ten** the observed value lies **below** that median. The only combination that leaves the cloud is Bitcoin at T+2, one out of ten: as much as one expects by chance. The rank correlation does not distinguish a market from a synthetic series with the same spectrum.

Proportionality, however, has a second formulation, more severe and more informative: the **slope of the relationship in logarithmic coordinates**. A pure diffusion gives a slope of one half; full proportionality gives a slope of one. The null value, however, is neither one nor the other: it must be measured, and on the surrogates it is $0.61-0.83$.

The observed slopes lie between 0.79 and 1.17 against a cloud at $0.61-0.83$: **they exceed the surrogate median in ten combinations out of ten**, and fall in the upper 1–2% in five. The consistency of sign over ten independent measurements is worth as much here as the significance of the individual ones.

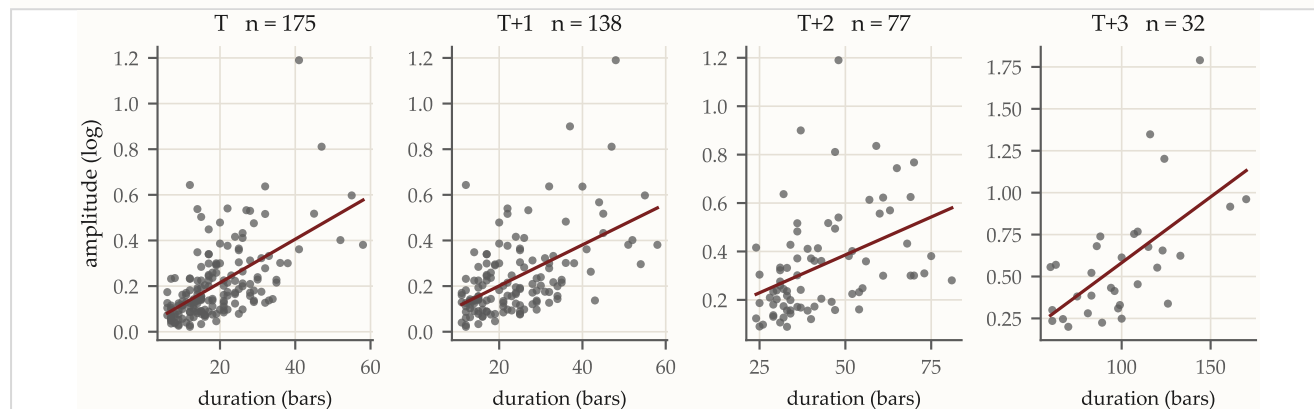


Figure 10: Duration and amplitude of every cycle, level by level, on Bitcoin. The line is the regression on the points of that level: the relationship exists and the slope is positive at every scale. It is the **slope** that is the quantity being measured, not the correlation coefficient, which with these numbers of cycles would be high almost by construction.

6.5.1. THE BULL/BEAR ASYMMETRY OF DURATION

What remains to be closed is the piece that proportionality carries with it: do cycles that close bullish last longer than those that close bearish? The descriptive measurement gave a ratio of $1.08-1.19$ depending on the market, and had never been compared with the surrogates.

The null hypothesis is not “ratio equal to one”: a series with positive drift produces longer bullish cycles even without structure, and how much longer must be measured. The comparison uses the elementary detector applied identically to the real series and to two hundred phase-randomized surrogates plus two hundred AAFT per level, and each level sits on the ruler that resolves it: from T+4 on the daily to T-7 on M1.

The verdict is **NOT PROVEN**, and the reason is not in the count. On the chain — twelve levels for three markets — **22 combinations out of 36** come out of the cloud against both surrogate families, and ten sit at the floor of the p -value that two hundred surrogates allow, 0.005 . But almost all of them come from the short levels: **4 out of 15** from T+4 to T, **18 out of 21** from T-1 down, none on the daily and on H2. Aggregated over the rulers the share is 0.611 , a whisker above the confirmation threshold (0.60): one combination fewer would give NOT PROVEN, and two of those that come out have a p -value of 0.045 . Read on the daily alone it is zero, that is, **FALSIFIED**. The outcome depends on how the rulers are aggregated, the rule declared beforehand did not say, and an aggregation is not chosen after seeing.

There is also a stronger reason not to read the 22 out of 36 as a confirmation. Both nulls — phase randomization and AAFT — are **time-reversible**: by construction they destroy any difference between the rise and the fall. A market whose returns fall faster than they rise comes out of the cloud anyway, even if its cycles have nothing asymmetric about them. With this null a confirmation is almost guaranteed; an exclusion, on the contrary, is informative, because the null errs in favor of the claim: on the daily and on H2, where no combination comes out, the asymmetry is absent even with a null that gives it away. And read as a symmetric question — “is there asymmetry?” — no combination comes out in the opposite direction. What the measurement supports is therefore narrower: **at the short levels, cycles that close up last longer than a reversible series with the same spectrum produces**. Whether this is a property of the cycles or of the returns will be decided by a null that preserves the asymmetry of returns and breaks the cyclical structure — a block shuffle —, declared with its criterion before running it.

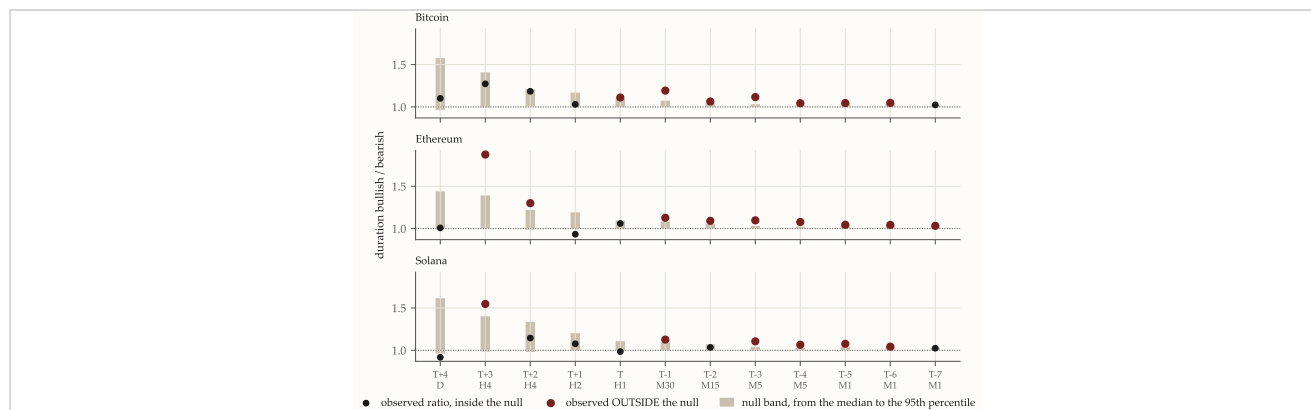


Figure 11: The ratio between the median duration of bullish cycles and that of bearish cycles, level by level on the three markets, alongside the band that the same detector produces on noise; under each level, the ruler that resolves it, and the line at one is symmetry. A point is marked outside the band only if it is so against **both** surrogate families, which is the criterion of the claim card. From T-1 down almost all come out, on the daily and on H2 none — and the null, being time-reversible, cannot tell an asymmetry of the cycles from one of the returns.

Verdict PARTIAL on proportionality. The amplitude-duration relationship is **stronger than diffusion**, and in this sense Hurst's Principle of Proportionality holds; but the measure that shows it is the slope, not the rank correlation, which is indistinguishable from chance. A caution: the comparison with the surrogates uses an elementary detector, because the proprietary one cannot be applied to synthetic series without exposing its implementation. And the bull/bear asymmetry, measured on the chain, is **not proven**: it comes out of the noise at the short levels, against a null that cannot exclude it.

6.6. Translation as a predictor of polarity

The position of the peak within the cycle — right or left translation — predicts the net polarity of the cycle. It is the quantitative and falsifiable substitute for the concept of the inverse cycle (§ 8.1).

Asset	Cycles	Accuracy	Empirical base rate	Uplift
Bitcoin	452	74.6%	56.2%	+18.4 points
Ethereum	447	74.3%	56.2%	+18.1 points
Solana	300	77.3%	52.7%	+24.7 points

Table 7: Translation against cycle polarity, on the descriptive snapshot (§ 5.7). The empirical base rate is the fraction of cycles that close bullish: it is the correct comparison, because the structural bullish bias of cryptocurrencies makes a random 50% too generous a benchmark.

Accuracy is stable across the three markets and stays 18–25 points above the trivial forecast “all cycles close bullish”. The signal can be computed once the cycle has closed and its verification uses no future data. On a more recent snapshot and with the same criterion the accuracy is 76.0%, 74.9% and 76.1% on 454, 454 and 314 cycles: the number moves by a point or two with the series, the order of magnitude does not. Outside the daily timeframe the signal holds on samples an order of magnitude larger (§ 6.13).

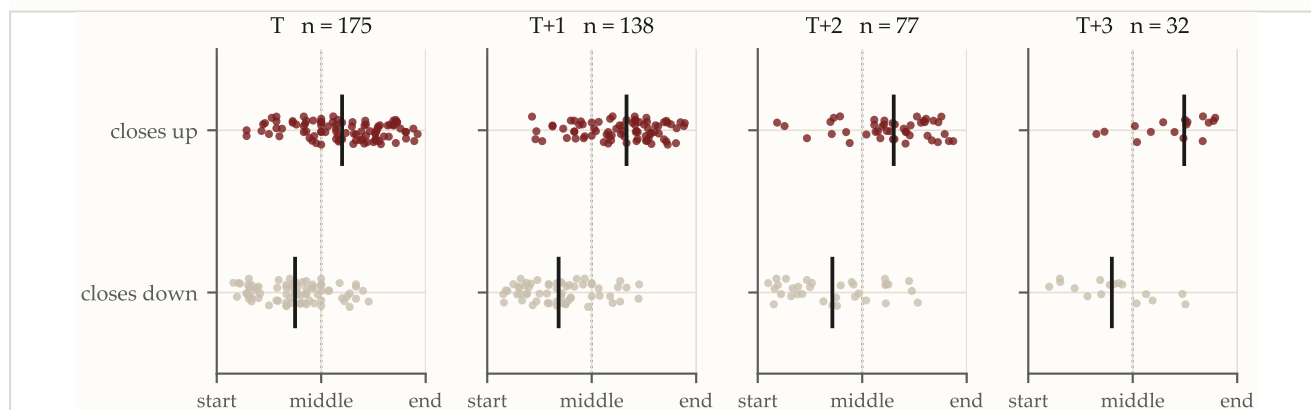


Figure 12: Where the peak falls within the cycle, separating the cycles that close up from those that close down. The vertical bars are the medians of the two groups: their distance is the signal.

Verdict CONFIRMED. With two caveats in reading, and the second is important.

The first: translation describes the polarity of the cycle *that is closing*, not of the next one. It is a measure of context, not a forecast.

The second: the *position of the peak within the cycle*, taken alone as a description of the market, **carries no information**. Compared with the null measured on the same detector — 199 surrogates for each of two generators — the observed median position falls **inside** the cloud in 14 decisive pairs out of 15, and no pair leaves it against both generators. It is geometry of the delimitation, like the short-cycle tail and the deviation from the nominal. The clearest case is Bitcoin at T+4: observed 0.5952 against a null center of 0.595. What remains, and does not depend on any null, are two things: that the canonical **constant** of 0.62 does not describe the position measured at the levels from T to T+4, and that translation **as a predictor of the polarity of the same cycle** exceeds the base rate. The description falls, the relationship holds.

6.6.1. DOES THE CYCLICAL CONTEXT REDUCE VARIANCE?

It is the premise of § 1, the thesis on which the whole framework rests, and it must be put to the test like everything else. Doing so requires an outcome, a context and a metric that is not the mean — because the premise speaks of **variance**. Outcome: the net logarithmic return of the child cycle, from trough to trough. Context: in which half of the parent cycle the child opens — structural and independent of the outcome, but **retrospective**, and the caveat is written below. Metric: the relative reduction of the interquartile range.

And it requires the right null, because **any partition reduces the conditional dispersion**, even a random one: without permuting the context labels, every positive result would have been arithmetic.

Verdict CONFIRMED on seven rulers out of nine. Over 23 evaluable cells and 83,531 child cycles the reduction of the interquartile range is positive in 22 and exceeds its own permutation null in 21; the rule declared beforehand — at least two markets above the null — is satisfied on H4, H2, H1, M30, M15, M5 and M1. On the Daily one cell out of two evaluable, on the Weekly none. The only negative cell is H4/Bitcoin, where the structural partition reduces the dispersion **less** than a random one (-4.5% , $p = 0.834$).

The caveat defines what the verdict means: the context is retrospective. The half of the parent cycle is computed with its **closing** trough, which at the moment the child opens has not yet occurred. The result therefore says that *the position within the parent carries information about the dispersion of the outcome* — an **upper** bound on what the structure can give to whoever decides — and not that that reduction is available in real time. The causal variant, with the half estimated from the nominal period instead of from the realized parent, is measurable with the same infrastructure and has not been run.

And the size is not uniform: from 11% to 18% on H2, H4 and M15, but between 1.6% and 3.5% on M5 and M1, where the samples number tens of thousands of cycles and the band of the null narrows to a few thousandths. There the significance comes from the sample, not from the size.

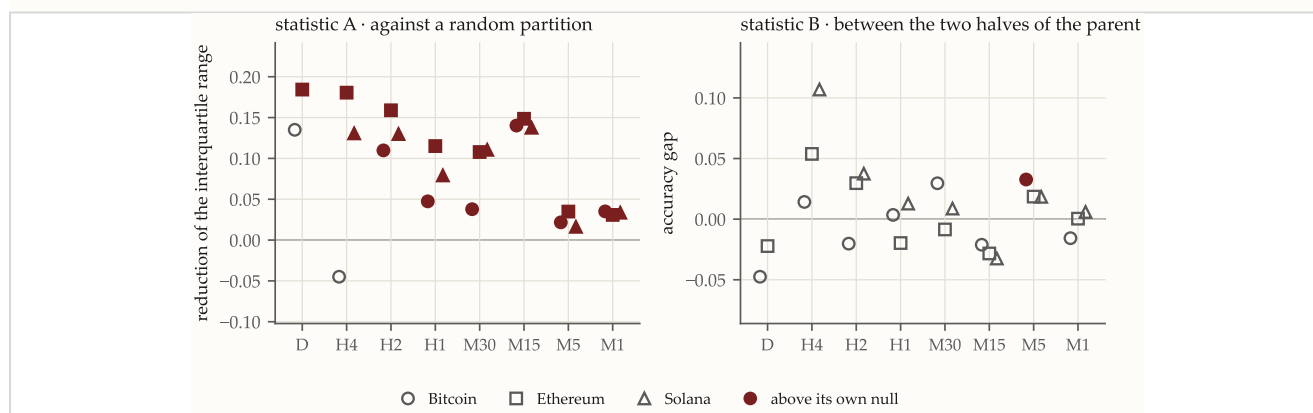


Figure 13: The two statistics, cell by cell over all rulers. On the left the reduction of the interquartile range: **any** partition reduces the dispersion, so the filled markers denote the cells that exceed their **own** permutation null, not zero. On the right the accuracy gap of the translation predictor between the two halves of the parent. The two are looked at together because **they have swapped places**: on the Daily alone the second held, over twenty-three cells the first holds.

It must also be said that the test is coarser than the premise: it uses a binary context and a signal of the same level, not of a lower level as the premise requires. A more faithful test wants a finer-grained signal — DSP,

order flow — conditioned by the structure, and that signal does not yet exist in measurable form in this work.

6.7. Conditional and unconditional swing (Elico64)

The Elico64 corpus distinguishes the **conditional swing** — a sub-cycle that closes negative in the first half of the parent cycle, which requires confirmation — from the **unconditional swing**, in the second half, which would be self-reliable.

Counted as the corpus defines it, the unconditional swing closes the parent cycle below its own internal peak in **100% of cases**, on all three markets. A hundred percent on a sample of that size is a number that invites one to stop, and it is exactly where one must not stop: **the control group is missing** (§ 5.4). The count concerns the parent cycles *with* a swing and says nothing about those *without*. Until that comparison exists, the 100% is compatible with two opposite worlds — one in which the swing selects the cycles that will close below the peak, and one in which **all** cycles close below their own internal peak and the swing selects nothing.

The condition, isolated. In the production module, unconditional branch, success is

$$\text{success} = [\text{low}(e) < \text{low}(\text{argmax}_{s \leq i \leq e} \text{high}(i))]]$$

where e is the **closing trough** of the parent cycle — a trough by construction, because the detector closes cycles on troughs — and the argmax is the bar of the internal peak. The question that condition poses is: “does the closing trough lie below the low of the highest bar of the cycle?”. Phrased this way it is almost self-evident.

The marginal, measured on the chain of levels that the production modules actually use:

Market	Parent cycles with swing	Without swing	Rate with	Rate without	Fisher p
Bitcoin	54	15	100.0%	100.0%	1.000
Ethereum	58	15	100.0%	100.0%	1.000
Solana	39	10	100.0%	100.0%	1.000

Table 8: The success condition of the unconditional swing, counted also on the parent cycles that do not have the swing: the first round, on the daily.

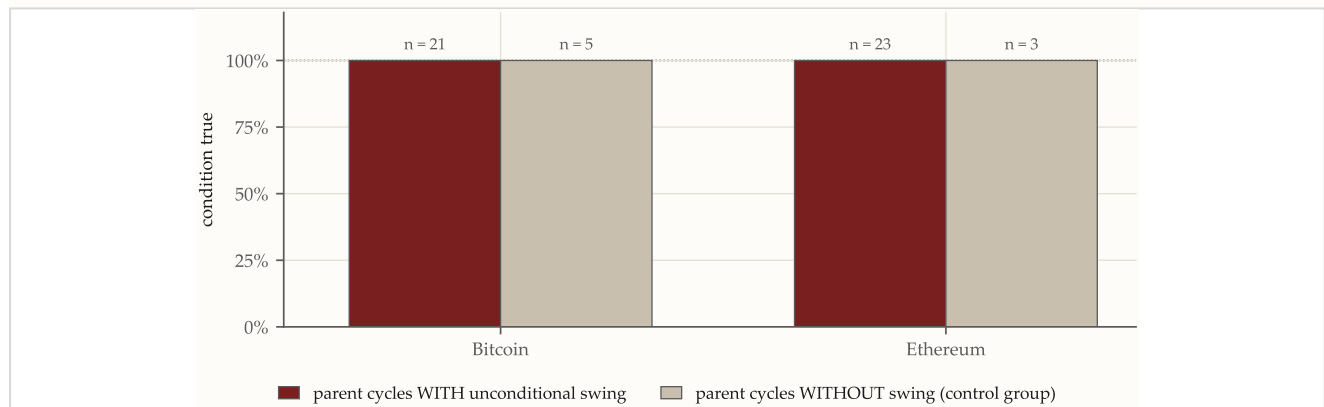


Figure 14: The condition counted on the two groups, ruler by ruler on the run and market by market: filled the group the swing selects, hollow the one it does not select, and under each ruler the parent cycles of the two groups. The two points stay together near one hundred percent; where the control group drops below the tautology threshold — on H2 and M30 — the gap stays below the tenth that the claim card requires for a confirmation.

Verdict: TAUTOLOGICAL — the claim is FALSIFIED as a signal. On the daily of the first round the condition is true in **all 191 parent cycles** measured, with and without swing: 151 with, 40 without, rate 1.000 everywhere, zero gap, Fisher $p = 1.000$ on all three markets. There is nothing to select. The tautology threshold was declared beforehand — rate on the cycles *without* swing ≥ 0.95 on all evaluable markets — and it is reached on three markets out of three.

On the run, read ruler by ruler, the tautology threshold is reached on five rulers out of seven. On H2 and M30 the group without swing drops to 91–94% on some market, and in three cases the gap is also significant (Fisher p 0.009 and 0.025 on M30, 0.043 on M5); but it stays small, and never reaches the tenth

the claim card required for a confirmation: the maximum is 0.091. A caution: the swing ties the parent to its third-degree sub-cycle, and the run measures it within each ruler, without the bridge that would put each member on the ruler that resolves it; the claim card keeps the outcome suspended on this point. In neither of the two available readings does the swing discriminate.

It remains true that the label is not retroactive: the swing is recognized before the parent closes, and in this it differs from the inverse cycle. But **being verifiable ex ante is not enough for a signal that does not discriminate**, and the two defects are independent: the inverse cycle is tautological in its *definition*, the unconditional swing is so in its *success condition*.

The conditional branch is not affected: its rates of 57–79% do not saturate, and the comparison between conditional and unconditional remains a comparison between two groups that assumes no base rate.

The rule that follows, and it holds for every conditional rate in this work: a rate measured on those that have the condition must always be accompanied by its **marginal**. Without it, a perfect number measures the condition and not the signal.

6.8. Violation of cycle tops (Elico64)

The corpus formalizes some signals based on the surpassing or the holding of levels of the previous cycle. There are four, and the marginal safeguard separates them sharply: two hold in the form in which the corpus states them, one holds only after being reformulated, one falls.

6.8.1. THE VIOLATION OF THE PREVIOUS PEAK, WITH ITS MARGINAL

The bullish signal — a sub-cycle that surpasses the peak of the previous one — requires two corrections before it can be read. The first is the **control group**: the peak hold is not a second signal, it is the marginal of the same comparison. The second is the **unit of count**: in the production module the event is counted *per pair of sub-cycles*, and a parent cycle with five sub-cycles produces up to three events that share the same outcome. They are not independent observations.

Market	Parent cycles	Base rate	With event	Without event	Gap	Fisher <i>p</i>
Bitcoin	73 (51 / 22)	53.4%	68.6%	18.2%	+50.5	0.0001
Ethereum	77 (60 / 17)	62.3%	76.7%	11.8%	+64.9	< 0.0001
Solana	49 (42 / 7)	67.4%	76.2%	14.3%	+61.9	0.0032

Table 9: Violation of the previous peak, counted per parent cycle and with the marginal alongside, read on H4 — the highest ruler on which all three markets have evaluable groups. No set inclusion: verified false on three markets out of three.

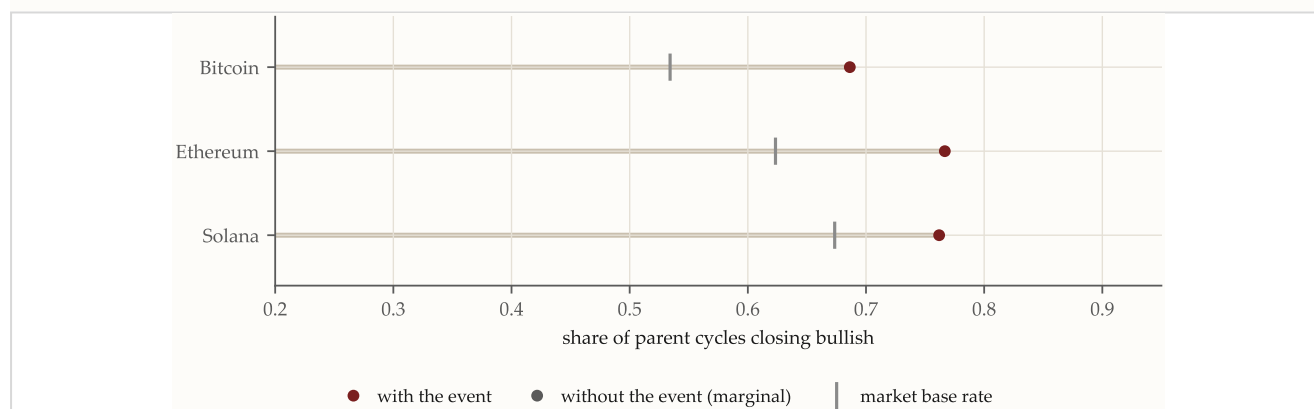


Figure 15: Read on H4, like the table. The signal is the **distance** between the two points; the vertical tick is the market's base rate, that is, what one would obtain without looking at anything. Covering the light point — the group that the event does not select — leaves a number that cannot be interpreted.

Verdict CONFIRMED, on three markets out of three and on six rulers out of nine. The signal is there and it is not an artifact: the set of events does not contain the positive outcomes by construction, and the control group closes bullish in 12–18% of cases against a base rate of 53–67%.

The edge case that this section used to record is gone, and the reason must be stated: in the reading on the daily two *p*-values lay two thousandths apart, on either side of the threshold, and the defensible reading was that the signal discriminated sharply on one market and marginally on the other two. On

the ruler that resolves the level the parent cycles go from 22 to 73 on Bitcoin and from 14 to 49 on Solana, and the three p -values lie between 0.0032 and one ten-thousandth. It was a fact of the sample, not of the market. On the daily the verdict remains NOT PROVEN, and for lack of cycles: two markets out of three do not have large enough groups.

And the peak hold is FALSIFIED as an autonomous signal: it is the marginal of this comparison, not an additional piece of information. Counting it as a separate signal means counting the same datum twice.

6.8.2. THE TWO BEARISH SIGNALS, WITHOUT THE INCLUSION

The corpus states two bearish signals: the break of the cycle's starting trough and the break of the trough of the previous cycle. The first, counted as the corpus defines it, contains a set defect that must be seen before the numbers.

The event is recorded if there **exists** a bar $k \in (s, e]$ with $\text{low}(k) < \text{low}(s)$, and the outcome is $\text{low}(e) < \text{low}(s)$. But if the cycle closes bearish, then $k = e$ already satisfies the condition: **every bearish cycle is an event by construction**. The set of events contains the whole set of positive outcomes — the implication is verified true in 100% of parent cycles on all three markets — and the resulting rate cannot fall below the base rate of bearish cycles.

The version that removes the inclusion requires the break to fall **within the first 75% of the duration**, and compares the share of bearish cycles between those that break early and those that do not break early: two groups both non-empty, neither of which contains by construction the outcomes of the other.

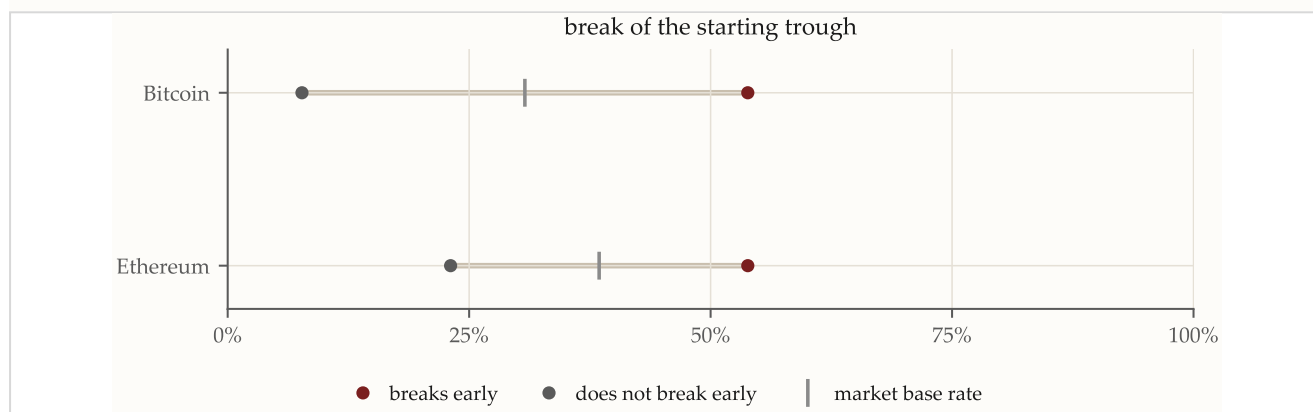


Figure 16: The two signals in discriminating form, read on H4 like the table: the share of bearish cycles among those that break early, among those that do not break, and the market's base rate. The distance between the two points is the signal; the tick is what one would obtain without looking at anything.

Verdict CONFIRMED on three markets out of three, for both signals, and on seven rulers out of nine.

The early break of the starting trough raises the share of bearish cycles by 13–23 points above the base rate; its **absence** lowers it by 19–33. The gap between the two groups is 32, 56 and 53 points; the previous-trough signal is worth 33, 45 and 51 points, and there the inclusion is absent: verified **false** on all three markets. Outside H4 the result does not change: from H2 down to the minute both signals discriminate on all three markets, with samples ranging from about a hundred to ten thousand parent cycles. Only the weekly and the daily remain out, and for lack of cycles, not because of the outcome.

The reformulation improves the claim in two ways. It is a **stronger** signal than the original form declared, and it is **two-sided**: the absence of an early break is an equally clear bullish indication — 13–22% of bearish cycles against a base rate of 42–51% — and in the corpus's formulation it does not appear at all.

The two sections, taken together, say something about the marginal safeguard: the same check, applied in the same way to neighboring claims of the same corpus, withdraws one, scales one down and strengthens two. **It is not a tool for demolition.**

6.9. Violated cycle troughs: what separates, and what does not anticipate

Two different questions fall on the same object — the violation of a level below the trough of a cycle — and current practice treats them together. They were measured separately, with the criterion frozen before the code, and they give two opposite outcomes.

6.9.1. THE STATIC LEVEL SEPARATES, THE ASCENDING TREND LINE DOES NOT

The first question is whether “the trough of a cycle is not violated” distinguishes a true cycle trough from any point whatsoever. The test builds **false troughs** with a fixed seed and windows of the same duration, and compares the share of violations on the true and on the false ones: if the claim has content, the true ones must be violated **less**.

With the static level — the trough of the previous cycle, a horizontal line — the separation is there and it is wide: true troughs turn out to be violated in 67.0% of cases and false ones in 91.7%, over 2,003 points per group; the direction is the same in **21 cells out of 21**, in 18 the difference is significant, and 9 contexts out of 11 pass the declared threshold. The verdict is **confirmed**, with a caveat that must be read together with the number: 67% is high in absolute terms. The claim “a cycle trough is not violated” is true only in a **relative** sense, as a filter against chance, and not as a description of what troughs do.

With the ascending trend line — the line joining two rising troughs, the tool that teaching presents alongside the level and with the same weight — the separation **is not there**. True ones are violated in 46.3% of cases and false ones in 44.9%: the true ones are on the wrong side, the expected direction appears in 8 cells out of 21, no cell is significant, none of the 11 contexts passes. The verdict is **falsified**.

The asymmetry is the result, not the detail: two tools that tradition teaches together behave in opposite ways when they are set against the same control group, and half of the teaching rests on the half that does not discriminate.

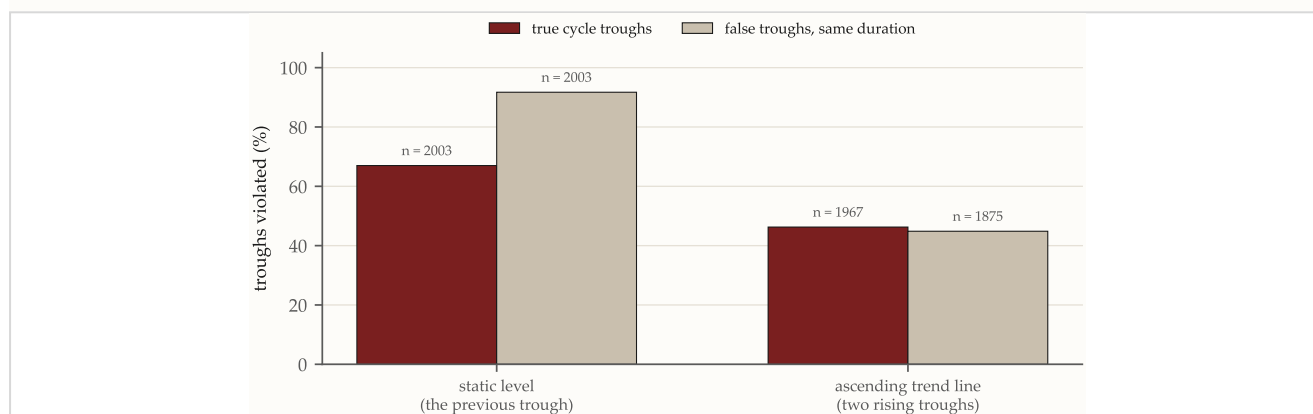


Figure 17: The share of violations on true cycle troughs and on false troughs built with a fixed seed and windows of the same duration, for the two tools. With the static level the two bars are twenty-five points apart in the expected direction; with the ascending trend line they touch, and the true ones are on the wrong side. The denominator is written above each bar: a share without its denominator is not verifiable.

6.9.2. THE VIOLATION DOES NOT ANTICIPATE THE CONFIRMATION OF THE TROUGH

The second question is operational: does the violation of a level **anticipate** the moment at which the trough is declared? Seven declarers — five, plus two placebos that act as a control — were put to the test on 21 contexts, each with its own base rate and with a **shape placebo** — a window of the same geometry placed where the event is not — because a window oriented in time produces by itself part of the advantage.

None of the seven passes: the share of contexts that meet the criterion is **zero** for all seven, and the overall verdict is **does not anticipate**. The way in which they fail, however, is not the same, and it is the part that indicates where to look. **VTL_SUB** is **precise**: it beats its own base rate in 15 cells out of 17 and exceeds the shape placebo in 11, that is, when it speaks it is right more often than chance. But it **does not cover**: in no cell does it reach half of the events. **NEST** is the opposite: it covers all 17 cells and does not beat the shape in 12 out of 17.

The two measurements of this section have their own runner and do not go through the construction described in § 5.6.1: their relationships are measured under the restriction declared there, ten complete cells out of seventeen.

6.10. Bayer tongues: encoding, calibration, discriminating power and predictive value

The tripartite detector (§ 3.3) classifies the connecting cycles. The breakdown is consistent across the three independent markets:

Calibration. Bayer’s canonical threshold — duration below half the mean — is not suited to cryptocurrencies: the real duration distribution requires an empirical calibration per asset and timeframe. That the classical

threshold must be recalibrated is a result and is reported; the threshold values, the percentiles and the procedure remain proprietary.

The discriminating power lies in the price step. Detection has two steps, and it is worth asking which of the two carries the information. Isolating step 1 — the brevity condition alone, with a causal threshold set at the lower quartile of the durations already observed at that level — and measuring whether short cycles are followed by a rebound more often than normal cycles:

Brevity alone carries no information about the rebound. The two-step criterion remains valid as a construction, but its discriminating power lies entirely in the price step: brevity selects the candidates, it does not classify them.

6.10.1. PREDICTIVE VALUE, JUDGED WITHOUT CIRCULARITY

A caution follows that must be carried all the way through: confirming a tongue through the subsequent rebound is **circular**, because the rebound is also what the price step measures. The predictive value must be measured with a design that does not use the rebound as confirmation.

Circularity is broken by a **blind interval**. The label is assigned by the causal detector as always; then one waits until the window the detector used has entirely passed — the maximum between 30 bars and half the nominal period of the level — and only from there does the return start to count, measured over a horizon equal to the nominal period. The control group consists of the troughs of the same level that the detector did **not** label as tongue, evaluated with the same rule. Without a control group a high rate would say nothing: in the crypto sector almost everything is bullish over long horizons.

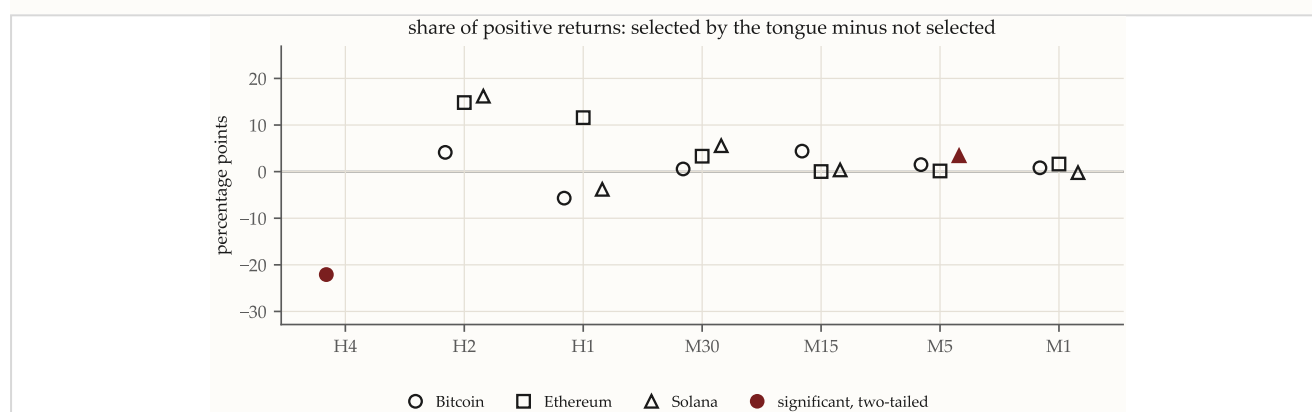


Figure 18: How much the share of positive returns of the cycles **selected** by the tongue exceeds that of the non-selected cycles, in percentage points, for every ruler and every market. Zero is the null hypothesis. Nineteen cells: two significant, in opposite directions, and a dispersion that **contracts as the sample grows** — up to sixteen points on H2 and H1, where the groups count hundreds of troughs, within three and a half points on M5 and M1, where they count thousands. It is the signature of noise averaging out.

Verdict NOT PROVEN, on 9,270 tongue troughs against 55,539 control troughs, nineteen evaluable cells out of twenty-seven. Only two significant cells, and in **opposite directions**: on H4/Bitcoin the tongue does clearly **worse** than the control — 44.6% positive returns against 66.7%, with a control of only 24 troughs on a bullish stretch — and on M5/Solana slightly better, 49.6% against 46.2%. The verdict is not **FALSIFIED** because the rule required that **no** cell be significant; it is not **CONFIRMED** because it required at least two significant markets with the same sign in favor, and there is one.

The ruler on which the measurement was born can no longer be evaluated, and this is the only honest way to report it: on the Daily the tongue label is almost universal at the only two remaining levels, and a control group of one or two troughs is not a control group.

This does not affect what the document states about the detection: that it is causal, verified non-repainting, dense and consistent across independent markets. These are properties of the **detection**. The defensible position is that tongue detection is a **verified descriptive tool, not a predictor**.

6.11. Cross-asset convergence, and the high levels over a century of history

The detector is run independently on the three markets. The convergence of the median durations, measured as the percentage range across the three assets, is strong at the short levels and seems to loosen at the high levels, where the sample shrinks.

Level	Range (3 crypto)	Range (6 indices, long history)
T	6.1%	11.1%
T+1	5.1%	5.0%
T+2	3.0%	8.6%
T+3	4.1%	8.9%
T+4	28.8%	9.2%
T+5	28.3%	9.7%

Table 10: Cross-asset convergence of the median durations. The right-hand column is measured on six equity indices with up to ninety-nine years of daily history.

The 28% at the high levels was small-sample noise, not a weakening of the principle. On the three crypto markets the cycles of level T+4 and T+5 number between four and twenty-one per market, and with that sample the dispersion across markets cannot be interpreted. On the six indices, where the sample exists, the range at T+4 and T+5 is **9.2% and 9.7%** against 8.6% and 8.9% at the middle levels: **there is no jump**. Hurst's cyclic commonality holds at the top too, as soon as it is given a sample.

The sample at the high levels was a limit of history, and now it has a number. Level T+5 goes from 4–6 cycles on crypto to **18–53 on the indices** and exceeds the measurability threshold — twenty cycles, inherited from the dip table of § 6.3 — on **five assets out of six**. The detector does not saturate: it produces troughs at the pace the level's duration allows, and the limit was the series'.

T+7 no, not even with a century. On the S&P 500, which starts in December 1927, fifteen come out. With a measured median between 1,166 and 1,799 bars, twenty cycles would require between **93 and 143 years of trading sessions**. T+7 is not an unmeasured level: it is a level that no existing dataset can measure, and it must be said so.

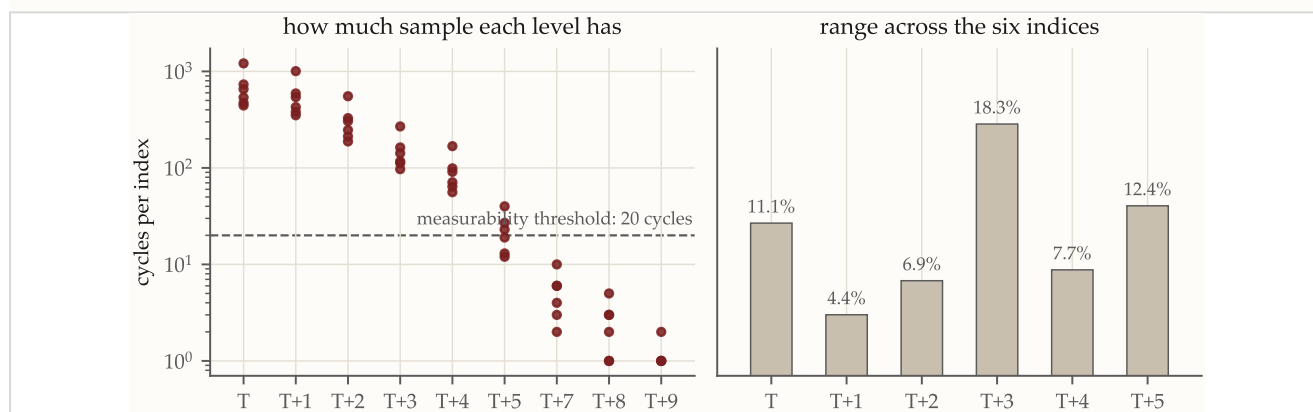


Figure 19: On the left, how many cycles each level has on each of the six long-history indices (logarithmic scale): the dashed line is the declared measurability threshold. On the right, the range of the median durations across the six indices, level by level. The left panel says **where the sample exists**; the right panel says that, where it exists, the convergence **does not loosen** moving up the levels — the jump to 28% in Table 10 is small-sample noise, not a weakening of the principle.

Verdict CONFIRMED. Cyclic commonality holds at the short levels on the three crypto markets (3.0–6.1% range) and at the high levels on the six indices (9.2–9.7%). Three markets with different microstructure, capitalization and history show median durations within a few percentage points of one another, and the same holds over a century of equities. It is the strongest evidence in favor of the Hurstian principle, and it does not depend on any calibration: the medians are direct measurements.

A declared limit, and it neither saves nor condemns the verdict: all the long-history measurements use the catalog of nominal periods calibrated on cryptocurrencies, applied to six equity indices. A catalog per asset class does not exist yet.

6.12. Admissible polarity sequences: a position effect

The Elico64 corpus postulates that the sub-cycles of a parent cycle can take only certain polarity sequences. It is the most ambitious hypothesis of the Italian strand. On the daily timeframe alone the sample is not enough — twenty-three four-part parent cycles across the three markets — and the natural path is to move down in

timeframe, where cycles with four sub-cycles should abound. The test, repeated on H4 and H1, says two things, and the second is far more important than the first.

Second: the rule is geometry of the detector. Verified by enumeration, the set of the eight admissible four-part sequences — four “index” and four “inverse” — coincides exactly with the set of sequences in which **the first and the last sub-cycle have opposite polarity**. The six admissible three-part ones coincide with “not all three equal”.

This opens a problem. The detector closes every parent cycle on a trough: the first sub-cycle starts from that trough and tends to rise, the last arrives at the next trough and tends to fall. The measurement confirms it unambiguously: in first position the sub-cycle is positive in 83–92% of cases, in last position in 13–22%.

Three references are therefore needed, not one.

Timeframe	<i>n</i>	Admissible rate	Measured baseline	<i>p</i> flow null	<i>p</i> position null
Daily	23	69.6%	49.4%	0.049	1.000
H4	127	73.2%	49.2%	0.0005	0.711
H1	686	75.1%	49.5%	0.0005	0.998

Table 11: Four-part structures: the observed rate against its nulls. The measured baseline uses the share of positive sub-cycles actually observed in place of the theoretical 50%.

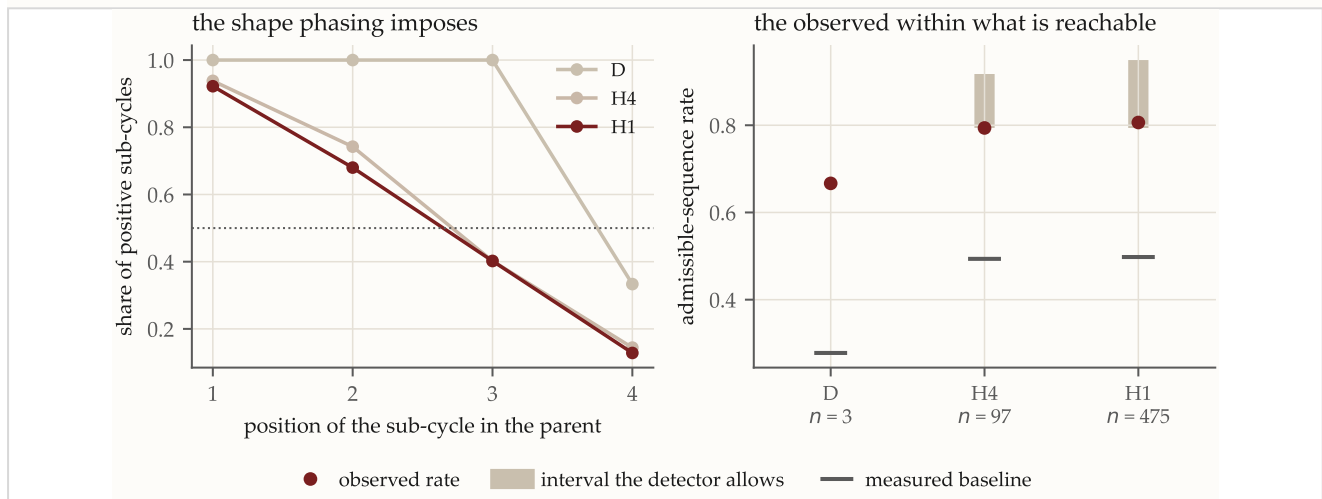


Figure 20: On the left the share of positive sub-cycles by position within the parent cycle: the almost obligatory shape that phasing imposes — it opens rising, it closes falling. On the right the rate of admissible sequences within the interval that shape allows: on the daily the observed value sits exactly at the bottom of the interval, that is, the market produces fewer admissible sequences than the detector permits.

The **naïve baseline** of 50% holds only if the polarities are fair; they are not, and the **measured** baseline drops to 49.2–49.5%. Against that, the observed rate looks high. The **flow permutation null** — which reshuffles all polarities together, preserving the overall marginal — confirms it: $p = 0.0005$ on H4 and H1, that is, the minimum declarable with two thousand reshuffles. With that null the point would have closed with a brilliant confirmation.

The **position null** instead reshuffles each column separately: it preserves the marginal of each position in the parent cycle and destroys only the dependence between different positions. It is the only one of the three that preserves the confounder. Against it the p -value is 0.711 on H4, 0.998 on H1 and 1.000 on the Daily: **the observed rate never exceeds it**.

And that 1.000 must be read carefully, because it is not a draw. Since the admissible four-part set coincides with “first and last of opposite polarity”, with the position marginals fixed the rate depends only on the overlap a between the positive sub-cycles in first and in last position: it equals $(k_1 + k_L - 2a)/n$, and can only range over a closed and narrow interval. Measured: **0.696–0.957 on the Daily, 0.701–0.858 on H4, 0.739–0.911 on H1**. On the Daily the observed value sits exactly at the **bottom** of that interval — every positive final sub-cycle belongs to a parent that starts positive — so $p = 1.000$ does not mean “no evidence” but “evidence in the opposite direction: the market produces fewer admissible sequences than the detector allows”.

On H4 the whole interval that the detector allows is fifteen percentage points wide. The 70–75% that reads as a regularity could not, on that timeframe, fall below 70% whatever the market did: the number was almost

obligatory before looking at the data.

Verdict NOT PROVEN, and the substantive reason matters more than the formal one: there is nothing to confirm. Against the null that preserves the confounder the p -value never falls below 0.7.

The three-part structure reads even better. The admissible set is “not all three equal”, the measured baseline is 0.720–0.741 and the observed rate 0.709–0.757, on samples from 74 to 2,129 sequences. There is not even the apparent gap observed with four parts: the three-part rule carries no information at any timeframe, and this does **not** depend on the sample.

Note of procedural honesty. The position null was added after seeing the first outcome, when the confounder was recognized. It is a more severe test, not a more permissive one: it can only remove support from the claim, never add any. The value of the first null remains recorded alongside the second and no threshold was touched.

The effect belongs to the detector at every scale tried, and now the scales are nine. The same measurement, with the same thresholds and the same null, was repeated on all the timeframes that the floor admits, down to the minute. The minimum sample to vote — one hundred four-part sequences, declared beforehand and not retouched — leaves out Weekly, Daily and H4, which bring zero, three and ninety-seven: **six vote**.

Two readings, and the second is the stronger. The first: on **six voting timeframes out of six** the verdict is **NOT PROVEN**, and the p -value against the null that preserves the position marginal never falls below 0.5 — on samples ranging from 172 to 3,220 sequences, that is, up to **a thousand times** that of the daily. The sample is no longer an available extenuating circumstance.

The second: **the confounder does not change with scale.** The share of positive sub-cycles in first position is 0.92–0.94 on all seven intraday timeframes and 0.11–0.14 in last position. The obligatory shape of the parent cycle — it opens rising, it closes falling — is the same from four hours to one minute, and with it the interval that shape **allows** stays identical. And the measured baseline converges to **0.500** as the sample grows: a position effect that reproduces unchanged over four orders of magnitude of resolution is a property of the detector, not an accident of one timeframe.

What remains of Elico64 on this point: nothing, in the form of the “admissible sequence”. There remains a useful observation that the corpus does not formulate but that the numbers support — a parent cycle, as the detector delimits it, has an almost obligatory shape: it opens rising and closes falling. Calling “admissible” the sequences that respect that shape is a correct description of the geometry and an empty prediction.

6.13. Outside the daily timeframe: four claims on three scales

A single timeframe for the statistical claims is a limitation, and it must be removed where possible. Four claims measurable without surrogates were re-measured on Weekly, H4 and H1 for the three markets, with the same pipeline and with the thresholds already declared — not new ones: that the measured durations do not match the nominal (A1), that some adjacent levels are not distinct (A2), that translation predicts polarity (A3), that there is a short-cycle tail growing with the level (A4).

A1 · durations deviate from the nominal	holds	no	no	1/3
A2 · two adjacent levels are not distinct	holds	holds	holds	3/3
A3 · translation predicts polarity	holds	holds	holds	3/3
A4 · a short-cycle tail exists	holds	holds	holds	3/3
	W	H4	H1	

Figure 21: The four load-bearing claims, cell by cell, on the three timeframes outside the daily. The filled boxes do not all sit in one column — it is not a scale that holds them up — and the **two empty ones sit on the same row**: A1, the only one that falls, and it falls precisely where the ruler resolves the levels of its own chain.

Verdict CONFIRMED, three claims out of four hold outside the Daily — the threshold declared before the measurement required at least three. The clearest case is translation, which holds in **nine cases out of nine**: on H4 it is 72.5–75.7% with 16.6–21.2 points of uplift over the measured base rate, on samples from 3,073 to 4,658 cycles per market; on H1 it is 72.0–73.8% with 17.5–21.4 points, on samples of up to 19,623 cycles. It is the same regularity as in § 6.6, measured with two orders of magnitude more data.

6.14. The widest test: thirty-seven instruments, five statistics

Everything above concerns three cryptocurrencies. It is the sample on which the work was born, and it is a narrow sample: three series of the same class, correlated with one another, within a decade. The question that closes the protocol is therefore the most uncomfortable one — **does the cycle structure that the detector extracts differ from chance, on a wide sample?** — and it must be posed with the most severe null available.

The design. Thirty-seven instruments in six classes (cryptocurrencies, stocks, indices, commodities, currencies, bonds), on Daily, with the same chain of levels and **five** different statistics. Fifty phase-randomized surrogates and fifty AAFT per instrument, with the production detector applied **identically** to the real and to the synthetic series; sign test on the 182 instrument/level pairs.

The outcome, over 1,820 comparisons. Against the phase surrogates the five statistics give 85–101 favorable pairs out of 182 — around half — with p between 0.079 and 0.832. Against the AAFT, which also preserve the distribution of returns and are the severe null, the picture worsens: 75–89 out of 182, p between 0.644 and 0.993, that is, **at half or below**, often in the *opposite* direction to the expected one. The pairs that exceed the threshold on the single comparison are between 2 and 13 out of 182: the order of magnitude that chance produces. **No significant cell in any of the six classes.**

Verdict: FALSIFIED. On a wide sample and with the severe null, the cycle structure extracted by the detector **does not differ** from the one that the same detector extracts from series with the same spectrum and the same tails. The measurable advantage on three cryptocurrencies against the phase surrogates alone — consistent in sign on ten pairs out of twelve — **does not survive** either the widening or the more severe surrogate.

A limit that must be declared, and it does not save the verdict. The catalog of nominal periods used is the one calibrated on cryptocurrencies, applied to all six classes: part of the failure may be a wrong catalog for stocks, indices and commodities — looking for cycles at the wrong periods finds no structure even where there is some. But the cryptocurrencies are **inside** the sample with **their** catalog, and even there no cell is significant.

6.15. The starting axiom: is the quaternary a binary of binaries?

The last claim to be tested belongs to no historical school: it is the theoretical starting point of this work. **The structural alphabet is {2, 3}, and the quaternary is not a third category: it is a binary of binaries.** The first part — the prevalence of twos and threes in the counts — was measured against the surrogates: against the

phase-randomized ones a few comparisons leave the cloud, against the AAFT — the more severe null — none does, and over thirty-seven instruments no structural statistic exceeds it (§ 6.14). The strong part must be tested on its own.

The test is direct and has no parameters. A parent cycle with four sub-cycles has three internal troughs. In the “by four” reading they are equivalent boundaries; in the “two by two” reading the middle one is the boundary between the two binaries, that is, a trough of higher degree, and on price it must be **the lowest of the three**. Under pure indifference the lowest is one of the three with probability one third.

On 557 cycles with four sub-cycles — aggregated over four rulers (D, H4, M5, M1) and three markets — the central trough is the lowest in 7 cases: 1.3%, against the 33.3% that indifference requires. The gap is thirty-two points, and on its own it does not yet say whose it is.

And the sample says where it comes from: 73% of the quaternaries arise on M5 and 27% on M1, that is, on the low inverse levels (T-3 ... T-7) that only the fine rulers resolve; the seven central cases **all** come from M5. With rulers stopping at M15 the analyzable pair would be T-1/T on H1: the scope of the measurement must be read accordingly, it concerns the quaternary **where the rulers resolve it**, and those are the fine ones.

Why so little. The null measured with the same detector on the surrogates has **median 0.0%** and ninety-fifth percentile **2.3%** against the phase surrogates and **2.1%** against the AAFT: the observed 1.3% lies **inside** it for both families (empirical p 0.176 and 0.189). On these series the detector alone produces the same share. The reason was predicted before measuring and it is the geometry documented in § 6.12 — the first sub-cycle rises in 83–92% of cases, the last falls — which pushes the central trough to be the **highest** of the three. If the quaternary were really a binary of binaries, it should be visible **despite** the detector’s geometry, not thanks to it. The market, with the geometry removed, is silent on that point — and it remains to ask whether the detector could have made it speak.

Would the test have seen a binary of binaries? A negative outcome counts only if the test has the power to see the phenomenon (§ 5.5), and here the question can be asked directly: on synthetic series with an **implanted** binary of binaries — three placements of the period, 288, 432 and 576 bars, and increasing intensities of the central trough — the same production procedure is run. On the true troughs of the synthetic series the central trough is the lowest in 64% of the parents at half the minimum intensity and in 90% at the minimum intensity: the phenomenon is there and it is marked. **The procedure does not return it.** It locks onto the intermediate level and reads it as two binaries; at the minimum intensity and above it produces no parent with four children, and where the effect is weak or absent it produces at most six per trial — 110 in all over 360 trials, never with the central trough the lowest. A trial **detects** the structure when the procedure returns at least the thirty four-child parents needed to vote: at the minimum intensity detections are **zero in thirty trials**, in each of the three placements. The simulation reproduces the M5 ruler, from which 73% of the market quaternaries come.

And the rest of the alphabet, which nobody had ever counted. The same pass censuses **all** the cardinalities, not only the fourth. Out of 26,017 parent cycles: 13,284 contain no trough of the level below, and of the 12,733 that have an internal structure **9,450 are binary (74.2%), 2,651 ternary (20.8%), 557 quaternary (4.4%) and 67 quinary (0.5%)**; eight have six or more. The alphabet {2, 3} covers **95.0%** of cases.

Two readings, and they must be kept separate. The **ternary** is not distributed like indifference: out of 2,651 parents the lowest trough falls 1,199 times in the first position and 1,452 in the second, a deviation from the uniform of 0.048 with $p = 0.0005$ — the minimum declarable with two thousand draws. **But that deviation does not get out of the detector’s cloud:** on the surrogates the median is 0.125, that is, two and a half times the observed value. The imbalance is real and it does not belong to the market: it is what the detector produces even when the structure is not there. The **quinary exists**, and with sixty-seven cases it is no longer a matter of sample: its asymmetry is 0.440 against a median of 0.524 on noise — **it lies inside the cloud, and lower than the noise itself**. It is not a new unit of the alphabet.

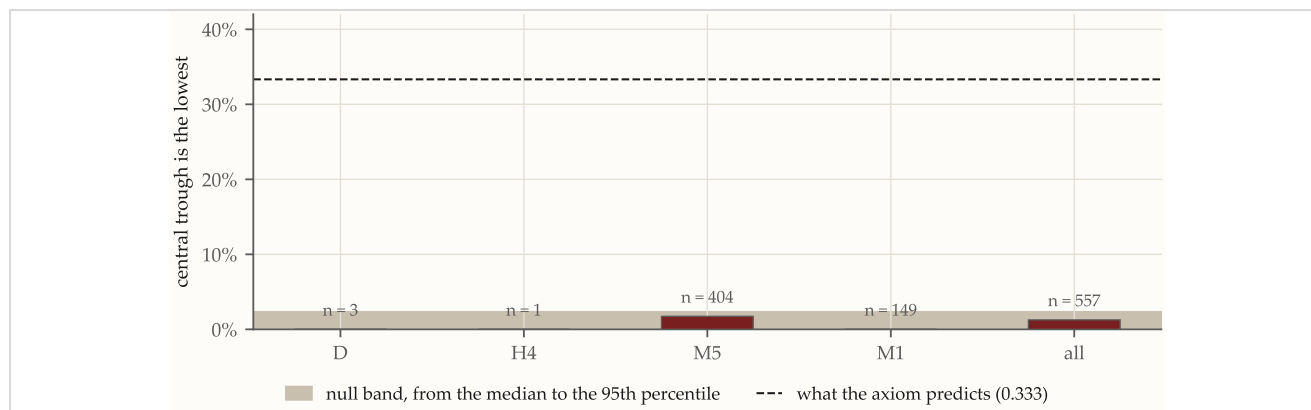


Figure 22: The share of four-part cycles in which the central internal trough is the lowest, against what the axiom predicts and against the detector's band on noise. The axiom asks for one third, the market gives 1.3% — and the detector's band on noise reaches 2.3%, that is, it contains it.

Verdict NOT PROVEN. The strong part of the axiom cannot be decided with this instrument: the market gives a lowest central trough in 1.3% of the cycles, within the detector's band, and the detector does not recognize a binary of binaries even where it is implanted. The alphabet {2, 3} remains a useful description of the counts; whether the quaternary is a binary of binaries, this work does not know.

It remains declared that the form tested is **one** reading of “by four equals two by two”: the one through the depth of the trough that separates the two binaries. Another — the symmetry of the durations between the two halves — is conceivable and has not been tested. Deciding either first requires a detector that, on the synthetic series, returns the structure placed in it.

6.16. The value of the idea, measured honestly

A cyclical idea has value if its causal detection produces, over at least one complete market cycle, a number of signals sufficient for inference and stable over time. The dataset covers the 2020–2021 bull, the 2022 bear, the 2023–2024 recovery and the 2025–2026 phase.

Three pieces of evidence support the informative value, and they must be weighed for what they are.

1. **Consistency of the taxonomy across independent markets** (the tripartite encoding above): the same criterion produces comparable breakdowns on series with different microstructure. It is a clue that the detected structure is not a calibration artifact, not a proof.
2. **A strong-value cyclical signal in the same period**: translation at 74–77% with an uplift of 18–25 points over the base rate (Table 7), replicated on H4 and H1 on samples an order of magnitude larger.
3. **Temporal stability**: the confirmed labels do not change as time progresses (the non-repainting gate above). It is a necessary requirement for a signal to have decision value.

What this section does not demonstrate. It does not demonstrate profitability: no number here is the result of an operational strategy. And the direct predictive value of the tongues, measured with a non-circular design, **is not there** (§ 6.10.1). The correct formulation is: detection is **dense, stable and consistent across markets**, and this makes it a context tool; what has measured directional value is translation, not the label.

7. Scope, limitations and reproducibility

7.1. What is demonstrated and what is not

In scope — non-repainting demonstrated. The coded Bayer tongues, on which the claims concerning detection rest, are causal and verified non-repainting: five asset/timeframe combinations out of five, walk-forward on the real pipeline, determinism verified. The confirmed historical labels do not change; the only moving entity is the last tongue with its window still open. The multi-timeframe resampling of the atomic data is also certified by an automated test.

Out of scope — not the object of the claims. The complete system includes cyclical composite components — structural anchoring, windowed spectral discovery, signal fusion — which are not the object of the claims

of this document. Any residual doubts about repainting in those components do not touch the results presented here, which rest on the confirmed tongues and on the cycle statistics derived from the causal detector. They remain out of scope for methodological honesty: a claim is made only where the property is demonstrated.

7.2. Limitations

1. **Dependence on the detector.** Durations, nesting and the short-cycle tail depend on the operational definition of a cycle. Convergence across three independent markets is a safeguard against *calibration* artifacts, **not** against *measurement* error: a wrong measurement converges wherever it is applied, because the series pass through the same procedure. Convergence across markets is a test of the **generality** of the phenomenon, never of the correctness of the instrument, and using it as the latter is one of the easiest ways of proving oneself right.
2. **The ruler.** The historical setup measures on the daily timeframe, which does not resolve the shortest levels of the chain: the numbers for T, T+1 and T+2 measured there describe a different object from the one they describe on the timeframe that resolves them (§ 5.6). **Polarity** measurements are robust to a change of ruler; those of **duration and dispersion** are not. The document keeps both readings and declares which one it is using.
3. **Multiple comparisons.** The work conducts many tests. They should be read as convergent cross-asset evidence rather than as isolated discoveries, and where only one market out of three rejects, this document declares it instead of reporting it as a confirmation.
4. **Sample at the very high levels.** T+7 and above are not measurable with any existing dataset: they require between 93 and 143 years of trading sessions (§ 6.11). It is no longer “insufficient sample”, it is a quantified requirement.
5. **The catalog by asset class.** All the long-history measurements and the test on 37 instruments use the catalog of nominal periods calibrated on cryptocurrencies. A catalog that calibrates itself on the asset does not exist yet, and it is the limitation that weighs most on the extended negative results.
6. **The cause of the short tail.** § 6.3.1 establishes that the tail belongs to the detector, not that it is known why. The mechanical hypothesis — the tail arises from the gap between nominal period and true duration — is testable: a calibrated catalog should reduce it if that is the cause.
7. **Intraday non-repainting.** The gate covers five combinations; a denser check shows that the robustness is not uniform across all intraday combinations. No result of this document rests on it, and it remains declared as open work.
8. **Asymmetry between confirmations and refutations.** A refutation on one market falsifies a universal claim; a confirmation holds for that market, in that period, with that detector. The two columns of the verdict table do not carry the same weight, and that is the reason why the confirmations are formulated with more caution.
9. **The instrument is not redistributed, and this bounds replicability.** The calibration thresholds, the parameter values and the tuning procedure of the cycle detector belong to a commercial product and are not published. The consequence for verifiability is precise, and it is stated here rather than left in a footnote. The results are **verifiable**: every verdict in this document comes with its evidence file, which carries the procedure, the dataset with the fingerprint of the prices, the null hypothesis **with its measured value**, the thresholds fixed before execution and the numbers cell by cell, so a third party can recheck each number and contest each verdict. The results are **not regenerable from scratch** by a third party, because the instrument that produced them is not shipped. This is a deliberate trade — protection of the product against full replicability — and not a gap in the experimental design; it is the same renunciation already declared in § 7.3, named here as a limitation of the work so that the reader can weigh it as one.

7.3. Reproducibility and independent verification

Two levels are distinguished.

Verification of the results (provided). The supporting material contains, **for every verdict declared in this document**, the evidence file it comes from: procedure, dataset with the fingerprint of the prices, **null hypothesis with its measured value**, thresholds fixed before execution, and the numbers cell by cell. The attached manifest carries their SHA-256 fingerprints.

The selection is not discretionary: the bundle is built **from the same table** used to verify that the verdicts on the page coincide with those of the claim cards, so a row added to the document adds a file to the material,

and a verdict that changes makes the check fail before delivery.

With that material a third party can recheck every number in this document without accessing the proprietary code — and it is the criterion that comes **before** commercial protection: if the reader cannot rerun the test on the delivered output and contest the verdict, it is not scientific work.

Repetition of the protocol (replicable). The methodology can be described in executable form on one's own data: determinism with fixed seeds and pinned dependencies; walk-forward with a prefix invariance criterion; cross-asset external validation; comparison with phase-randomized and AAFT surrogates for **every** claim of the form “this regularity exists”; a marginal next to every conditional rate; a power check before reading a negative outcome; and every level measured where it resolves. An analyst can apply the expected cycle window (§ 3.2), the two-step causal criterion (§ 3.3) and confirmation at window close (§ 3.4) to their own data with their own calibration, obtaining a conceptually equivalent system.

Only the calibration values remain reserved — window tolerance, brevity percentile, thresholds of the price criterion, band of the intermediate category, fusion weights — and the code.

8. Refutations, confirmations and operational substitutes

Claim	Source	Test	Verdict
The scale of levels is complete	tradition	own troughs and bidirectional overlap, against its null	FALSIFIED
Fixed harmonic nesting by two	Hurst	ratios between median durations and observed structures	FALSIFIED
Harmonicity of periods	Hurst	distance of the peaks from the rungs of a scale, with verified power; intraday it holds only against a null that is too wide	FALSIFIED
Cyclic commonality	Hurst	range of variation across markets, on crypto and on a century of equities	CONFIRMED at the short and high levels
Synchronicity of troughs	Hurst	overlap between levels	CONFIRMED
Translation as a predictor of polarity	Hurst	position of the peak against the net polarity of the cycle, with base rate	CONFIRMED : describes the closing cycle, not the next one
Amplitude-duration proportionality	Hurst	rank and log-log slope, against surrogates	PARTIAL : the slope yes, the rank no
Bull/bear asymmetry of duration	this work	ratio between bullish and bearish durations, against surrogates	NOT PROVEN : comes out at the short levels, against a reversible null that cannot exclude it
Duration equal to the nominal period	all three	median against nominal, on D, W, H4 and H1	FALSIFIED
The deviation from the nominal is a market fact	this work	median/nominal against a cloud of series without cycles	FALSIFIED : inside the cloud in 28 cells out of 37
The position of the peak is a market fact	this work	median position against measured null, two generators	FALSIFIED : inside the cloud in 14 pairs out of 15
The cycle structure is distinguishable from chance	this work	five statistics, 37 instruments, 182 pairs, against AAF	FALSIFIED : no significant cell
Local stationarity	Ehlers	drift of the dominant period against surrogates with fixed spectrum	NOT PROVEN , with its own test
Bimodality of durations	this work	Hartigan's dip	FALSIFIED
Short tail as a market fact	this work	observed share against surrogate cloud	FALSIFIED : 1 combination out of 33 on the chain
Regime as driver of the tail	this work	exogenous and causal regime, in terciles	NOT PROVEN
Ternary level skip	this work	children of binary parents against children of ternary parents	FALSIFIED
Sequential mean-reversion	inference from Hurst	first-order autocorrelation	FALSIFIED in the strong form

Claim	Source	Test	Verdict
Operational inverse cycle	Elico64	logical form of the definition	TAUTOLOGICAL
Canonical brevity threshold	Bayer, via Elico64	application to the real distribution	FALSIFIED
Unconditional swing	Elico64	conditional rate against its marginal, Fisher exact	TAUTOLOGICAL : true in all 191 parent cycles of the daily; on the chain the gap never reaches 0.10
Violation of the previous peak (D>P)	Elico64	rate per parent cycle, with marginal and base rate	CONFIRMED on 3 markets out of 3
Break of the starting trough	Elico64	form without inclusion, against the group that does not break	CONFIRMED 3 out of 3
Break of the trough of the previous cycle	Elico64	same	CONFIRMED 3 out of 3
Peak hold (Hold)	Elico64	frequency against base rate	FALSIFIED : it is the marginal of D>P
Predictive value of the tongues	this work	forward return after blind interval, against controls	NOT PROVEN
Admissible polarity sequences	Elico64	rate against three nulls, one of which preserves the position	NOT PROVEN : position effect
Context reduces decision variance	premise § 1	IQR reduction against permutation of the context, on nine rulers	CONFIRMED on 7 rulers out of 9 — but the context is retrospective
The structural alphabet is {2,3}, the quaternary is a binary of binaries	this work	depth of the trough that would separate the two binaries, with a power check on synthetic series	NOT PROVEN : the detector does not return the implanted structure
The 80% overlap threshold is informative	this work	observed overlap against the null of the same detector, each pair on the ruler that owns its members	NOT PROVEN : complete coverage — 11 pairs out of 11, 33 rows across the three markets — and margin in all of them; the blocker is the worst margin, 0.096 on T-7/T-6 of M1, below the declared 0.10
The cause of the T+3 discrepancy is the duration filter	this work	three readings with the filter off and on, on the ruler that resolves the level	FALSIFIED — on H4 the filter cuts nothing
Measuring a level on the timeframe that resolves it changes the numbers	this work	four statistics, Daily against the owner timeframe, two controls	CONFIRMED : 12 discordant pairs out of 12
The daily cycle separates from the half-daily	this work	validity rule on the pair, three resolutions	CONFIRMED : ratio 2.34–2.38 where it is resolved
The four claims of the Daily hold outside the Daily	this work	remeasurement on W, H4 and H1 with the same thresholds	CONFIRMED 3 out of 4 — A1 falls where the ruler resolves
The sample at the high levels is an insurmountable limit	this work	count of cycles on six long-history indices	FALSIFIED for T+5, which becomes measurable

Claim	Source	Test	Verdict
The cycle trough is not violated (static level)	Elico64	violations on true troughs against false troughs of equal duration	CONFIRMED : 67.0% against 91.7%, 9 contexts out of 11
The cycle trough is not violated (ascending trend line)	Elico64	same test, same control group	FALSIFIED : 46.3% against 44.9%, 0 contexts out of 11
The violation of a level anticipates the confirmation of the trough	this work	seven declarers, with base rate and shape placebo	FALSIFIED : 0 contexts out of 21

Table 12: Verdict table. The rows with source “this work” are claims of the author, and should be read together with the “Source” column beside each row; the operational substitute of each, where it exists, is in the section that deals with it.

The count, row by row on the table: thirty-eight claims, of which **seventeen falsified** and **two tautological** — nineteen refuted in all —, **eleven confirmed**, **seven not proven** and **one partial**. The count is not written by hand: it is computed by a check that reads the table, so that a row added to the document cannot leave it behind.

8.1. On the inverse cycle

The inverse cycle is Elico64’s distinctive contribution to the Italian strand: every cycle would admit a symmetric reading, “index” and “inverse”. The refutation here is **epistemological, not empirical**: the operational definition “a cycle that closes bearish is inverse” is tautological, because every bearish cycle is by definition inverse, and the assertion “the inverse predicts the decline” is not falsifiable ex ante — the label is assigned only a posteriori. The phenomenon that the notion attempts to describe is not denied: it is observed that its formulation is not testable. Cyclical translation, instead, has a quantitative formulation, is falsifiable and reaches 74–77% accuracy with 18–25 points of uplift (§ 6.6): it replaces intuition without abandoning it.

8.2. On Hurst’s assumptions

The balance is mixed, and should be read as such. Harmonicity as a fixed ratio and the complete scale of levels fall. Cyclic commonality — at the short levels on cryptocurrencies and at the high levels on a century of equities — and the synchronicity of troughs hold. Proportionality holds in the form of a slope and falls in the form of rank correlation. It is remarkable, for a framework built on equity data from the 1960s and applied here to a market that did not exist, how much of it survives: Hurst’s claims are the most exposed because they are the most precise, and it is exactly this that makes them useful.

9. Conclusions

Cycle analysis of markets has survived for over fifty years as a set of doctrinal traditions. This work shows that it can be converted into a research program by applying a Popperian protocol to a multi-asset dataset, with a causal and deterministic methodology verified by automated tests, and with four requirements that cycle analysis has almost never set itself: the null value of every measurement must be **measured**; every conditional rate must be accompanied by its own **marginal**; a negative outcome counts only if the test would have the **power** to see the phenomenon; and every level must be measured on the **timeframe that resolves it**.

The balance is clear-cut and should be read with the asymmetry of § 7.2. Among the **refuted** claims are the complete scale of levels, fixed harmonic nesting, the harmonicity of periods on the daily ruler — with a test whose power is verified —, duration equal to the nominal period, sequential mean-reversion, the ternary level skip, the canonical brevity threshold, peak hold and, **by tautology**, two claims of the Italian school: the inverse cycle, tautological in its definition, and the unconditional swing, tautological in its success condition. **Confirmed**, with cross-asset convergence, are cyclic commonality — at the short levels on cryptocurrencies and at the high levels on a century of equities —, the synchronicity of troughs, translation as a predictor of polarity and, in the form that eliminates the set-theoretic inclusion, the three signals of violation of cycle tops. **Not proven** remain local stationarity, the association between duration and regime, the bull/bear asymmetry, the predictive value of the tongues, the admissible polarity sequences, the eighty-percent overlap threshold and the quaternary as a binary of binaries. The premise of § 1 on variance reduction is **confirmed on seven rulers out of nine**, but with the context measured a posteriori: it says that the position within the parent cycle carries information about the dispersion of the outcome, not that this information is available to the one who decides.

Of the regularities that this work could claim as undocumented by the previous schools, **two** remain, and neither of them concerns the market: the **absence of two steps** of the nominal scale — now replicated on a century of equities — and the **cross-asset convergence** of durations, which extends to the high levels as soon as they are given a sample. Two of a different nature are added, because they concern the **instrument** and not the market: the almost obligatory shape of the parent cycle under phasing, and the short tail as a property of the procedure.

The fallen candidates are more instructive than those that remained. The **bimodality** of durations — which would have implied two populations of cycles — is not there. The **short-cycle tail** is there, it is stable across markets and, within each ruler, grows with the level, and **it does not belong to the market**: on the chain of rulers one asset/level combination out of thirty-three comes out of the surrogate cloud, and otherwise the observed value lies above or below the synthetic median with no direction. The same holds for the deviation from the nominal period and for the position of the peak within the cycle. They were the results we had the most interest in keeping, and they fell for the same reason the protocol exists: **the null value is measured**. And for the twin reason the **axiom from which the work had started** remains open: a negative outcome counts only if the test could have seen the phenomenon, and this detector does not recognize a binary of binaries even where it is implanted.

If there is a result that this work delivers, it is not an indicator and it is not a regularity. It is a procedure, and it is simple enough to fit in four lines:

A rate without its marginal is not evidence. A statistic without its measured null is not evidence. A null that does not preserve the geometry of the instrument measures the instrument; one that preserves what is under discussion measures nothing. And a negative outcome does not count until it has been shown that the test would recognize the phenomenon if it were there.

To which a fifth line must be added, the most recent and perhaps the most costly: **a level measured where it does not resolve is another level**, and the name it bears does not say so.

Applied to one's own work, this procedure has a cost. The cost is the price of the fact that, in the end, what remains standing truly remains standing — and it is the only thing a research work can honestly promise. That almost nothing of what remains standing is of one's own making is the way a research program declares that it has worked, not that it has failed.

The natural direction of development is already indicated by the limitations: a catalog of periods calibrated by asset class, the systematic measurement of every level on the timeframe that resolves it — the troughs of the five intraday levels now exist (§ 5.6.1), but the statistics that this document reports for those levels remain those read on the daily — and a test of conditional value with a signal of finer grain than the structural one.

References

- Hurst, J. M. (1970). *The Profit Magic of Stock Transaction Timing*. Prentice-Hall.
- Hurst, J. M. (1973–1976). *Cyclitec Cycles Course*.
- Ehlers, J. F. (2001). *Rocket Science for Traders*. Wiley.
- Ehlers, J. F. (2013). *Cycle Analytics for Traders*. Wiley.
- Bayer, G. (1948). *The Egg of Columbus*.
- Hamilton, J. D. (1989). "A New Approach to the Economic Analysis of Nonstationary Time Series and the Business Cycle". *Econometrica*, 57(2), 357–384.
- Lunde, A. and Timmermann, A. (2004). "Duration Dependence in Stock Prices". *Journal of Business & Economic Statistics*, 22(3).
- Claessens, S., Kose, M. A. and Terrones, M. E. (2011). "Financial Cycles: What? How? When?". IMF Working Paper WP/11/76.
- Bartels, J. (1932). "Random Fluctuations, Persistence and Quasi-Persistence in Geophysical and Cosmical Periodicities". *Terrestrial Magnetism*, 37(3).
- Hartigan, J. A. and Hartigan, P. M. (1985). "The Dip Test of Unimodality". *Annals of Statistics*, 13(1), 70–84.
- Theiler, J., Eubank, S., Longtin, A., Galdrikian, B. and Farmer, J. D. (1992). "Testing for Nonlinearity in Time Series: the Method of Surrogate Data". *Physica D*, 58(1–4), 77–94.
- Schreiber, T. and Schmitz, A. (2000). "Surrogate Time Series". *Physica D*, 142(3–4), 346–382.
- Goertzel, G. (1958). "An Algorithm for the Evaluation of Finite Trigonometric Series". *American Mathematical Monthly*, 65(1).
- Holm, S. (1979). "A Simple Sequentially Rejective Multiple Test Procedure". *Scandinavian Journal of Statistics*, 6(2), 65–70.
- Migliorino. *Cicli di Borsa, Battleplan, Timing*. Teaching material. Historical author of the battleplan as a visual tool of cyclical context, and of the numbering of levels as a relative scale: it is from here that the cycle structure becomes a unit of measurement of time (Italian strand, phase 1).

16. Sartorelli, E. *Analisi Ciclica e Gann sul Bitcoin*. Extension to the crypto market and to the integration with Gann, conducted with an engineering approach: explicit definitions, repeatable procedure and rules that do not change when the market changes. It is the contribution of the Italian strand on which this work has been able to work best, because it is the one formulated precisely enough to be put to the test.
17. Elico64. *Appunti di Analisi Ciclica — Basi 1-10*. Primary source of the inverse cycle, of the polarity rules, of the definitions of conditional and unconditional swing and of the violation of cycle tops (Italian strand, phase 2).
18. Marini. Didactic formalization of the Italian school of cycle analysis; takes up and disseminates the Elico64 corpus.
19. Popper, K. R. (1959). *The Logic of Scientific Discovery*. Hutchinson, London. (Expanded English translation of *Logik der Forschung*, Vienna, 1934; Routledge Classics edition.)
20. Yule, G. U. (1927). "On a Method of Investigating Periodicities in Disturbed Series, with Special Reference to Wolfer's Sunspot Numbers". *Philosophical Transactions of the Royal Society of London, Series A*, 226, 267–298.
21. Slutsky, E. (1937). "The Summation of Random Causes as the Source of Cyclic Processes". *Econometrica*, 5(2), 105–146.
22. Fisher, R. A. (1929). "Tests of Significance in Harmonic Analysis". *Proceedings of the Royal Society of London, Series A*, 125(796), 54–59.
23. Brock, W., Lakonishok, J. and LeBaron, B. (1992). "Simple Technical Trading Rules and the Stochastic Properties of Stock Returns". *The Journal of Finance*, 47(5), 1731–1764.
24. White, H. (2000). "A Reality Check for Data Snooping". *Econometrica*, 68(5), 1097–1126.
25. Sullivan, R., Timmermann, A. and White, H. (1999). "Data-Snooping, Technical Trading Rule Performance, and the Bootstrap". *The Journal of Finance*, 54(5), 1647–1691.
26. Romano, J. P. and Wolf, M. (2005). "Stepwise Multiple Testing as Formalized Data Snooping". *Econometrica*, 73(4), 1237–1282.
27. Harvey, C. R., Liu, Y. and Zhu, H. (2016). "...and the Cross-Section of Expected Returns". *The Review of Financial Studies*, 29(1), 5–68.
28. Harvey, C. R. (2017). "Presidential Address: The Scientific Outlook in Financial Economics". *The Journal of Finance*, 72(4), 1399–1440.
29. Bailey, D. H. and López de Prado, M. (2014). "The Deflated Sharpe Ratio: Correcting for Selection Bias, Backtest Overfitting and Non-Normality". *The Journal of Portfolio Management*, 40(5), 94–107.
30. Bailey, D. H., Borwein, J., López de Prado, M. and Zhu, Q. J. (2017). "The Probability of Backtest Overfitting". *Journal of Computational Finance*, 20(4), 39–69.
31. Urquhart, A. (2016). "The Inefficiency of Bitcoin". *Economics Letters*, 148, 80–82.
32. Nadarajah, S. and Chu, J. (2017). "On the Inefficiency of Bitcoin". *Economics Letters*, 150, 6–9.
33. Brauneis, A. and Mestel, R. (2018). "Price Discovery of Cryptocurrencies: Bitcoin and Beyond". *Economics Letters*, 165, 58–61.
34. Caporale, G. M. and Plastun, A. (2019). "The Day of the Week Effect in the Cryptocurrency Market". *Finance Research Letters*, 31, 258–269.