

Analisi ciclica falsificazionista delle criptovalute

Un protocollo di validazione popperiana per la teoria ciclica dei mercati, con verifica empirica su criptovalute, indici azionari e trentasette strumenti in sei classi

Luigi Garone · Ricercatore indipendente · Cryptoverso — Ricerca ciclica

Settembre 2026 · Licenza CC BY-NC-ND 4.0 · Corrispondenza: luigi@cryptoverso.net

Confine di divulgazione. Tesi, metodo, ipotesi nulle, regole di verdetto e risultati statistici sono pubblicati e verificabili; il dataset e i verdetti in formato JSON sono forniti per la verifica indipendente, così che il lettore possa rieseguire i test sull'output consegnato e contestare i verdetti. Restano proprietari il rilevatore ciclico e la sua calibrazione — soglie, valori dei parametri, procedura di taratura, regole operative e codice — che il documento **descrive** senza consegnare.

Abstract

L'analisi ciclica dei mercati produce, da oltre mezzo secolo, framework concettuali ricchi ma raramente sottoposti a falsificazione empirica. Le tre scuole storiche — classica americana (Hurst), italiana (battleplan e ciclo inverso) e DSP quantitativo (Ehlers) — vengono adottate come dottrine alternative, quasi mai testate in modo congiunto su un campione statisticamente rilevante. Questo lavoro formalizza un protocollo popperiano che mette alla prova ciascuna affermazione fondante delle tre scuole su un dataset multi-asset — Bitcoin ed Ethereum dal 2017, Solana dal 2020, esteso a trentasette strumenti in sei classi e a sei indici azionari con fino a novantanove anni di storia giornaliera — su alcune migliaia di cicli gerarchici fasati.

La metodologia è causale per costruzione e deterministica; entrambe le proprietà sono formalizzate in test automatici e dimostrate empiricamente, in particolare sul modulo di rilevazione delle lingue di Bayer, verificato non-repainting in modalità walk-forward su cinque combinazioni asset/timeframe. A questa impalcatura il protocollo aggiunge quattro requisiti, che sono la parte del lavoro che sopravvive ai singoli risultati: **il valore nullo di una misura va misurato, non assunto**, con serie surrogate che conservano lo spettro di potenza e distruggono la struttura temporale; **un tasso condizionale va accompagnato dalla propria marginale**, cioè dalla stessa condizione contata sui casi che l'evento non seleziona; **un esito negativo non vale finché non si è mostrato che il test riconoscerebbe il fenomeno se ci fosse**; e **ogni livello va misurato sul timeframe che lo risolve**, perché il timeframe giornaliero non separa i livelli più brevi della gerarchia e misurarli lì cambia i numeri di cui si parla.

I risultati confutano le asserzioni canoniche più esposte: il nesting armonico fisso di Hurst e il principio di armonicità — quest'ultimo con un test la cui potenza è verificata su una scala sintetica vera, dove sul giornaliero nessuno dei mercati con abbastanza picchi risulta più vicino ai pioli di un insieme di periodi estratti a caso dalla stessa griglia —, la mean-reversion sequenziale delle durate, la soglia canonica di Bayer per i cicli di raccordo, la tenuta del massimo, e — per tautologia — due affermazioni della scuola italiana: il ciclo inverso, la cui definizione operativa non è falsificabile ex ante, e lo swing incondizionato, la cui condizione di successo è vera quasi sempre anche senza lo swing — in tutti i cicli padre del giornaliero, e in almeno nove su dieci su ogni righello. Confermano, con significatività statistica e convergenza cross-asset, la comunanza ciclica delle durate — forte ai livelli brevi e, su storia lunga, altrettanto forte ai livelli alti —, la traslazione ciclica come predittore di polarità, la violazione del massimo precedente e, in una forma riscritta che elimina un'inclusione insiemistica, i due segnali ribassisti del corpus italiano.

Lo stesso metro si applica alle regolarità che il lavoro potrebbe rivendicare come proprie, e lì è più severo. La **coda dei cicli brevi**, misurata contro cento surrogate a fase randomizzata e cento AAFT per mercato con ogni livello sul righello che lo risolve, **esce dalla nube in una combinazione asset/livello su trentatré**, sotto la soglia di falsificazione: è ciò che il rilevatore produce, non ciò che il mercato fa. Lo stesso vale per lo scarto dal periodo nominale e per la posizione del massimo dentro il ciclo. E l'**assioma di partenza** del lavoro, per cui una struttura quaternaria sarebbe un binario di binari, **non si può decidere** con questo strumento: il

minimo che separerebbe i due binari è il più basso dei tre interni nell'1,3% dei 557 cicli misurati, dentro la banda che lo stesso rilevatore produce sul rumore; e su serie sintetiche in cui il binario di binari è impiantato, il rilevatore non lo restituisce mai in numero sufficiente per votare.

L'apporto è triplice: un protocollo falsificabile e riproducibile applicato in modo uniforme a tutto il corpus, senza riguardo all'origine dell'asserzione; la dimostrazione di causalità e non-repainting della componente analitica centrale; e la distinzione sistematica fra ciò che il **mercato** produce e ciò che produce il **rilevatore**, che a più riprese cambia il segno di una conclusione — fino a togliere di mezzo, fra le regolarità che questo lavoro avrebbe potuto rivendicare come proprie, quelle che sembravano più solide.

1. Introduzione e motivazione

L'idea che i prezzi finanziari oscillino attorno a periodicità ricorrenti accompagna l'analisi tecnica dagli anni Trenta (Wheeler, Dewey, Bayer) e raggiunge formulazione matura con J. M. Hurst (*The Profit Magic of Stock Transaction Timing*, 1970) attraverso il Nominal Model e i sette principi cardine. Il filone italiano si articola in due fasi: una prima (Migliorino, Sartorelli) introduce il battleplan come strumento di contesto e le **strutture cicliche** come unità di misura del tempo, con il nesting a due e a tre, integrate con Gann sui mercati equity e crypto; una seconda (Elico64, poi ripresa in chiave didattica da Marini) formalizza in chiave operativa il ciclo inverso, le regole di polarità, le definizioni di swing condizionato e incondizionato e la violazione dei top ciclici.

Un debito va detto qui, perché riguarda il metodo prima che i risultati. Di tutto ciò che questo lavoro ha potuto mettere alla prova, la parte più testabile viene dalla scuola italiana: l'idea che la struttura ciclica sia un'**unità di misura** — e non una figura da riconoscere a occhio — è ciò che rende possibile contare, confrontare e falsificare. Sartorelli, in particolare, porta a quell'idea un'impostazione da ingegnere: definizioni esplicite, procedura ripetibile, estensione a un mercato nuovo senza cambiare le regole per farle tornare. E il corpus di **Elico64** arriva a un grado di formalizzazione — il ciclo inverso, la polarità, lo swing condizionato e incondizionato, la violazione dei top — che qui è stato possibile tradurre in test **senza doverlo interpretare**: sono definizioni operative, scritte da chi non disponeva dei mezzi per verificarle su decine di migliaia di cicli. Metterle alla prova è, in buona parte, il contenuto di questo lavoro.

Una dottrina si può falsificare solo se qualcuno l'ha formulata con abbastanza precisione, e su questo terreno il merito non è di chi misura: è di chi ha scritto in modo abbastanza chiaro da poter essere smentito. Sei delle affermazioni qui falsificate non sarebbero nemmeno state formulabili senza quel lavoro. Parallelamente J. F. Ehlers porta nel mondo finanziario il DSP classico: trasformata di Hilbert, MESA, SuperSmoother, autocorrelation periodogram.

Tre scuole, tre vocabolari, tre dottrine. Nessuna delle tre — nei testi fondativi e nella pratica di mercato — sottopone le proprie affermazioni cardine a un test empirico falsificabile su un campione di cicli statisticamente rilevante e su un asset moderno. Il mercato delle criptovalute offre l'occasione per chiudere questa lacuna: operatività continua, alta volatilità, storia digitale completa al minuto dal 2017, struttura ciclica visivamente riconoscibile, abbondanza di asset correlati per la validazione esterna.

Il lavoro persegue tre obiettivi: falsificare in modo formale (ipotesi → test → valore p → verdetto) le affermazioni cardine delle tre scuole; sintetizzare ciò che resiste in un protocollo causale, multi-livello e adattivo al regime; documentare le regolarità empiriche emerse nel processo, **inclusi i risultati negativi e le confutazioni delle proprie ipotesi**.

Premessa. La struttura ciclica visualizzata sul prezzo non è un sistema predittivo deterministico: è uno strumento di contesto che riduce la varianza decisionale fornendo all'analista un riferimento quantitativo della posizione nel ciclo. L'eventuale predittività emerge da segnali di livello inferiore condizionati dal contesto strutturale, non dalla sola sovrapposizione ciclica.

Questa premessa è essa stessa un'affermazione, e come tale è stata messa alla prova: il § 6.6 ne riporta l'esito, che la sostiene su sette righe su nove — con una riserva che ne definisce la portata, perché il contesto misurato è noto solo a posteriori.

2. Stato dell'arte: tre scuole a confronto

Scuola	Riferimento	Focus	Approccio	Limite chiave
Classica americana	J. M. Hurst (1970)	Nominal Model, sette principi, nesting armonico	Top-down deterministico	Nesting fisso mai testato fuori dall'azionario USA
Italiana, fase 1	Migliorino, Sartorelli	Battleplan, strutture a due e a tre, integrazione con Gann	Bottom-up visuale, manualistica	Battleplan come sistema meccanico difficile da falsificare
Italiana, fase 2	Elico64 (ripreso da Marini)	Ciclo inverso, polarità, swing, violazione dei top	Bottom-up con regole formalizzate	Definizione tautologica di «ciclo inverso»
DSP quantitativo	J. F. Ehlers (2001, 2013)	Filtri adattativi, frequenza istantanea	Signal processing	Assume stazionarietà locale; cieco al contesto strutturale
Sintesi (questo lavoro)	—	Fusione e falsificazione	Data-driven popperiano	—

Mapa delle tre scuole storiche e proposta di sintesi falsificazionista.

Ciascuna porta un limite proprio, ed è il limite a rendere le sue affermazioni interessanti da mettere alla prova: in Hurst il nesting fisso, mai testato fuori dall'azionario statunitense; nella prima fase italiana un battleplan difficile da falsificare perché non meccanizzato; nella seconda una definizione tautologica di «ciclo inverso»; in Ehlers una stazionarietà locale assunta e cieca al contesto strutturale.

Dai tre corpora sono state estratte le affermazioni operative più caratterizzanti e ridotte a un insieme di ipotesi falsificabili, ciascuna corredata da test statistico, soglia decisionale fissata a priori e campione minimo. A queste si aggiungono le ipotesi specifiche del corpus Elico64 sullo swing e sulle violazioni dei top ciclici, e le ipotesi proprie di questo lavoro. Il protocollo è applicato in modo uniforme: **le ipotesi dell'autore sono sottoposte allo stesso rigore di quelle delle scuole storiche**, e nella tavola dei verdetti (§ 8) compaiono con lo stesso dettaglio, comprese quelle cadute. La colonna «Fonte» di quella tavola riporta il conto separato per origine, ed è la riga più scomoda del documento.

2.1. Lavori correlati: la letteratura statistica a cui questo protocollo risponde

Le affermazioni messe alla prova qui vengono da tradizioni di praticanti, ma le obiezioni a cui devono sopravvivere vengono da una letteratura statistica che questo lavoro assume come metro.

Cicli che non ci sono. L'obiezione più antica a qualunque lettura ciclica di una serie di prezzi è che un processo puramente stocastico può sembrare periodico. Yule (1927) mostrò che un processo autoregressivo del secondo ordine mosso da shock casuali produce una pseudo-periodicità visibile — l'effetto che in seguito ha preso il nome suo e di Slutsky (1937) — e Fisher (1929) diede il test di significatività classico per l'ordinata massima di un periodogramma contro il rumore bianco. L'approccio a surrogati usato qui (Theiler et al. 1992; Schreiber e Schmitz 2000) è il discendente computazionale di quel test, ed è preferito per una ragione di metodo più che pratica: il test g di Fisher assume una nulla di rumore bianco gaussiano stazionario, un presupposto che le serie finanziarie, non stazionarie e a code pesanti, violano, mentre i surrogati a fase randomizzata e AAFIT fissano la nulla sullo spettro empirico della serie effettivamente studiata.

Molti test, un solo dataset. Un protocollo che mette alla prova decine di affermazioni su una manciata di mercati è esposto esattamente al guasto che la letteratura finanziaria documenta da due decenni: White (2000) ha formalizzato il bootstrap reality check per il data snooping, Sullivan, Timmermann e White (1999) lo hanno applicato alle regole di trading tecnico trovando che la migliore regola nel campione non sopravvive fuori campione, Harvey, Liu e Zhu (2016) hanno sostenuto che la ricerca accumulata sulla sezione trasversale dei rendimenti richiede soglie ben più severe di quelle convenzionali, e Harvey (2017) ne ha fatto una questione istituzionale. Questo lavoro risponde all'obiezione con soglie fissate prima dell'esecuzione, con la correzione di Holm (1979) dentro ogni famiglia dichiarata, e con i risultati negativi riportati con lo stesso standard di evidenza di quelli positivi. Due scelte meritano di essere dichiarate invece che date per scontate. La prima: Holm e non la procedura a passi di Romano e Wolf (2005), che lo domina in potenza sotto dipendenza: il metodo a passi richiede di ricampionare la distribuzione congiunta dell'intera famiglia di statistiche, mentre qui le nuvole di surrogati sono generate per mercato e per ipotesi, quindi la distribuzione congiunta non è disponibile. Il costo è la potenza, e lo si dichiara: sotto Holm un verdetto NON PROVATA può essere un falso negativo, ed è precisamente per questo che ogni esito negativo di questo documento è

accompagnato da un controllo di potenza. La seconda: non si usano né lo Sharpe ratio deflazionato (Bailey e López de Prado 2014) né la probabilità di overfitting del backtest (Bailey, Borwein, López de Prado e Zhu 2017), e la loro assenza non è una svista: entrambi correggono la selezione di una strategia fra molte in base alla sua statistica di performance, mentre l'oggetto messo alla prova qui è un rilevatore misurato contro una nulla, non una strategia misurata dal suo Sharpe ratio. Se il quadro dovesse in seguito produrre segnali operativi con una statistica di performance associata, quelle correzioni diventerebbero necessarie e la loro assenza andrebbe argomentata di nuovo.

Il mercato studiato. La tesi che le criptovalute siano informativamente inefficienti, e quindi un luogo plausibile in cui cercare struttura sfruttabile, fu avanzata in modo sistematico per la prima volta da Urquhart (2016) e subito contestata da Nadarajah e Chu (2017), che arrivarono alla conclusione opposta sugli stessi dati con una trasformazione di potenza dei rendimenti. Quello scambio è citato qui di proposito e per intero: è la dimostrazione più chiara, dentro la stessa letteratura sulle criptovalute, che il verdetto segue il test, ed è la ragione per cui questo lavoro misura il valore nullo di ogni statistica invece di assumerlo. Brauneis e Mestel (2018) mostrano che l'efficienza non è uniforme nell'universo delle criptovalute, ed è parte della ragione per cui il test più ampio qui copre trentasette strumenti in sei classi invece di un solo asset. Caporale e Plastun (2019) mettono alla prova una periodicità di calendario — un effetto giorno della settimana — sulle criptovalute; l'oggetto è diverso da quello studiato qui, una regolarità di calendario esogena invece di cicli endogeni multiscala, e la distinzione va fatta esplicitamente perché le due cose si confondono facilmente.

Il precedente di metodo più vicino. Brock, Lakonishok e LeBaron (1992) misero alla prova regole a media mobile e a trading range su novant'anni di dati del Dow Jones contro quattro modelli nulli dichiarati esplicitamente — random walk, AR(1), GARCH-M, EGARCH — stimati per bootstrap. Quel disegno, modelli nulli dichiarati in anticipo e misurati invece che assunti, è quello che questo protocollo segue; ciò che si aggiunge qui è il requisito di potenza prima che un esito negativo conti, il marginale riportato accanto a ogni tasso condizionato, e la regola per cui ogni livello si misura sul timeframe che lo risolve.

La cornice epistemologica è Popper (1959), ed è presa alla lettera e non come ornamento: un'affermazione guadagna il verdetto CONFERMATA solo sopravvivendo a un test che avrebbe potuto confutarla, e il documento distingue FALSIFICATA da NON PROVATA dall'inizio alla fine.

3. Principio del metodo

Confine di divulgazione. Questa sezione descrive il principio concettuale del metodo: che cosa fa il sistema, su quali fonti teoriche si fonda e con quale logica causale. Non espone soglie di calibrazione, valori di parametri, la procedura di calibrazione, il set delle regole decisionali né codice. Metodologia e risultati statistici sono invece interamente documentati, e il dataset è fornito per la verifica indipendente.

3.1. Architettura concettuale: tre livelli

Il metodo separa il problema in tre livelli, e la separazione è essa stessa un contributo concettuale: **L1**, struttura ciclica — dove siamo nella gerarchia dei cicli; **L2**, timing DSP — quando si apre e si chiude il ciclo corrente; **L3**, conferma di order flow, predisposto come estensione. Ciascun livello risponde a una domanda diversa e usa evidenza diversa: è questo che rende il problema meccanizzabile.

3.2. Gerarchia ciclica e finestra ciclica attesa

Il primo livello risponde alla domanda «dove siamo adesso?». Il prezzo viene decomposto in una gerarchia di cicli annidati, dal ciclo dominante di lungo periodo fino ai sotto-cicli minori, secondo il principio di comunanza ciclica e di sincronicità dei minimi della tradizione hurstiana. La scoperta dei periodi candidati avviene per via spettrale, con un criterio che non assume stazionarietà sull'intera finestra; il tracciamento del ciclo dominante e della sua fase istantanea segue il DSP causale di Ehlers.

La **finestra ciclica attesa** è la nozione che rende applicabile tutto il resto ed è definibile simbolicamente. Sia P_L il periodo nominale del livello L e t_0 il trough precedente confermato a quel livello. La fase ciclica corrente è

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

che vale 0 al trough, $\approx 0,5$ a metà ciclo e ≈ 1 alla chiusura attesa. Un nuovo trough di livello L è atteso nella finestra

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

dove τ è la tolleranza relativa che assorbe l'elasticità ciclica. La struttura della finestra — ancoraggio al trough precedente e ampiezza proporzionale al periodo nominale — è sufficiente a riprodurre il concetto; il valore di τ e il criterio di ammissione del periodo nominale per asset restano proprietari, così come i criteri di promozione e rimozione dei livelli e le tolleranze di allineamento fra livelli.

3.3. Detection causale delle lingue: la matematica del criterio

Una **lingua di Bayer** è un ciclo anomalmente corto che funge da raccordo: segnala una presa di liquidità con ripartenza, oppure un ciclo di transizione fra due strutture maggiori. Il sistema la rileva con un classificatore basato su regole — non un modello opaco — e causale: ogni valutazione usa solo dati fino all'istante corrente. La logica è replicabile a livello di famiglie di criteri, articolate in due passi.

Passo 1 — brevità nella finestra. Il ciclo è candidato lingua se la sua durata sta sotto un **percentile basso** della distribuzione di durata osservata *fino a t* per quel livello e quell'asset. **Passo 2 — criterio di prezzo causale.** Un ciclo corto da solo non basta: il criterio, di stampo Dow e Migliorino, chiede un rally di ampiezza sufficiente dal minimo candidato e il superamento dei massimi recenti entro la finestra.

In simboli, sia D la durata del ciclo candidato misurata fra due trough consecutivi: il ciclo è candidato se $D < \delta_{\text{short}}(L)$, con δ_{short} stimato in modo adattivo e causale sul livello L e sull'asset. **La famiglia del criterio — soglia data da un percentile basso, stimato causalmente — è ciò che il lettore deve replicare; il valore del percentile è riservato**, ed è uno dei punti in cui il confine di divulgazione morde. La differenza dalla soglia canonica di Bayer, che fissa la brevità a una frazione costante della media, è sostanziale e viene verificata al § 6.10.

Entrambe le condizioni del passo 2 sono causali per costruzione: guardano barre già chiuse, e la finestra entro cui cercano il superamento è determinata dal livello del ciclo candidato, non dal suo esito.

Dei due passi, uno solo porta informazione. Il passo 1, isolato dal passo 2, **non ne porta alcuna** sulla ripartenza (§ 6.10): il potere discriminante del criterio sta interamente nel passo di prezzo. Conta per chi voglia replicare il metodo, e conta anche come avvertimento — la brevità serve a **selezionare** i candidati, non a **classificarli**, e usarla da sola darebbe un criterio che sembra funzionare perché non decide nulla.

3.4. Conferma a finestra chiusa

Una lingua si conferma quando la sua finestra di rilevazione si chiude, esattamente come uno swing si conferma a fine ciclo. Da quel momento la sua etichetta storica non cambia più. L'unica entità mobile è l'ultima lingua a finestra ancora aperta, che per costruzione può aggiornarsi: è un comportamento atteso e non un repaint causale. Tutti i risultati di questo lavoro si basano esclusivamente sulle lingue confermate.

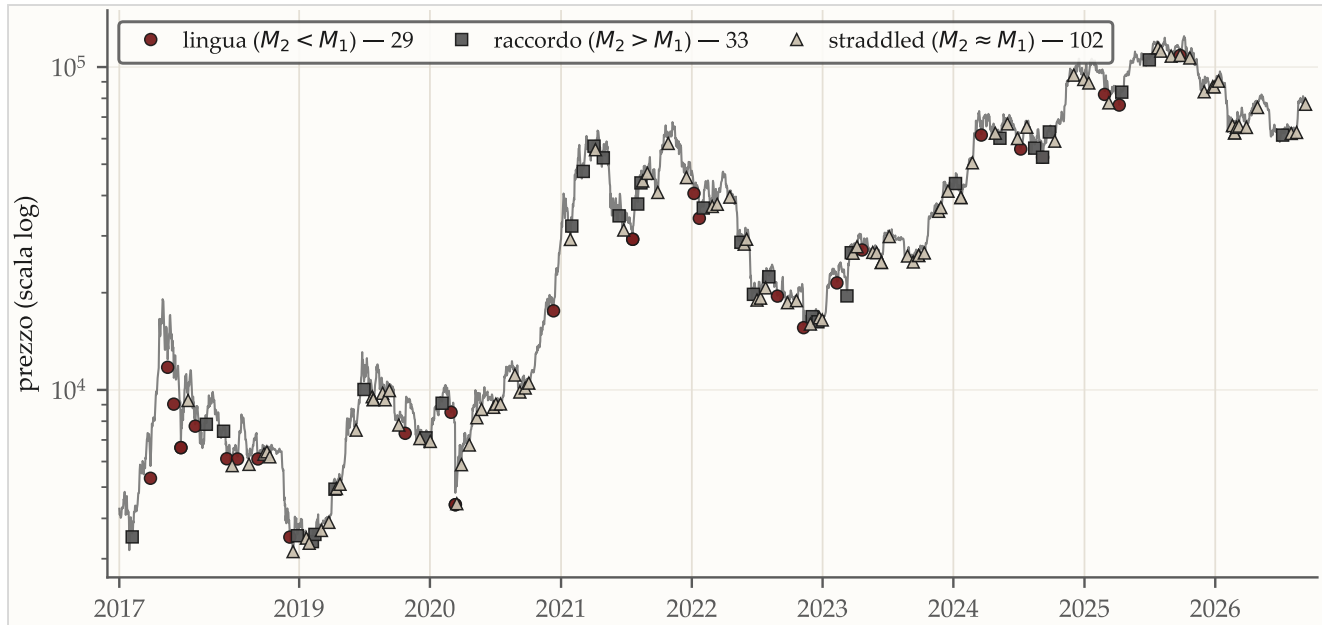


Fig. 1: Le lingue confermate dal rilevatore causale su nove anni di Bitcoin, in scala logaritmica. Ogni marcatore è una lingua al proprio minimo, colorata per categoria secondo la relazione fra il minimo del ciclo corto e quello del precedente. È l'unica figura di questo documento che mostra l'uscita del sistema invece di una statistica su di essa, e serve a una verifica che nessun numero fa: la rilevazione è **densa e distribuita** sul toro 2020–2021, sull'orso 2022 e sulla ripresa successiva, non concentrata in un regime.

4. Nomenclatura ciclica unificata

Le tre scuole usano nomi diversi per gli stessi oggetti. La scala relativa, di uso italiano, è adottata qui come vocabolario, con le durate **misurate** accanto ai periodi nominali di listino. L'origine della scala non è arbitraria e vale ricordarlo: **T sta per trading cycle**, il livello su cui si opera, ed è il termine con cui Migliorino introduce la numerazione. La scala è relativa nelle distanze, non nell'origine. La scala **non si ferma al ciclo di dieci giorni**: scende fino al ciclo giornaliero e sotto, e la colonna del timeframe che risolve ogni livello dice perché il giornaliero non basta a misurarli.

Livello	Nome	Nominale	Risolto su	Durata misurata
T+9	—	~15 anni	W	fuori dalla storia disponibile
T+8	Hyper	~8 anni	W	1 durata per mercato: un intervallo, non una distribuzione
T+7	Super	~4 anni	W	1–2 durate per mercato
T+6	—	480 g (duplicato di T+5)	D	non prodotto da nessun righello
T+5	Long	480 g	D	Tabella 2
T+4	Major	160 g	D	Tabella 2
T+3	Minor	90 g	H4	Tabella 2
T+2	Short	45 g	H4	Tabella 2
T+1	Swift	20 g	H2	Tabella 2
T	Base	10 g	H1	Tabella 2
T-1	Micro	5 g	M30	Tabella 2
T-2	Nano	2,5 g	M15	Tabella 2
T-3	giornaliero	1 g	M5	Tabella 2
T-4	metà-giornaliero	~12 h	M5	Tabella 2
T-5	—	~7 h	M1	Tabella 2
T-6	—	~5 h	M1	Tabella 2
T-7	—	~3 h	M1	Tabella 2
T-8	—	~1,5 h	nessuno	nessun righello ammesso lo risolve

Tab. 1: Nomenclatura ciclica unificata, dal ciclo pluriennale a quello di tre ore, con il **righello che risolve** ogni livello. Le durate misurate stanno in una tavola sola — **Tabella 2** — e non sono ripetute qui: la durata di un ciclo letta dove il rilevatore non lo separa è la durata di un altro oggetto, quindi un livello senza il suo righello accanto non è un numero confrontabile. Ai due livelli in alto la durata non compare perché **con una o due durate non c'è una distribuzione, c'è un intervallo**: su nove anni di storia un ciclo da quattro anni entra due volte, e la terza volta che compare su Ethereum esiste solo perché il rilevatore ha accettato due cicli più corti del nominale. T-5, T-6 e T-7 hanno minimi propri ma **non sono distinti fra loro** (§ 5.6.1).

La colonna «risolto su» non è una scelta redazionale: la calcola la stessa regola di produzione che il § 5.6 descrive, cercando il timeframe in cui il livello cade dentro la finestra di risoluzione e più vicino al suo centro geometrico. È la ragione per cui una nomenclatura completa **deve** avere quella colonna: senza, il lettore assume che i numeri di ogni riga siano confrontabili fra loro, e non lo sono.

La scala conta **diciotto gradini**, da T+9 (~15 anni) a T-8 (~1,5 ore), e non sono misurabili tutti dallo stesso timeframe: **T+6 non è mai prodotto** su nessuno dei tre mercati, T+7, T+8 e T+9 non hanno campione sufficiente sul solo campione crypto (§ 6.11), T-8 non ha un righello ammesso che lo risolva, e sotto il ciclo di dieci giorni il Daily non separa più — lì la misura va fatta dove il livello si risolve (§ 5.6). Restano **tredici livelli misurati**, ed è la tavola che segue.

Il listino descrive bene le durate misurate, e lo scarto ha un verso. Misurata dove il livello si risolve, la mediana sta vicino al nominale su tutta la scala breve — a T vale 9,5 giorni contro 10, a T-1 4,7–4,9 contro 5, a T+1 18,8–19,8 contro 20 — e lo scarto è quasi sempre **verso il basso**, di pochi punti percentuali. Il gradino si allarga solo a T+4, dove la mediana sta fra il 10% e il 24% sotto il nominale. È il verso che ci si aspetta da un rilevatore che chiude i cicli sui minimi e impone una durata minima: taglia le code lunghe più di quanto allunghi quelle corte. Non è una conferma del listino come modello del mercato — è la constatazione che il **nostro** rilevatore, letto sul righello giusto, restituisce le durate del catalogo entro pochi punti.

Livello	Righello	Nominale	Bitcoin	Ethereum	Solana	Scarto
T-7	M1	0,13 g	0,13 (4.870)	0,13 (4.931)	0,13 (4.816)	0% ... -3%
T-6	M1	0,20 g	0,18 (3.438)	0,18 (3.433)	0,18 (3.394)	-8% ... -9%
T-5	M1	0,30 g	0,27 (2.318)	0,27 (2.298)	0,27 (2.281)	-9% ... -10%
T-4	M5	0,50 g	0,48 (6.613)	0,48 (6.617)	0,48 (4.407)	-4% ... -5%
T-3	M5	1,00 g	0,94 (3.297)	0,94 (3.276)	0,95 (2.175)	-5% ... -6%
T-2	M15	2,50 g	2,37 (1.318)	2,38 (1.317)	2,43 (873)	-3% ... -5%
T-1	M30	5 g	4,81 (648)	4,74 (664)	4,94 (433)	-1% ... -5%
T	H1	10 g	9,46 (333)	9,46 (328)	9,58 (218)	-4% ... -5%
T+1	H2	20 g	19,79 (158)	18,83 (162)	19,25 (105)	-1% ... -6%
T+2	H4	45 g	39,67 (73)	44,08 (70)	41,42 (50)	-2% ... -12%
T+3	H4	90 g	88,17 (35)	96,83 (35)	90,17 (23)	-2% ... +8%
T+4	D	160 g	144,0 (21)	121,0 (21)	137,0 (13)	-10% ... -24%
T+5	D	480 g	495,0 (5)	481,0 (5)	— (4)	+0,2% ... +3%

Tab. 2: Le durate **accertate dal sistema**, ciascuna misurata sul righello che risolve il suo livello. Tredici livelli su diciotto, dal minuto al giorno. Le mediane sono in **giorni**; fra parentesi il numero di cicli su cui la mediana è calcolata, perché una mediana senza il suo denominatore non è verificabile. Lo scarto è fra la mediana e il nominale del listino, sui tre mercati. La riga di Solana a T+5 è vuota perché quattro durate non fanno una distribuzione, e la misura lo dichiara invece di riportare un numero. Ogni valore viene da **CR-009**, che porta le stesse durate **in barre del proprio timeframe** con il nominale convertito nella stessa unità: a T, per esempio, 227 barre orarie contro 240.

I sei livelli in basso entrano ora, e sono la resa concreta del pavimento sceso al minuto: fino alla corsa precedente T-2 e i cinque sotto non avevano un righello ammesso e non venivano misurati. Il loro scarto dal nominale sta fra lo zero e il -10%, cioè nella stessa fascia dei livelli alti, su campioni di migliaia di cicli invece che di decine.

T+6 non compare. La procedura non produce alcun livello in quell'intervallo di periodi: fra T+5, che sui tre mercati sta fra 368 e 495 barre, e T+7, che su Ethereum vale 1.153 barre, non c'è gradino intermedio con minimi propri e sostegno spettrale.

4.1. Il livello T+3, e perché la sua misura era doppia

Per il solo livello T+3 il materiale di supporto ha contenuto a lungo **due misure incompatibili** prodotte dalla stessa gerarchia: 28/24/19 cicli con mediane 92,5 / 104 / 89 barre da un modulo, 38/37/28 cicli con mediane 75 / 72 / 75 da altri due. Gli altri livelli coincidevano. Una discordanza di questo tipo non si risolve scegliendo la serie più comoda: si isola la causa.

La misura la isola in due passi. **Primo:** si leggono le durate di T+3 dalla stessa gerarchia per le vie dei tre moduli, con i filtri di livello disattivati. **Secondo:** si riaccende il filtro e si misura di quanto sposta conteggio e mediana. Entrambi i passi girano su ciascuno dei nove righelli, e il verdetto si legge su **H4** — il righello che risolve T+3 (§ 5.6).

Sul righello che risolve il livello le letture coincidono cifra per cifra: 35 cicli e mediana 529,0 barre su Bitcoin, 35 e 581,0 su Ethereum, 23 e 541,0 su Solana — in giorni 88,2 / 96,8 / 90,2, le stesse durate di **Tabella 2**. Non esiste una seconda causa da cercare: nessuna divergenza di gerarchia, nessun minimo diverso, nessun conteggio diverso. Con un'avvertenza che la misura dichiara da sé: due delle tre vie sono **identiche per costruzione** — entrambe la differenza fra minimi consecutivi — quindi i testimoni distinti sono due, non tre.

E il filtro, su quel righello, non sposta niente: 35 cicli restano 35, 529,0 barre restano 529,0, e così sugli altri due mercati. La soglia è **60 barre** e la mediana di T+3 su H4 vale **529**: un pavimento di sessanta barre sotto una distribuzione che vive attorno a cinquecento non incontra nulla da tagliare.

La causa c'è, ed è più profonda di un parametro. La doppia misura nasceva sul **Daily**, dove un ciclo di T+3 dura una novantina di barre e un pavimento di sessanta ne taglia una parte; sul righello che risolve il livello lo stesso pavimento è invisibile. Non era un disaccordo fra moduli e non era soltanto un parametro applicato senza dichiararlo: era **un livello misurato su un righello che non lo risolve**, cioè la tesi che § 5.6 misura per conto proprio. La serie da usare è quella non filtrata, letta su H4; l'altra va

ritirata, e le conclusioni che dipendevano da T+3 — la convergenza cross-asset a quel livello e il rapporto di nesting T+3 → T+4 — non sono più provvisorie.

La regola di verdetto, applicata alla lettera, falsifica. Per dichiarare *causa accertata* era stata fissata a priori una riduzione di almeno il 20% dei cicli e un aumento di almeno 15 barre della mediana **su tutti e tre i mercati**; la regola dichiarava falsificata l'ipotesi se nessun mercato riproduceva l'effetto. Su H4 la riduzione è **zero**: il verdetto formale è **FALSIFICATA**. Va letto per quello che dice: non che il filtro fosse innocente, ma che **su quel righello non può agire**. La regola, scritta quando la misura viveva su un righello solo, non distingue *l'effetto non c'è* da *l'effetto non può esserci*, e la distinzione non si recupera riscrivendola dopo aver visto il numero. La domanda storica — il filtro spiegava la doppia misura **sul Daily?** — non è più computabile: là la lettura che passa per l'insieme dei livelli non produce più alcun ciclo di T+3, e resta solo la differenza grezza fra minimi (32 / 34 / 24 cicli, mediane 98,5 / 102,5 / 88,0 giorni).

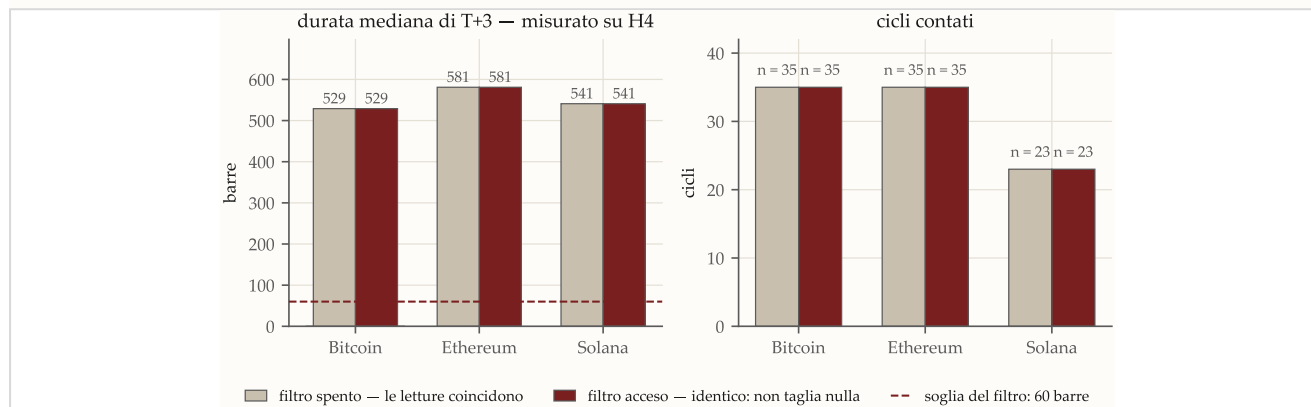


Fig. 2: Il livello T+3 sul righello che lo risolve, letto per tre vie con il filtro spento — le letture coincidono, stesso conteggio e stessa mediana — e le stesse grandezze con il filtro acceso, che sono le medesime: a cinquecento barre di mediana un pavimento di sessanta non taglia nulla. La doppia misura del materiale precedente non era un disaccordo fra moduli, era un livello letto su un righello che non lo risolve.

5. Metodologia di validazione

5.1. Dataset e copertura

Voce	Valore
Mercati primari	BTCUSDT, ETHUSDT e SOLUSDT, dal 2017 (Solana dal 2020), granularità da un minuto a settimanale
Prova cross-asset larga	37 strumenti in sei classi: criptovalute, azioni, indici, commodities, valute, obbligazionario (§ 6.14)
Storia lunga	6 indici azionari: ^GSPC 24.775 barre (dal dicembre 1927), ^N225 15.149, ^IXIC 14.001, ^FTSE 10.769, ^GDAXI 9.770, ^DJI 8.719 (§ 6.11)
Timeframe delle affermazioni	Daily per l'impianto storico; nove righelli — W, D, H4, H2, H1, M30, M15, M5 e M1 — per le rimisure sul timeframe che risolve il livello (§ 5.6) e per la verifica cross-timeframe (§ 6.13)
Dato atomico	OHLCV a un minuto da fonte exchange, partizionato Parquet; ogni timeframe superiore è derivato dal minuto con resampler certificato no-repainting
Taglio dei dati	ultimo scambio della corsa che ha rifatto trenta carte in una lettura sola: 5 settembre 2026 (§ 5.7)

Tab. 3: Dataset di riferimento. Il numero di barre di ciascuna misura, con l'impronta dei prezzi, è dichiarato nel proprio file di evidenza e riassunto al § 5.7.

5.2. Protocollo popperiano

- 1. Causalità.** Ogni filtro, feature e segnale produce output al tempo t usando esclusivamente dati disponibili fino a t . Un test automatico di *prefix invariance* verifica che l'output al tempo t non cambi aggiungendo barre successive. È il presupposto irrinunciabile: un risultato calcolato con dati futuri produce falsa confidenza e invalida ogni validazione.

2. **Falsificabilità.** Ogni ipotesi è espressa con un test statistico, una soglia decisionale fissata a priori e un campione minimo. I verdeti ammessi sono CONFERMATA, FALSIFICATA, PARZIALE e NON PROVATA. L'ultimo verdetto è necessario e non è una sfumatura: «smentito» e «non sostenuto» sono **due esiti diversi**, e confonderli fa passare per confutazione ciò che è soltanto assenza di prova — o, peggio, per conferma ciò che è soltanto assenza di smentita.
3. **Le soglie prima dei numeri.** Ogni misura dichiara le proprie soglie di verdetto nel modulo che la calcola, **prima** di essere eseguita, e non le ritocca dopo. In più di un caso questo produce un verdetto più severo di quanto i numeri suggerirebbero: dove accade, il documento lo dichiara apertamente come fragilità, con scritto in chiaro il valore che avrebbe cambiato l'esito (§ 4.1, § 6.3). È l'unico modo di mantenere la disciplina senza nascondere informazione.
4. **Determinismo e riproducibilità.** Semi fissi e dipendenze bloccate: esecuzioni ripetute producono output identici. Ogni numero pubblicato vive in un file di misura versionato che contiene il commit del codice, il seme, l'elenco dei dataset e l'impronta SHA-256 dei prezzi. **Nessun numero compare in questo documento se non esiste in un file di misura**, e i file sono consegnati come materiale di supporto.
5. **Validazione esterna.** Convergenza cross-asset come test di generalità, e analisi walk-forward per la stabilità temporale delle classificazioni.

5.3. Il valore nullo si misura: il ruolo dei surrogati

È il requisito che regge tutto il resto, e va enunciato prima dei risultati perché nessuno di essi avrebbe senso senza: **un valore di riferimento si misura, non si assume.**

Molte affermazioni dell'analisi ciclica hanno la forma «la grandezza X e la grandezza Y sono legate». Il test ingenuo confronta la misura osservata con il valore che avrebbe sotto indipendenza. Il confronto è però quasi sempre sbagliato, perché **il rilevatore di cicli induce da solo una relazione**: qualunque procedura che chiuda un ciclo su un minimo produce, anche su rumore, cicli lunghi mediamente più ampi dei brevi.

Il valore nullo corretto si ottiene con **serie surrogate**: repliche sintetiche che conservano alcune proprietà della serie vera e ne distruggono altre.

- **Surrogati a fase randomizzata:** si calcola la trasformata di Fourier della serie, si sostituiscono le fasi con valori casuali uniformi e si antitrasforma. Conservano esattamente lo spettro di potenza — cioè le periodicità — e distruggono la posizione temporale degli estremi.
- **Surrogati AAFT** (*amplitude adjusted Fourier transform*): conservano anche la distribuzione dei valori della serie originale, e sono quindi più severi.

Il protocollo è: si applica **lo stesso rilevatore** ai dati veri e a un centinaio o due di repliche per mercato e per tipo, si misura la stessa quantità, e si dichiara conferma solo se il valore osservato cade fuori dalla nube dei surrogati.

Una nulla, però, non è sempre la stessa cosa di un'altra, e la scelta decide il segno del risultato. Vale in entrambe le direzioni, e il documento incontra tutti e due i casi:

Una nulla che distrugge anche la geometria del rilevatore misura il rilevatore, non il mercato. Una nulla che conserva proprio ciò di cui si discute non misura niente.

Il § 6.12 è il primo caso — rimescolare tutte le polarità insieme distrugge la posizione nel ciclo padre, che è una proprietà del rilevatore, e produce un valore p di 0,0005 dove la nulla che conserva la posizione dà 0,71. Il § 6.2 è il secondo: surrogati che conservano lo spettro conservano anche i picchi, cioè esattamente l'insieme di cui si discute la disposizione, e su quella domanda non hanno alcun potere.

Una nota sui valori p empirici. Quelli stimati per ricampionamento — surrogati, permutazioni, bootstrap — sono riportati con lo stimatore $(k + 1)/(n + 1)$ e non come la frazione grezza k/n : con n estrazioni il minimo dichiarabile è $1/(n + 1)$, e scrivere 0,000 dichiarerebbe una precisione che il campionamento non ha. Dove invece una soglia era stata fissata a priori sul **quantile empirico**, quella grandezza resta la regola di decisione e il valore p valido le compare accanto nell'evidenza.

5.4. La marginale accanto al condizionale

Il secondo requisito riguarda i tassi di successo, ed è quello che questo lavoro ha imparato a proprie spese su due affermazioni che aveva considerato concluse:

Ogni tasso di successo va accompagnato dal tasso della stessa condizione sui casi che l'evento NON seleziona. Se i due coincidono, l'evento non seleziona nulla e il tasso non è evidenza, qualunque valore abbia.

È la forma generale della critica che il § 8.1 muove al ciclo inverso di Elico64: una definizione che si applica solo a posteriori non è falsificabile ex ante. Applicata allo swing incondizionato mostra che l'evento non seleziona (§ 6.7); applicata alla rottura del minimo di partenza mostra che il numero pubblicabile era gonfiato da un'**inclusione insiemistica**, cioè da un insieme di eventi che contiene per costruzione tutti gli esiti positivi (§ 6.8). Lo stesso presidio, applicato allo stesso modo a due affermazioni vicine dello stesso corpus, ne ritira una e ne rafforza l'altra: non è uno strumento per demolire.

5.5. Un esito negativo vale solo se il test avrebbe potere

Il terzo requisito è simmetrico al secondo e chiude la porta all'altro modo di sbagliare:

Un test che non falsifica mai e uno che falsifica sempre sono ugualmente inservibili. Prima di leggere un esito negativo va verificato che il test riconoscerebbe il fenomeno se ci fosse, su un campione sintetico costruito per averlo.

Il § 6.2 lo esegue esplicitamente, e senza quel controllo il suo verdetto non avrebbe valore: la prima formulazione del test di armonicità dava un esito negativo che non distingueva «l'armonicità non c'è» da «questo test non la vedrebbe comunque». La stessa logica governa i controlli di identità e di inadeguatezza del § 5.6.

5.6. Ogni livello sul timeframe che lo risolve

Il quarto requisito è il più recente e il più invasivo, perché non riguarda una singola misura ma il **righello** con cui sono state prese quasi tutte.

Un livello ciclico ha un periodo nominale espresso in unità di tempo, non in barre. Su un dato timeframe quel periodo diventa un numero di barre N , e N decide che cosa il rilevatore può vedere: sotto una certa risoluzione due livelli adiacenti non sono più due rilevatori diversi, sono lo stesso rilevatore chiamato con due nomi. La griglia di ammissibilità che ne risulta si **fattorizza** in due assi indipendenti, ed è un fatto misurato e non un'assunzione: il numero di cicli C che un livello possiede è **invariante al timeframe** a parità di mercato — scarto relativo massimo 0,0015 su 54 coppie mercato/livello, contro una soglia dichiarata dell'1% — mentre la risoluzione N dipende solo dal timeframe. **Scendere di timeframe non regala campione: aumenta soltanto la risoluzione con cui quel campione è visto.**

Il controllo del criterio è il gruppo dei timeframe in cui un tetto sulle barre servite morde: lì la storia è una frazione dichiarata di quella disponibile e C deve cadere nella stessa proporzione. Cade infatti di 0,70–0,80 su tutti e diciotto i livelli di ciascun mercato, due ordini di grandezza sopra la soglia: il criterio discrimina, e l'invarianza non è saturazione.

Da qui una **regola del proprietario**: il timeframe che possiede un livello è quello in cui N cade dentro la finestra ideale e più vicino al suo centro geometrico. La scelta si fa sull'asse N e **solo** su quello, perché C è invariante e non può discriminare.

Timeframe	Livelli che possiede	Livello più basso che separa davvero
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

Tab. 4: Chi possiede che cosa, col pavimento sceso al minuto. Diciassette livelli su diciotto hanno un proprietario secondo la regola dell'asse N; **tredici** vengono effettivamente misurati, ed è la tavola delle durate. Fuori restano T+6 — che la procedura non produce, si veda sotto — e i tre livelli alti T+7, T+8 e T+9, per i quali nove anni di storia non bastano. La regola è «sul timeframe che lo **risolve**», non «sul più centrale»: cambia l'insieme dei candidati, non il criterio di scelta.

Il tetto si è alzato, e con esso la griglia. Col limite precedente di 100.000 barre nessun righello sotto l'ora era ammesso e sei livelli su diciotto erano dichiarati non risolvibili *da questa griglia*; col tetto a un milione i quattro righelli finì entrano e quei livelli trovano il proprio. **Non erano livelli non misurabili: erano livelli che quella griglia non risolveva**, ed è la distinzione che il resto del documento usa.

La conseguenza non è teorica, ed è misurata. Le stesse quattro statistiche di un livello — quanti cicli, durata mediana, coefficiente di variazione delle durate, quota di cicli ribassisti — calcolate sul Daily e sul timeframe che quel livello lo risolve, con la **stessa** catena di produzione e sulla **stessa** istantanea di prezzo, **discordano in 12 coppie su 12**. Le soglie di concordanza erano dichiarate prima ed erano generose (10% sulla mediana, 0,05 sul CV, 10% sulla quota, rapporto fra 0,80 e 1,25 sul numero di cicli). Due controlli delimitano il criterio dai due lati: il Daily contro sé stesso concorda su 12 coppie su 12 — il criterio non spara a vuoto — e il Daily contro il Weekly, dove quei livelli valgono fra 2 e 13 barre, discorda 12 volte su 12 — il criterio non è saturo.

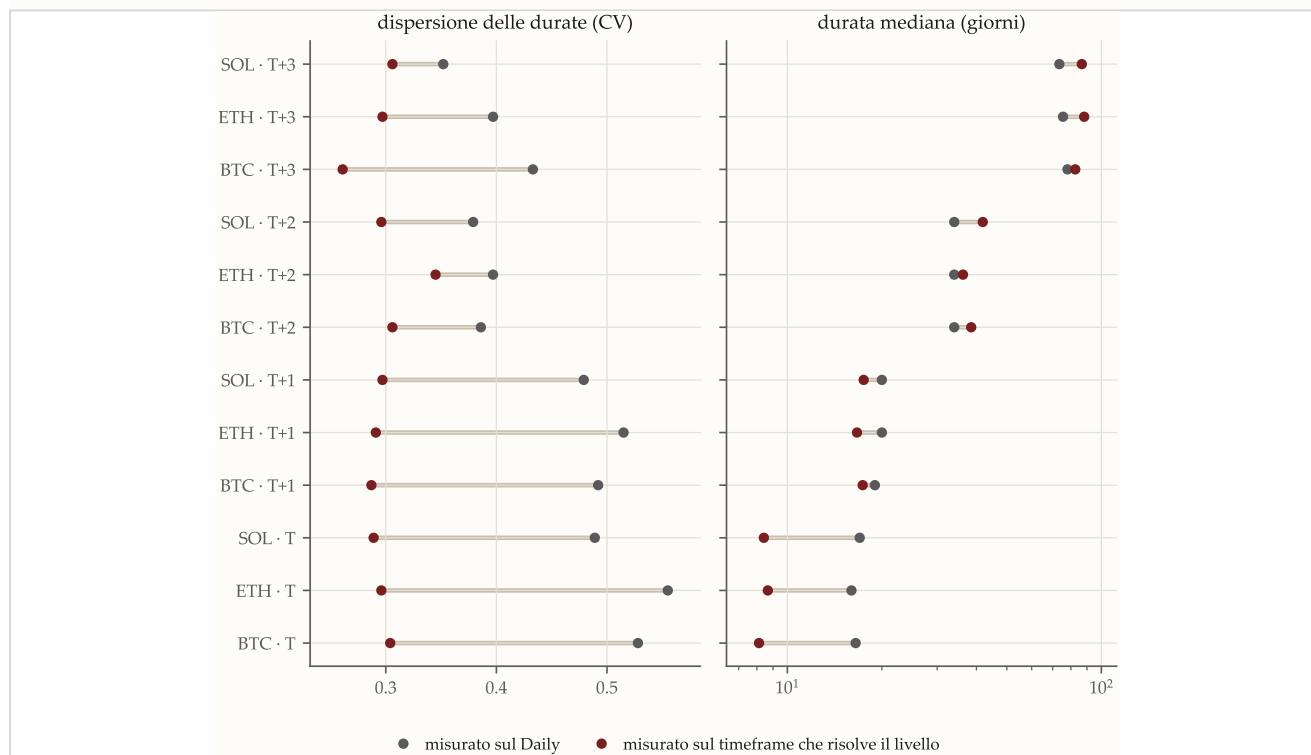


Fig. 3: Le stesse quattro statistiche, misurate sul Daily e sul timeframe che possiede il livello. A sinistra il coefficiente di variazione delle durate, che è più alto sul Daily in **dodici coppie su dodici**; a destra la durata mediana in giorni, dove il livello T cambia di un fattore due. La quota di cicli ribassisti, che non compare perché non discorda mai, è la statistica robusta al cambio di righello.

Il coefficiente di variazione misurato sul Daily è gonfio, e in una direzione sola. Il rapporto fra il timeframe proprietario e il Daily ha mediana 0,61: misurato dove è risolto, lo stesso livello ha una dispersione delle

durate **circa il 40% inferiore**. Questo tocca direttamente la lettura dei cicli come «quasi-fissi con elasticità del venti-quaranta per cento»: una parte di quella dispersione non è elasticità del mercato, è risoluzione del righello. Quanta, questa misura non lo dice — dice che il righello ne produce da solo abbastanza da spostare il numero di un terzo.

Il livello T misurato sul Daily non è T. La sua mediana sul Daily vale 16,5, 16,0 e 17,0 giorni sui tre mercati, contro 8,1, 8,7 e 8,4 sul timeframe che lo risolve: un rapporto fra **1,85 e 2,03**, cioè un fattore due, e il numero di cicli più che raddoppia (da 176, 175 e 114 a 374, 365 e 245). Sedici giorni non è il nominale di T, che vale dieci: è quello di T+1, che vale venti. Sul Daily il livello T dispone di dieci barre, che la griglia marca come non risolvibili, e il rilevatore fa la cosa prevedibile — trova la struttura più vicina che quelle barre permettono di vedere. È lo stesso fatto che il § 6.1 registra come «T e T+1 non sono distinti», visto dal lato della causa invece che da quello del sintomo.

I quattro numeri di questo confronto **non vanno letti insieme a quelli di Tabella 2**, dove lo stesso T su H1 vale 9,46 giorni: quella tavola misura con la catena di produzione e sull'istantanea corrente, questo confronto con i filtri di livello disattivi e su un'istantanea precedente. Ciò che la misura afferma è il **rapporto fra i due righelli**, che è calcolato dentro la stessa catena; l'ordine di grandezza assoluto va preso dalla tavola.

Che cosa NON segue da qui. Non segue che i numeri del timeframe proprietario siano quelli «veri». Segue che sono **diversi**, e che la differenza non è rumore: è sistematica nel verso, su tre mercati e quattro livelli. Quale delle due misure descriva meglio il mercato è una domanda sul mercato e chiede un criterio esterno, che questo lavoro non ha e non finge di avere. Segue anche, molto concretamente, che **le misure di polarità sono robuste al cambio di timeframe e quelle di durata no** — la quota di cicli ribassisti non discorda mai, con scarto massimo 0,083 contro una soglia di 0,10 — e questo dice quali risultati del § 6 il righello tocca e quali no.

Il documento, di conseguenza, tiene i numeri del Daily e quelli del timeframe risolvente: i primi come braccio di controllo, i secondi come misura del livello. Dove il livello è alto — da T+4 in su — il Daily è il proprietario, e la questione non si pone.

5.6.1. LA REGOLA VALE PER LE PROPRIETÀ, NON ANCORA PER LE RELAZIONI

La regola del proprietario è ben definita per le **proprietà di un livello** — quanti cicli ha, quanto durano, quanto sono disperse le durate — perché quelle grandezze appartengono a un livello solo, e quel livello ha un righello. Non è definita, allo stesso modo, per le **relazioni fra due livelli**: un legame fra un ciclo e il suo sotto-ciclo non appartiene a nessuno dei due righelli, e quando i due membri sono risolti da timeframe diversi la regola, presa alla lettera, non dice che cosa fare.

Le misure di questo documento che riguardano relazioni — il legame fra un livello e il suo sotto-livello — **non risolvono** il problema: lo **dichiarano**. Cella per cella si registra se anche il sotto-livello risulti risolto sullo stesso timeframe del padre, e i verdetti vengono riportati anche sulle sole coppie complete. Sulla corsa del 7 settembre 2026 le celle con un sotto-livello sono diciassette, e le coppie complete sono **dieci**: nelle altre sette il sotto-livello sarebbe risolto altrove, e la cella entra nel conto generale ma non in quello ristretto.

È una restrizione, non un ponte, e va letta come tale: dove i due membri capitano sullo stesso righello la relazione è misurata come si deve; dove non capitano, la relazione o è misurata sul righello del padre — e allora il sotto-livello è visto sotto la propria risoluzione — oppure esce dal conto ristretto. La costruzione che misura **ciascun membro sul proprio righello**, allineando i minimi sul tempo assoluto invece che sugli indici, è implementata e in corso di validazione; le misure di relazione riportate qui non ci passano ancora, e per questo la restrizione è dichiarata con il suo numero invece di essere taciuta.

5.6.2. IL CICLO GIORNALIERO NON È IL TIMEFRAME DAILY

Un caso vale la pena di essere isolato, perché è quello in cui il nome inganna di più. Nel catalogo, il **ciclo giornaliero** è il livello con periodo nominale di un giorno — T-3 nella scala relativa — e il metà-giornaliero è quello da mezza giornata, T-4. Il timeframe chiamato *Daily* non è il loro: possiede T+4, T+5 e T+6, cioè cicli da 160 a 480 giorni, e su di lui il ciclo giornaliero vale **una barra**. L'unico timeframe che porta T-3 dentro la finestra ideale della risoluzione è M5, dove vale 288 barre contro un centro geometrico di 282,8.

Ma i cinque conteggi non sono cinque livelli, ed è il risultato più importante della misura. Applicando ai minimi appena prodotti lo stesso criterio del § 6.1 — sovrapposizione bidirezionale sopra l'ottanta per cento, soglia dichiarata e non ritoccata — i tre livelli più brevi **non sono distinti fra loro**: su M1 T-7 e T-6 hanno insieme di minimi **identici** e T-6 e T-5 stanno a 0,949. La lettura difendibile è quindi più stretta di quella che

il conteggio suggerisce: sotto il pavimento del Daily i minimi propri esistono, e i livelli distinti sono **due** — T-3 e T-4 — più un terzo oggetto che il catalogo chiama con tre nomi.

Resta dichiarato il limite del campione: su M5 e M1 lo store copre una **finestra** e non l'intera storia, e i conteggi sono quelli di quella finestra.

5.7. Le istantanee di prezzo, dichiarate

I dati sono vivi: la barra in corso si aggiorna e ogni tanto ne arriva una nuova. Un lavoro composto di decine di misure eseguite in tempi diversi non poggia perciò su **una** istantanea di prezzo, e questo documento non finge il contrario.

Che cosa è congelato, e che cosa no. Congelati sono i **prezzi**: la serie ha un ultimo scambio ed è dichiarato. Alla data di questa redazione lo store contiene, per i tre mercati primari, gli scambi **fino al 5 settembre 2026**. Non sono congelate le **misure**: nessuna misura è stata rifatta per allineare la propria istantanea a quella delle altre, perché rifare una misura per farla coincidere con un'altra significa scegliere l'istantanea guardando il risultato.

Che cosa viene dichiarato al posto di una data unica. Ogni file di evidenza porta, per ogni mercato e timeframe che usa, il numero di barre e l'**impronta SHA-256 dei prezzi**: due misure poggiano sulla stessa istantanea se e solo se le due impronte coincidono, e la verifica non richiede fiducia. Alla redazione le evidenze consegnate si distribuiscono su diciassette istantanee distinte; la **maggioritaria** — quella della corsa che ha rifatto trenta carte in una sola lettura dei prezzi — vale 3.308 barre giornaliere su Bitcoin e raccoglie 29 evidenze su 53, e le successive valgono 3.124, 3.299 e 1.894 barre. Le tavole descrittive dei § 4 e § 6 che vengono dall'impianto storico poggiano sullo snapshot del 15 giugno 2026.

Cinque carte citate stanno fuori da quella maggioritaria per costruzione, e vanno nominate: **CR-021** (storia lunga sugli indici), **CR-041**, **CR-046**, **CR-049** e **LG-004** si rifanno con un runner proprio, quindi con un'invocazione propria e un'istantanea propria. Non è un difetto che si possa correggere rilanciando la raccolta — nessuna corsa le comprende — ed è la ragione per cui il controllo di coerenza le esclude dal confronto invece di chiedere loro l'impossibile. Ma l'esclusione va **detta qui**, perché i loro numeri finiscono in pagina accanto agli altri: dove questo può ingannare, il paragrafo lo dichiara sul posto (§ 5.6).

5.8. Gate di non-repainting e determinismo

Poiché il modulo lingue è la componente analitica centrale, la sua causalità è verificata con un gate dedicato, su pipeline reale e in modalità walk-forward.

Il metodo è walk-forward sulla pipeline reale: si aggiungono progressivamente barre e si verifica che le etichette **già confermate** non mutino. **Cinque combinazioni su cinque passano** — i tre mercati su Daily più Bitcoin su H1 e H4 — con esecuzioni identiche a parità di dati e zero etichette mature mancanti o in eccesso.

Il limite si dichiara qui e vale per tutto il documento: il gate copre il **Daily**, dove è misurato l'impianto storico del § 6, e due timeframe intraday su un mercato solo. La proprietà è dimostrata dove serve ai risultati, non ovunque; una verifica più densa, con la latenza di conferma misurata invece che assunta, mostra che la tenuta non è uniforme su tutte le combinazioni intraday, e questo resta dichiarato come lavoro aperto (§ 7.2).

6. Risultati statistici

I numeri di questa sezione provengono dai file di misura versionati, ciascuno con la propria istantanea di prezzo dichiarata (§ 5.7), e sono ricalcolabili dal materiale di supporto.

6.1. Quanti livelli esistono davvero

Prima di chiedersi se i cicli si annidino per due o per tre, occorre sapere quanti livelli esistano. Un livello esiste, sui dati, se ha minimi propri e un periodo distinguibile da quello dei vicini. Il criterio di falsificazione è dichiarato prima: **un livello i cui minimi coincidono con quelli di un vicino in oltre l'ottanta per cento dei casi in entrambe le direzioni non è un livello distinto**.

La bidirezionalità è la parte che conta. In una gerarchia sincronizzata ogni minimo di un livello alto è per costruzione anche minimo di tutti i livelli sotto: la sovrapposizione dal padre verso il figlio vale sempre il cento per cento e non distingue nulla. L'informazione sta nella direzione opposta — quanti minimi propri ha il figlio — e se anche quella supera la soglia, il figlio non è un livello ma un altro nome per lo stesso oggetto.

Un livello risulta doppione del vicino solo quando lo si legge su un righello che non lo risolve. Misurata con ogni livello sul timeframe che lo risolve — e i figli prestati dai loro, secondo il ponte del § 5.6.1 — la scala non contiene nessuna coppia non autonoma: **trentotto coppie adiacenti, zero marcate**, con sovrapposizione massima 0,71. La stessa gerarchia letta tutta sul righello di base, che è la forma di prima, ne marca **quaranta su centottantaquattro**; e ognuna di quelle quaranta ha almeno un membro che quel righello non risolve. Il conto dai due lati:

	marcate doppione	non marcate	totale
il righello risolve entrambi i livelli	0	45	45
almeno un livello sotto la sua risoluzione	40	99	139

Tab. 5: Le coppie adiacenti di tre mercati su nove righelli, secondo il criterio dichiarato e secondo una proprietà che non guarda i minimi — se il righello risolve o no ciascun membro. Test esatto di Fisher, verso dichiarato prima: $p = 2,6 \cdot 10^{-6}$.

Le due condizioni non sono simmetriche. «Il righello non lo risolve» è **necessaria e non sufficiente** — quaranta coppie su centotrentanove — mentre il verso opposto non ha eccezioni: quando il righello risolve entrambi i membri, la coppia non è mai marcata. E le due nubi non si toccano, 0,27–0,71 contro 0,81–1,00: la soglia cade nel vuoto fra le due, e alzarla a 0,85 non cambia la conclusione. **Col pavimento sceso al minuto il risultato si è rafforzato senza cambiare forma**: le coppie con entrambi i membri risolti passano da ventuno a quarantacinque e restano zero le marcate, mentre il valore p scende di due ordini di grandezza.

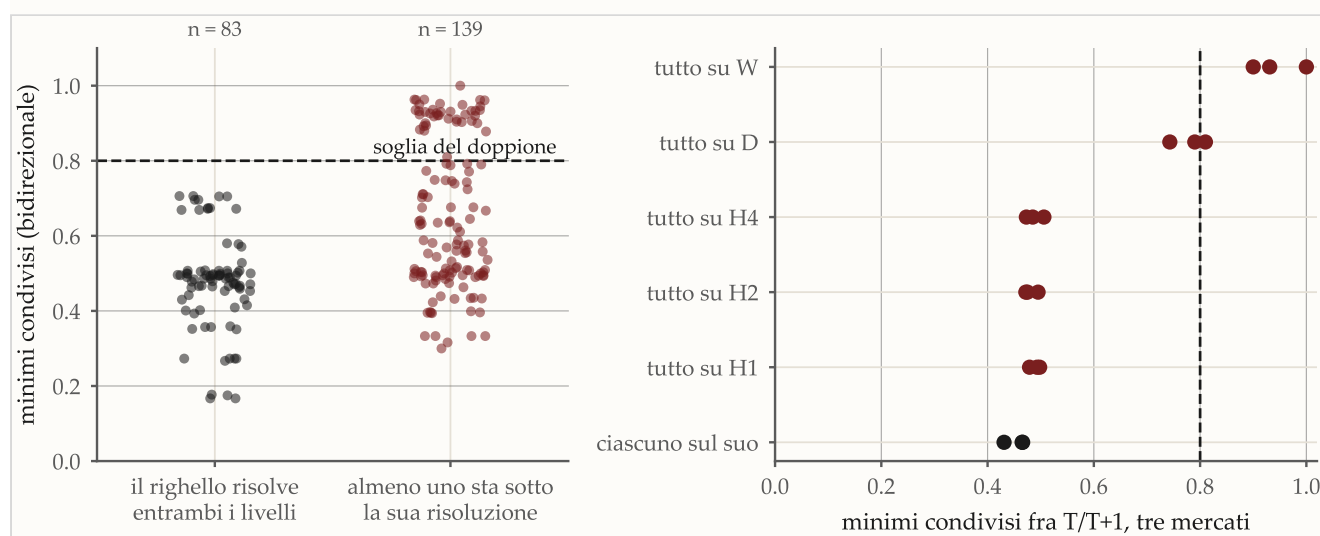


Fig. 4: A sinistra tutte le coppie adiacenti misurate, separate secondo una proprietà che non guarda i minimi: le due nubi non si toccano e la soglia cade nel vuoto fra loro. A destra la coppia T/T+1, letta su ciascun righello e sui tre mercati — sopra la soglia dove il righello non risolve nessuno dei due, alla metà dove ne risolve almeno uno.

Il caso di T e T+1 è la dimostrazione più netta. La stessa coppia, sugli stessi tre mercati e con la stessa funzione, dà 0,90–1,00 sul settimanale, 0,74–0,81 sul giornaliero — che non risolve né l'uno né l'altro — e crolla a 0,47–0,51 su ogni righello che ne risolva almeno uno, fino a 0,43–0,47 nella lettura a ponte. Il rapporto fra le durate mediane si muove insieme: 1,24–1,35 sul giornaliero, **1,99–2,10** con ciascun livello sul proprio. Il due che la gerarchia richiede c'è, e compare appena si smette di leggerlo sul righello sbagliato.

E non è un fatto delle criptovalute. Sui sei indici azionari a storia lunga, dove ai livelli alti il campione è dieci volte più grande — fino a 251 e 64 cicli su S&P 500, contro 22 e 6 su Bitcoin — nessuna delle quindici coppie con entrambi i membri risolti risulta doppione, in nessuna delle due convenzioni di durata.

Ma la soglia dell'ottanta per cento resta da confrontare con un valore nullo. È una soglia ereditata, e finché non la si misura contro ciò che il rilevatore produce sul rumore resta una scelta. Il controllo è lo stesso di tutti gli altri: la funzione di produzione applicata alla gerarchia vera e a quella che lo stesso rilevatore produce su cento surrogati a fase randomizzata e cento AAFT per mercato. Il primo giro l'ha misurata sul giornaliero, ed è la tavola che segue.

Il verso della conclusione, e la sua portata. Sul giornaliero T e T+1 sono indistinguibili, e lì il rilevatore non li separa nemmeno sul rumore: la soglia dell'ottanta per cento viene superata dal 98–100% dei surrogati, e l'osservato sta perfino sotto la mediana della nube. Su quel righello, quindi, T+1 non va offerto: non perché il mercato non lo contenga, ma perché il giornaliero non ha la risoluzione per

vederlo. Dove il righello ce l'ha — da H2 in giù — gli stessi due livelli risultano distinti, con rapporto 2,0 e sovrapposizione dimezzata. E sulla catena della corsa le trentatré righe che un righello possiede per intero hanno tutte margine, la peggiore di **0,096** — Ethereum su T-7/T-6 di M1, con Solana a 0,097 e Bitcoin a 0,099. Il verdetto resta NON PROVATA, ma **non per la copertura**: in nessuna delle trentatré il nullo supera la soglia, e il blocco è quel margine, che sta sotto lo 0,10 dichiarato prima della misura. È una distinzione che conta, perché la copertura si può ampliare misurando di più, mentre un margine sotto la soglia è un esito.

Il criterio di validità di un livello va dunque dichiarato **per timeframe**, come il catalogo dei periodi: la stessa soglia, sugli stessi livelli, è vuota sul settimanale, indecisa sul giornaliero e ha 29–32 punti di margine su H4 e H1.

T+6 non esiste e T+8 non è misurabile su crypto. Il primo è un duplicato del catalogo — ha il periodo nominale di T+5 — e nessun righello lo produce (§ 4); il secondo coincide con T+7 su Bitcoin, dove entrambi hanno due soli minimi. Sugli indici a storia lunga T+8 conta invece fra quattro e nove cicli sul settimanale: è un limite di storia, non di scala.

Verdetto. Dei nove gradini della scala ereditata, nessuno risulta un doppione del vicino quando lo si misura sul righello che lo risolve. Due non esistono come livelli autonomi **sul timeframe giornaliero** (T+1 e T+6) e due non hanno campione su crypto (T+7 e T+8). La scala va usata come vocabolario, non come inventario: i livelli vanno dichiarati per mercato **e per timeframe**, non presupposti — e la loro autonomia va misurata dove si risolvono, o si misura il righello invece del mercato.

6.2. Il nesting: rapporti misurati e principio di armonicità

Hurst afferma che i periodi delle onde stanno fra loro in rapporti armonici costanti, tipicamente due; la scuola italiana ammette anche il tre. È l'affermazione più esposta della tradizione, e si misura in tre modi indipendenti.

Secondo modo: quanti sotto-cicli contiene di fatto ogni ciclo padre.

Le due misure concordano, e i file di evidenza ne portano la tavola cella per cella. **La struttura a due è la più frequente a ogni livello fino a T+4 su tutti e tre i mercati** (49–75% dei cicli padre) e il rapporto fra durate resta fra 1,62 e 2,27 fino al passo T+3 → T+4. Il rapporto si allontana da due solo all'ultimo gradino misurabile, T+4 → T+5 (2,42–4,11), dove i cicli padre sono **quattro o cinque per mercato** e Solana mostra ancora la struttura a due.

La forma che ha potere: la scala, non le coppie. Se i periodi sono armonici, esiste una **fondamentale** P_0 tale che i picchi cadono vicino ai pioli di una scala — binaria ($P_0 \cdot 2^k$, il raddoppio di Hurst) o intera ($P_0 \cdot n$, i rapporti interi piccoli). La statistica è la distanza media dei picchi dal piolo più vicino, in ottave, minimizzata sulla fondamentale; la nulla è un insieme di periodi estratti **dalla stessa griglia del rilevatore**, con la stessa cardinalità. E il **controllo di potenza viene prima del risultato**: la stessa catena su una serie sintetica con componenti a 20, 40, 80 e 160 barre **riconosce la scala** — distanza 0,0000 ottave sulla scala intera, battendo il caso nel 100% dei confronti, e 0,0233 sulla binaria, nel 99,85%.

Sul giornaliero, dove la misura è nata, la scala non c'è. Bitcoin batte il caso nel **92,95%** dei confronti ed Ethereum nel **93,6%**, contro una soglia dichiarata del 95%; sulla scala binaria entrambi stanno **sotto** la mediana del caso, Bitcoin più in basso di nove estrazioni su dieci. Solana produce tre picchi soli e non entra nel conteggio. Sui sei righelli intraday la scala intera sembra dare l'esito opposto, e il verdetto qui sotto dice perché non è una conferma.

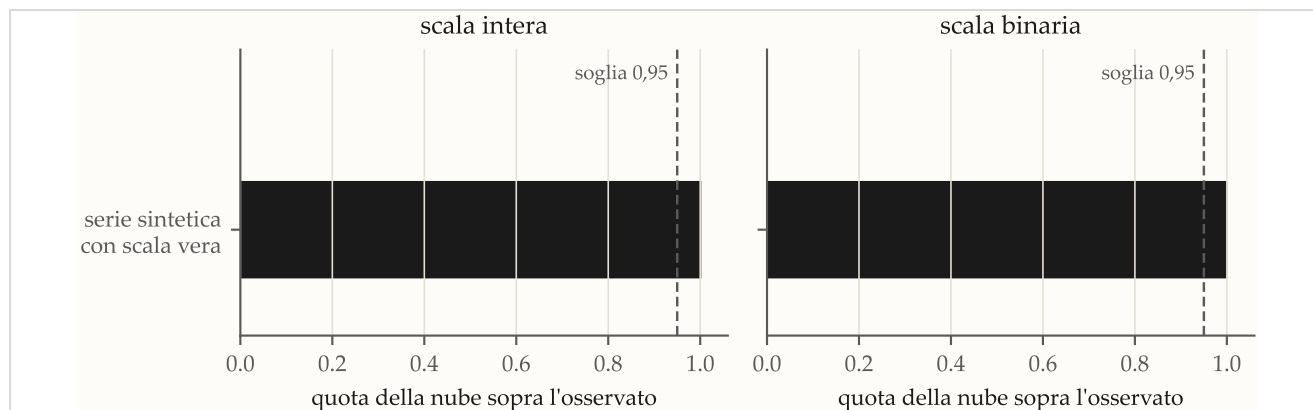


Fig. 5: Dove cade la distanza osservata dentro la nube nulla, per mercato e per scala. La grandezza è la quota di nube che sta sopra l'osservato: siccome la statistica misura una distanza dai pioli, più la barra è lunga più quel mercato è armonico. La riga tratteggiata è la soglia dichiarata prima. In nero il **controllo di potenza**: la stessa catena su una serie sintetica costruita con componenti a 20, 40, 80 e 160 barre.

Un'obiezione alla nulla, misurata invece che discussa. L'estrazione libera dalla griglia può produrre periodi adiacenti, cosa che un rilevatore a prominenza non fa mai, e due periodi vicini cadono quasi sempre presso lo stesso piolo: la nulla libera potrebbe risultare più armonica del dovuto. Imponendo alla nulla la stessa separazione minima dei picchi osservati, Bitcoin scende dal 92,95% al 92,3% ed Ethereum sale dal 93,6% al 95,8%, sopra la soglia. Questa nulla è dichiarata di robustezza e non entra nel verdetto; se entrasse, il giornaliero passerebbe da falsificato a non provato — un mercato su due — e non a confermato. Va scritto, perché i due verdetti li separa mezzo punto.

Sui righelli intraday la scala intera sembra rovesciare l'esito, ed è la nulla. La stessa misura — stessa nulla, stessa soglia — sui sei righelli da H4 a M1 rimisurati con questo pavimento trova da sedici a ventiquattro picchi per mercato, e la scala intera batte il caso in **diciotto contesti su diciotto**; la binaria in nessuno. Prima di leggerlo come una conferma c'è una spiegazione concorrente da togliere: su quei righelli nessun picco sta sotto le 53 barre, mentre la nulla estrae da tutta la griglia fra 10 e 500 — quasi metà dei nodi sta sotto quella soglia — e in ottave i pioli di una scala intera si infittiscono al crescere del periodo. Una nulla che occupa una banda che i picchi non occupano fa sembrare più armonico chi sta in alto. La misura che separa le due letture è la stessa con la nulla ristretta alla banda dei picchi, dichiarata prima di eseguirla: la scala intera resta sopra la soglia su **un righello solo su sei** (M1, due mercati su tre), e anche il quasi-passaggio del giornaliero sparisce — Bitcoin scende dal 92,95% al 49,35%, Ethereum dal 93,6% al 40,15%. Il diciotto su diciotto era geometria della nulla.

I quattro numeri di questo capoverso vengono dalla misura con la nulla ristretta **rifatta il 12 settembre** sull'evidenza rimisurata, di cui l'esito dichiara l'impronta: questa volta l'impronta corrisponde al file da cui è stata letta. Il denominatore è **sei** e non sette perché il pavimento della rimisura non comprende M5 — un righello che nella misura precedente stava già fra quelli non confermati, e la cui uscita infatti non muove nessun verdetto: righello per righello i risultati sono identici. **L'esito però cade sulla soglia esatta**: la regola dichiarata chiede cinque righelli non confermati, e sono cinque su sei. Un solo righello in più dalla parte della conferma avrebbe reso l'esito indeterminato. È detto qui perché un verdetto che passa per un'unità non è un verdetto robusto, ed è più onesto dichiararlo che lasciarlo dedurre.

E resta una scelta che nessuna misura può arbitrare, ed è la più importante di questa sezione: la nulla estrae dalla **griglia** del rilevatore, non da un continuo. Con l'estrazione log-uniforme continua la scala intera risultava confermata su entrambi i mercati. La griglia è la scelta corretta — l'osservato può cadere solo lì, e una nulla che vive altrove misura la griglia invece dell'armonicità — ma l'esito di questa misura poggia su un argomento di principio, non su un margine numerico.

Verdetto. Il **nesting armonico fisso** è falsificato come legge: il rapporto misurato non è costante e la struttura a due, pur dominante, non supera il 75% dei cicli padre a nessun livello. E il **principio di armonicità** è **FALSIFICATO sul giornaliero**, non semplicemente non sostenuto: con un test la cui potenza è verificata su una scala vera, **zero mercati su due** risultano più armonici di un insieme di periodi estratti a caso dalla stessa griglia. Il raddoppio non passa su nessuno dei sei righelli misurati; la scala intera, che con questa nulla passa su tutti quelli intraday, con la nulla ristretta alla banda dei picchi passa su uno solo: il vantaggio era della nulla, non dei mercati. La differenza fra i due verdetti non è formale — *non sostenuto* significa che non si sa, *falsificato* significa che si è guardato con uno strumento capace di vedere. Il sostituto operativo non è una seconda legge, ma una misura: **il rapporto va stimato**

per coppia di livelli e per mercato, e sulla scala esaminata resta vicino a due fin dove il campione consente di misurarlo.

E non c'è un fulcro. L'idea che il rapporto valga due sotto un certo livello e tre sopra — un «nesting bimodale» con un gradino di passaggio — è la prima cosa che i numeri suggeriscono a occhio, e i rapporti misurati e i conteggi delle strutture non la sostengono: il passaggio a tre si osserva solo all'ultimo gradino, dove il campione non basta a distinguerlo dal rumore.

6.3. La distribuzione delle durate: un modo solo e una coda di cicli brevi

L'affermazione sotto esame. Che le durate cicliche siano **bimodali**: a ogni livello un gruppo di cicli corti e uno di cicli lunghi, due popolazioni distinte invece di una sola. È un'affermazione attraente, perché una distribuzione a due gobbe suggerisce due regimi e quindi una regola operativa che li distingue.

Il test, e perché proprio questo. La domanda « quanti modi ha questa distribuzione? » ha un test suo, il **dip di Hartigan**, che misura la distanza massima fra la ripartizione empirica e la più vicina distribuzione unimodale. Non va confuso con un test di **uniformità**: una statistica di Kolmogorov–Smirnov contro la rampa uniforme risponde a « questa distribuzione è piatta? », e una distribuzione a un modo solo — che è tutto tranne che piatta — la massimizza. Le due domande hanno risposte opposte sullo stesso dato, ed è una confusione facile da fare: il dip di Hartigan vive per costruzione in $[0, 1/4]$, quindi **un valore fuori da quell'intervallo è la spia che si sta guardando un'altra statistica**.

Il dip di Hartigan sulle serie di durate: due rigetti su trentanove, e cadono nello stesso posto. Misurato con ogni livello sul righello che lo risolve, il test copre 39 combinazioni asset/livello su nove righelli. Trentasette non rigettano l'unimodalità, con p fra 0,064 e 0,988. Le due che rigettano stanno **sullo stesso livello e sullo stesso righello** — T-4 su M5 — e il terzo mercato, sulla stessa cella, sta al bordo.

La fragilità va detta qui, perché tocca il verdetto. La metrica dichiarata prima della misura era « rigetto dell'unimodalità su almeno una combinazione asset/livello »: sul solo Daily, con 18 combinazioni, nessuna rigettava e il verdetto FALSIFICATA seguiva dalla lettera del criterio. Con 39 combinazioni due rigettano, e **la lettera del criterio è soddisfatta**. Due fatti stanno dall'altra parte, e sono dichiarati senza essere usati per scegliere: con 39 test al 5% due rigetti sono quanti se ne attendono per caso, e con la correzione di Holm — soglia 0,00128 per il più piccolo — **nessuno** rigetta. Ma i due non sono sparsi: cadono sulla stessa cella (T-4 su M5) di due mercati, col terzo a 0,052. La soglia non è stata ritoccata e il verdetto non è stato cambiato: resta FALSIFICATA con questa fragilità scritta, e la domanda che ne segue — se a quel livello convivano davvero due popolazioni di durate, o se sia la delimitazione a produrle, come per la coda dei brevi — si risolve con la stessa misura sui surrogati dello stesso rilevatore.

Che cosa resta, ed è misurabile. La distribuzione ha un modo solo e una **coda di cicli brevi**, la cui consistenza è stabile fra mercati e cresce con il livello. La soglia — 0,65 della durata nominale — è dichiarata prima, e i due gruppi sono costruiti da una regola nota, non trovati nei dati.

Livello	Bitcoin	Ethereum	Solana
T	3,0%	6,1%	—
T+1	18,8%	19,0%	16,8%
T+2	32,9%	35,0%	35,3%

Tab. 6: Quota di cicli che si chiudono sotto lo 0,65 della durata nominale del proprio livello.

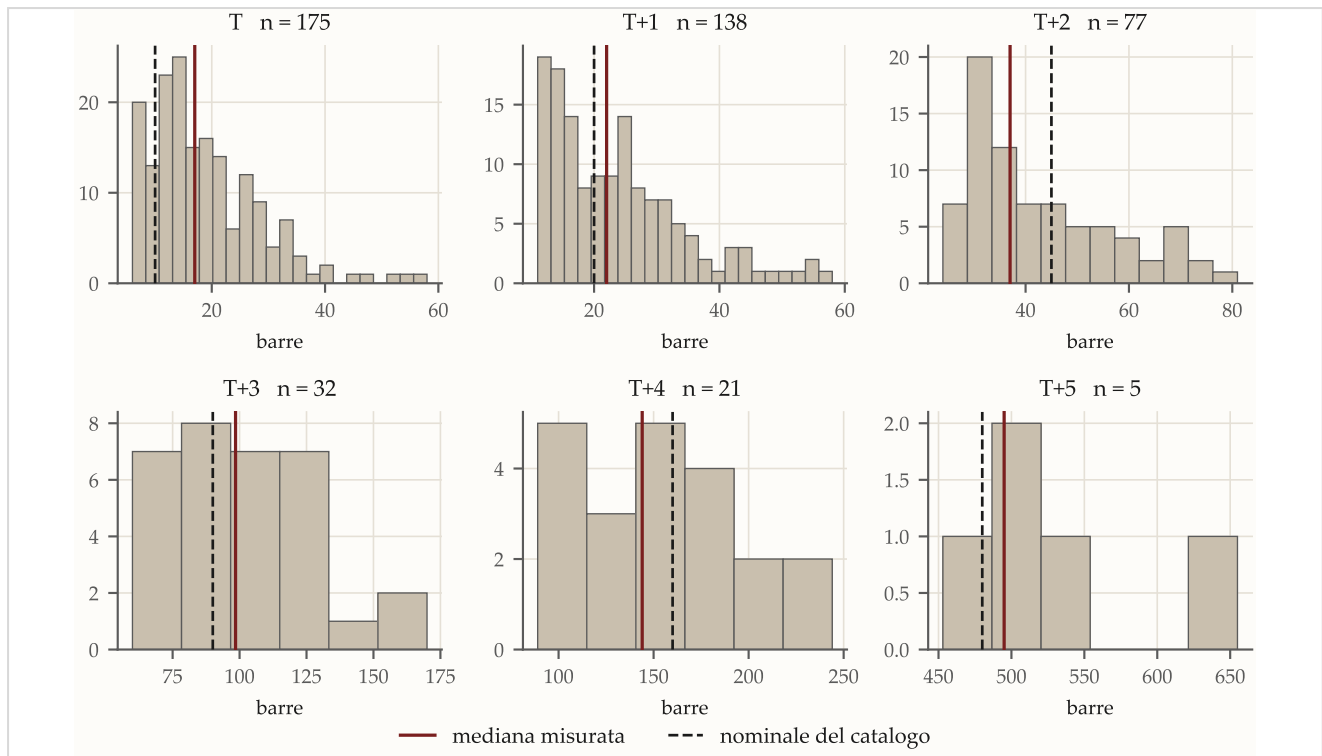


Fig. 6: La distribuzione delle durate di ogni livello, con la mediana misurata e il nominale del catalogo. Un modo solo e una coda a sinistra: nessuna delle sei mostra i due gruppi separati che la bimodalità richiedeva.

Verdetto. La bimodalità delle durate è **FALSIFICATA**: non c'è in trentasette celle su trentanove, e dove il campione è quello del giornaliero non c'è a nessun livello. La distribuzione ha **un modo solo con una coda di cicli brevi**, e la quota di quella coda cresce con il livello in modo ordinato su tutti e tre i mercati. **Le due eccezioni sono nominate e non assorbite**: stanno sulla stessa cella T-4/M5 di due mercati, il terzo è al bordo, e la fragilità che ne segue per la lettera del criterio è dichiarata sopra. Resta anche lo **scarto sistematico** fra mediana misurata e periodo nominale (§ 4), che è però una proprietà del **sistema di misura** e non del mercato: si veda § 6.3.2.

6.3.1. LA CODA DEI BREVI È DEL MERCATO O DEL RILEVATORE?

La tavola **Tabella 6** descrive un fatto ordinato: la quota di cicli brevi è stabile fra mercati e cresce con il livello. È il genere di regolarità che si è tentati di rivendicare come scoperta — e la domanda che il § 5.3 impone di fare **prima** di rivendicarla è quale sia il suo valore nullo.

Si misura applicando il rilevatore, con gli argomenti di produzione, a cento surrogati a fase randomizzata e cento AAFT per mercato, con ogni livello sul righello che lo risolve, e confrontando la quota osservata con la nube che ne esce. Le soglie sono dichiarate prima dell'esecuzione: la coda è del mercato se esce dalla nube in almeno il 60% delle combinazioni asset/livello, è del rilevatore se ne esce in non più del 10%.

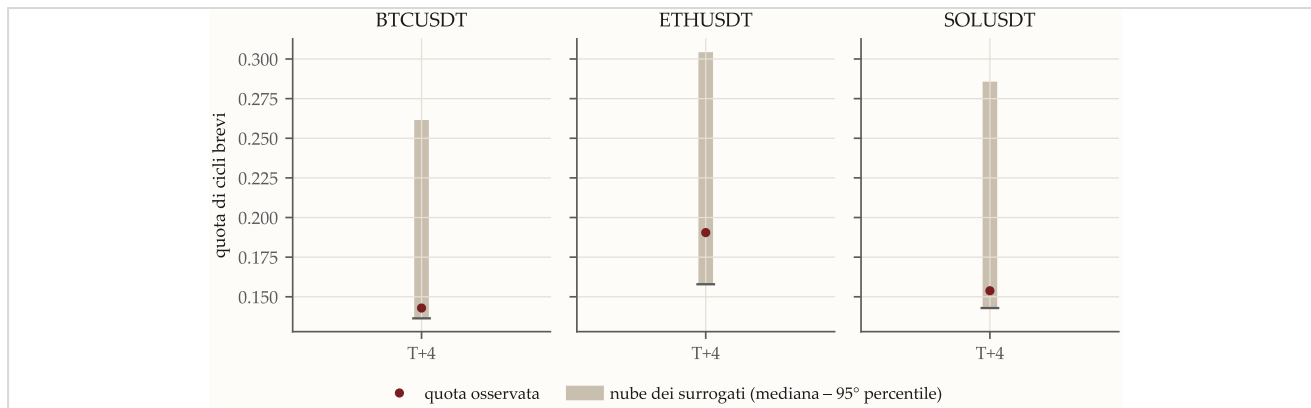


Fig. 7: La quota osservata di cicli brevi (punto) contro la nube prodotta dallo stesso rilevatore su serie senza struttura temporale (banda: dalla mediana al 95° percentile), livello per livello sui tre mercati, con sotto ogni livello il righello che lo risolve. Fra le due famiglie di surrogati — fase randomizzata e AAF — è disegnata la **più severa**, cioè quella con la banda più alta. Un punto su trentatré esce sopra la propria banda — Ethereum a T-1, sul righello M30 —; gli altri stanno dentro o sotto, senza un verso.

L'esito. Sulla catena — undici livelli per tre mercati, trentatré combinazioni — ne esce dalla nube contro entrambe le famiglie di surrogati **una**: Ethereum a T-1 su M30, con quota 0,1551 contro un 95° percentile di 0,1367 (fase) e 0,1328 (AAF). La quota è $1/33 = 0,030$, sotto la soglia di falsificazione dichiarata. Su sei dei sette righelli valutabili non ne esce nessuna; solo M30 ne ha una, e letto da solo darebbe NON PROVATA. Bitcoin a T-5 su M1 esce contro **una sola** famiglia, e il verdetto di M1 è FALSIFICATA; Bitcoin a T-4 su M5 **non è misurato**, perché M5 è fuori dal pavimento di questa corsa: è un dato che manca, non un'eccezione. L'osservato sta sotto la mediana della famiglia più severa — la più severa è quella con la banda più alta, come in **Figura 7** — in 18 combinazioni su 33: il mercato non produce sistematicamente né più né meno cicli brevi del sintetico. La definizione va dichiarata perché il conteggio ci poggia: prendendo invece la famiglia con la mediana più alta le combinazioni sarebbero 19. Il primo giro, sul solo giornaliero, dava 0 su 15 e l'osservato sotto la mediana a T+2 e T+3 su tutti e tre i mercati; ma il giornaliero quei due livelli non li risolve, e sul righello che li risolve il verso sparisce.

Verdetto: FALSIFICATA. La coda dei cicli brevi **non è una proprietà del mercato**. È ciò che il rilevatore produce comunque: esce dalla nube una volta su trentatré, e per il resto sta sopra o sotto la mediana dei surrogati senza un verso.

È il risultato che questo lavoro avrebbe avuto più interesse a tenere — era elencato fra le regolarità che le scuole precedenti non documentano — ed è la ragione per cui il valore nullo si misura invece di assumerlo: **l'ipotesi nulla sbagliata non fa sbagliare il numero, fa sbagliare la conclusione**. Spiega anche perché le altre tre ipotesi fallivano: se la coda non è un fenomeno di mercato, nessuna causa di mercato può spiegarla.

La coda resta una **descrizione corretta dell'output**, utile a chi calibra, e le quote di **Tabella 6** sono quelle giuste. Esce dall'elenco delle regolarità di mercato ed entra in quello delle proprietà della procedura di phasing.

6.3.2. LO SCARTO DAL NOMINALE: STESSA DOMANDA, STESSA RISPOSTA

Che le durate mediane stiano **sotto** il periodo nominale del catalogo è, in questo lavoro, il fatto ripetuto più spesso: la nomenclatura del § 4 lo mostra a tutti i livelli bassi, e viene naturale leggerlo come un difetto del catalogo ereditato — durate pensate per l'azionario degli anni Sessanta, applicate a un mercato che non esisteva.

La domanda è però la stessa della coda dei brevi, e va posta con lo stesso strumento: il rapporto $r = \text{mediana}(\text{durata})/P_{\text{nominale}}$ è quello che il rilevatore produrrebbe **su una serie che di cicli non ne ha affatto**? La nulla primaria è una nube di cammini con incrementi gaussiani a media nulla e dispersione stimata dai log-rendimenti dell'asset — niente ciclo, niente deriva — con la sola chiusura da entrambe le parti, perché un sintetico non ha massimi e minimi di barra e confrontarlo con quelli reali misurerebbe la presenza degli estremi di barra.

Su sei indici azionari a storia lunga e tre crypto, l'osservato cade **dentro** la nube senza cicli in **28 celle su 37** (quota 0,757; la soglia dichiarata prima era 0,667). E dove non ci cade — T+2 e T+3, sei celle su nove ciascuno

— il verso è l'**opposto** di quello che rafforzerebbe la tesi di mercato: l'osservato è **più alto** del nullo (0,733–0,800 contro mediane di 0,544–0,600), cioè il mercato produce cicli **più lunghi** del rumore, non più corti.

Verdetto: lo scarto dal nominale è del rilevatore. Il reticolo esiste — le durate cadono davvero sotto il catalogo — ma a imporlo sono i pavimenti e le proporzioni del rilevatore, non una proprietà dei mercati. Il catalogo nominale resta da rivedere; **la ragione per rivederlo non è più «i mercati sono più veloci»**, ed è una distinzione che cambia chi deve fare il lavoro.

6.3.3. LA STAZIONARIETÀ LOCALE DI EHLERS, CON UN TEST PROPRIO

L'assenza di bimodalità sarebbe stata tentante da usare contro Ehlers, e non si può: la bimodalità, se ci fosse stata, avrebbe smentito la stazionarietà locale, ma la sua assenza non la conferma. Serve un test proprio, ed è stato costruito.

Va costruito sulla forma che Ehlers usa davvero. Non «la durata dei cicli è costante», che nessuno sostiene, ma: **dentro la finestra di lavoro di un filtro adattativo il periodo dominante è abbastanza stabile da poter essere trattato come costante.** La forma operativa è la deriva del periodo dominante fra finestre scorrevoli — finestre da 750 barre con passo 100, roofing filter sui log-prezzi, spettro di Goertzel su una griglia log-spaziata di 60 periodi fra 10 e 250 barre, statistica MAD/mediana dei periodi dominanti — confrontata con 100 surrogati a fase randomizzata e 100 AAF, che sono a spettro stazionario per costruzione.

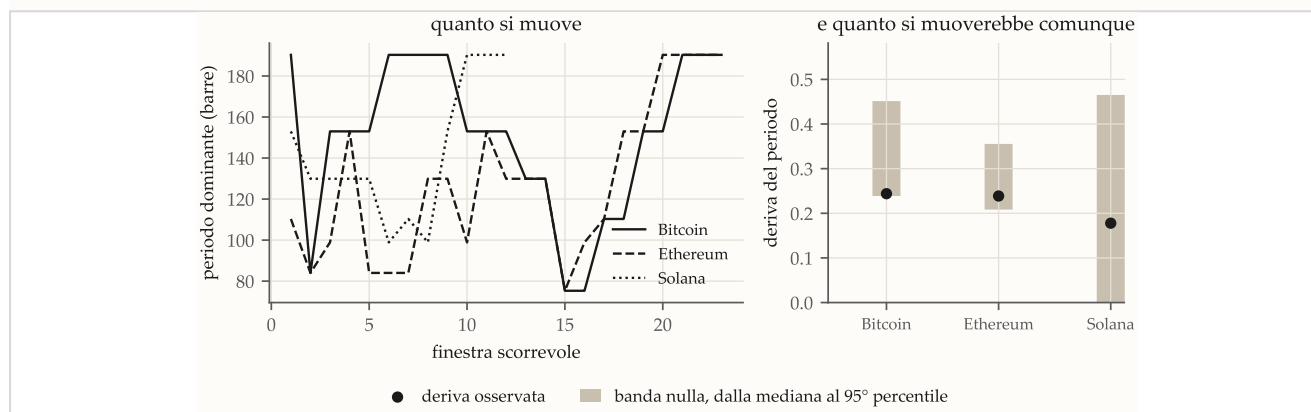


Fig. 8: In alto il periodo dominante misurato finestra per finestra sul giornaliero: si muove, e molto. In basso la deriva messa accanto alla banda che lo stesso rilevatore produce su serie senza struttura, righello per righello, e per ogni righello Bitcoin, Ethereum e Solana da sinistra. La domanda non è se il periodo si muova — si muove — ma se si muova più di quanto si muoverebbe comunque: sul giornaliero no; sui righelli intraday la deriva sta sotto la nube in 14 casi su 21, dentro in 6 e sopra una volta sola, Bitcoin su M5. Il colore è il verdetto della carta per quel mercato; la deriva è una statistica a valori discreti, e un punto marcato può stare sul bordo della banda.

Verdetto NON PROVATA sul giornaliero, CONFERMATA su cinque righelli su nove. Sul giornaliero nessuno dei tre mercati esce dalla nube, né sopra né sotto: la misura non distingue il mercato da un processo a spettro stazionario, e il verdetto è NON PROVATA. Ma la stessa misura è stata rifatta su nove righelli, ed è lì che il quadro si forma: su H2, H1, M30, M5 e M1 la deriva osservata sta **sotto** la nube in due o tre mercati su tre — il periodo dominante è più stabile di quanto lo stesso rilevatore produca su serie senza struttura — e la regola dichiarata prima (due mercati su tre sotto la nube) dà **CONFERMATA**. Su H4 e M15 un mercato solo, quindi NON PROVATA; sul settimanale la serie non è abbastanza lunga.

Il verso conta più del conteggio, ed è favorevole a Ehlers: in nessun righello e in nessun mercato la deriva sta **sopra** la nube. Non esiste un solo contesto in cui il periodo dominante derivi più di un processo stazionario — che è ciò che falsificherebbe l'ipotesi su cui i filtri adattativi sono costruiti. Dove la misura ha campione (i righelli intraday, dove le finestre sono centinaia invece di venti) la stazionarietà locale non solo regge: risulta **più stretta** del caso.

Il limite, dichiarato: la scala su cui la misura ha più potenza è anche quella su cui il filtro lavora, quindi il risultato dice che *dentro la finestra di lavoro* l'assunzione tiene — non che il periodo dominante di un mercato sia stabile nel lungo.

Fragilità dichiarata. La griglia dei periodi è discreta, quindi la statistica atterra su un reticolo di valori possibili e su due delle sei nubi la mediana dei surrogati è esattamente zero — oltre metà dei surrogati ha periodo dominante costante su tutte le finestre. Con una griglia più fitta la distribuzione nulla sarebbe più

liscia e i valori p potrebbero spostarsi. È la fragilità principale di questa misura. Se ne aggiunge una seconda, di natura diversa: un roofing filter ha un transitorio lungo, e uno spettro calcolato su una finestra troppo corta ne porta dentro l'assestamento invece del segnale. Le finestre usate qui sono lunghe abbastanza da renderlo trascurabile, ma la stessa misura su finestre corte non sarebbe leggibile.

6.4. La durata dipende dal regime di mercato?

L'ipotesi è dell'autore: i cicli brevi appartenerebbero a fasi di mercato diverse da quelle dei cicli lunghi, con verso **invertito** rispetto all'intuizione — compressione in laterale, allungamento in trend. È l'ipotesi che, fra quelle proprie, aveva l'aria di reggere meglio.

La prima formulazione confrontava la volatilità dei cicli brevi con quella dei cicli lunghi, cioè **costruiva il regime a partire dai cicli che voleva spiegare**. Qui il regime esiste prima dei cicli. Due assi, entrambi funzione dei soli prezzi: la volatilità realizzata a 30 barre e la forza della tendenza a 90 barre, $|\log(P_t/P_{t-90})|$. Entrambi trasformati in terzili **in modo causale** — al tempo t i due quantili di taglio sono calcolati sulla sola storia fino a t , con un periodo di riscaldamento e nessun percentile calcolato su tutto il campione e applicato all'indietro. Ogni ciclo riceve il terzile del proprio minimo di apertura ed è classificato breve se dura meno di 0,65 del proprio nominale — la soglia già dichiarata al § 6.3, non ritoccata.

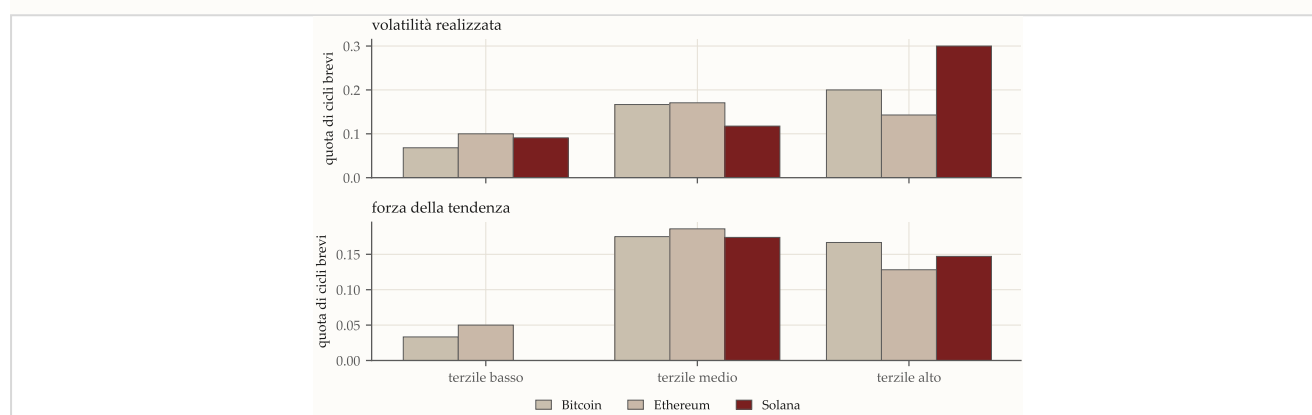


Fig. 9: La quota di cicli brevi per terzile di regime, sui due assi esogeni e sui tre mercati, misurata su H_4 . Le colonne salgono insieme — il verso è lo stesso in sei celle su sei — ma la differenza fra la prima e la terza colonna non supera mai il rumore: nessuna delle sei è significativa. È l'immagine di un'ipotesi che ha la forma giusta e non ha la forza.

Verdetto NON PROVATA. Su H_4 nessuna delle sei combinazioni è significativa: i valori p stanno fra 0,17 e 0,65. La misura è stata ripetuta su tutti e nove i righelli — 48 celle in tutto — e il quadro non cambia: **tre celle su 48** scendono sotto 0,05, meno di quante il caso ne produca alla stessa soglia, e nessun righello arriva alla regola dichiarata. Due righelli (H_2 e il minuto) danno **FALSIFICATA**, cinque **NON PROVATA**, e sul settimanale e sul giornaliero il campione non basta.

Una cosa però va detta, e non è un risultato: su H_4 il verso è lo stesso in **sei celle su sei** — più regime, più cicli brevi — e lo stesso accade sul righello a cinque minuti. Sugli altri cinque righelli il verso si divide (due o quattro celle su sei). Una concordanza di segno senza una sola cella significativa non sostiene l'ipotesi: la si registra perché sarebbe disonesto tacerla, e perché indica dove guardare se un giorno la si vorrà rimisurare con una statistica che usi il segno invece di ignorarlo.

6.5. Ampiezza e durata: la pendenza, non la correlazione

L'ipotesi, originale nella sua formulazione quantitativa, è che i cicli si comportino come una fisarmonica: ampiezza maggiore, durata maggiore. È coerente con la letteratura sulla *duration dependence* (Lunde e Timmermann, 2004; Claessens, Kose e Terrones, 2011) e con il Principio di Proporzionalità di Hurst.

La correlazione di rango fra le due grandezze, misurata sui cicli del sistema, è forte e replicata: ρ di Spearman +0,724 su Bitcoin (452 cicli), +0,765 su Ethereum (445) e +0,725 su Solana (300). **Il numero è solido; l'inferenza che se ne trae d'istinto no.** Confrontare quel ρ con l'ipotesi nulla di indipendenza — il riflesso naturale — è qui l'ipotesi sbagliata (§ 5.3): **il rilevatore induce la correlazione anche sul rumore**, quindi il valore atteso sotto il caso non è zero, e va misurato. Confrontato con la nube dei surrogati, il quadro cambia.

Il confronto, su dieci combinazioni asset/livello: la mediana dei surrogati vale +0,554 — +0,660, cioè **quanto l'osservato o più**, e su **otto combinazioni su dieci** il valore osservato sta **sotto** quella mediana.

L'unica combinazione che esce dalla nube è Bitcoin a T+2, una su dieci: quanto ci si attende per caso. La correlazione di rango non distingue un mercato da una serie sintetica con lo stesso spettro.

La proporzionalità ha però una seconda formulazione, più severa e più informativa: la **pendenza della relazione in coordinate logaritmiche**. Una diffusione pura dà pendenza un mezzo; la proporzionalità piena dà pendenza uno. Il valore nullo, però, non è né l'uno né l'altro: va misurato, e sui surrogati vale 0,61–0,83.

Le pendenze osservate stanno fra 0,79 e 1,17 contro una nube a 0,61–0,83: **superano la mediana dei surrogati su dieci combinazioni su dieci**, e cadono nell'1–2% superiore su cinque. La coerenza del segno su dieci misure indipendenti vale qui quanto la significatività delle singole.

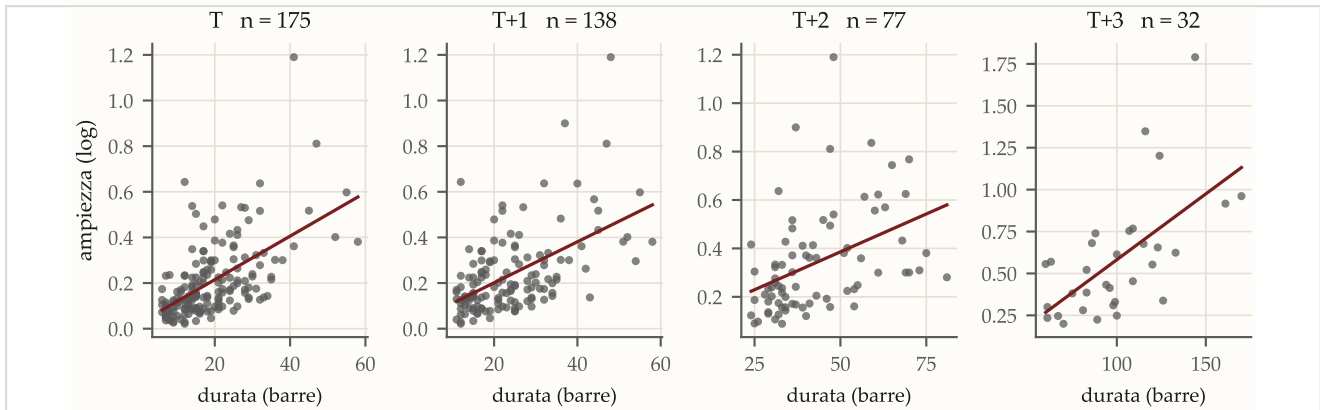


Fig. 10: Durata e ampiezza di ogni ciclo, livello per livello, su Bitcoin. La retta è la regressione sui punti di quel livello: la relazione esiste e la pendenza è positiva a ogni scala. È la **pendenza** la grandezza che si misura, non il coefficiente di correlazione, che con questi numeri di cicli sarebbe alto quasi per costruzione.

6.5.1. L'ASIMMETRIA BULL/BEAR DELLA DURATA

Resta da chiudere il pezzo che la proporzionalità si porta dietro: i cicli che chiudono rialzisti durano più di quelli che chiudono ribassisti? La misura descrittiva dava un rapporto di 1,08–1,19 secondo il mercato, e non era mai stata confrontata con i surrogati.

L'ipotesi nulla non è «rapporto uguale a uno»: una serie con deriva positiva produce cicli rialzisti più lunghi anche senza struttura, e quanto più lunghi va misurato. Il confronto usa il rilevatore elementare applicato identico alla serie vera e a duecento surrogati a fase randomizzata più duecento AAFT per livello, e ogni livello sta sul righello che lo risolve: da T+4 sul giornaliero a T-7 su M1.

Il verdetto è **NON PROVATA**, e la ragione non sta nel conteggio. Sulla catena — dodici livelli per tre mercati — escono dalla nube, contro entrambe le famiglie di surrogati, **22 combinazioni su 36**, e dieci stanno al pavimento del valore p che duecento surrogati consentono, 0,005. Ma escono quasi tutte dai livelli brevi: **4 su 15** da T+4 a T, **18 su 21** da T-1 in giù, nessuna sul giornaliero e su H2. Aggregata sui righelli la quota vale 0,611, a un soffio dalla soglia di conferma (0,60): una combinazione in meno darebbe NON PROVATA, e due di quelle che escono hanno un valore p di 0,045. Letta sul solo giornaliero vale zero, cioè FALSIFICATA. L'esito dipende da come si aggregano i righelli, la regola dichiarata prima non lo diceva, e un'aggregazione non si sceglie dopo aver visto.

C'è poi una ragione più forte per non leggere il 22 su 36 come una conferma. Le due nulle — fase randomizzata e AAFT — sono **reversibili nel tempo**: distruggono per costruzione ogni differenza fra la salita e la discesa. Un mercato i cui rendimenti scendono più in fretta di quanto salgano esce dalla nube comunque, anche se i suoi cicli non hanno niente di asimmetrico. Con questa nulla una conferma è quasi garantita; un'esclusione, al contrario, è informativa, perché la nulla sbaglia in favore dell'affermazione: sul giornaliero e su H2, dove nessuna combinazione esce, l'asimmetria non c'è nemmeno con una nulla che la regala. E letta come domanda simmetrica — «esiste asimmetria?» — nessuna combinazione esce nel verso opposto. Ciò che la misura sostiene è quindi più stretto: **ai livelli brevi i cicli che chiudono in rialzo durano più di quanto produca una serie reversibile con lo stesso spettro**. Se sia una proprietà dei cicli o dei rendimenti lo dirà una nulla che conservi l'asimmetria dei rendimenti e rompa la struttura ciclica — un rimescolamento a blocchi —, dichiarata con il suo criterio prima di eseguirla.

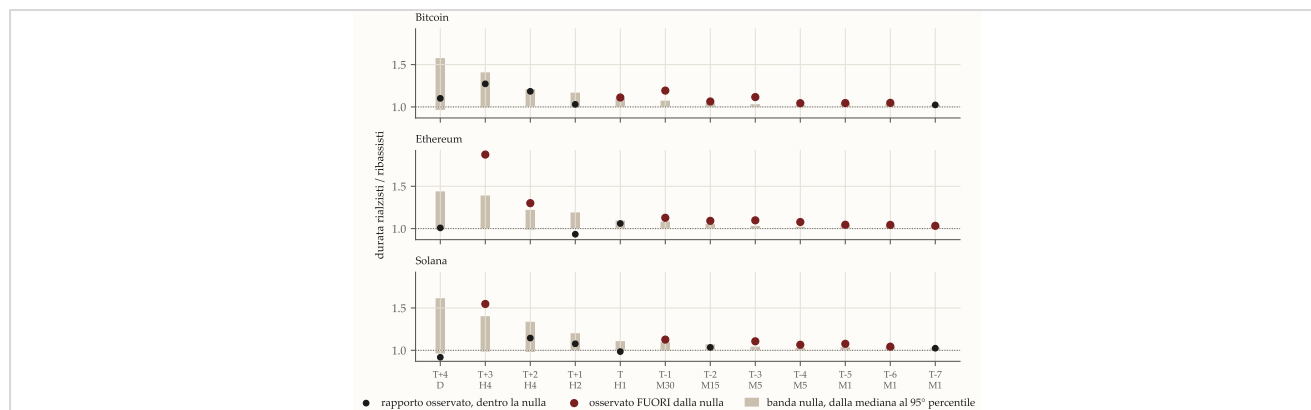


Fig. 11: Il rapporto fra la durata mediana dei cicli rialzisti e quella dei ribassisti, livello per livello sui tre mercati, con accanto la banda che lo stesso rilevatore produce sul rumore; sotto ogni livello, il righello che lo risolve, e la riga a uno è la simmetria. Un punto è marcato fuori dalla banda solo se lo è contro **entrambe** le famiglie di surrogate, che è il criterio della carta. Da T-1 in giù escono quasi tutti, sul giornaliero e su H2 nessuno — e la nulla, reversibile nel tempo, non sa distinguere un'asimmetria dei cicli da una dei rendimenti.

Verdetto PARZIALE sulla proporzionalità. La relazione ampiezza-durata è **più forte della diffusione**, e in questo senso il Principio di Proporzionalità di Hurst regge; ma la misura che lo mostra è la pendenza, non la correlazione di rango, che è indistinguibile dal caso. Una cautela: il confronto con i surrogate usa un rilevatore elementare, perché quello proprietario non è applicabile a serie sintetiche senza esporne l'implementazione. E l'asimmetria bull/bear, misurata sulla catena, è **non provata**: esce dal rumore sui livelli brevi, contro una nulla che non la sa escludere.

6.6. La traslazione come predittore di polarità

La posizione del massimo entro il ciclo — traslazione a destra o a sinistra — predice la polarità netta del ciclo. È il sostituto quantitativo e falsificabile del concetto di ciclo inverso (§ 8.1).

Asset	Cicli	Accuratezza	Base-rate empirica	Uplift
Bitcoin	452	74,6%	56,2%	+18,4 punti
Ethereum	447	74,3%	56,2%	+18,1 punti
Solana	300	77,3%	52,7%	+24,7 punti

Tab. 7: Traslazione contro polarità del ciclo, sullo snapshot descrittivo (§ 5.7). La base-rate empirica è la frazione di cicli che chiudono rialzisti; è il confronto corretto, perché il bias rialzista strutturale delle criptovalute rende il 50% casuale un riferimento troppo generoso.

L'accuratezza è stabile fra i tre mercati e resta 18–25 punti sopra la previsione banale «tutti i cicli chiudono rialzisti». Il segnale è calcolabile a ciclo chiuso e la sua verifica non usa dati futuri. Su un'istantanea più recente e con lo stesso criterio l'accuratezza vale 76,0%, 74,9% e 76,1% su 454, 454 e 314 cicli: il numero si muove di un punto o due con la serie, l'ordine di grandezza no. Fuori dal timeframe giornaliero il segnale regge su campioni di un ordine di grandezza più grandi (§ 6.13).

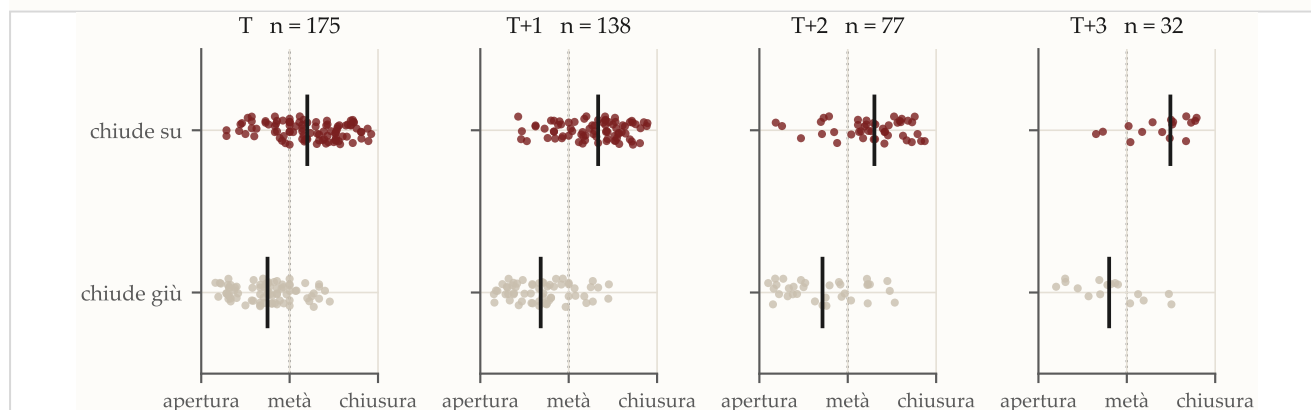


Fig. 12: Dove cade il massimo dentro il ciclo, separando i cicli che chiudono in rialzo da quelli che chiudono in ribasso. Le barre verticali sono le mediane dei due gruppi: è la loro distanza il segnale.

Verdetto CONFERMATA. Con due cautele di lettura, e la seconda è importante.

La prima: la traslazione descrive la polarità del ciclo *in corso di chiusura*, non del successivo. È misura di contesto, non previsione.

La seconda: la *posizione del massimo dentro il ciclo*, presa da sola come descrizione del mercato, **non porta informazione**. Confrontata con la nulla misurata sullo stesso rilevatore — 199 surrogati per ciascuno di due generatori — la posizione mediana osservata cade **dentro** la nube in 14 coppie decisive su 15, e nessuna coppia ne esce contro entrambi i generatori. È geometria della delimitazione, come la coda dei brevi e lo scarto dal nominale. Il caso più netto è Bitcoin a T+4: osservato 0,5952 contro un centro di nulla di 0,595. Ciò che resta, e non dipende da alcuna nulla, sono due cose: che la **costante** canonica del 0,62 non descriva la posizione misurata ai livelli da T a T+4, e che la traslazione **come predittore della polarità dello stesso ciclo** superi la base-rate. La descrizione cade, la relazione regge.

6.6.1. IL CONTESTO CICLICO RIDUCE LA VARIANZA?

È la premessa del § 1, la tesi che regge l'intero impianto, e va messa alla prova come tutto il resto. Farlo richiede un esito, un contesto e una metrica che non sia la media — perché la premessa parla di **varianza**. Esito: il rendimento logaritmico netto del ciclo figlio, da minimo a minimo. Contesto: in quale metà del ciclo padre il figlio si apre — strutturale e indipendente dall'esito, ma **retrospettivo**, e la riserva è scritta sotto. Metrica: la riduzione relativa dello scarto interquartile.

E richiede la nulla giusta, perché **qualunque partizione riduce la dispersione condizionata**, anche una casuale: senza permutare le etichette di contesto, ogni risultato positivo sarebbe stato aritmetica.

Verdetto CONFERMATA su sette righelli su nove. Su 23 celle valutabili e 83.531 cicli figlio la riduzione dello scarto interquartile è positiva in 22 e supera la propria nulla per permutazione in 21; la regola dichiarata prima — almeno due mercati sopra la nulla — è soddisfatta su H4, H2, H1, M30, M15, M5 e M1. Sul Daily una cella su due valutabili, sul Weekly nessuna. L'unica cella negativa è H4/Bitcoin, dove la partizione strutturale riduce la dispersione **meno** di una casuale ($-4, 5\%$, $p = 0,834$).

La riserva definisce che cosa il verdetto significa: il contesto è retrospettivo. La metà del ciclo padre si calcola con il suo minimo di **chiusura**, che al momento in cui il figlio si apre non è ancora avvenuto. Il risultato dice quindi che *la posizione dentro il padre porta informazione sulla dispersione dell'esito* — un limite **superiore** a quanto la struttura possa dare a chi decide — e non che quella riduzione sia disponibile in tempo reale. La variante causale, con la metà stimata dal periodo nominale invece che dal padre realizzato, è misurabile con la stessa infrastruttura e non è stata eseguita.

E l'entità non è uniforme: dall'11% al 18% su H2, H4 e M15, ma fra l'1,6% e il 3,5% su M5 e M1, dove i campioni contano decine di migliaia di cicli e la banda della nulla si stringe a pochi millesimi. Lì la significatività viene dal campione, non dall'entità.

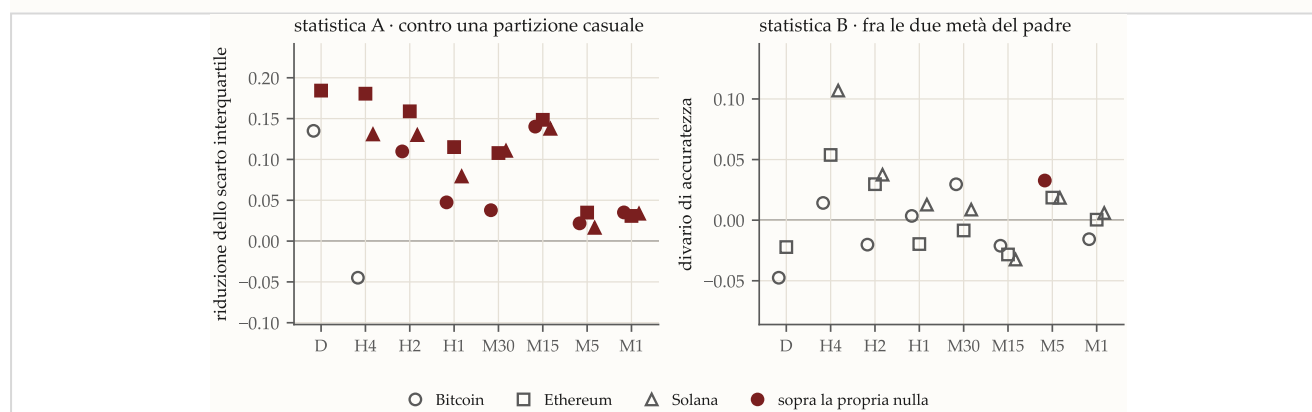


Fig. 13: Le due statistiche, cella per cella su tutti i righelli. A sinistra la riduzione dello scarto interquartile: **qualunque** partizione riduce la dispersione, quindi il pieno segna le celle che superano la **propria** nulla per permutazione, non lo zero. A destra il divario di accuratezza del predittore di traslazione fra le due metà del padre. Le due si guardano insieme perché **si sono scambiate di posto**: sul solo Daily reggeva la seconda, su ventitré celle regge la prima.

Va anche detto che il test è più grossolano della premessa: usa un contesto binario e un segnale di pari livello, non di livello inferiore come la premessa richiede. Un test più fedele vuole un segnale a grana più fine — DSP, order flow — condizionato dalla struttura, e quel segnale non esiste ancora in forma misurabile in questo lavoro.

6.7. Swing condizionato e incondizionato (Elico64)

Il corpus Elico64 distingue lo **swing condizionato** — un sotto-ciclo che chiude negativo nella prima metà del ciclo padre, che richiede conferma — dallo **swing incondizionato**, nella seconda metà, che sarebbe auto-affidabile.

Contato come il corpus lo definisce, lo swing incondizionato chiude il ciclo padre sotto il proprio massimo interno nel **100% dei casi**, su tutti e tre i mercati. Un cento per cento su un campione di quella taglia è un numero che invita a fermarsi, ed è esattamente dove non ci si deve fermare: **manca il gruppo di controllo** (§ 5.4). Il conteggio riguarda i cicli padre *con* swing e non dice nulla su quelli *senza*. Finché quel confronto non c'è, il 100% è compatibile con due mondi opposti — uno in cui lo swing seleziona i cicli che chiuderanno sotto il massimo, e uno in cui **tutti** i cicli chiudono sotto il proprio massimo interno e lo swing non seleziona niente.

La condizione, isolata. Nel modulo di produzione, ramo incondizionato, il successo è

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

dove e è il **minimo di chiusura** del ciclo padre — un minimo per costruzione, perché il rilevatore i cicli li chiude sui minimi — e argmax è la barra del massimo interno. La domanda che quella condizione pone è: «il minimo di chiusura sta sotto il minimo della barra più alta del ciclo?». Formulata così è quasi evidente.

La marginale, misurata sulla catena di livelli che i moduli di produzione usano davvero:

Mercato	Cicli padre con swing	Senza swing	Tasso con	Tasso senza	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

Tab. 8: La condizione di successo dello swing incondizionato, contata anche sui cicli padre che lo swing non ha: il primo giro, sul giornaliero.

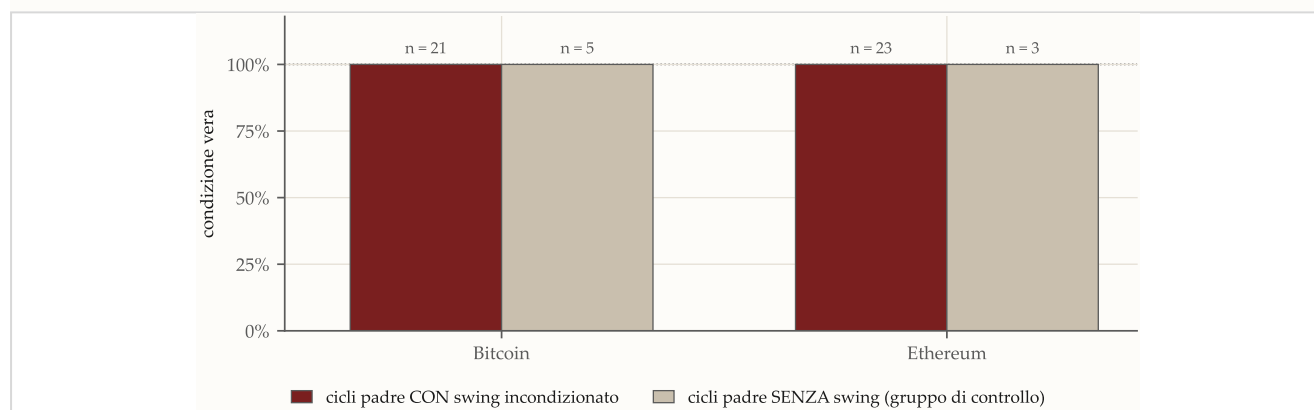


Fig. 14: La condizione contata sui due gruppi, righello per righello sulla corsa e mercato per mercato: pieno il gruppo che lo swing seleziona, vuoto quello che non seleziona, e sotto ogni righello i cicli padre dei due gruppi. I due punti stanno insieme vicino al cento per cento; dove il gruppo di controllo scende sotto la soglia di tautologia — su H2 e M30 — lo scarto resta sotto il decimo che la carta chiede per una conferma.

Verdetto: TAUTOLOGICO — l'affermazione è FALSIFICATA come segnale. Sul giornaliero del primo giro la condizione è vera in **tutti i 191 cicli padre** misurati, con e senza swing: 151 con, 40 senza, tasso 1,000 ovunque, scarto nullo, Fisher $p = 1,000$ su tutti e tre i mercati. Non c'è nulla da selezionare. La soglia di tautologia era dichiarata prima — tasso sui cicli *senza* swing $\geq 0,95$ su tutti i mercati valutabili — ed è raggiunta su tre mercati su tre.

Sulla corsa, letta righello per righello, la soglia di tautologia è raggiunta su cinque righelli su sette. Su H2 e M30 il gruppo senza swing scende al 91–94% in qualche mercato, e in tre casi lo scarto è anche significativo (Fisher p 0,009 e 0,025 su M30, 0,043 su M5); ma resta piccolo, e non arriva mai al decimo che la carta chiedeva per una conferma: il massimo è 0,091. Una cautela: lo swing lega il padre al suo sotto-ciclo di terzo grado, e la corsa lo misura dentro ogni righello, senza il ponte che metterebbe ciascun membro sul righello che lo risolve; la carta tiene l'esito sospeso su questo punto. In nessuna delle due letture disponibili lo swing discrimina.

Resta vero che l'etichetta non è retroattiva: lo swing si riconosce prima della chiusura del padre, e in questo differisce dal ciclo inverso. Ma **essere verificabile ex ante non basta a un segnale che non discrimina**, e i due difetti sono indipendenti: il ciclo inverso è tautologico nella *definizione*, lo swing incondizionato lo è nella *condizione di successo*.

Il ramo condizionato non è toccato: i suoi tassi del 57–79% non saturano, e il confronto fra condizionato e incondizionato resta un confronto fra due gruppi che non assume alcun base rate.

La regola che ne discende, e vale per ogni tasso condizionale di questo lavoro: un tasso misurato su chi ha la condizione va sempre accompagnato dalla sua **marginale**. Senza, un numero perfetto misura la condizione e non il segnale.

6.8. Violazione dei top ciclici (Elico64)

Il corpus formalizza alcuni segnali fondati sul superamento o sulla tenuta di livelli del ciclo precedente. Sono quattro, e il presidio della marginale li separa in modo netto: due reggono nella forma in cui il corpus li enuncia, uno regge solo dopo essere stato riformulato, uno cade.

6.8.1. LA VIOLAZIONE DEL MASSIMO PRECEDENTE, CON LA SUA MARGINALE

Il segnale rialzista — un sotto-ciclo che supera il massimo del precedente — richiede due correzioni prima di essere leggibile. La prima è il **gruppo di controllo**: la tenuta del massimo non è un secondo segnale, è la marginale dello stesso confronto. La seconda è l'**unità di conto**: nel modulo di produzione l'evento è contato *per coppia di sotto-cicli*, e un ciclo padre con cinque sotto-cicli produce fino a tre eventi che condividono lo stesso esito. Non sono osservazioni indipendenti.

Mercato	Cicli padre	Base rate	Con evento	Senza evento	Scarto	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

Tab. 9: Violazione del massimo precedente, contata per ciclo padre e con la marginale accanto, letta su H4 — il righello più alto su cui tutti e tre i mercati hanno gruppi valutabili. Nessuna inclusione insiemistica: verificata falsa su tre mercati su tre.

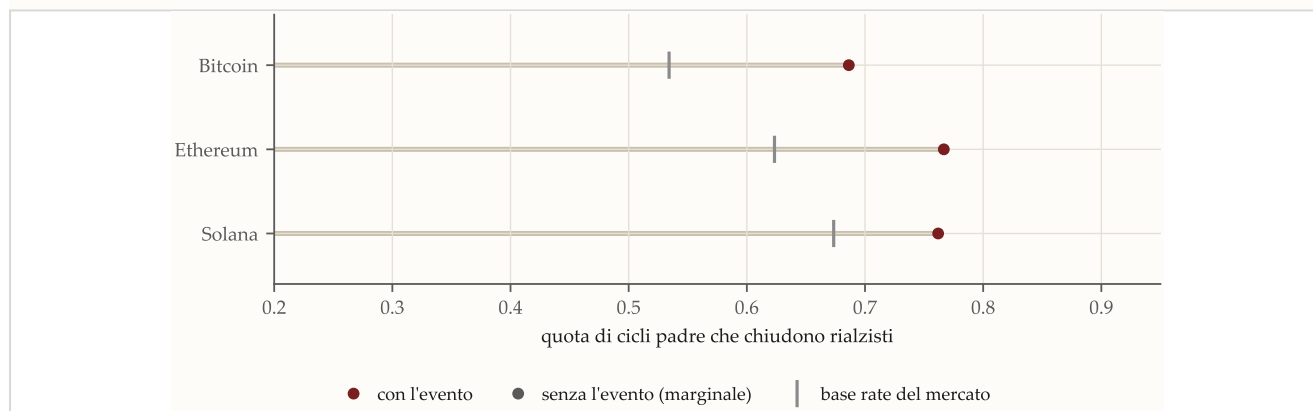


Fig. 15: Letta su H4, come la tavola. Il segnale è la **distanza** fra i due punti; il trattino verticale è il base rate del mercato, cioè quello che si otterrebbe senza guardare nulla. Coprendo il punto chiaro — il gruppo che l'evento non seleziona — resta un numero che non si può interpretare.

Verdetto CONFERMATA, su tre mercati su tre e su sei righelli su nove. Il segnale c'è e non è un artefatto: l'insieme degli eventi non contiene per costruzione gli esiti positivi, e il gruppo di controllo chiude rialzista nel 12–18% dei casi contro un base rate del 53–67%.

Il caso al bordo che questa sezione registrava non c'è più, e la ragione va detta: nella lettura sul giornaliero due valori *p* stavano a due millesimi l'uno dall'altro ai lati della soglia, e la lettura difendibile era che il segnale discriminasse in modo netto su un mercato e marginale sugli altri due. Sul righello che risolve il livello i cicli padre passano da 22 a 73 su Bitcoin e da 14 a 49 su Solana, e i tre valori *p* stanno fra 0,0032 e un decimillesimo. Era un fatto del campione, non del mercato. Sul giornaliero il verdetto resta NON PROVATA, e per mancanza di cicli: due mercati su tre non hanno gruppi abbastanza grandi.

E la tenuta del massimo è FALSIFICATA come segnale autonomo: è la marginale di questo confronto, non un'informazione in più. Contarla come segnale separato significa contare due volte lo stesso dato.

6.8.2. I DUE SEGNALI RIBASSISTI, SENZA L'INCLUSIONE

Il corpus enuncia due segnali ribassisti: la rottura del minimo di partenza del ciclo e la rottura del minimo del ciclo precedente. Il primo, contato come il corpus lo definisce, contiene un difetto di insieme che va visto prima dei numeri.

L'evento è registrato se **esiste** una barra $k \in (s, e]$ con $\text{low}(k) < \text{low}(s)$, e l'esito è $\text{low}(e) < \text{low}(s)$. Ma se il ciclo chiude ribassista, allora $k = e$ soddisfa già la condizione: **ogni ciclo ribassista è per costruzione un evento**. L'insieme degli eventi contiene per intero l'insieme degli esiti positivi — l'implicazione è verificata vera nel 100% dei cicli padre su tutti e tre i mercati — e il tasso che ne risulta non può scendere sotto il base rate dei ribassisti.

La versione che elimina l'inclusione chiede che la rottura cada **entro il primo 75% della durata**, e confronta la quota di ribassisti fra chi rompe presto e chi non rompe presto: due gruppi entrambi non vuoti, nessuno dei quali contiene per costruzione gli esiti dell'altro.

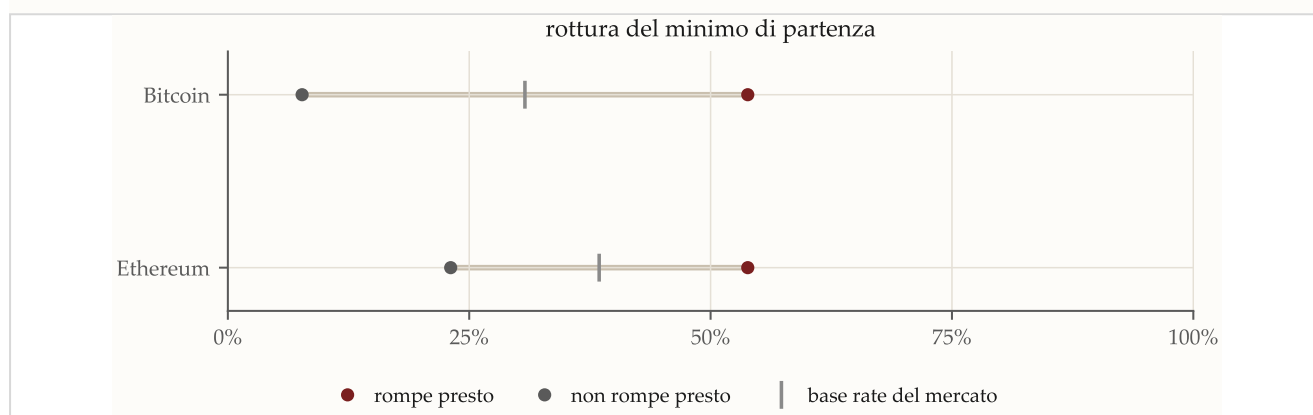


Fig. 16: I due segnali nella forma discriminante, letti su H4 come la tavola: la quota di cicli ribassisti fra chi rompe presto, fra chi non rompe, e il base rate del mercato. La distanza fra i due punti è il segnale; il trattino è ciò che si otterrebbe senza guardare nulla.

Verdetto CONFERMATA su tre mercati su tre, per entrambi i segnali, e su sette righelli su nove. La rottura precoce del minimo di partenza alza la quota di cicli ribassisti di 13–23 punti sopra il base rate; la sua **assenza** la abbassa di 19–33. Lo scarto fra i due gruppi vale 32, 56 e 53 punti; il segnale del minimo precedente vale 33, 45 e 51 punti, e lì l'inclusione non c'è: verificata **falsa** su tutti e tre i mercati. Fuori da H4 il risultato non cambia: da H2 fino al minuto entrambi i segnali discriminano in tutti e tre i mercati, con campioni che vanno da un centinaio a diecimila cicli padre. Solo il settimanale e il giornaliero restano fuori, e per mancanza di cicli, non per esito.

La riformulazione migliora l'affermazione in due modi. È un segnale **più forte** di quanto la forma originale dichiarasse, ed è **a due facce**: l'assenza di rottura precoce è un'indicazione rialzista altrettanto netta — 13–22% di cicli ribassisti contro un base rate del 42–51% — e nella formulazione del corpus non compare affatto.

Le due sezioni, prese insieme, dicono qualcosa sul presidio della marginale: lo stesso controllo, applicato allo stesso modo a affermazioni vicine dello stesso corpus, ne ritira una, ne ridimensiona una e ne rafforza due. **Non è uno strumento per demolire.**

6.9. I minimi ciclici violati: che cosa separa, e che cosa non anticipa

Due domande diverse cadono sullo stesso oggetto — la violazione di un livello sotto il minimo di un ciclo — e la pratica corrente le tratta insieme. Sono state misurate separatamente, con il criterio congelato prima del codice, e danno due esiti opposti.

6.9.1. IL LIVELLO STATICO SEPARA, LA TREND LINE ASCENDENTE NO

La prima domanda è se «il minimo di un ciclo non viene violato» distingue un minimo ciclico vero da un punto qualunque. Il test costruisce **falsi minimi** con seme fisso e finestre della stessa durata, e confronta la quota di violazioni sui veri e sui falsi: se l'affermazione ha contenuto, i veri devono essere violati **meno**.

Con il livello statico — il minimo del ciclo precedente, una retta orizzontale — la separazione c'è ed è larga: i minimi veri risultano violati nel 67,0% dei casi e i falsi nel 91,7%, su 2.003 punti per gruppo; la direzione è la stessa in **21 celle su 21**, in 18 la differenza è significativa, e 9 contesti su 11 superano la soglia dichiarata. Il verdetto è **confermato**, con un'avvertenza che va letta insieme al numero: il 67% è alto in senso assoluto. L'affermazione «un minimo ciclico non viene violato» è vera solo in senso **relativo**, come filtro contro il caso, e non come descrizione di ciò che i minimi fanno.

Con la trend line ascendente — la retta che unisce due minimi crescenti, lo strumento che la didattica presenta accanto al livello e con lo stesso peso — la separazione **non c'è**. I veri sono violati nel 46,3% dei casi e i falsi nel 44,9%: i veri stanno dalla parte sbagliata, la direzione attesa compare in 8 celle su 21, nessuna cella è significativa, nessuno degli 11 contesti passa. Il verdetto è **falsificato**.

L'asimmetria è il risultato, non il dettaglio: due strumenti che la tradizione insegna insieme si comportano in modo opposto quando vengono messi contro lo stesso gruppo di controllo, e metà della didattica poggia sulla metà che non discrimina.

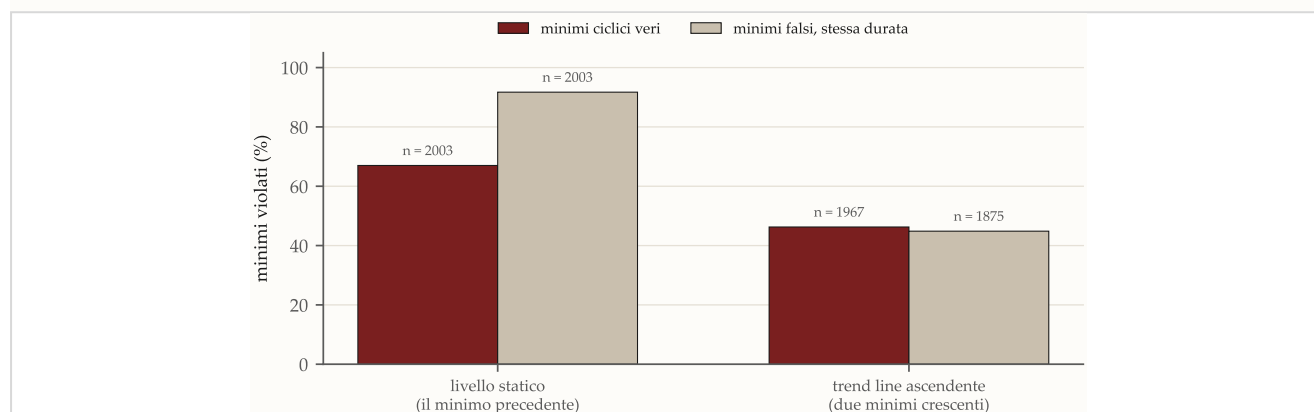


Fig. 17: La quota di violazioni sui minimi ciclici veri e su minimi falsi costruiti con seme fisso e finestre della stessa durata, per i due strumenti. Con il livello statico le due barre distano venticinque punti nella direzione attesa; con la trend line ascendente si toccano, e i veri stanno dalla parte sbagliata. Il denominatore è scritto sopra ogni barra: una quota senza il suo denominatore non è verificabile.

6.9.2. LA VIOLAZIONE NON ANTICIPA LA CONFERMA DEL MINIMO

La seconda domanda è operativa: la violazione di un livello **anticipa** il momento in cui il minimo viene dichiarato? Sette dichiaratori — cinque, più due placebo che fanno da controllo — sono stati messi alla prova su 21 contesti, ciascuno con il proprio tasso di base e con un **placebo di forma** — una finestra della stessa geometria piazzata dove l'evento non c'è — perché una finestra orientata nel tempo produce da sola una parte del vantaggio.

Nessuno dei sette passa: la quota di contesti che superano il criterio è **zero** per tutti e sette, e il verdetto complessivo è **non anticipa**. Il modo in cui falliscono, però, non è lo stesso, ed è la parte che indica dove guardare. **VTL_SUB** è **preciso**: batte il proprio tasso di base in 15 celle su 17 e supera il placebo di forma in 11, cioè quando parla ha ragione più spesso del caso. Ma **non copre**: in nessuna cella arriva alla metà degli eventi. **NEST** è il contrario: copre tutte e 17 le celle e non batte la forma in 12 su 17.

Le due misure di questa sezione hanno un runner proprio e non passano dalla costruzione descritta nel § 5.6.1: le loro relazioni sono misurate sotto la restrizione dichiarata lì, dieci celle complete su diciassette.

6.10. Lingue di Bayer: codifica, calibrazione, potere discriminante e valore predittivo

Il rilevatore tripartito (§ 3.3) classifica i cicli di raccordo. La ripartizione è coerente fra i tre mercati indipendenti:

Calibrazione. La soglia canonica di Bayer — durata sotto la metà della media — non è adatta alle criptovalute: la distribuzione di durata reale richiede una calibrazione empirica per asset e timeframe. Che la soglia classica vada ricalibrata è un risultato e viene riportato; i valori-soglia, i percentili e la procedura restano proprietari.

Il potere discriminante sta nel passo di prezzo. La detection ha due passi, e vale la pena chiedersi quale dei due porti l'informazione. Isolando il passo 1 — la sola condizione di brevità, con soglia causale posta al quarto inferiore delle durate già osservate a quel livello — e misurando se i cicli brevi siano seguiti da una ripartenza più spesso dei cicli normali:

La sola brevità non porta informazione sulla ripartenza. Il criterio a due passi resta valido come costruzione, ma il suo potere discriminante è interamente nel passo di prezzo: la brevità seleziona i candidati, non li classifica.

6.10.1. IL VALORE PREDITTIVO, GIUDICATO SENZA CIRCULARITÀ

Ne segue una cautela che va portata fino in fondo: confermare una lingua tramite la ripartenza successiva è **circolare**, perché la ripartenza è anche ciò che il passo di prezzo misura. Il valore predittivo va misurato con un disegno che non usi la ripartenza come conferma.

La rottura della circolarità è un **intervallo cieco**. L'etichetta viene assegnata dal rilevatore causale come sempre; poi si aspetta che la finestra che il rilevatore ha usato sia interamente passata — il massimo fra 30 barre e metà del periodo nominale del livello — e solo da lì comincia a contare il rendimento, misurato su un orizzonte pari al periodo nominale. Il gruppo di controllo sono i minimi dello stesso livello che il rilevatore **non** ha etichettato lingua, valutati con la stessa regola. Senza gruppo di controllo un tasso alto non direbbe nulla: sul comparto crypto quasi tutto è rialzista su orizzonti lunghi.

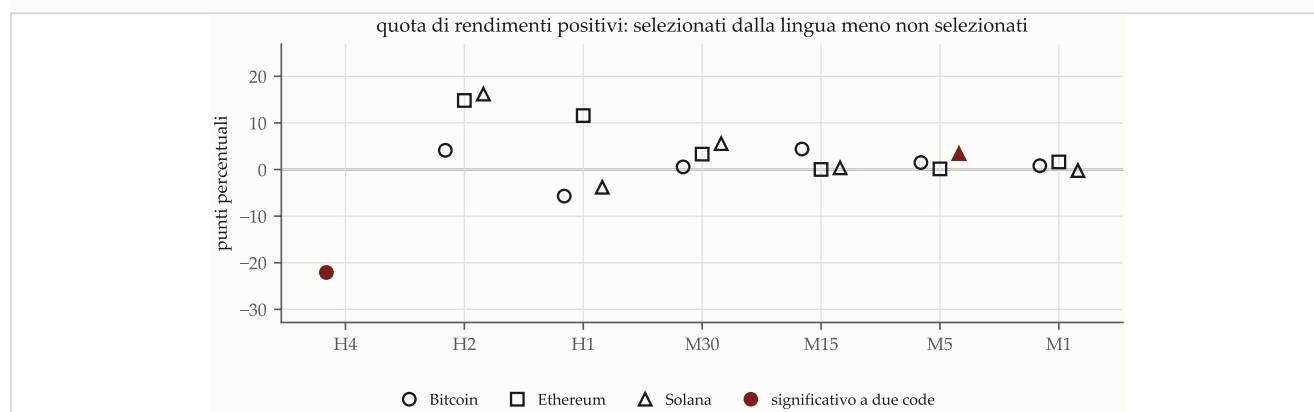


Fig. 18: Di quanto la quota di rendimenti positivi dei cicli **selezionati** dalla lingua supera quella dei cicli non selezionati, in punti percentuali, per ogni righello e ogni mercato. Lo zero è l'ipotesi nulla. Diciannove celle: due significative, in direzioni opposte, e la dispersione che **si contrae man mano che il campione cresce** — fino a sedici punti su H2 e H1, dove i gruppi contano centinaia di minimi, entro tre punti e mezzo su M5 e M1, dove ne contano migliaia. È la firma del rumore che si media.

Verdetto NON PROVATA, su **9.270 minimi lingua** contro **55.539 di controllo**, diciannove celle valutabili su ventisette. Due sole celle significative, e in **direzioni opposte**: su H4/Bitcoin la lingua fa nettamente **peggio** del controllo — 44,6% di rendimenti positivi contro 66,7%, con un controllo di soli 24 minimi su un tratto rialzista — e su M5/Solana un poco meglio, 49,6% contro 46,2%. Il verdetto non è **FALSIFICATA** perché la regola chiedeva che **nessuna** cella lo fosse; non è **CONFERMATA** perché chiedeva almeno due mercati significativi con lo stesso segno a favore, e ce n'è uno.

Il righello su cui la misura era nata non è più valutabile, ed è l'unico modo onesto di riportarlo: sul Daily l'etichetta lingua è quasi universale ai due soli livelli rimasti, e un gruppo di controllo di uno o due minimi non è un gruppo di controllo.

Questo non tocca ciò che il documento afferma sulla rilevazione: che sia causale, verificata non-repainting, densa e coerente fra mercati indipendenti. Sono proprietà della **rilevazione**. La posizione difendibile è che la rilevazione delle lingue sia uno strumento **descrittivo verificato, non un predittore**.

6.11. Convergenza cross-asset, e i livelli alti su un secolo di storia

Il rilevatore viene eseguito indipendentemente sui tre mercati. La convergenza delle durate mediane, misurata come campo di variazione percentuale fra i tre asset, è forte sui livelli brevi e sembra allentarsi ai livelli alti, dove il campione cala.

Livello	Campo di variazione (3 crypto)	Campo di variazione (6 indici, storia lunga)
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%

Tab. 10: Convergenza cross-asset delle durate mediane. La colonna di destra è misurata su sei indici azionari con fino a novantanove anni di storia giornaliera.

Il 28% ai livelli alti era rumore da campione piccolo, non un indebolimento del principio. Sui tre mercati crypto i cicli di livello T+4 e T+5 sono fra quattro e ventuno per mercato, e con quel campione la dispersione fra mercati non è interpretabile. Sui sei indici, dove il campione c'è, il campo di variazione a T+4 e T+5 vale **9,2% e 9,7%** contro 8,6% e 8,9% ai livelli medi: **non c'è salto**. La comunanza ciclica di Hurst regge anche in alto appena le si dà campione.

Il campione ai livelli alti era un limite di storia, e ora ha un numero. Il livello T+5 passa da 4–6 cicli su crypto a **18–53 sugli indici** e supera la soglia di misurabilità — venti cicli, ereditata dalla tavola del dip del § 6.3 — su **cinque asset su sei**. Il rilevatore non satura: produce minimi al ritmo che la durata del livello consente, e il limite era della serie.

T+7 no, e nemmeno con un secolo. Sull'S&P 500, che parte dal dicembre 1927, ne escono quindici. Con una mediana misurata fra 1.166 e 1.799 barre, per venti cicli servono fra **93 e 143 anni di sedute**. T+7 non è un livello non misurato: è un livello che nessun dataset esistente può misurare, e va detto così.

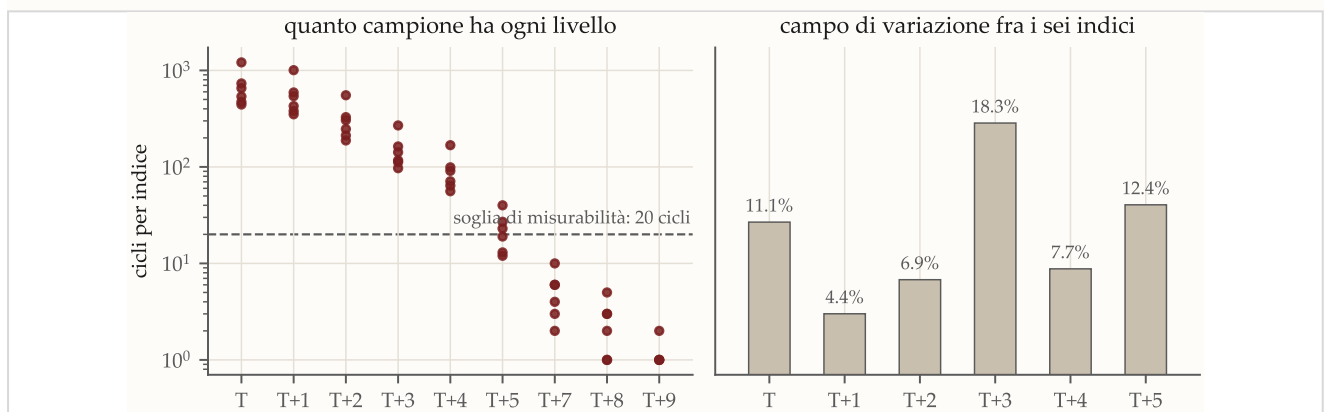


Fig. 19: A sinistra, quanti cicli ciascun livello possiede su ciascuno dei sei indici a storia lunga (scala logaritmica): la riga tratteggiata è la soglia di misurabilità dichiarata. A destra, il campo di variazione delle durate mediane fra i sei indici, livello per livello. Il pannello di sinistra dice dove il campione c'è; quello di destra dice che, dove c'è, la convergenza **non si allenta** salendo di livello — il salto al 28% della **Tabella 10** è rumore da campione piccolo, non un indebolimento del principio.

Verdetto CONFERMATA. La comunanza ciclica regge ai livelli brevi sui tre mercati crypto (3,0–6,1% di campo di variazione) e ai livelli alti sui sei indici (9,2–9,7%). Tre mercati con microstruttura, capitalizzazione e storia diverse mostrano durate mediane entro pochi punti percentuali l'una dall'altra, e la stessa cosa vale su un secolo di azionario. È l'evidenza più forte a favore del principio hurstiano, e non dipende da alcuna calibrazione: le mediane sono misure dirette.

Un limite dichiarato, e non salva né condanna il verdetto: tutte le misure di storia lunga usano il catalogo dei periodi nominali tarato su criptovalute, applicato a sei indici azionari. Un catalogo per classe di asset non esiste ancora.

6.12. Sequenze di polarità ammesse: un effetto di posizione

Il corpus Elico64 postula che i sotto-cicli di un ciclo padre possano assumere solo alcune sequenze di polarità. È l'ipotesi più ambiziosa del filone italiano. Sul solo timeframe giornaliero il campione non basta — ventitré cicli padre a quattro tempi sui tre mercati — e la strada naturale è scendere di timeframe, dove i cicli a

quattro sotto-cicli dovrebbero abbondare. Il test, rifatto su H4 e H1, dice due cose, e la seconda è molto più importante della prima.

Secondo: la regola è geometria del rilevatore. Verificato per enumerazione, l'insieme delle otto sequenze ammesse a quattro tempi — quattro «indice» e quattro «inverso» — coincide esattamente con l'insieme delle sequenze in cui **il primo e l'ultimo sotto-ciclo hanno polarità opposta**. Le sei ammesse a tre tempi coincidono con «non tutte e tre uguali».

Questo apre un problema. Il rilevatore chiude ogni ciclo padre su un minimo: il primo sotto-ciclo parte da quel minimo e tende a salire, l'ultimo arriva al minimo successivo e tende a scendere. La misura lo conferma senza ambiguità: in prima posizione il sotto-ciclo è positivo nell'**83–92%** dei casi, in ultima nel **13–22%**.

Servono perciò tre riferimenti, non uno.

Timeframe	<i>n</i>	Tasso ammesse	Baseline misurata	<i>p</i> nulla sul flusso	<i>p</i> nulla per posizione
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

Tab. 11: Strutture a quattro tempi: il tasso osservato contro le sue nulle. La baseline misurata usa la quota di sotto-cicli positivi effettivamente osservata al posto del 50% teorico.

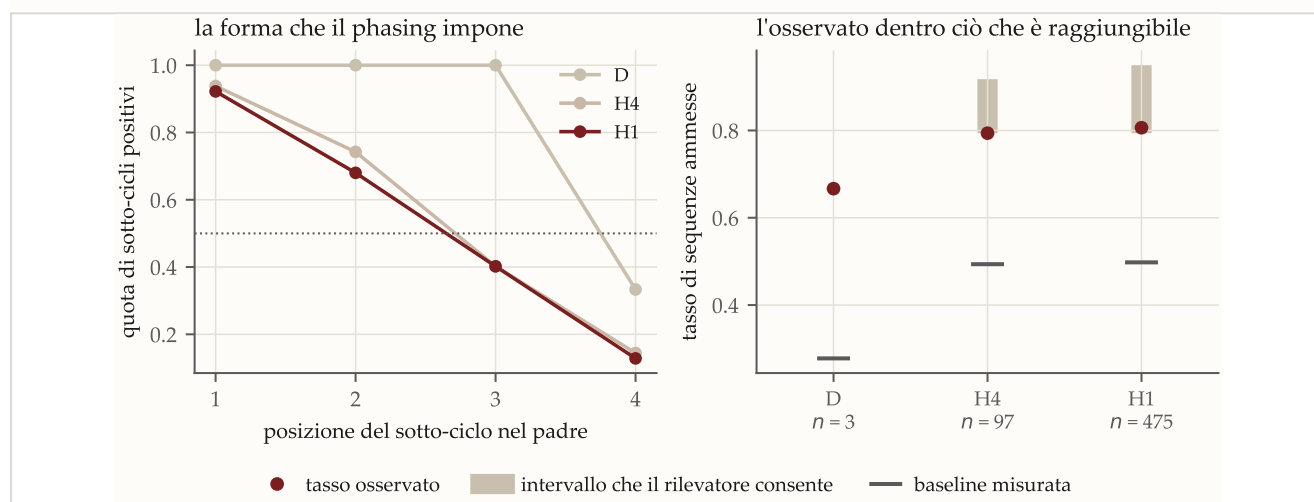


Fig. 20: A sinistra la quota di sotto-cicli positivi per posizione dentro il ciclo padre: la forma quasi obbligata che il phasing impone — si apre salendo, si chiude scendendo. A destra il tasso di sequenze ammesse dentro l'intervallo che quella forma consente: sul giornaliero l'osservato sta esattamente sul fondo dell'intervallo, cioè il mercato produce meno sequenze ammesse di quante il rilevatore ne permetta.

La **baseline ingenua** del 50% vale solo se le polarità sono eque; non lo sono, e la baseline **misurata** scende a 49,2–49,5%. Contro quella, il tasso osservato sembra alto. La **nulla per permutazione del flusso** — che rimescola tutte le polarità insieme, conservando la marginale complessiva — lo conferma: $p = 0,0005$ su H4 e H1, cioè il minimo dichiarabile con duemila rimescolamenti. Con quella nulla il punto si sarebbe chiuso con una conferma brillante.

La **nulla per posizione** rimescola invece ogni colonna separatamente: conserva la marginale di ciascuna posizione nel ciclo padre e distrugge la sola dipendenza fra posizioni diverse. È l'unica delle tre che conserva il confondente. Contro di lei il valore p vale 0,711 su H4, 0,998 su H1 e 1,000 sul Daily: **il tasso osservato non la supera mai**.

E quel 1,000 va letto con attenzione, perché non è un pareggio. Poiché l'insieme ammesso a quattro tempi coincide con «primo e ultimo di polarità opposta», a marginali di posizione fissate il tasso dipende dalla sola sovrapposizione a fra i sotto-cicli positivi in prima e in ultima posizione: vale $(k_1 + k_L - 2a)/n$, e può percorrere solo un intervallo chiuso e stretto. Misurato: **0,696–0,957 sul Daily, 0,701–0,858 su H4, 0,739–0,911 su H1**. Sul Daily l'osservato sta esattamente sul **fondo** di quell'intervallo — ogni sotto-ciclo finale positivo appartiene a un padre che inizia positivo — quindi $p = 1,000$ non significa «nessuna evidenza» ma «evidenza nel verso opposto: il mercato produce meno sequenze ammesse di quante il rilevatore ne consenta».

Su H4 l'intero intervallo che il rilevatore consente è largo quindici punti percentuali. Il 70–75% che si legge come regolarità non poteva, su quel timeframe, scendere sotto il 70% qualunque cosa facesse il mercato: il numero era quasi obbligato prima di guardare i dati.

Verdetto NON PROVATA, e la ragione sostanziale conta più di quella formale: non c'è niente da confermare. Contro la nulla che conserva il confondente il valore p non scende mai sotto 0,7.

La struttura a tre tempi si legge ancora meglio. L'insieme ammesso è «non tutte e tre uguali», la baseline misurata vale 0,720–0,741 e il tasso osservato 0,709–0,757, su campioni da 74 a 2.129 sequenze. Non c'è nemmeno lo scarto apparente che si osserva a quattro tempi: la regola a tre tempi non porta informazione a nessun timeframe, e questo **non** dipende dal campione.

Nota di onestà procedurale. La nulla per posizione è stata aggiunta dopo aver visto il primo esito, quando il confondente è stato riconosciuto. È un test più severo, non più permissivo: può solo togliere sostegno all'affermazione, mai aggiungerne. Il valore della prima nulla resta registrato accanto al secondo e nessuna soglia è stata toccata.

L'effetto è del rilevatore a ogni scala provata, e adesso le scale sono nove. La stessa misura, con le stesse soglie e la stessa nulla, è stata ripetuta su tutti i timeframe che il pavimento ammette, fino al minuto. Il campione minimo per votare — cento sequenze a quattro tempi, dichiarato prima e non ritoccato — lascia fuori Weekly, Daily e H4, che ne portano zero, tre e novantasette: **votano in sei**.

Due letture, e la seconda è la più forte. La prima: su **sei timeframe votanti su sei** il verdetto è **NON PROVATA**, e il valore p contro la nulla che conserva la marginale di posizione non scende mai sotto 0,5 — su campioni che vanno da 172 a 3.220 sequenze, cioè fino a **mille volte** quello del giornaliero. Il campione non è più un'attenuante disponibile.

La seconda: **il confondente non cambia con la scala**. La quota di sotto-cicli positivi in prima posizione vale 0,92–0,94 su tutti e sette i timeframe intraday e 0,11–0,14 in ultima. La forma obbligata del ciclo padre — si apre salendo, si chiude scendendo — è la stessa dai quattro ore al minuto, e con essa resta identico l'intervallo che quella forma **consente**. E la baseline misurata converge a **0,500** man mano che il campione cresce: un effetto di posizione che si riproduce invariato su quattro ordini di grandezza di risoluzione è una proprietà del rilevatore, non un accidente di un timeframe.

Che cosa resta di Elico64 su questo punto: niente, nella forma della «sequenza ammessa». Resta una osservazione utile che il corpus non formula ma che i numeri sostengono — un ciclo padre, così come il rilevatore lo delimita, ha una forma quasi obbligata: si apre salendo e si chiude scendendo. Chiamare «ammesse» le sequenze che rispettano quella forma è una descrizione corretta della geometria e una previsione vuota.

6.13. Fuori dal timeframe giornaliero: quattro affermazioni su tre scale

Un solo timeframe per le affermazioni statistiche è un limite, e va tolto dove si può. Quattro affermazioni misurabili senza surrogati sono state rimisurate su Weekly, H4 e H1 per i tre mercati, con la stessa pipeline e con le soglie già dichiarate — non nuove: che le durate misurate non corrispondano al nominale (A1), che alcuni livelli adiacenti non siano distinti (A2), che la traslazione predica la polarità (A3), che esista una coda di cicli brevi crescente con il livello (A4).

A1 · le durate si scostano dal nominale	regge	no	no	1/3
A2 · due livelli adiacenti non si distinguono	regge	regge	regge	3/3
A3 · la traslazione predice la polarità	regge	regge	regge	3/3
A4 · esiste una coda di cicli brevi	regge	regge	regge	3/3
	W	H4	H1	

Fig. 21: Le quattro affermazioni portanti, cella per cella, sui tre timeframe fuori dal giornaliero. Le caselle piene non stanno tutte su una colonna — non è una scala a reggerle — e le **due vuote stanno sulla stessa riga**: A1, la sola che cade, e cade proprio dove il righello risolve i livelli della propria catena.

Verdetto CONFERMATA, tre affermazioni su quattro reggono fuori dal Daily — la soglia dichiarata prima della misura ne chiedeva almeno tre. Il caso più netto è la traslazione, che regge su **nove casi su nove**: su H4 vale 72,5–75,7% con 16,6–21,2 punti di uplift sulla base rate misurata, su campioni da 3.073 a 4.658 cicli per mercato; su H1 vale 72,0–73,8% con 17,5–21,4 punti, su campioni fino a 19.623 cicli. È la stessa regolarità del § 6.6, misurata con due ordini di grandezza di dati in più.

6.14. La prova più larga: trentasette strumenti, cinque statistiche

Tutto ciò che precede riguarda tre criptovalute. È il campione su cui il lavoro è nato, ed è un campione stretto: tre serie della stessa classe, correlate fra loro, dentro un decennio. La domanda che chiude il protocollo è perciò la più scomoda — **la struttura ciclica che il rilevatore estrae si distingue dal caso, su un campione largo?** — e va posta con la nulla più severa disponibile.

Il disegno. Trentasette strumenti in sei classi (criptovalute, azioni, indici, commodities, valute, obbligazionario), su Daily, con la stessa catena di livelli e **cinque** statistiche diverse. Cinquanta surrogati a fase randomizzata e cinquanta AAFT per strumento, con il rilevatore di produzione applicato **identico** al vero e al sintetico; test dei segni sulle 182 coppie strumento/livello.

L'esito, su 1.820 confronti. Contro i surrogati a fase le cinque statistiche danno 85–101 coppie favorevoli su 182 — attorno alla metà — con p fra 0,079 e 0,832. Contro gli **AAFT**, che conservano anche la distribuzione dei rendimenti e sono la nulla severa, il quadro peggiora: 75–89 su 182, p fra 0,644 e 0,993, cioè **a metà o sotto**, spesso nella direzione *opposta* a quella attesa. Le coppie che superano la soglia sul singolo confronto sono fra 2 e 13 su 182: l'ordine di grandezza che il caso produce. **Nessuna cella significativa in nessuna delle sei classi.**

Verdetto: FALSIFICATA. Su un campione largo e con la nulla severa, la struttura ciclica estratta dal rilevatore **non si distingue** da quella che lo stesso rilevatore estrae da serie con lo stesso spettro e le stesse code. Il vantaggio misurabile su tre criptovalute contro i soli surrogati a fase — coerente di segno su dieci coppie su dodici — **non sopravvive** all'allargamento né al surrogato più severo.

Un limite che va dichiarato, e non salva il verdetto. Il catalogo dei periodi nominali usato è quello tarato su criptovalute, applicato a tutte e sei le classi: parte del fallimento può essere un catalogo sbagliato per azioni, indici e commodities — cercare cicli ai periodi sbagliati non trova struttura nemmeno dove c'è. Ma le criptovalute sono **dentro** il campione con il **loro** catalogo, e anche lì nessuna cella è significativa.

6.15. L'assioma di partenza: il quaternario è un binario di binari?

L'ultima affermazione da mettere alla prova non appartiene a nessuna scuola storica: è il punto di partenza teorico di questo lavoro. **L'alfabeto strutturale** è $\{2, 3\}$, e il **quaternario non è una terza categoria: è un binario di binari**. La prima parte — la prevalenza di due e di tre nei conteggi — è stata misurata contro i surrogati: contro quelli a fase randomizzata qualche confronto esce dalla nube, contro gli AAFT — la nulla

più severa — nessuno, e su trentasette strumenti nessuna statistica di struttura la supera (§ 6.14). La parte forte va provata a sé.

Il test è diretto e non ha parametri. Un ciclo padre con quattro sotto-cicli ha tre minimi interni. Nella lettura «per quattro» sono confini equivalenti; nella lettura «due per due» quello di mezzo è il confine fra i due binari, cioè un minimo di grado superiore, e sul prezzo deve essere **il più basso dei tre**. Sotto pura indifferenza il più basso è uno dei tre con probabilità un terzo.

Su **557 cicli con quattro sotto-cicli** — aggregati su quattro righelli (D, H4, M5, M1) e tre mercati — il minimo centrale è il più basso in **7 casi: l'1,3%**, contro il 33,3% che l'indifferenza richiede. Lo scarto è di trentadue punti, e da solo non dice ancora di chi sia.

E il campione dice da dove viene: il 73% dei quaternari nasce su M5 e il 27% su M1, cioè sui livelli inversi bassi (T-3 ... T-7) che solo i righelli fini risolvono; i sette casi centrali vengono **tutti** da M5. Con righelli fermi a M15 la coppia analizzabile sarebbe T-1/T su H1: la portata della misura va letta così, riguarda il quaternario **dove i righelli lo risolvono**, e sono quelli fini.

Perché così poco. La nulla misurata con lo stesso rilevatore sui surrogati ha **mediana 0,0%** e novantacinquesimo percentile **2,3%** contro i surrogati a fase e **2,1%** contro gli AAF: l'1,3% osservato ci sta **dentro** su entrambe le famiglie (p empirico 0,176 e 0,189). Su queste serie il rilevatore da solo produce la stessa quota. La ragione era prevista prima di misurare ed è la geometria documentata al § 6.12 — il primo sotto-ciclo sale nell'83–92% dei casi, l'ultimo scende — che spinge il minimo centrale a essere il più **alto** dei tre. Se il quaternario fosse davvero un binario di binari, dovrebbe vedersi **nonostante** la geometria del rilevatore, non grazie a essa. Il mercato, toltà la geometria, è muto su quel punto — e resta da chiedere se il rilevatore avrebbe potuto farlo parlare.

Il test avrebbe visto un binario di binari? Un esito negativo vale solo se il test ha il potere di vedere il fenomeno (§ 5.5), e qui la domanda si può porre direttamente: su serie sintetiche con un binario di binari **impiantato** — tre collocazioni del periodo, 288, 432 e 576 barre, e intensità crescenti del minimo centrale — si fa girare la stessa procedura di produzione. Sui minimi veri del sintetico il minimo centrale è il più basso nel 64% dei padri a metà dell'intensità minima e nel 90% all'intensità minima: il fenomeno c'è ed è marcato. **La procedura non lo restituisce.** Aggancia il livello intermedio e lo legge come due binari; all'intensità minima e sopra non produce nessun padre a quattro figli, e dove l'effetto è debole o assente ne produce al più sei per prova — 110 in tutto su 360 prove, mai con il minimo centrale più basso. Una prova **rileva** la struttura quando la procedura restituisce almeno i trenta padri a quattro figli che servono per votare: all'intensità minima le rilevazioni sono **zero in trenta prove**, in ciascuna delle tre collocazioni. La simulazione riproduce il righello M5, da cui viene il 73% dei quaternari di mercato.

E il resto dell'alfabeto, che nessuno aveva mai contato. La stessa passata censisce **tutte** le cardinalità, non solo la quarta. Su 26.017 cicli padre: 13.284 non contengono nessun minimo del livello sotto, e dei 12.733 che hanno struttura interna **9.450 sono binari (74,2%), 2.651 ternari (20,8%), 557 quaternari (4,4%) e 67 quinari (0,5%)**; otto ne hanno sei o più. L'alfabeto {2, 3} copre il **95,0%** dei casi.

Due letture, e vanno tenute separate. Il **ternario** non si distribuisce come l'indifferenza: su 2.651 padri il minimo più basso cade 1.199 volte nella prima posizione e 1.452 nella seconda, uno scarto dalla uniforme di 0,048 con $p = 0,0005$ — il minimo dichiarabile con duemila estrazioni. **Ma quello scarto non esce dalla nube del rilevatore:** sui surrogati la mediana vale 0,125, cioè il doppio e mezzo dell'osservato. Lo squilibrio è reale e non è del mercato: è ciò che il rilevatore produce anche quando la struttura non c'è. Il **quinario esiste**, e con sessantasette casi non è più questione di campione: la sua asimmetria vale 0,440 contro una mediana di 0,524 sul rumore — **sta dentro la nube, e più in basso del rumore stesso**. Non è un'unità nuova dell'alfabeto.

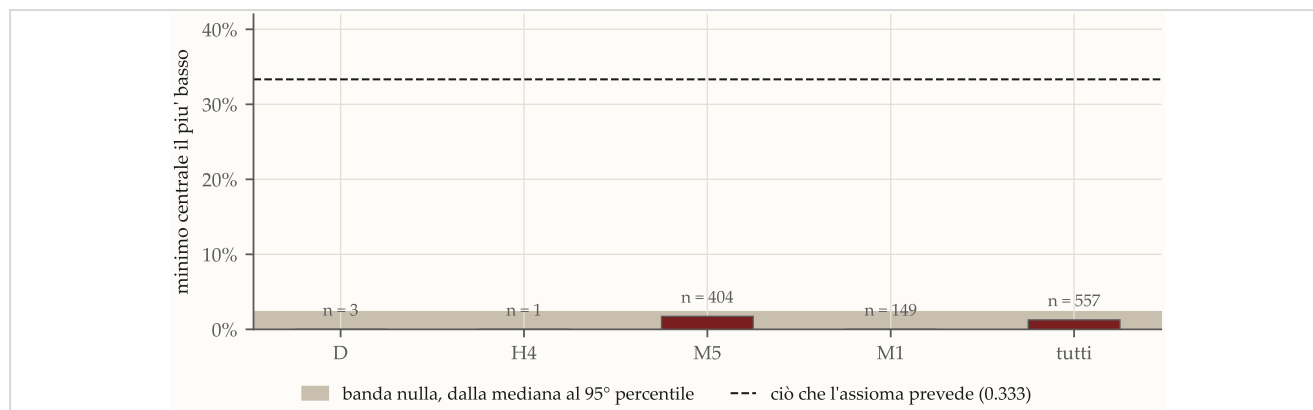


Fig. 22: La quota di cicli a quattro tempi in cui il minimo interno centrale è il più basso, contro ciò che l'assioma prevede e contro la banda del rilevatore sul rumore. L'assioma chiede un terzo, il mercato dà l'1,3% — e la banda del rilevatore sul rumore arriva al 2,3%, cioè lo contiene.

Verdetto NON PROVATA. La parte forte dell'assioma non si può decidere con questo strumento: il mercato dà un minimo centrale più basso nell'1,3% dei cicli, dentro la banda del rilevatore, e il rilevatore non riconosce un binario di binari nemmeno dove è impiantato. L'alfabeto {2, 3} resta una descrizione utile dei conteggi; se il quaternario sia un binario di binari, questo lavoro non lo sa.

Resta dichiarato che la forma testata è **una** lettura di «per quattro uguale due per due»: quella attraverso la profondità del minimo che separa i due binari. Un'altra — la simmetria delle durate fra le due metà — è concepibile e non è stata provata. Deciderle entrambe chiede prima un rilevatore che, sul sintetico, restituisca la struttura che vi si è messa.

6.16. Il valore dell'idea, misurato con onestà

Un'idea ciclica ha valore se la sua rilevazione causale produce, su almeno un ciclo di mercato completo, un numero di segnali sufficiente per l'inferenza e stabile nel tempo. Il dataset copre il toro 2020–2021, l'orso 2022, la ripresa 2023–2024 e la fase 2025–2026.

Tre evidenze sostengono il valore informativo, e vanno pesate per quello che sono.

1. **Coerenza della tassonomia fra mercati indipendenti** (la codifica tripartita qui sopra): lo stesso criterio produce ripartizioni confrontabili su serie con microstruttura diversa. È un indizio che la struttura rilevata non sia un artefatto di calibrazione, non una dimostrazione.
2. **Un segnale ciclico a valore forte nello stesso periodo:** la traslazione al 74–77% con uplift di 18–25 punti sulla base-rate ([Tabella 7](#)), replicata su H4 e H1 su campioni di un ordine di grandezza più grandi.
3. **Stabilità temporale:** le etichette confermate non mutano al progredire del tempo (il cancello di non-repainting qui sopra). È requisito necessario perché un segnale abbia valore decisionale.

Che cosa questa sezione non dimostra. Non dimostra redditività: nessun numero qui è il risultato di una strategia operativa. E il valore predittivo diretto delle lingue, misurato con un disegno non circolare, **non c'è** (§ 6.10.1). La formulazione corretta è: la rilevazione è **densa, stabile e coerente fra mercati**, e questo la rende uno strumento di contesto; ciò che ha valore direzionale misurato è la traslazione, non l'etichetta.

7. Ambito, limiti e riproducibilità

7.1. Ciò che è dimostrato e ciò che non lo è

In ambito — non-repainting dimostrato. Le lingue di Bayer codificate, su cui poggiano le affermazioni relative alla rilevazione, sono causali e verificate non-repainting: cinque combinazioni asset/timeframe su cinque, walk-forward su pipeline reale, determinismo verificato. Le etichette storiche confermate non cambiano; l'unica entità mobile è l'ultima lingua a finestra aperta. Anche il ricampionamento multi-timeframe del dato atomico è certificato da test automatico.

Fuori ambito — non oggetto delle affermazioni. Il sistema completo comprende componenti di composito ciclico — ancoraggio strutturale, discovery spettrale a finestra, fusione dei segnali — che non sono oggetto delle affermazioni di questo documento. Eventuali dubbi residui di repainting su tali componenti non

toccano i risultati qui presentati, che poggiano sulle lingue confermate e sulle statistiche cicliche derivate dal rilevatore causale. Restano fuori ambito per onestà metodologica: un'affermazione si fa solo dove la proprietà è dimostrata.

7.2. Limiti

1. **Dipendenza dal rilevatore.** Durate, nesting e coda dei brevi dipendono dalla definizione operativa di ciclo. La convergenza su tre mercati indipendenti è un presidio contro l'artefatto di *calibrazione*, **non** contro l'errore di *misura*: una misura sbagliata converge ovunque la si applichi, perché le serie passano dalla stessa procedura. La convergenza fra mercati è un test di **generalità** del fenomeno, mai di correttezza dello strumento, e usarla come secondo è uno dei modi più facili di darsi ragione da soli.
2. **Il righello.** L'impianto storico misura sul timeframe giornaliero, che non risolve i livelli più brevi della catena: i numeri di T, T+1 e T+2 misurati li descrivono un oggetto diverso da quello che descrivono sul timeframe che li risolve (§ 5.6). Le misure di **polarità** sono robuste al cambio di righello; quelle di **durata e dispersione** no. Il documento tiene entrambe le letture e dichiara quale sta usando.
3. **Confronti multipli.** Il lavoro conduce molti test. Sono da leggere come evidenza convergente cross-asset più che come scoperte isolate, e dove un solo mercato su tre rigetta, questo documento lo dichiara invece di riportarlo come conferma.
4. **Campione ai livelli molto alti.** T+7 e oltre non sono misurabili con alcun dataset esistente: servono fra 93 e 143 anni di sedute (§ 6.11). Non è più «campione insufficiente», è un fabbisogno quantificato.
5. **Il catalogo per classe di asset.** Tutte le misure di storia lunga e la prova su 37 strumenti usano il catalogo dei periodi nominali tarato su criptovalute. Un catalogo che si tari sull'asset non esiste ancora, ed è il limite che pesa di più sui risultati negativi allargati.
6. **La causa della coda corta.** Il § 6.3.1 stabilisce che la coda è del rilevatore, non che si sappia perché. L'ipotesi meccanica — la coda nasce dallo scarto fra periodo nominale e durata vera — è testabile: un catalogo tarato dovrebbe ridurla se è quella la causa.
7. **Il non-repainting intraday.** Il gate copre cinque combinazioni; una verifica più densa mostra che la tenuta non è uniforme su tutte le combinazioni intraday. Nessun risultato di questo documento vi poggia, e resta dichiarato come lavoro aperto.
8. **Asimmetria fra conferme e confutazioni.** Una confutazione su un mercato falsifica un'affermazione universale; una conferma vale per quel mercato, in quel periodo, con quel rilevatore. Le due colonne della tavola dei verdetti non hanno lo stesso peso, ed è la ragione per cui le conferme sono formulate con più cautela.
9. **Lo strumento non viene redistribuito, e questo limita la replicabilità.** Le soglie di calibrazione, i valori dei parametri e la procedura di taratura del rilevatore ciclico appartengono a un prodotto commerciale e non sono pubblicati. La conseguenza per la verificabilità è precisa, e la si dichiara qui invece di lasciarla in una nota. I risultati sono **verificabili**: ogni verdetto di questo documento viene con il suo file di evidenza, che porta la procedura, il dataset con l'impronta dei prezzi, l'ipotesi nulla **con il suo valore misurato**, le soglie fissate prima dell'esecuzione e i numeri cella per cella, così che un terzo possa ricontrollare ogni numero e contestare ogni verdetto. I risultati **non sono rigenerabili da zero** da un terzo, perché lo strumento che li ha prodotti non viene consegnato. È uno scambio voluto — tutela del prodotto contro replicabilità piena — e non una lacuna del disegno sperimentale; è la stessa rinuncia già dichiarata nel § 7.3, nominata qui come limite del lavoro perché il lettore possa pensarla come tale.

7.3. Riproducibilità e verifica indipendente

Sono distinti due piani.

Verifica dei risultati (fornita). Il materiale di supporto contiene, **per ogni verdetto dichiarato in questo documento**, il file di evidenza da cui proviene: procedura, dataset con l'impronta dei prezzi, **ipotesi nulla con il suo valore misurato**, soglie fissate prima dell'esecuzione, e i numeri cella per cella. Il manifesto allegato ne porta le impronte SHA-256.

La selezione non è discrezionale: il bundle è costruito **dalla stessa tabella** con cui si verifica che i verdetti in pagina coincidano con quelli delle carte di misura, quindi una riga aggiunta al documento aggiunge un file al materiale, e un verdetto che cambia fa fallire il controllo prima della consegna.

Con quel materiale un terzo può ricontrollare ogni numero di questo documento senza accedere al codice proprietario — ed è il criterio che viene **prima** della tutela commerciale: se il lettore non può rieseguire il test sull'output consegnato e contestare il verdetto, non è un lavoro scientifico.

Ripetizione del protocollo (replicabile). La metodologia è descrivibile in modo eseguibile su dati propri: determinismo con semi fissi e dipendenze bloccate; walk-forward con criterio di prefix invariance; validazione esterna cross-asset; confronto con surrogati a fase randomizzata e AAFT per **ogni** affermazione di forma «questa regolarità esiste»; marginale accanto a ogni tasso condizionale; controllo di potenza prima di leggere un esito negativo; e ogni livello misurato dove si risolve. Un analista può applicare la finestra ciclica attesa (§ 3.2), il criterio causale a due passi (§ 3.3) e la conferma a finestra chiusa (§ 3.4) ai propri dati con la propria calibrazione, ottenendo un sistema concettualmente equivalente.

Restano riservati soltanto i valori di calibrazione — tolleranza della finestra, percentile di brevità, soglie del criterio di prezzo, banda della categoria intermedia, pesi della fusione — e il codice.

8. Confutazioni, conferme e sostituti operativi

Affermazione	Fonte	Test	Verdetto
La scala dei livelli è completa	tradizione	minimi propri e overlap bidirezionale, contro il suo nullo	FALSIFICATA
Nesting armonico fisso a due	Hurst	rapporti fra durate mediane e strutture osservate	FALSIFICATO
Armonicità dei periodi	Hurst	distanza dei picchi dai pioli di una scala, con potenza verificata; sull'intraday regge solo contro una nulla troppo larga	FALSIFICATA
Comunanza ciclica	Hurst	campo di variazione fra mercati, su crypto e su un secolo di azionario	CONFERMATA ai livelli brevi e alti
Sincronicità dei minimi	Hurst	sovrapposizione fra livelli	CONFERMATA
Traslazione come predittore di polarità	Hurst	posizione del massimo contro polarità netta del ciclo, con base-rate	CONFERMATA : descrive il ciclo in chiusura, non il successivo
Proporzionalità ampiezza-durata	Hurst	rango e pendenza log-log, contro surrogati	PARZIALE : la pendenza sì, il rango no
Asimmetria bull/bear della durata	questo lavoro	rapporto fra durate rialziste e ribassiste, contro surrogati	NON PROVATA : esce sui livelli brevi, contro una nulla reversibile che non la sa escludere
Durata pari al periodo nominale	tutte e tre	mediana contro nominale, su D, W, H4 e H1	FALSIFICATA
Lo scarto dal nominale è un fatto di mercato	questo lavoro	mediana/nominale contro nube di serie senza cicli	FALSIFICATA : dentro la nube in 28 celle su 37
La posizione del massimo è un fatto di mercato	questo lavoro	posizione mediana contro nulla misurata, due generatori	FALSIFICATA : dentro la nube in 14 coppie su 15
La struttura ciclica si distingue dal caso	questo lavoro	cinque statistiche, 37 strumenti, 182 coppie, contro AAFT	FALSIFICATA : nessuna cella significativa
Stazionarietà locale	Ehlers	deriva del periodo dominante contro surrogati a spettro fermo	NON PROVATA , con test proprio
Bimodalità delle durate	questo lavoro	dip di Hartigan	FALSIFICATA
Coda dei brevi come fatto di mercato	questo lavoro	quota osservata contro nube di surrogati	FALSIFICATA : 1 combinazione su 33 sulla catena
Regime come driver della coda	questo lavoro	regime esogeno e causale, in terzi	NON PROVATO
Salto di livello ternario	questo lavoro	figli di padri binari contro figli di padri ternari	FALSIFICATO
Mean-reversion sequenziale	inferenza da Hurst	autocorrelazione di ordine uno	FALSIFICATA nella forma forte

Affermazione	Fonte	Test	Verdetto
Ciclo inverso operativo	Elico64	forma logica della definizione	TAUTOLOGICO
Soglia canonica di brevità	Bayer, via Elico64	applicazione alla distribuzione reale	FALSIFICATA
Swing incondizionato	Elico64	tasso condizionale contro la sua marginale, Fisher esatto	TAUTOLOGICO : vero in tutti i 191 cicli padre del giornaliero; sulla catena lo scarto non arriva mai a 0,10
Violazione del massimo precedente (D>P)	Elico64	tasso per ciclo padre, con marginale e base rate	CONFERMATA su 3 mercati su 3
Rottura del minimo di partenza	Elico64	forma senza inclusione, contro il gruppo che non rompe	CONFERMATA 3 su 3
Rottura del minimo del ciclo precedente	Elico64	idem	CONFERMATA 3 su 3
Tenuta del massimo (Hold)	Elico64	frequenza contro base-rate	FALSIFICATA : è la marginale di D>P
Valore predittivo delle lingue	questo lavoro	rendimento forward dopo intervallo cieco, contro controlli	NON PROVATO
Sequenze di polarità ammesse	Elico64	tasso contro tre nulle, una delle quali conserva la posizione	NON PROVATA : effetto di posizione
Il contesto riduce la varianza decisionale	premessa § 1	riduzione dell'IQR contro permutazione del contesto, su nove righelli	CONFERMATA su 7 righelli su 9 — ma il contesto è retrospettivo
L'alfabeto strutturale è {2,3}, il quaternario è un binario di binari	questo lavoro	profondità del minimo che separerebbe i due binari, con controllo di potenza su sintetico	NON PROVATA : il rilevatore non restituisce la struttura impiantata
La soglia dell'80% di overlap è informativa	questo lavoro	overlap osservato contro il nullo dello stesso rilevatore, ogni coppia sul righello che ne possiede i membri	NON PROVATA : copertura completa — 11 coppie su 11, 33 righe sui tre mercati — e margine in tutte; il blocco è il margine peggiore, 0,096 su T-7/T-6 di M1, sotto lo 0,10 dichiarato
La causa della discordanza T+3 è il filtro di durata	questo lavoro	tre letture a filtro spento e acceso, sul righello che risolve il livello	FALSIFICATA — su H4 il filtro non taglia nulla
Misurare un livello sul timeframe che lo risolve cambia i numeri	questo lavoro	quattro statistiche, Daily contro proprietario, due controlli	CONFERMATA : 12 coppie discordi su 12
Il ciclo giornaliero si separa dal metà-giornaliero	questo lavoro	regola di validità sulla coppia, tre risoluzioni	CONFERMATA : rapporto 2,34–2,38 dove è risolto
Le quattro affermazioni del Daily reggono fuori dal Daily	questo lavoro	rimisura su W, H4 e H1 con le stesse soglie	CONFERMATA 3 su 4 — A1 cade dove il righello risolve
Il campione ai livelli alti è un limite invalicabile	questo lavoro	conteggio dei cicli su sei indici a storia lunga	FALSIFICATA per T+5, che diventa misurabile
Il minimo ciclico non viene violato (livello statico)	Elico64	violazioni sui minimi veri contro minimi falsi di pari durata	CONFERMATA : 67,0% contro 91,7%, 9 contesti su 11

Affermazione	Fonte	Test	Verdetto
Il minimo ciclico non viene violato (trend line ascendente)	Elico64	stesso test, stesso gruppo di controllo	FALSIFICATA : 46,3% contro 44,9%, 0 contesti su 11
La violazione di un livello anticipa la conferma del minimo	questo lavoro	sette dichiaratori, con tasso di base e placebo di forma	FALSIFICATA : 0 contesti su 21

Tab. 12: Quadro dei verdeti. Le righe di fonte «questo lavoro» sono affermazioni dell'autore, e vanno lette con la colonna «Fonte» accanto a ogni riga; il sostituto operativo di ciascuna, dove esiste, è nella sezione che la tratta.

Il conto, riga per riga sulla tavola: trentotto affermazioni, di cui diciassette falsificate e due tautologiche — diciannove confutate in tutto —, **undici confermate, sette non provate e una parziale**. Il conto non è scritto a mano: lo calcola un controllo che legge la tavola, così che una riga aggiunta al documento non possa lasciarlo indietro.

8.1. Sul ciclo inverso

Il ciclo inverso è il contributo distintivo di Elico64 al filone italiano: ogni ciclo ammetterebbe una lettura simmetrica, «indice» e «inverso». La confutazione qui è **epistemologica, non empirica**: la definizione operativa «ciclo che chiude ribassista è inverso» è tautologica, perché ogni ciclo ribassista è per definizione inverso, e l'asserzione «l'inverso predice il ribasso» non è falsificabile ex ante — l'etichetta si assegna solo a posteriori. Non si nega il fenomeno che la nozione tenta di descrivere: si osserva che la sua formulazione non è testabile. La traslazione ciclica ha invece formulazione quantitativa, è falsificabile e raggiunge il 74–77% di accuratezza con 18–25 punti di uplift (§ 6.6): sostituisce l'intuizione senza abbandonarla.

8.2. Sulle assunzioni di Hurst

Il bilancio è misto, e va letto come tale. Cadono l'armonicità come rapporto fisso e la scala completa dei livelli. Reggono la comunanza ciclica — ai livelli brevi sulle criptovalute e ai livelli alti su un secolo di azionario — e la sincronicità dei minimi. La proporzionalità regge in forma di pendenza e cade in forma di correlazione di rango. È notevole, per un impianto costruito su dati azionari degli anni Sessanta e applicato qui a un mercato che non esisteva, quanta parte resista: le affermazioni di Hurst sono le più esposte perché sono le più precise, ed è esattamente questo a renderle utili.

9. Conclusioni

L'analisi ciclica dei mercati è sopravvissuta per oltre cinquant'anni come insieme di tradizioni dottrinali. Questo lavoro mostra che può essere convertita in programma di ricerca applicando un protocollo popperiano a un dataset multi-asset, con metodologia causale e deterministica verificata da test automatici, e con quattro requisiti che l'analisi ciclica non si è quasi mai posta: il valore nullo di ogni misura va **misurato**; ogni tasso condizionale va accompagnato dalla propria **marginale**; un esito negativo vale solo se il test avrebbe il **potere** di vedere il fenomeno; e ogni livello va misurato sul **timeframe che lo risolve**.

Il bilancio è netto e va letto con l'asimmetria del § 7.2. Fra le affermazioni **confutate** ci sono la scala completa dei livelli, il nesting armonico fisso, l'armonicità dei periodi sul giornaliero — con un test la cui potenza è verificata —, la durata pari al periodo nominale, la mean-reversion sequenziale, il salto di livello ternario, la soglia canonica di brevità, la tenuta del massimo e, **per tautologia**, due affermazioni della scuola italiana: il ciclo inverso, tautologico nella definizione, e lo swing incondizionato, tautologico nella condizione di successo. Sono **confermate**, con convergenza cross-asset, la comunanza ciclica — ai livelli brevi sulle criptovalute e ai livelli alti su un secolo di azionario —, la sincronicità dei minimi, la traslazione come predittore di polarità e, nella forma che elimina l'inclusione insiemistica, i tre segnali di violazione dei top ciclici. Restano **non provate** la stazionarietà locale, l'associazione fra durata e regime, l'asimmetria bull/bear, il valore predittivo delle lingue, le sequenze di polarità ammesse, la soglia dell'ottanta per cento di overlap e il quaternario come binario di binari. La premessa del § 1 sulla riduzione della varianza è **confermata su sette righelli su nove**, ma con il contesto misurato a posteriori: dice che la posizione dentro il ciclo padre porta informazione sulla dispersione dell'esito, non che quell'informazione sia disponibile a chi decide.

Delle regolarità che questo lavoro potrebbe rivendicare come non documentate dalle scuole precedenti ne restano **due**, e nessuna delle due riguarda il mercato: l'**assenza di due gradini** della scala nominale — ora replicata su un secolo di azionario — e la **convergenza cross-asset** delle durate, che si estende ai livelli alti appena le si dà campione. Se ne aggiungono due di natura diversa, perché riguardano lo **strumento** e non il

mercato: la forma quasi obbligata del ciclo padre sotto phasing, e la coda corta come proprietà della procedura.

Le candidate cadute sono più istruttive di quelle rimaste. La **bimodalità** delle durate — che avrebbe implicato due popolazioni di cicli — non c'è. La **coda dei cicli brevi** c'è, è stabile fra mercati e, dentro ogni righetto, cresce con il livello, e **non è del mercato**: sulla catena dei righeggi ne esce dalla nube dei surrogati una combinazione asset/livello su trentatré, e per il resto l'osservato sta sopra o sotto la mediana del sintetico senza un verso. Lo stesso vale per lo scarto dal periodo nominale e per la posizione del massimo dentro il ciclo. Erano i risultati che avevamo più interesse a tenere, e sono caduti per la stessa ragione per cui il protocollo esiste: **il valore nullo si misura**. E per la ragione gemella **l'assioma da cui il lavoro era partito** resta aperto: un esito negativo vale solo se il test avrebbe potuto vedere il fenomeno, e questo rilevatore non riconosce un binario di binari nemmeno dove è impiantato.

Se c'è un risultato che questo lavoro consegna, non è un indicatore e non è una regolarità. È una procedura, ed è abbastanza semplice da stare in quattro righe:

Un tasso senza la sua marginale non è evidenza. Una statistica senza la sua nulla misurata non è evidenza. Una nulla che non conserva la geometria dello strumento misura lo strumento; una che conserva ciò di cui si discute non misura niente. E un esito negativo non vale finché non si è mostrato che il test riconoscerebbe il fenomeno se ci fosse.

A cui va aggiunta una quinta riga, che è la più recente e forse la più costosa: **un livello misurato dove non si risolve è un altro livello**, e il nome che porta non lo dice.

Applicata al proprio lavoro, questa procedura costa. Il costo è il prezzo del fatto che, alla fine, ciò che resta in piedi resta in piedi davvero — ed è l'unica cosa che un lavoro di ricerca può onestamente promettere. Che di quello che resta in piedi quasi niente sia farina del proprio sacco è il modo in cui un programma di ricerca dichiara di aver funzionato, non di aver fallito.

La direzione naturale di sviluppo è già indicata dai limiti: un catalogo dei periodi tarato per classe di asset, la misura sistematica di ogni livello sul timeframe che lo risolve — i minimi dei cinque livelli infra-giornalieri esistono ora (§ 5.6.1), ma le statistiche che questo documento riporta per quei livelli restano quelle lette sul giornaliero — e un test del valore condizionato con un segnale a grana più fine di quello strutturale.

Riferimenti

- Hurst, J. M. (1970). *The Profit Magic of Stock Transaction Timing*. Prentice-Hall.
- Hurst, J. M. (1973–1976). *Cyclite 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. e Timmermann, A. (2004). «Duration Dependence in Stock Prices». *Journal of Business & Economic Statistics*, 22(3).
- Claessens, S., Kose, M. A. e 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. e Hartigan, P. M. (1985). «The Dip Test of Unimodality». *Annals of Statistics*, 13(1), 70–84.
- Theiler, J., Eubank, S., Longtin, A., Galdrikian, B. e Farmer, J. D. (1992). «Testing for Nonlinearity in Time Series: the Method of Surrogate Data». *Physica D*, 58(1–4), 77–94.
- Schreiber, T. e 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*. Materiale didattico. Autore storico del battleplan come strumento visuale di contesto ciclico, e della numerazione dei livelli come scala relativa: è da qui che la struttura ciclica diventa un'unità di misura del tempo (filone italiano, fase 1).
- Sartorelli, E. *Analisi Ciclica e Gann sul Bitcoin*. Estensione al mercato crypto e all'integrazione con Gann, condotta con impostazione ingegneristica: definizioni esplicite, procedura ripetibile e regole che non cambiano quando cambia il mercato. È il contributo del filone italiano su cui questo lavoro ha potuto lavorare meglio, perché è quello formulato in modo abbastanza preciso da poter essere messo alla prova.
- Elico64. *Appunti di Analisi Ciclica — Basi 1-10*. Fonte primaria del ciclo inverso, delle regole di polarità, delle definizioni di swing condizionato e incondizionato e della violazione dei top ciclici (filone italiano, fase 2).
- Marini. Formalizzazione didattica della scuola italiana di analisi ciclica; riprende e ridiffonde il corpus Elico64.

19. Popper, K. R. (1959). *The Logic of Scientific Discovery*. Hutchinson, London. (Traduzione inglese ampliata di *Logik der Forschung*, Vienna, 1934; edizione Routledge Classics.)
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.