

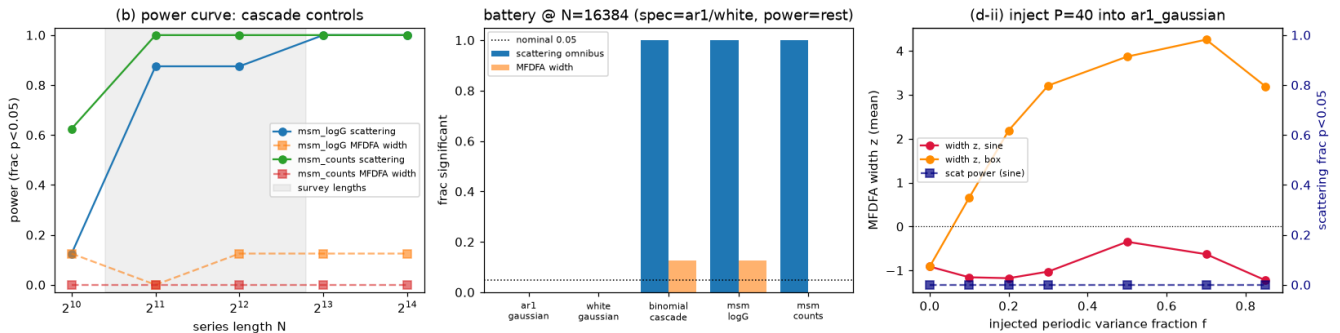
The width and the null

The measured power and specificity of the standard multifractality test, a mechanism for its failure, and a statistic that prices its own failure modes

Evan Tabak Atlas (ORCID [0009-0007-7374-2338](https://orcid.org/0009-0007-7374-2338))

Draft — methods preprint. Target venue: *arXiv first (physics.data-an / astro-ph.IM)*, then *Chaos or Physical Review E*. Every quantitative claim traces to a named, immutable results record: in this repository ([out/v22/](#) the scattering battery, [out/v7/](#) the width null audit, [out/v21/](#) the event-time surrogate honesty layer, [out/v33/](#) the survey-length specificity control, [out/v34/](#) the σ -parameterisation of the power claim), in *Fenosoma's ledger* (record v3, the same-day physiological replication), or in the *Fenopan meta-repo* (records v5/v6, the two in-house limits on the scattering omnibus; *CONSTITUTION.md* §6.1/§6.2/§10). Figures regenerate deterministically from the committed records: [out/v22/p34_scattering_battery.png](#)

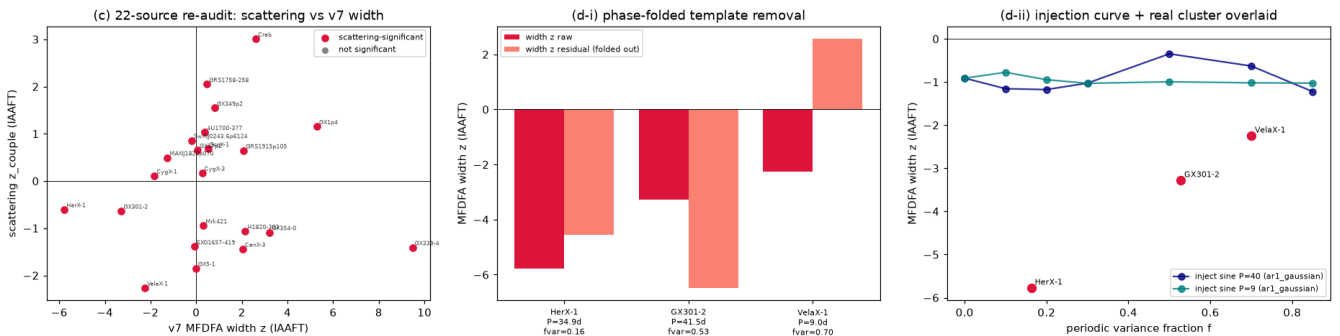
Record v22 (P-3.4): scattering-moment null audit -- anchor battery (committed before source verdicts)



[out/v22/p34_scattering_battery.png](#)

(specificity/power/length curves) and [out/v22/p34_scattering_audit.png](#)

Record v22 (P-3.4): the 22-source re-audit and the negative-z diagnosis



[out/v22/p34_scattering_audit.png](#)

(the 22-source re-audit and the negative-z diagnosis).

Priority note (updated 2026-08-31; this note previously read “no DOI and no priority date”): the Zenodo deposit is published — [10.5281/zenodo.22180477](https://zenodo.org/record/22180477) went live 2026-08-31 and establishes the priority date. The *t_{MF}* robustness line (Mangalam, Likens & Kelty-Stephen 2025) and one-command IAAFT attribution toolkits (MF-toolkit, *arXiv:2604.16257*, April 2026) remain active. *arXiv* posting is an owner action (endorsement).

Abstract

The field-standard test for multifractality asks whether the width of a series' singularity spectrum — measured by multifractal detrended fluctuation analysis (MFDFA), the $h(-4) - h(+4)$ or $\Delta\alpha$ width — exceeds the width of iteratively-amplitude-adjusted-Fourier-transform (IAAFT) surrogates that share its power spectrum and marginal distribution. A width that does not exceed the surrogate band is routinely read as evidence that the process is monofractal, linear, or “not a cascade.” **We show, on a pre-registered known-truth battery, that this test has near-zero measured detection power on its own canonical positive control — the binomial multiplicative cascade.** In record v22 the MFDFA width flags multifractality on 0.125 of binomial-cascade realisations (and 0.00 under the earlier v7 protocol), where a pre-registered wavelet-scattering-moment omnibus flags 1.00 on the identical IAAFT nulls at matched specificity (0.00 on AR(1) and white-noise monofractal controls at the battery's gap-free length; 0.045 / 0.036 — at the 2/51 rank budget, not elevated — when those same controls are pushed through the real sources' own concatenate-and-censor masks at survey lengths, record v33). The gap was replicated the same day, byte-identical nulls, on physiological data (Fenosoma record v3).

The failure is not a bug but an **incoherent pairing** of statistic and null. The width is a second-order, time-symmetric functional; IAAFT hands back the marginal and spectrum that carry the width; and the structure IAAFT actually destroys — fourth-order cross-scale coupling and time-asymmetry — is precisely what the width cannot measure. The test's sensitivity and the null's destructiveness are orthogonal. The consequence is a **pincer**: the width reads too *high* when nothing multiplicative is there (finite size, trends, isolated singularities) and too *low* when something is (marginal-encoded, fourth-order, time-asymmetric cascade structure). A statistic biased in both directions depending on the truth is not a measurement instrument.

We state precisely what this does and does not overturn. The **converse inference** — “width did not exceed its IAAFT null, therefore linear / not a cascade” — dies as uninformative. **Likelihood-based cascade results and width-as-biomarker classifications survive untouched**; positive width-vs-IAAFT rejections become *under-claims*. We engage the live t_{MF} robustness counterclaim directly, reading its own reported split (51.31% positive, 27.07% significant-negative, over 7,000 generations pooled across additive and multiplicative cascade types) as a two-sided instability that a one-sided literature reads one-sidedly, and citing the same group's independent, same-month derivation of the low jaw's mechanism (Mangalam & Kelty-Stephen, PRE 114, 024215) as allied prior art. Finally — because a paper whose demand of the field is “measure your statistic's power and specificity before you claim” cannot exempt its own replacement statistic — we price the scattering omnibus's two independently-measured limits (a genuine high-side excess on isolated singularities; a from-scratch specificity that requires scale-separated leverage) and declare its operating envelope. The proposed contribution is a **reporting standard**, not a new hammer: a measured specificity-and-power certificate on the canonical control, committed before the data, travelling with every multifractality verdict.

Neither jaw of the pincer is new, and we say so on page one: the low side — that cascade-driven series can have *narrower* spectra than their surrogates — was published inside the target literature by Lee & Kelty-Stephen (2017), and the high side — that isolated singularities artefactually broaden spectra — by Oświecimka et al. (2020). **The two jaws were known separately; what was missing was the price tag.** Our contribution is the **quantification**: a measured power/specificity certificate at a matched operating point on the canonical control, and a replacement statistic whose own failure modes ship in the same certificate. That certificate is now a **surface rather than a point**: across a stochastic-multiplier σ -axis the width's detection rate runs from 0.042 to 0.875 while the omnibus holds 0.958 – 1.000 (record v34), so what voids the converse inference is not a uniformly dead test but one whose power swings twenty-fold with a property of the generator no analyst observes.

1. Introduction

1.1 One test, everywhere

Across physics, astrophysics, physiology, finance, and hydrology, a single recipe has become the default test of whether a time series is “genuinely multifractal” as opposed to trivially wide-spectrumed: compute the singularity spectrum by MF DFA (Kantelhardt et al. 2002), summarise it by a width ($\Delta\alpha$), or the generalised-Hurst contrast $h(q_-) - h(q_+)$, and compare that width to an ensemble of surrogates that preserve the linear structure — usually IAAFT surrogates (Schreiber & Schmitz 1996), which iteratively match both the power spectrum and the amplitude distribution of the original. If the original width exceeds the surrogate band at $p < 0.05$, the multifractality is declared “nonlinear” or “genuine”; if it does not, the series is declared monofractal, linear-with-fat-tails, or “not a cascade.” The comparison is frequently distilled into a one-sample statistic, the multifractal-nonlinearity t-statistic t_{MF} (§6).

This paper is about the second half of that recipe — the null-not-exceeded read. Our claim is narrow, quantitative, and, we argue, load-bearing for a large literature: **as a detector of the canonical positive control, the width-vs-IAAFT test has almost no power, and it fails asymmetrically, so a null-not-exceeded result carries no information about whether the process is a cascade.** The narrowness is deliberate and now measured on both sides: on a stochastic-multiplier cascade whose width is *not* marginal-encoded the same test recovers substantial power (§6, record v34). What kills the converse inference is that its power swings between those two regimes according to a property of the generator the analyst never sees.

1.2 Prior art, cited on page one

The qualitative core of this result is not new, and it was published *inside* the target literature. **Lee & Kelty-Stephen (2017)** showed that cascade-driven series can generate singularity spectra *narrower* than their IAAFT surrogates ($w_{orig} < w_{surr}$), and explicitly warned against reading a narrower-than-surrogate width as monofractal or linear, urging a bidirectional interpretation. That paper is *allied* with this measurement, not a target of it; framing the present result as a discovery would be ungenerous re-discovery. Our delta is the **quantification and the price tag**: a measured power/specificity certificate at a matched operating point on the canonical binomial cascade, plus a replacement statistic whose own failure modes are measured and shipped in the same certificate (§7–§8) rather than left for a referee to find.

The methodological demand is not new either — it is nineteen years old. Wendt & Abry (2007) built bootstrap multifractality hypothesis tests on wavelet leaders in six design declinations and assessed their significances, powers and p -values by Monte Carlo on synthetic processes. “Measure your statistic’s power before you claim” was therefore met, properly, for a competing statistic, in the signal-processing literature, in 2007. What was never done is the same exercise for the **MF DFA-width + IAAFT pairing that the applied literature actually runs** — and, as this paper reports, when it is finally done the answer is near-zero. That is the honest framing of our contribution, and it is stronger than a novelty claim on the demand: the standard exists, it is old, the applied default was simply never held to it.

The high jaw is published too, and by a different group. Oświęcimka et al. (2020) showed that signals harbouring only **isolated singularities** — with no cascade-like hierarchy — artefactually give rise to broad multifractal spectra under both DFA-family methods and wavelet leaders, and warned in those words of “a real risk of incorrectly inferring multifractality.” That is the same mechanism our own §7 admits about the scattering

omnibus, published six years earlier for the width. In the same tradition, Drożdż et al. (2023) argue that genuine multifractality requires long-range temporal correlations and that fat tails alone broaden the spectrum only when such correlations are present, with short heavy-tailed series a known trap for shuffle-based surrogate tests; Kluszczyński et al. (2025) carry the decomposition further. We cite all of this as allied, and it sharpens the framing of the pincer: **both jaws were already known separately — Lee & Kelty-Stephen (2017) on the low side, Oświęcimka et al. (2020) on the high side. What was missing was the price tag,** and a statistic that ships one.

The test under audit, cited in its canonical statement. The MFDEFA-width + surrogate recipe entered the behavioural and physiological literature through Kelty-Stephen, Lane, Bloomfield & Mangalam (2022/2023), the methods tutorial for “multifractal analysis and surrogate data production for resolving when multifractality entails nonlinear changes over time”; the same group’s mechanism-side work (Mangalam & Kelty-Stephen 2024a, PRE **109**, 064212; 2024b, *Physica A*, on additivity suppressing multifractal nonlinearity) is the line the §6 counterclaim sits in. Auditing a recipe without citing its canonical methods statement would be the exact failure mode this paper is about.

1.3 Why it matters now, and why it is bounded

Two developments make the calibration urgent. First, one-command IAAFT attribution pipelines with synthetic-validation suites are shipping — MF-toolkit (Mendez et al., arXiv:2604.16257, April 2026) validates against fBm, a **binomial multiplicative cascade**, heavy-tailed transforms and controlled crossovers, and ships an IAAFT module and a source-identification demonstration, yet reports **no detection rate, no power and no false-positive rate** anywhere: its source-identification result is qualitative. The number this paper measures was computable, with that toolkit’s own components, and uncomputed, as of April 2026. Second, a live “robustness” line (Mangalam, Likens & Kelty-Stephen 2025, §6) argues that t_{MF} is a robust cascade estimator, and its own reported numbers, read with their pool stated exactly, show a two-sided instability rather than robustness.

We are equally careful about what the result is *not*. It does not say multifractal analysis is worthless; it says the *converse inference* from a null result is void, while the biomarker use and the likelihood-based use survive (§5). And it does not claim the field is uniformly wrong: the existence test we measure (width vs IAAFT: is anything nonlinear there?) is not the same as the three-way *attribution* decomposition (original vs shuffle vs FT vs IAAFT: fat tails vs long-range correlation vs nonlinearity) that the largest cross-domain targets run; the power of *that* decomposition on the canonical cascade has been measured by nobody, including this program, and is deferred to a companion study (Appendix B). This paper’s stratum is the width-vs-IAAFT existence test and its one-sample t_{MF} distillate.

1.4 Contributions

1. **A mechanism (§3).** The incoherent-pairing account: why a second-order, time-symmetric statistic tested against a spectrum-and-marginal-preserving null has structurally near-zero power on a fourth-order, time-asymmetric, marginal-encoded cascade.
2. **The pincer (§4).** The two-sided bias, with the high side (finite size, trends, isolated singularities — Oświęcimka et al. 2020) and the low side (Lee & Kelty-Stephen 2017) both drawn from the published critique tradition, and this measurement placed beside them as the price tag neither carries. Note explicitly what is *not* claimed: neither jaw is our discovery, and the demand that a test report its power is Wendt & Abry’s (2007), met there for wavelet leaders and never met for the MFDEFA-width + IAAFT pairing (§1.2).

3. **The two-domain battery (§5).** Pre-registered specificity and power curves for the width and for a wavelet-scattering-moment omnibus, on IAAFT nulls, at survey-matched lengths; the same-day independent replication; the reconciliation of the two width-power columns.
4. **What dies, survives, strengthens (§5.4, §11).** The converse inference dies; the likelihood and biomarker uses survive; positive rejections become under-claims. (The scope-honesty of the 22-source re-audit is §5.4; the full dies/survives/strengthens treatment is §11.)
5. **Direct engagement with the t_MF counterclaim (§6),** reading its formula first, not paraphrased, its own reported split read with its pool stated exactly, the estimator mismatch named before a referee names it, and the group's same-month sublinear-multifractality mechanism cited as allied.
6. **The scattering statistic's own high side (§7)** — isolated singularities produce a genuine positive omnibus excess with no cascade present — stated before a referee states it, with the decomposition that bounds it.
7. **Reproducibility of the battery (§8):** an independent in-house reimplementations's two readings — the scale-separation strength and the protocol-sensitivity weakness — reported both ways.
8. **The branching-ratio rule (§9),** the allied §6.1 result for which this paper is the program's named vehicle, with its four independent in-house reproductions and a dated correction to an earlier prior-art claim.

The proposed reporting standard is collected in §10.

2. The statistics and the null

All widths and scattering vectors in this paper are measured on the same **log-rate, concatenated-good-ticks** series (the masked convention of record v7) so that “the same ruler” is applied to model, control, and data, and every claim names the **amplitude domain** on which it is measured (Standing Rule 8: altitude travels with the number — the statistics here constrain the *amplitude* distribution of scale-wise fluctuations, not event times, which are the domain of §9 and record v21).

The width. MFDFA (Kantelhardt et al. 2002) partitions the integrated series into windows at scale s , removes a local polynomial trend, and forms the q -th order fluctuation function

$$F_q(s) = \left(\frac{1}{2N_s} \sum_{v=1}^{2N_s} \left| F^{(2)}(v,s) \right|^{q/2} \right)^{1/q}, \quad F_q(s) \sim s^{h(q)},$$

from which the generalised Hurst exponents $h(q)$ and the singularity spectrum follow. The width statistic is the two-point contrast $h(q_-) - h(q_+)$ (here $h(-4) - h(+4)$, v7's pinned moments) or, equivalently, the Hölder-exponent support $\Delta\alpha$ of the spectrum. Two structural facts about this statistic drive the entire paper and are proved in §3: $F^{(2)}(v,s)$ is a *variance* (second-order in the fluctuations) and is *invariant under time reversal*.

The null. IAAFT (Schreiber & Schmitz 1996, reimplemented — Standing Rule 4, `fenocosm/surrogates.py`) alternately imposes the original series' amplitude distribution and its Fourier amplitude spectrum, converging to a surrogate that preserves *both* the marginal and the power spectrum (the full second-order autocorrelation) while randomising the Fourier phases. It is the standard “nonlinearity” null: what survives IAAFT is, by construction, everything a linear-Gaussian process with the same spectrum and the same (possibly heavy-tailed) marginal could produce.

The stronger statistic (record v22). The replacement is a pre-registered wavelet-scattering-moment vector (`fenocosm/scattering.py`; Bruna, Mallat, Bacry & Muzy 2015; Morel et al. 2025 — reimplemented, nothing vendored). On the same standardized log-rate series it computes, over an analytic Morlet bank (one wavelet per octave, $J = \min(8, \lfloor \log_2 N \rfloor - 3)$ octaves):

- **S1(j)** first-order scattering $\text{mean}|x*\psi_j|$ and the sparsity ratio $\rho(j) = (E|W_j|)^2/E|W_j|^2$ (Gaussian value $\pi/4 \approx 0.785$);
- **S2n(j₁,j₂)** normalized second-order scattering, $j_2 > j_1$;
- **C(j₁,j₂)** wavelet-modulus correlation $\text{corr}(|x*\psi_{\{j_1\}}|, |x*\psi_{\{j_2\}}|)$ — the cross-scale, **fourth-order** coupling IAAFT does not preserve;
- **τ(j)** a phase/modulus time-asymmetry coefficient, **odd under time reversal** (proved and gate-tested), hence ≈ 0 for any IAAFT/FT surrogate.

A single aggregate-excess rule scores the whole vector: per-coefficient z against the surrogate ensemble, omnibus $A(y) = \text{mean}_d z_d^2$ with **symmetric leave-one-out** normalisation (real and every surrogate each scored against the other pool members, so real is exchangeable with a surrogate under the null and the rank test is exactly calibrated), global $p = (1 + \#\{A_{\text{surr}} \geq A_{\text{real}}\})/(N+1)$. With $n_{\text{surr}} = 50$, $p = 0.020$ is the resolution floor (the real aggregate exceeds all 50 surrogates), and a directional summary z_{couple} (mean z over the C coefficients; positive = more cross-scale coupling than the null) carries the sign the omnibus does not. **The battery and this rule are pinned in `protocol.py` before any real source is read** (Standing Rule 2 / the anchors-before-claims discipline).

3. The mechanism: an incoherent pairing

The central claim of the mechanism is one sentence: **the width measures a class of structure that IAAFT preserves, and IAAFT destroys a class of structure the width cannot measure, so the pairing of this statistic with this null has a detection power set by their overlap — which is near zero for the canonical cascade.**

Make it precise in three steps.

(i) The width is a second-order, time-symmetric functional. The MFDFA spectrum is built entirely from $F^2(v, s)$, the detrended *variance* of the integrated series in window v at scale s . Re-weighting by the moment order q and tracking the scaling in s produces the $h(q)$ curve, but every input to that curve is a second-order quantity of the fluctuations — the width is a functional of the *distribution of local variances across scales*, not of any genuinely higher-order ($\geq$ fourth) cross-scale correlation except insofar as such correlation leaves a footprint in the marginal of those local variances. Moreover $F^2(v, s)$ is a sum of squared detrended residuals within a window and is therefore **invariant under reversing the series in time**. Consequently the width has *structurally zero* sensitivity to time-asymmetry: a process and its time-reverse have the same MFDFA width, exactly.

(ii) IAAFT returns exactly the width-bearing part of a cascade. A binomial multiplicative cascade is multifractal, but on a finite record the dominant carrier of its measured spectral *width* is its heavy-tailed marginal together with its linear correlation structure — both of which IAAFT reproduces by construction. This is visible directly in record v7's battery: the binomial cascade has by far the largest raw width (0.745) of any control, yet its IAAFT excess is only $z = +0.49$ (frac significant 0.00) — the surrogate, having matched the marginal and the spectrum, carries a width almost as large as the data's. The width the cascade *has* is handed straight back by

the surrogate. (This is exactly the caveat v7 flagged as an honest power ceiling: a null IAAFT excess is not proof of monofractality; the discriminating signature v7 could still see was the collapse from a large *shuffle* excess — marginal only — to a small IAAFT excess.)

The words “the dominant carrier ... is its heavy-tailed marginal” are a *premise*, not a decoration, and record v34 turns it into a measured dial. Hold the cascade construction fixed and randomise only the multiplier *magnitude* — moving along a stochastic-multiplier σ -axis on which the binomial is the two-point member — and the width-bearing structure leaves the marginal for the arrangement. The width’s power against IAAFT tracks that move, from $1/24$ at the marginal-encoded end to $21/24$ at the other (§6). Step (ii) holds exactly where its premise holds, which is where the canonical control sits.

(iii) What IAAFT destroys, the width cannot see. IAAFT randomises the Fourier phases. What that destroys is the higher-order phase organisation: the fourth-order cross-scale modulus coupling $C(j_1, j_2)$ and the time-asymmetry τ . These are precisely the coordinates on which a genuine multiplicative cascade differs from its linear-Gaussian shadow — and precisely the coordinates the width does not measure (by (i)). The information the null removes lands in the *null space of the statistic*. Sensitivity and destruction are orthogonal.

That orthogonality is the “incoherent pairing.” It is measurable. The scattering vector was built to occupy exactly the destroyed coordinates: the C coefficients are fourth-order in the series (a product of two wavelet moduli, each already a modulus of a linear filter), so IAAFT — a second-order null — does not preserve them, and τ is odd under time reversal, so it vanishes for any phase-randomised surrogate. On the binomial cascade the omnibus recovers full power with $z_{\text{couple}} \approx +13.8$ (§5), reading the coupling the width is blind to. **Measured, not assumed: v7’s power ceiling was a property of the width statistic, not of the IAAFT null** — the “null-audit power” question the founding program left open, resolved.

A note on scope of the mechanism (developed fully in §7): the same argument implies that IAAFT is the *wrong null for any process whose discriminating structure is fourth-order or time-asymmetric* — which includes not only cascades but frozen deterministic geometries. That symmetry is why the replacement statistic has a high side of its own, and why the honest contribution is a reporting standard rather than a new default statistic.

4. The pincer

The width fails in **both** directions, and which way it fails depends on the truth it is meant to detect. This two-sidedness is the headline.

The high jaw (width reads too high when nothing multiplicative is there). This side is the older critique tradition and is not our finding; we place it here for completeness because a two-sided bias is the point.

- *Finite size.* At small N the MFDFA width is positively biased: short, fat-tailed, or linearly-correlated series produce a non-zero width with no multiplicative structure (Zhou 2009; the finance attribution consensus). A raw width is never claim-bearing (Standing Rule 2 / canon F2).
- *Trends and smooth periodicity.* A smooth, high-variance-fraction modulation can drive the width-vs-IAAFT z either sign depending on the modulation’s *shape*. Record v22’s injection curve measures this: a **smooth (sinusoidal)** modulation drives the width- z **negative** (the regular modulation reads monofractal while its

phase-scrambled IAAFT surrogate carries more width), while a **sharp (box/flare)** modulation drives it **positive** (its harmonics add width to the data). The sign of a width anomaly can be an artifact of an observation-process modulation, not of the process.

- *Isolated singularities.* A single sharp, isolated feature (a step, a corner, an aperiodic deterministic turn) produces a genuine scaling excess against a phase-scrambling null with no cascade present. This is the sharpest form of the high jaw, and — because it afflicts scattering statistics too — we treat it in full in §7.

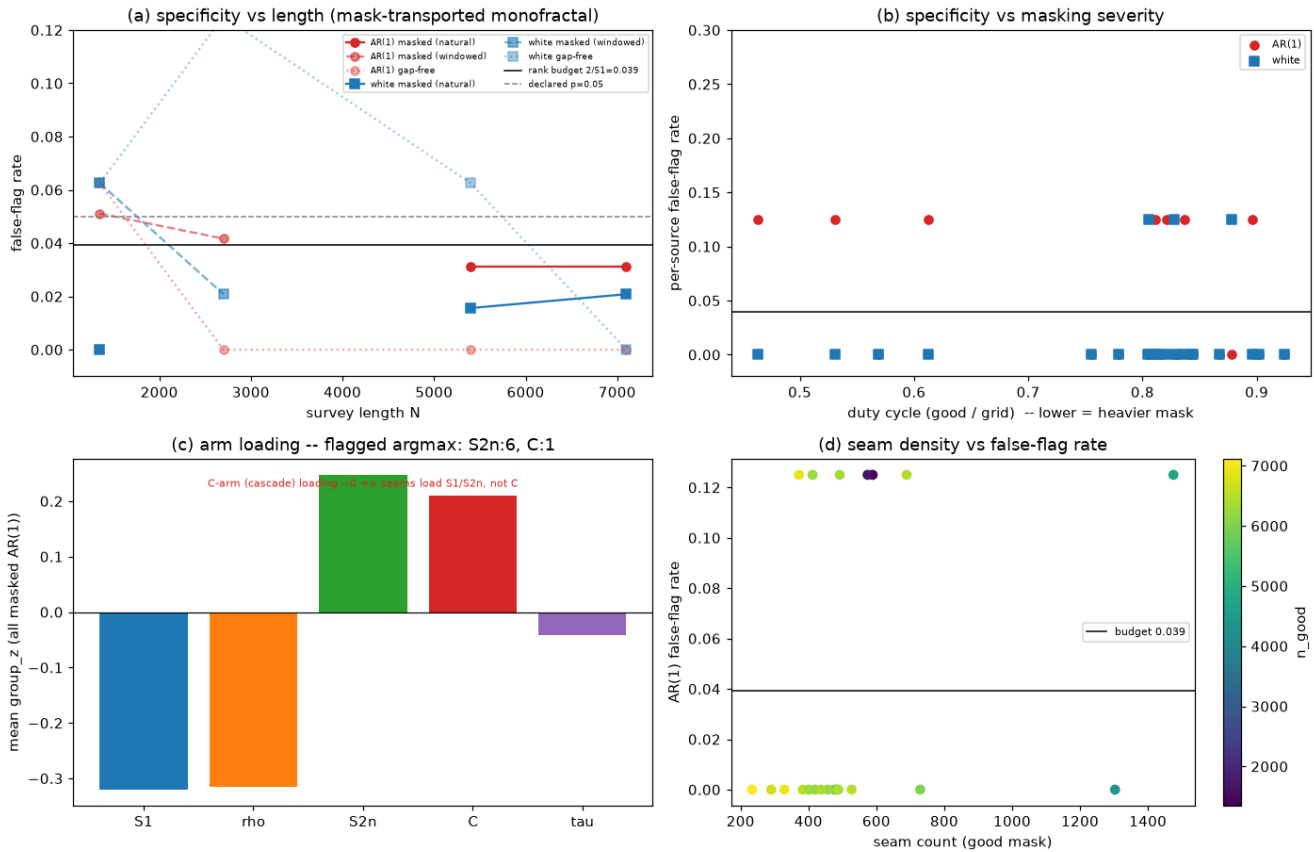
The low jaw (width reads too low when something multiplicative is there). This side is the present measurement (§3, §5). On the canonical binomial cascade the width flags multifractality on `0.125` of realisations (v22 protocol) or `0.00` (v7 protocol) — at 8 seeds, not distinguishable from zero power (§5.3). The cascade’s genuine multiplicative structure is marginal-encoded and fourth-order, and the width, tested against IAAFT, sees none of it.

A statistic that is positively biased on linear artifacts and negatively biased on genuine cascades does not have a “detection threshold” one can correct for: the same numerical width can arise from a finite-size linear series (false positive) or hide a genuine cascade (false negative). The width is a *feature*, not a *measurement of multifractality*. §5 shows what a statistic occupying the destroyed coordinates does on the same battery.

5. The two-domain battery

(Records: `out/v22/` — the scattering battery and the 22-source re-audit; `out/v7/` — the width null audit whose columns v22 quotes unchanged; `out/v33/` — the mask-transported monofractal specificity control (the 22/22’s survey-length, concatenate-and-censor specificity cell, §5.4). Figures: `p34_scattering_battery.png`, `p34_scattering_audit.png`, `p35_maskspec.png`

Record v33 (P-3.5): mask-transported monofractal specificity -- AR(1) 0.045 / white 0.036 (budget 0.039); C-arm 0.009; SPEC GREEN



p35_maskspec.png

. Independent replication: Fenosoma record v3.)

5.1 Specificity and power on known truth

The anchor battery uses known-truth controls, 8 seeds each, $N = 16384$, IAAFT null. The scattering columns are record v22; the width column is v7's, quoted unchanged under the append-only rule.

Table 1 — Specificity and power on the pre-registered control battery (record v22; `z_couple` is the directional cross-scale summary; the width column is v7's).

control	reads as	scattering frac sig	mean <code>z_couple</code>	v7 MFDEFA-width frac sig
AR(1) Gaussian (monofractal)	specificity	0.00	+0.6	0.00
white Gaussian (monofractal)	specificity	0.00	+0.0	0.00
binomial cascade (multifractal)	power	1.00	+13.8	0.00 ← the blind spot
<code>msm_logG</code> (cascade latent)	power	1.00	+3.5	0.38
<code>msm_counts</code> (Poisson-observed cascade)	power	1.00	+5.1	0.00

Two readings, both measured:

- **Specificity is clean and the calibration is exact.** The monofractal linear controls show no scattering excess (AR(1) 0.00, white 0.00; median omnibus $p \approx 0.32$, mildly conservative). The stronger statistic does not manufacture a cascade where a linear-Gaussian process with the same spectrum and

marginal suffices.

- **Power is full where the width had little or none.** All three cascade controls are detected on every seed, including the binomial cascade whose multifractality is marginal-encoded and invisible to the width under IAAFT ($z_{\text{couple}} \approx +13.8$). The single genuine-cascade control on which the width had any power at all, `msm_logG` , is the exception that proves the rule: its multifractality is *not* marginal-encoded, so the width caught 0.38 of it — and the scattering omnibus catches all of it. That “exception” is the whole mechanism in miniature, and §6 turns it into a controlled dial: hold the cascade construction fixed, randomise only the multiplier *magnitude*, and the width’s power against IAAFT rises with how much of the multifractality has left the marginal (record v34).

5.2 The power curve at survey-matched lengths

The failure is not a finite-size artifact of any single length. Across the whole survey band (the Swift/BAT sources this battery was built for are $\approx 1,350\text{--}7,100$ points), the width never acquires usable power while the scattering omnibus saturates by $N \approx 2,048\text{--}8,192$.

Table 2 — Power vs record length (record v22; frac significant).

N	msm_logG scattering	msm_logG width	msm_counts scattering	msm_counts width
1024	0.12	0.12	0.62	0.00
2048	0.88	0.00	1.00	0.00
4096	0.88	0.12	1.00	0.00
8192	1.00	0.12	1.00	0.00
16384	1.00	0.12	1.00	0.00

The width stays ≤ 0.12 at every length tested — at 8 seeds per cell, not distinguishable from zero power (§5.3) — for both genuine-cascade controls. There is no length at which “run more data” rescues the width-vs-IAAFT test on a Poisson-observed cascade.

5.3 The reconciliation of the two width-power columns

Two width-power figures appear in the v22 materials, and they are different objects; any citation must name which protocol its number comes from (FINDINGS reconciliation note, 2026-08-21, no record edited).

- The **v7-quoted** column (binomial cascade 0.00 , `msm_logG` 0.38) is computed under v7’s protocol, lengths, and surrogate budget, and is reproduced byte-for-byte in v22 so the width columns match v7 exactly.
- The **v22 freshly-computed** battery (`p34_scattering_battery.json`) reports `width_frac_sig` = 0.125 for both cascade controls under v22’s pre-registered protocol (8 seeds, $N = 16,384$, $n_{\text{surr}} = 50$) — a 1 flag in 8 seeds that is not distinguishable from zero power at this seed count.

The two quantities in that last clause must be kept apart, and an earlier draft of this section (and the corresponding row of Appendix A) ran them together. The decision rule is $\text{width_p_excess} = (1 + \#\{w_{\text{surr}} \geq w_{\text{obs}}\}) / (n_{\text{surr}} + 1) < 0.05$, so with $n_{\text{surr}} = 50$ **ranks 1–2 of 51 clear $p < 0.05$ ($1/51 = 0.0196$ and $2/51 \approx 0.0392$, $3/51 = 0.0588$ does not), and the per-realisation false-positive rate is therefore $2/51 \approx 0.0392$ — that is the floor. The quantity $1 - (49/51)^8 \approx 0.274$ is something else: it is the**

probability of seeing **at least one** false flag across 8 independent seeds, i.e. the binomial p -value of the observed $1/8$ under the null of zero power. So the correct statement is: the per-realisation false-positive rate is 0.0392 ; observing 1 flag in 8 seeds has binomial $p \approx 0.274$; 0.125 is therefore not distinguishable from zero power at this seed count.

Dated correction (2026-08-30). The 2026-08-25 correction paragraph this replaces was itself wrong by a factor of two: it read the rank budget as $1/51 = 0.0196$ and the dependent binomial as $1 - (50/51)^8 \approx 0.147$, while the repo's own code and records use the two-rank budget throughout — `p36_sigma.NOMINAL_BUDGET = 2/(n_surr + 1) = 0.0392`, written into `out/v34`'s `protocol.nominal_budget`, and quoted as $2/51 \approx 0.039$ by `out/v33`, by §5.4 below and by this paper's own abstract. Both numbers are restated above; the citation gate now binds $2/(n_surr + 1)$. The direction is safe — a *higher* floor makes the observed rates *less* distinguishable from zero power, so no conclusion moves — and it sharpens one: **the deeper panel's $1/24 = 0.042$ is then statistically indistinguishable from the floor itself** ($2/51 = 0.0392$; under that null, seeing at least one flag in 24 seeds has probability 0.62), so the width's measured rate on the canonical control is, to the resolution this design has, exactly the rate a statistic with *no* power would produce.

Both measurements support the same conclusion: **the width has no usable power on the canonical cascade controls**. The 0.00 and the 0.125 bracket the same substantive fact. A third figure for the same quantity appears in §6, where record `v34` re-runs the binomial control on a 24-seed panel at the identical operating point and reads $1/24 = 0.042$; that panel's leading eight seeds *are* these eight and reproduce this 0.125 exactly, so the deeper number is a tighter estimate of the same near-zero rate, not a fourth object. Any citation should still name its panel size along with its protocol.

5.4 The 22-source re-audit: a more sensitive instrument, and a broader question

Run on 22 public Swift/BAT light curves (same masks, freshly generated IAAFT + shuffle ensembles at the pinned seed), the width columns reproduce `v7` exactly (Her X-1 -5.77 , GX 339-4 $+9.48$, Vela X-1 -2.25). The stronger statistic detects far more:

- **22/22 sources are a significant scattering departure from the IAAFT null**, versus the width's 7/22 IAAFT-excess (both from the same run). The monofractal controls behind that claim read 0.00 at the battery's own $N = 16,384$ gap-free length, and **$0.045 / 0.036$ at these sources' survey lengths pushed through their own concatenate-and-censor masks** (record `v33`; at the $2/51$ rank budget, not elevated — see the scope note below, and read the 22/22 with it).
- **But the departure is diverse, not a wholesale multifractality verdict**. The dominant departing family is second-order intermittency $S2n$ (7 sources), time-asymmetry τ (6), and sparsity/secular change $S1$ (6); the **cross-scale multiplicative coupling C that a multifractal claim specifically needs dominates in only 3/22 sources**. So the omnibus answers a *broader* question (any higher-order structure) with far more power than the width; it does **not** recast the 7/22 width-collapse as “22/22 multifractal.”

Scope of the 22/22, stated as the standard this paper demands of others — and now measured (record `v33`).

The specificity controls behind that 0.00 (§5.1) are generated at $N = 16,384$ as continuous, gap-free, uncensored series. The 22 real sources run at $N \approx 1,350\text{--}7,100$ — $2.3\times$ to $12\times$ shorter — and reach the statistic through the `v7` preprocessing convention: presence-masked, **concatenated across observing gaps**, and one-sidedly censored at non-positive background-subtracted rates. Whether the omnibus's specificity survives *that* — the survey lengths and that preprocessing — is a distinct question from the $N = 16,384$ battery, and §7 of this paper is the reason to care: concatenation seams are isolated singularities, and isolated singularities are precisely

where the omnibus has a genuine high side. **So we measured it.** Record v33 ([out/v33/](#)) generates the *same* monofractal controls — AR(1) $\phi = 0.98$ and white Gaussian — on each source’s full daily grid and **pushes them through that source’s own presence-and-censor mask** (dropping exactly the days the v7 pipeline drops), producing a concatenated, seamed control of the source’s own length, then scores it by the v22 omnibus against its own IAAFT null ($n_surr = 50$, $SCAT_SEED = 0$ — unchanged, so the numbers are comparable to the battery’s 0.00). The result, at survey lengths $N \in \{1,350, 2,700, 5,400, 7,100\}$ and under the full mask-concatenate-and-censor pipeline:

- **Specificity holds.** With $n_surr = 50$ only ranks 1–2 of 51 clear $p < 0.05$, so the nominal false-flag budget is $2/51 \approx 0.039$. The mask-transported controls false-flag the omnibus at **$20/448 = 0.045$** (AR(1); binomial $p = 0.31$ against the budget — **not elevated**) and **$16/448 = 0.036$** (white; $p = 0.68$) — the survey-length, mask-transported analogue of the battery’s 0.00 , at the rank budget. The rate does **not** rise toward the shortest, most-heavily-concatenated lengths, nor with masking severity (the heaviest-masked duty-cycle tercile, ~760 seams, sits at 0.047 , no higher than the lightest); the one-sided censoring adds nothing over the observing gaps. **The Crab — the source flagged at the ensemble floor and carrying a 365-day concatenation gap — is 0/8 on both controls.**
- **The (budget-level) flags load $S2n$, not C — exactly the §7 isolated-singularity signature.** Of the natural-length AR(1) flags, $S2n$ (second-order intermittency) is dominant in 6 of 7 and C in 1; the cross-scale- C share is **$7/896$** pooled across both controls, sitting at its uniform-over-arms calibration baseline ($2/51 \div 5$, $p = 0.68$). **The seams do not manufacture a cross-scale-cascade reading** — so the C -dominant-in-3/22 sub-claim is seam-robust, and the diverse $S2n / \tau$ departure the omnibus reads on the real sources is not a preprocessing artifact.

The contrast with the real sources is the point: **all 22 real sources land on the two smallest achievable ranks** ($iaaft_scat_p \in \{0.020, 0.039\}$), including the most heavily concatenated — while the monofractal controls pushed through those *same* masks sit at the 0.045 budget. So the honest reading of the 22/22 is “**a significant, genuine departure from the IAAFT null under the v7 concatenation-and-censoring convention, whose monofractal specificity at survey lengths is measured and holds (record v33)**” — upgraded from the pre-registered caveat’s “not yet measured,” and *not* recast as a clean multifractality verdict (the departure is $S2n / \tau$, not the multiplicative cross-scale cascade). Nothing in §§5.1–5.3 — the power certificate, which is where this paper’s contribution lives — depended on it; the specificity of the 22/22 re-audit now does not either (measured 2026-08-26).

This is the crucial honesty of the result, and it cuts against over-reading our own instrument as hard as it cuts against the width: **v7’s conservative reading stands and is refined.** Most apparent multifractal *width* in these sources is linear-correlation artifact (the shuffle $\rightarrow$ IAAFT collapse; GRS 1915+105 falls from shuffle $z \approx +10.5$ to IAAFT $z \approx +2.1$); the residual structure the width could not see is real, but it is mostly time-irreversibility and intermittency, not the multiplicative cross-scale cascade. A statistic with power is not a licence to claim the thing it was not decomposed to isolate — a discipline the rest of this paper (§7) applies to the scattering omnibus itself.

5.5 The negative-z diagnosis (the sign the width gets wrong)

Record v22 also resolves the cluster of *negative* IAAFT excess v7 left undiagnosed in strongly-periodic HMXBs — the clearest single demonstration that the width’s *sign* is untrustworthy. Folding out Vela X-1’s 8.964-d eclipsing orbit **flips its width- z from -2.25 to $+2.56$** : its negative width was the smooth eclipsing periodicity (whose phase IAAFT scrambles, inflating the surrogate band), and beneath it lies genuine excess — Vela X-1

joins the excess group after de-orbiting. Her X-1's negative z is instead the detection-floor smoothness of its faint turn-off state (16% of kept points at $S/N < 1$; folding the 35-d cycle barely moves it, $-5.77 \rightarrow -4.56$), and GX 301-2's becomes *more* negative after de-flaring ($-3.28 \rightarrow -6.48$), an honestly open puzzle. Under the omnibus all three are a significant scattering departure ($p = 0.020$, the resolution floor): the negative width- z was never evidence of monofractality; it was the low-power two-point statistic's particular response to a smooth, high-variance-fraction modulation. A test whose sign flips when you remove a known orbital period is not reading a property of the accretion flow.

5.6 The independent same-day replication

The gap is not an artifact of this repository's implementation. On the same day, in an unrelated domain (physiological signals), an independent in-house program (**Fenosoma record v3**) ran the identical comparison against **byte-identical IAAFT nulls** and found the same collapse: the width's minority detection rising to full detection under the scattering omnibus (a $19/54 \rightarrow 54/54$ shift across its source panel), with the recovered departure **dominated by time-irreversibility** — the τ coordinate this paper identifies as one of the destroyed directions. Two skies, one day, byte-identical nulls, same verdict.

(Trace note, discharged 2026-08-25: the $19/54 \rightarrow 54/54$, byte-identical-nulls, and time-irreversibility figures are Fenosoma's sealed record v3, external to this repository. They have now been verified directly against `Fenosoma/out/v3/README.md`: the scattering aggregate is significant on 54/54 nsr2db records under both editing policies, against the x1 width's 19/54 (NN) and 21/54 (raw); the surrogates are byte-identical, regenerated from x1's pinned seeds; and the healthy-cohort excess is carried predominantly by the time-asymmetry family. The companion chf2db panel reads 29/29 scattering versus the width's 18/29 (NN). The earlier instruction to re-verify at submission is therefore discharged, not deferred.)

6. Direct engagement with the t_{MF} robustness counterclaim

A live line of work argues the opposite of this paper: that multifractal nonlinearity, distilled into the one-sample statistic t_{MF} , is a **robust** estimator of multiplicative cascade dynamics (Mangalam, Likens & Kelty-Stephen 2025, *Phys. Rev. E* **111**, 034126; arXiv:2312.05653). Because that claim and this one cannot both be read naively, we engage it directly and read its formula first.

Their statistic, quoted, not paraphrased. The multifractal-nonlinearity t-statistic compares the original series' singularity-spectrum width to the widths of its IAAFT surrogates and standardises the difference. In the lineage's own words (Bell et al. 2019, *Front. Physiol.* **10**:998, from Kelty-Stephen et al. 2013):

$$t_{MF} = (W_{MF} \text{ of original series} - \text{Average } W_{MF} \text{ of surrogate-data series}) / (\text{Standard error of } W_{MF} \text{ of surrogate-data series})$$

with the surrogate ensemble generated by IAAFT, and a significant *positive* t_{MF} (original outside the surrogate 95% interval) read as evidence of nonlinear interactions across scales. The robustness paper uses the same statistic in the form $t_{MF} = (\Delta\alpha_{Orig} - \langle \Delta\alpha_{Surr} \rangle) / \sigma_{\{\Delta\alpha, Surr\}}$.

Their own numbers, read exactly. The robustness paper simulates cascade generations and reports, verbatim, the distribution of t_{MF} :

“We modeled the 7 seven generations from the 9th to 15th generations of the 1000 cascade processes across additive and multiplicative cascades with five kinds of noises, yielding 7000 individual generations of random cascade processes. Out of these 7000, there were 3592 significant positive values of t_{MF} ... for **51.31% of the simulated generations.**”

“Within our sample of 7000 simulated generations, we found 1895 such significant negative t_{MF} , **27.07% of the total simulated generations.**”

The pool must be stated precisely, because the obvious reading of it is wrong. The 7,000 is **1,000 cascade processes $\times$ 7 generations (the 9th to the 15th)**, pooled across **additive and multiplicative** cascade types and five noise types — it is *not* 7,000 genuine multiplicative cascades. The authors themselves attribute the significant-negative tail particularly to the additive and “defractalizing” conditions, noting a raised likelihood of negative t_{MF} for random additive defractalizing cascades and connecting it to increasing additive-white-Gaussian-noise strength rather than to multiplicativity. We therefore do **not** read 27.07% as a wrong-sign rate on genuine cascades, and the argument we do make is weaker and true:

- **Pooled over conditions that include genuine multiplicative cascades, the flagship t_{MF} study reports the statistic firing significant-positive on barely half its own simulated generations (51.31%) and significant-negative on better than a quarter (27.07%), with $\approx 21.6\%$ null.** That is a two-sided instability, published by the statistic’s own proponents, in a literature that reads the statistic one-sidedly — significant-positive $\Rightarrow$ nonlinear cascade, null $\Rightarrow$ not one. Whatever the mixture’s composition, no decision rule with that distribution over a pool built entirely of cascade processes supports the converse inference.
- **The correctly-scoped number is the multiplicative-only subset, and it is not reported.** If the authors publish that split, it should be quoted here in place of the pooled figures. Our expectation is that it makes the argument *stronger*, not weaker — Lee & Kelty-Stephen (2017), from the same group, reach $0 < w_{orig} < w_{surr}$ on stochastic-multiplier cascades by raising the multiplier σ , so the low jaw is reachable inside the multiplicative family and is not an artifact of pooling additive conditions in. Until that subset exists we treat the pooled split as what it is: a two-sided instability of the pooled statistic, not a measured wrong-sign rate on cascades.

And the statistic is not the one this paper measured — said plainly, because it is the first reply a referee will make. Mangalam et al. estimate $\Delta\alpha$ with **Chhabra & Jensen’s direct method**, against **32 IAAFT surrogates**, with the decision rule $|t_{MF}| > 2.04$ ($df = 31$), on binomial-fracturing cascades with awGn/fGn noise terms at lengths $2^8 - 2^{14}$. This paper targets the **MFDEFA width** at a ~ 50 -surrogate rank operating point (§5.3). Estimator, null size and decision rule all differ, so nothing here is a direct refutation of their number — and our certificate does not claim to be one.

What makes MFDEFA the *right* target rather than a convenient one is published evidence from inside their own literature. Lee & Kelty-Stephen’s σ -sweep estimated w_{orig} two ways, by MFDEFA and by Chhabra–Jensen, on 200 simulated series at each σ from 0.1 to 1.1 in steps of 0.05, and found that increasing σ draws w_{orig} away from $w_{surr} < w_{orig}$ and towards $0 < w_{orig} < w_{surr}$ for both methods — **“but more for MF-DEFA.”** By the target literature’s own measurement the MFDEFA width is the more fragile of the two estimators, and it is

also the one the applied literature overwhelmingly runs. Pricing the fragile, widely-used estimator is the useful measurement; pricing Chhabra–Jensen as well is a natural extension and belongs in the adversarial collaboration offered below. One further scope note in the same spirit — and we measured it rather than leaving it a note. Their σ -sweep is over *stochastic*-multiplier cascades, while our canonical control is the *deterministic* binomial cascade: **one point on their axis**. Record v34 ([out/v34/](#)) puts three pre-registered points of their published grid ($\sigma = 0.30, 0.70, 1.10$ of $[0.10, 1.10]$ step 0.05) through the v22 protocol **unchanged** — same length $N = 16,384$, same MFDEFA scales and moments, same omnibus, same IAAFT null at $n_{\text{surr}} = 50$, $\text{SCAT_SEED} = 0$ — with the deterministic binomial re-run beside them on the same 24-seed panel as the in-run reference. That reference’s leading eight seeds *are* the v22 seeds on the v22 code path, and they reproduce the sealed battery’s $0.125 / 1.00$ exactly, so the σ cells and the certificate are the same measurement. The axis is stated exactly: **σ is the per-level standard deviation of the $\log_2$ split ratio**, under which the binomial is the *two-point* member of the same family at $\sigma_{\text{equiv}} = |\log_2 p - \log_2 (1-p)| = 1.222$ ($p = 0.3$), matched in the second moment.

The answer is not the one we pre-registered. We registered the hypothesis that the near-zero power would prove σ -generic, and named in advance the outcome that would scope this paper’s headline claim rather than confirm it. That is the outcome:

cell (24 realisations each)	width flags	mean width z	omnibus flags	mean z_couple	raw width
deterministic binomial, $p = 0.3$ (the canonical control)	1/24	+0.63	24/24	+13.9	0.760
stochastic, $\sigma = 0.30$	1/24	+0.08	23/24	+3.2	0.070
stochastic, $\sigma = 0.50$ (<i>post-hoc, added to locate the transition</i>)	9/24	+1.23	24/24	+2.9	0.198
stochastic, $\sigma = 0.70$	18/24	+2.73	24/24	+3.0	0.357
stochastic, $\sigma = 1.10$	21/24	+3.56	24/24	+2.3	0.712
stochastic, $\sigma = 1.222$ (the binomial’s own σ -equivalent)	18/24	+3.93	24/24	+2.4	0.899

The width’s detection rate against IAAFT is not σ -generic; it climbs the axis — $0.042, 0.375, 0.750, 0.875$ at $\sigma = 0.30, 0.50, 0.70, 1.10$ — while the scattering omnibus holds $0.958 - 1.000$ at every σ . And the decisive comparison is the last row against the first: at the binomial’s *own* σ -equivalent, and at a raw cascade width (0.899) that brackets the binomial’s own (0.760 , which falls between the $\sigma = 1.10$ and $\sigma = 1.222$ cells), the **stochastic** twin is flagged by the width $18/24$ of the time where the **deterministic** original is flagged $1/24$. So the invisibility is not a property of cascade strength, and not a property of the multiplicative family. It is a property of the *deterministic multiplier*.

And the record measures why, on the same surrogate path the verdict uses. An IAAFT-fidelity cell scores what fraction of the original’s MFDEFA width the surrogate ensemble hands back — §3(ii)’s premise, priced rather than asserted. **On the deterministic binomial the null reproduces 0.939 of the width; on every stochastic cell with a non-trivial width it reproduces about half ($0.465 - 0.552$).** That is the entire power difference in one number: where the null returns 94% of the statistic’s own quantity, the statistic has nothing left to measure. (The $\sigma = 0.30$ ratio is a quotient of two near-zero widths — that cascade’s raw width is 0.070 — and is not read.) The surrogate marginal is exact in every cell, and the spectral residual is ≤ 0.037 **on the pre-registered σ grid, 0.085 on the σ -equivalent anchor, and 0.351 on the deterministic reference** — so the null is best converged where the pre-registered grid is read, and an order of magnitude better there than on the reference.

Stated against this paper's own standard (dated correction, 2026-08-30). An earlier draft quoted the ≤ 0.037 bound over “the σ cells” without qualification. It covers the three *pre-registered grid* points only ($\sigma = 0.30 / 0.70 / 1.10$: $6.96e-5$, $2.53e-3$, $3.68e-2$). The $\sigma = 1.222$ **anchor** — the cell carrying the decisive $18/24$ against the binomial's $1/24$ — sits at 0.085 , $2.3\times$ the bound, and by the same directional argument this record makes just below (a shuffle-initialised, under-converged IAAFT ensemble is *under*-correlated and therefore *narrower*), that residual biases the anchor's own $18/24$ **upward**. The clean corroboration is the pre-registered $\sigma = 1.10$ grid cell: residual 0.037 , inside the bound, reading $21/24$ — a *higher* detection rate than the anchor's, on a better-converged null. So the σ -dependence conclusion does not rest on the anchor cell, and the fairness bound is quoted per-cell above rather than as one number over “the σ cells”.

Indeed the canonical control is where IAAFT converges **worst** in the panel, and because IAAFT is shuffle-initialised an under-converged ensemble sits *below* the target autocorrelation (acf_1 drift -0.194) — a less-correlated surrogate of a cascade is a *narrower* one, which biases the width test **towards** detection. It still flags $1/24$. **The near-zero power on the canonical control is conservative: it survives a null mis-converged in the direction that would favour it.**

Which is exactly what §3(ii) predicts, and the reason this result strengthens the mechanism instead of denting it. In the binomial every cell's multiplier magnitudes are identical — only the orientation is random — so a cell's local exponent is a function of its *value alone*: the multifractality is **maximally marginal-encoded**, and IAAFT, which returns the marginal exactly, hands the width straight back. Randomise the multiplier magnitude and that value $\leftrightarrow$ exponent correspondence breaks: part of the width moves out of the marginal and into the *arrangement*, above second order, where IAAFT scrambles it and the width can then see the difference. The omnibus's own coordinates say the same thing from the other side — on the binomial the departure is almost purely cross-scale coupling (z_couple **+13.9**, with $S1$, p and $S2n$ all *negative*), while on the σ -cells it is a milder C ($+2.3$ to $+3.2$) beside a rising second-order intermittency $S2n$ ($+0.5 \rightarrow +2.1$). §3's claim was never “the width cannot see cascades”; it was that the width cannot see a cascade *whose width is marginal-encoded*. v34 turns that premise into a measured dial and locates its boundary between $\sigma = 0.30$ and $\sigma = 0.70$.

What it does and does not change. It does not touch a number in §§5.1–5.3, where the certificate lives, and it does not weaken the headline as stated: on its own canonical positive control the test does have near-zero measured power, and that control is the one the applied literature uses. What it forecloses is the *over-general* reading — “the width cannot detect cascades” — which this paper has never claimed and now cannot be read as claiming. The converse inference dies by a better argument than uniformity: **the width's power against IAAFT is truth-dependent across a single cascade family, spanning 0.042 to 0.875 as a function of a property of the generator that no practitioner observes.** A test whose power varies twenty-fold with something you cannot measure cannot license “null $\Rightarrow$ not a cascade” — and that statement is now measured rather than assumed.

What we do not claim. We do **not** reproduce the *direction* Lee & Kelty-Stephen report. Their sweep moves w_orig from above its surrogate band towards $0 < w_orig < w_surr$ as σ rises; ours moves the mean width z the other way ($+0.08 \rightarrow +1.23 \rightarrow +2.73 \rightarrow +3.56$). Our σ -family is a stated, reproducible parameterisation of the stochastic-multiplier construction, not their generator, whose code we do not have — so the discrepancy is most plausibly a difference of family (or of estimator and decision rule: Chhabra–Jensen at 32 surrogates against MF DFA at a ~ 50 -surrogate rank point). **This cell prices our control family, and only ours.** Fixing the cascade construction and the σ convention jointly is precisely the first item the adversarial collaboration offered below would settle, and this result makes that offer more valuable, not less.

The disagreement, then, is not about their numbers — we accept them — but about what a distribution like that licenses. A statistic that fires significant-positive on about half of a pool built entirely of cascade processes and significant in the *opposite* direction on better than a quarter of it is not thereby shown robust. The t_{MF} framework is a fine *feature extractor*; the inference “significant $t_{MF} \Rightarrow$ cascade / null $t_{MF} \Rightarrow$ not a cascade” is what fails, in both directions.

An independent, same-month derivation of the low jaw’s mechanism. While this paper was in draft, the same authors published a mechanism for exactly the phenomenon our low jaw measures: Mangalam & Kelty-Stephen, “Constrained interactions and sublinear multifractality in multiscale systems,” *Phys. Rev. E* **114**, 024215 (published 17 August 2026). Their abstract states that various **nonmultiplicative constraints on the nonlinear correlations across scales of a multiplicative cascade might produce sublinear multifractality**, framed as analogues of the physical limits characteristic of homeostatic regulation and feedback control. That is a derivation, from the mechanism side, of *why a genuine cascade can present with suppressed multifractality* — the same object our certificate prices from the measurement side, obtained independently and in the same month. We cite it as **allied prior art**, exactly as we treat Lee & Kelty-Stephen (2017) in §1.2, and we state the delta plainly: they supply a mechanism by which the width can be narrow on a real cascade; we supply the **measured certificate** — the detection rate at a matched operating point on the canonical control, pre-registered, with a replacement statistic whose own failure modes ship beside it. Neither result contains the other. (That paper has no arXiv preprint; the sentence above is quoted from the APS abstract page and should be re-read against the published text before submission.)

Adversarial collaboration, offered here. The productive path is not duelling preprints. We offer an adversarial collaboration: a jointly pre-registered battery — agreed cascade constructions, agreed lengths and surrogate budgets, agreed decision rule — on which both t_{MF} and the scattering omnibus report power and specificity, with the sign convention and the $\Delta\alpha$ -vs- $h(q)$ width definition fixed in advance. The Kelty-Stephen group sits at the origin of this statistic and published the prior art that anticipates our low jaw (§1.2); the disagreement is narrow and empirically decidable. (Whether to open that collaboration is an owner decision — this paper only offers it.)

7. The scattering statistic’s own high side

A paper whose entire demand of the field is “measure your statistic’s power and specificity before you claim” cannot exempt its own replacement statistic. The scattering omnibus has a genuine high side, it is measured, and we state it before a referee does.

Isolated singularities produce a real positive omnibus excess with no cascade present. This is a **published** result before it is an in-house one, and the published statement comes first: Oświęcimka et al. (2020) found that both MF DFA and wavelet leaders are essentially unable to infer the non-multifractal nature of signals devoid of a cascade-like hierarchy of singularities — signals harbouring only isolated singularities “artefactually give rise to broad multifractal spectra, resembling those expected in the presence of a well-developed underlying multifractal structure,” so that “there is a real risk of incorrectly inferring multifractality due to isolated singularities.” Our own material adds an *enumerable-null* instance of the same mechanism for the scattering statistic. An independent in-house study (Fenopan record v6; CONSTITUTION.md §6.2) read a **unicursal labyrinth** — a frozen,

deterministic, aperiodic geometry with no multiplicative cascade anywhere in it — and found a scattering omnibus excess of $z = +7.8$, driven by the labyrinth's isolated turns (singularities). Two features make that instance decisive rather than anecdotal:

- The excess is **generic to 100% of the enumerable $M(12) = 1,828$ null** — the null there is *exactly enumerable* (every labyrinth of that class counted), not surrogate-approximated, so the positive is not a surrogate artifact.
- The **two-point width is correctly null** on the same object: the high side is not shared by the width here; it belongs to the scattering leverage.

The lesson is symmetric with §3's mechanism. IAAFT-class phase-scrambling nulls are the **wrong null for a frozen deterministic geometry**, just as they are the right null for a stochastic linear process — so *any* scaling statistic that occupies the destroyed higher-order coordinates (scattering included) will read an isolated singularity as an excess. **The pincer's high side is a property of scaling statistics against phase-scrambling nulls, not of the width alone.** Our own instrument inherits it.

What distinguishes the scattering vector is not immunity but decomposability. The omnibus does not merely say “something higher-order is here”; its aggregate-excess rule carries the per-family direction, and specifically the cross-scale-coupling arm C — the coordinate a *multiplicative* cascade must load (§3). That arm dominates in only 3/22 Swift/BAT sources (§5.4), which is exactly how the vector separates “some higher-order structure” (an isolated singularity, an intermittency burst, a time-asymmetry) from “the multiplicative cross-scale cascade a multifractal claim needs.” An isolated-singularity excess shows up in $S1/S2n$, not as a dominant C . **This is the statistic's declared operating envelope:** it is a certified detector of *higher-order departure from the IAAFT null*, and a *cascade* claim requires reading the C arm, not the bare omnibus p . Stated as a limitation so the certificate is complete, not buried.

This limit is published (Oświęcimka et al. 2020) and independently reproduced in-house on an enumerable null; the in-house half is dated, adversarially reviewed, and now program canon (CONSTITUTION.md §6.2/§10). Citing only the in-house record for a six-year-old published result would have been the exact sin this paper audits, and the correction is dated here (2026-08-25).

8. Reproducibility of the battery: two readings, both reported

A second independent in-house reimplement (Fenopan record v5) rebuilt the scattering battery from scratch and did **not** reproduce the full-vector result cleanly. It found the **full S2 omnibus not robustly specific from scratch**: it required a *restricted, well-separated-scale* leverage to hold specificity, and a genuinely **analytic** filterbank (a rectified carrier over-detected residual periodicity and had to be replaced). Both facts must be reported, and they admit two readings, which we give both ways because the certificate is only honest if it names its own fragility:

- **The strength reading.** Scale separation is *why* the restricted leverage works: separating the scales isolates the specifically-multiplicative coupling from the adjacent-band overlap that fires on *any* process. Under that reading, the mechanism is understood and the restriction is principled, not a fudge — a strength worth stating plainly.

- **The weakness reading.** The v22 full-vector result is *protocol-sensitive*: a from-scratch reimplementa-tion that does not adopt the scale-separated, analytic-filterbank choices can lose specificity. Under that reading the headline is more fragile than a single battery suggests — a weakness worth stating *first*.

The same reimplementa-tion also **retired, as empirically false, a claim it would have been convenient to make**: that an IAAFT surrogate of a cascade must **not** detect. For the binomial cascade the multiplicative leverage is largely marginal-encoded, so an IAAFT surrogate of a cascade **still detects**. The certification this paper rests on is therefore stated exactly, and never overstated: it is “**power on real cascades + specificity on genuine linear/rough negatives,**” never “IAAFT erases the cascade.” The distinction matters because the naive slogan (IAAFT destroys the cascade, so a surrogate is a clean negative) is the very over-reading §3 is at pains to avoid.

9. The branching-ratio rule and its in-house reproductions

This preprint is also the program’s named vehicle for an allied methodological rule that travels with the same discipline (CONSTITUTION.md §6.1, the Filimonov–Sornette rule): **a branching ratio is never evidence of contagion until a regime/background alternative has been raced and the bidirectional recovery gate has passed at the actual η** . A fitted Hawkes branching ratio $\hat{\eta}$ near 1 is routinely read as “near-critical self-excitation,” but a smooth, non-stationary background alone drives $\hat{\eta} \rightarrow 1$ as the estimator absorbs the trend.

Correction, dated 2026-08-25. An earlier draft of this section stated that a literature search found this “in print nowhere as a hard rule.” That was wrong, and the correction is made here rather than left to a referee. The trap is Filimonov & Sornette’s own published result — calibrating a Hawkes model on mixtures of Poisson processes with regime changes, true $\eta = 0$ by construction, returns $\hat{\eta}$ hovering around the critical value 1, and the residual analysis of those calibrations *cannot reject* Poisson residuals, so time-rescaling goodness-of-fit does not adjudicate it (2015, *Quantitative Finance* 15(8), §4.3–4.4). The size of the effect has been measured on real data with everything else held fixed: fitting the same tick data under backgrounds of increasing flexibility drops the average branching ratio monotonically from **0.74** (constant μ) to **0.41** (a Bayesian cubic-B-spline $\mu(t)$), so the estimate is partly a readout of how flexible the baseline was allowed to be (Omi, Hirata & Aihara 2017, *PRE* 96, 012303). And the failure is formal as well as empirical: Chen & Hall introduced a time-varying background precisely because a constant baseline degrades the analysis (2013), then proved that in the fully nonparametric case not every aspect of the model can be identified from a single observed sample path (2016). Seismology has argued the same question under the name “clustering vs. rate inhomogeneity” and answered it architecturally, by putting a variable background inside the model. **What is a program deliverable is therefore not the rule but the cross-domain measured confusion/recovery framing** — the same trap demonstrated in four unrelated arenas with a bidirectional recovery gate attached to each — and that is what the list below should be read as claiming.

Four independent in-house reproductions, in four unrelated domains:

1. **Crisis data** (Fenocrisis record v0): a self-excitation-free background reproduces a near-critical $\hat{\eta}$.
2. **This repository’s astrophysics** (record v21; CONSTITUTION.md §6.1). A *provably* self-excitation-free smooth solar cycle — a homogeneous-background world with zero Hawkes coupling by construction — reproduces **median $\eta = 0.956$** against a committed 0.957 , and even beats the cascade on the event-time BIC. So the “Hawkes beats the cascade” reading on that record is a rate-nonstationarity artifact: neither the

intensity-preserving nor the local-interval-shuffle null separates self-excitation from nonstationarity, because every generator preserves the trend, and the decisive evidence is null-reproduction. (The record's counts-arena headline is untouched; only the η reading moves.)

3. **Real ant-colony data** (Fenopan record v5): a self-excitation-free, intensity-preserving surrogate drives the naive η higher than the data, on 15/16 replicates.
4. **A real astrophysical event record** (this repository, record v29). On the ~50-year GOES/XRS flare record the fitted η runs to the estimator's ceiling at both thresholds (M1.0, 8,393 flares; C1.0, 64,996) and the zero-self-excitation smooth-solar-cycle intensity-preserving null reproduces it, while the informative local-interval-shuffle bands (median 0.886 / 0.976) place the observed value above a *short-range* null — some structure, but the dominant driver is a smooth cycle no surrogate separates from self-excitation.

Four domains — crisis, astrophysics on synthetic and on real records, and animal behaviour; synthetic-provable and real — is the claim's strength, and we state it that way. The rule is the event-time-domain sibling of this paper's amplitude-domain thesis: a summary statistic (η , or the width) read without a background-matched null and a power check at the actual N is not a measurement.

10. The proposed reporting standard

The constructive contribution is not “use scattering instead of the width.” The scattering omnibus has its own high side (§7) and a from-scratch specificity that must be earned with scale-separated leverage (§8); promoting it to a new default hammer would repeat the error this paper diagnoses. The contribution is a **standard for reporting any multifractality verdict that rests on a surrogate-null width test**:

1. **A measured specificity-and-power certificate on the canonical control, committed before the data.** Report the fraction of genuine binomial-cascade realisations the test flags (power) and the fraction of matched linear-Gaussian realisations it flags (specificity), at the study's own N , surrogate budget, and decision rule — *pinned before the real data are read* (the anchors-before-claims discipline).
2. **The width is a feature, not a verdict.** A raw width, or a shuffle-only test, is never claim-bearing (canon F2). A width excess over IAAFT is an *under-claim* (real structure the low-power statistic happened to catch); a width *deficit* is uninformative about linearity (the low jaw). Report the width as a biomarker/feature and let a powered statistic or an exact likelihood carry any mechanistic claim.
3. **Name the null's blind spots.** State that IAAFT is the wrong null for fourth-order/time-asymmetric structure (it cannot destroy what it must, §3) and for frozen deterministic geometry (it destroys what a phase-scrambling null must not, §7). A verdict inherits the null's blind spots.
4. **Decompose before attributing.** “Higher-order departure from IAAFT” is not “multiplicative cascade”; a cascade claim reads the cross-scale-coupling arm, not the bare omnibus (§7).

Adopted, the standard costs one synthetic battery per study and makes the converse inference (§11) unspeakable without a power number beside it.

This is the second half of a standard already being written from the other side. A parallel programme (Mendez, Mariani, Beccar-Varela, Tweneboah & Jaroszewicz, and collaborators) has published, within the last six months, a finite-size protocol showing previously reported broad singularity spectra in the 2D Ising model to be artifacts (arXiv:2603.04609), and an SST audit which finds that none of the indices studied are multifractal within record resolution and states plainly that “reported widths should not be interpreted without verifying both the sign

of $h'(q)$ and the stability of the scaling range,” revising their own earlier estimates downward (arXiv:2606.18901). That is a reporting-standard demand from the **high** side: what to check before believing a positive width. The standard proposed here is its two-sided completion — **what to check before believing a null** — and their MF-toolkit (arXiv:2604.16257) is much of the tooling that makes both runnable. We name it as an allied programme, not a rival one: the two demands are complementary and a study that meets only one is still under-specified.

11. Discussion and limitations

What dies. The **converse inference** — “the width did not exceed its IAAFT null, therefore the process is monofractal / linear / not a cascade” — is void as uninformative wherever it rests on a test with unmeasured, near-zero power on the canonical control. Record v34 sharpens the reason it is void: across a *single* cascade family the width’s power against IAAFT runs from 0.042 to 0.875 as a function of how completely the cascade’s width is marginal-encoded — a property of the generator that no practitioner observes (§6). The converse inference therefore fails not because the power is uniformly near zero, but because it is **truth-dependent by a factor of twenty in an unobservable direction**, which is the stronger and the measured statement. That inference is common; the demonstration exhibit is a study that read a non-significant width-vs-IAAFT difference at $N = 512$ ($p \approx 0.06-0.09$) as “the strength of nonlinear multifractality is low” — a mechanism inference from a test with no measurable power at that length (Arsac 2023). The *substantive* answer there may well be right; the *inference* is void, and that is the cleanest demonstration of the distinction.

What survives. (i) **Likelihood-based cascade results** are untouched — they never used the width as their discriminator; an exact-filter model comparison (this program’s MSM-Cox tournament) reads the cascade from the data’s likelihood, not from a surrogate width. (ii) **Width-as-biomarker classifications survive** — a feature needs no mechanism to classify; if a narrower spectrum reliably separates congestive-heart-failure from healthy records, that discrimination stands whether or not the width “means” multifractality. (iii) **Positive width-vs-IAAFT rejections become under-claims** — where the low-power width *did* exceed its null, the real structure is at least that large.

What strengthens. The scope-correction is generative: it identifies the destroyed coordinates (fourth-order coupling, time-asymmetry), hands them a powered statistic, and — critically — prices that statistic’s own failure modes in the same certificate. The result is a *reporting standard* the field can adopt at the cost of one battery.

Honest limitations of this preprint.

- **Existence, not attribution.** The measured deficit is on the width-vs-IAAFT *existence* test. The three-way *attribution* decomposition (original vs shuffle vs FT vs IAAFT) that the largest cross-domain targets run (Zhou 2009; Baruník et al. 2012; the rs-fMRI lineage) has a detection power on the canonical cascade that **nobody has measured, including this program**. Its blast radius is therefore *conjectured*, not measured, and a companion study must close it before any “the field is wrong” claim (Appendix B). This paper’s stratum is the existence test and its `t_MF` distillate.
- **Synthetic controls at a canonical operating point.** The power/specificity certificate is measured on binomial cascades, `msm_logG`, and `msm_counts` at a pinned operating point; it is a calibration, not a population survey. The 22-source re-audit is real data, but its verdict is deliberately conservative (3/22 cross-scale).

- **The replacement statistic is not a clean instrument** (§7–§8). Its high side (isolated singularities) and its from-scratch specificity fragility are real; the contribution is the standard, not the hammer.
- **The counterclaim engagement rests on the counterclaim’s own reported figures** (§6), and on a *different* estimator from the one we measured (they: Chhabra–Jensen with 32 IAAFT surrogates and $|t_{MF}| > 2.04$; here: the MFDFA width at a ~50-surrogate rank point). The exact t_{MF} decision convention (σ vs standard error; the $\Delta\alpha$ vs $h(q)$ width) should be fixed jointly in the offered adversarial collaboration, and the primary text re-read against the final draft before submission.
- **Both cells that were queued pre-submission are now measured, and they did not come back the same way.** (i) The length- and preprocessing-matched AR(1)/white specificity control for the scattering omnibus at $N \approx 1,350\text{--}7,100$ through the v7 mask-concatenate-and-censor pipeline, which scopes the 22/22, is **record v33**: the mask-transported controls false-flag at $0.045 / 0.036$ against a 2/51 rank budget, load $S2n$ rather than the cross-scale C , and specificity holds — §5.4 stands, upgraded (§5.4). (ii) The placement of our deterministic binomial control on Lee & Kelty-Stephen’s stochastic-multiplier σ -axis is **record v34**, and it went the other way: the width’s near-zero detection rate is **σ -specific, not σ -generic**, climbing from $1/24$ at $\sigma = 0.30$ to $21/24$ at $\sigma = 1.10$ while the omnibus holds $0.958\text{--}1.000$ across the axis (§6). Neither changes a number in §§5.1–5.3, where the certificate lives; the second changes what may be *generalised* from them, and it is reported here in the form the pre-registration named in advance rather than in the form we expected.

Priority and posture. This finding’s Zenodo deposit was published 2026-08-31 ([10.5281/zenodo.22180477](https://doi.org/10.5281/zenodo.22180477) — the DOI is live and the priority date is established; before that date this paragraph correctly read “no DOI and no priority date”), and several active lines sit adjacent to it: the t_{MF} robustness paper and its same-month sublinear-multifractality sequel (PRE **114**, 024215, 17 Aug 2026); one-command IAAFT-attribution toolkits shipping a binomial-cascade generator and an IAAFT module but no detection rate (MF-toolkit, arXiv:2604.16257, Apr 2026); and a parallel audit programme publishing reporting-standard demands from the *high* side (arXiv:2606.18901, revised 7 Aug 2026). The prior art is allied and cited on page one — Lee & Kelty-Stephen 2017 on the low jaw, Oświęcimka et al. 2020 on the high jaw, Wendt & Abry 2007 on the demand itself — and the contribution is scoped throughout as **quantification, never discovery**: the two genuinely unclaimed assets are the *measured certificate* at a matched operating point on the canonical control and the *scattering omnibus deployed as a hypothesis test against a surrogate null*, which no 2025–2026 work does. arXiv posting and any adversarial-collaboration outreach are owner actions.

12. References

(Bibliographic details verified where possible; the two engaged external claims — the t_{MF} robustness paper and the prior art — carry their DOIs so a referee can check the reading against the source.)

- **Bell, C. A., Carver, N. S., Zbaracki, J. A., & Kelty-Stephen, D. G. 2019**, “Non-linear Amplification of Variability Through Interaction Across Scales Supports Greater Accuracy in Manual Aiming: Evidence From a Multifractal Analysis With Comparisons to Linear Surrogates in the Fitts Task,” *Frontiers in Physiology*, 10, 998. doi:[10.3389/fphys.2019.00998](https://doi.org/10.3389/fphys.2019.00998) (the explicit t_{MF} formula quoted in §6).
- **Bruna, J., Mallat, S., Bacry, E., & Muzy, J.-F. 2015**, “Intermittent process analysis with scattering moments,” *Annals of Statistics*, 43(1), 323–351. doi:[10.1214/14-AOS1276](https://doi.org/10.1214/14-AOS1276)

- **Chen, F., & Hall, P. 2013**, “Inference for a nonstationary self-exciting point process with an application in ultra-high frequency financial data modeling,” *Journal of Applied Probability*, 50(4), 1006–1024. doi:[10.1239/jap/1389370096](https://doi.org/10.1239/jap/1389370096) (§9: the time-varying background, introduced because a constant one degrades the analysis).
- **Chen, F., & Hall, P. 2016**, “Nonparametric estimation for self-exciting point processes — a parsimonious approach,” *Journal of Computational and Graphical Statistics*, 25(1), 209–224. doi:[10.1080/10618600.2014.1001491](https://doi.org/10.1080/10618600.2014.1001491) (§9: the formal single-sample-path identifiability failure).
- **Drożdż, S., Kwapień, J., Oświęcimka, P., et al. 2023**, “Genuine multifractality in time series is due to temporal correlations,” *Physical Review E*, 107, 034139 (arXiv:2211.00728). doi:[10.1103/PhysRevE.107.034139](https://doi.org/10.1103/PhysRevE.107.034139) (§1.2: fat tails broaden the spectrum only when correlations are present; short heavy-tailed series are a trap for shuffle-based surrogate tests).
- **Filimonov, V., & Sornette, D. 2015**, “Apparent criticality and calibration issues in the Hawkes self-excited point process model: application to high-frequency financial data,” *Quantitative Finance*, 15(8), 1293–1314. doi:[10.1080/14697688.2015.1032544](https://doi.org/10.1080/14697688.2015.1032544) (§9: true $n = 0$ calibrating to $n \approx 1$, with Poisson residuals that cannot reject it — the published statement of the rule §9 travels with).
- **Kantelhardt, J. W., Zschiegner, S. A., Koscielny-Bunde, E., Havlin, S., Bunde, A., & Stanley, H. E. 2002**, “Multifractal detrended fluctuation analysis of nonstationary time series,” *Physica A*, 316, 87–114. doi:[10.1016/S0378-4371\(02\)01383-3](https://doi.org/10.1016/S0378-4371(02)01383-3)
- **Kelty-Stephen, D. G., Lane, E., Bloomfield, L., & Mangalam, M. 2022/2023**, “Multifractal test for nonlinearity of interactions across scales in time series,” *Behavior Research Methods*, 55, 2249–2282. doi:[10.3758/s13428-022-01866-9](https://doi.org/10.3758/s13428-022-01866-9) (the canonical methods statement of the test audited here — the tutorial that propagated the MFDEFA-width + surrogate recipe into the behavioural literature).
- **Kelty-Stephen, D. G., Palatinus, K., Saltzman, E., & Dixon, J. A. 2013**, “A tutorial on multifractality, cascades, and interactivity for empirical time series in ecological science,” *Ecological Psychology*, 25(1), 1–62. doi:[10.1080/10407413.2013.753804](https://doi.org/10.1080/10407413.2013.753804)
- **Kluszczyński, R., Drożdż, S., Kwapień, J., Stanisław, T., & Wątorek, M. 2025**, “Disentangling sources of multifractality in time series,” *Mathematics*, 13, 205 (arXiv:2501.08898). doi:[10.3390/math13020205](https://doi.org/10.3390/math13020205)
- **Lee, H., & Kelty-Stephen, D. G. 2017**, “Cascade-Driven Series with Narrower Multifractal Spectra Than Their Surrogates: Standard Deviation of Multipliers Changes Interactions across Scales,” *Complexity*, 2017, 7015243. doi:[10.1155/2017/7015243](https://doi.org/10.1155/2017/7015243) (prior art; the qualitative low jaw, cited on page one).
- **Mangalam, M., & Kelty-Stephen, D. G. 2024a**, “Multifractal perturbations to multiplicative cascades promote multifractal nonlinearity with asymmetric spectra,” *Physical Review E*, 109, 064212. doi:[10.1103/PhysRevE.109.064212](https://doi.org/10.1103/PhysRevE.109.064212)
- **Mangalam, M., & Kelty-Stephen, D. G. 2024b**, “Additivity suppresses multifractal nonlinearity due to multiplicative cascade dynamics,” *Physica A*. doi:[10.1016/j.physa.2024.129651](https://doi.org/10.1016/j.physa.2024.129651) (an earlier statement, from the same group, that additivity suppresses the very signal the width test looks for — relevant to §6’s reading of the pooled t_{MF} split).
- **Mangalam, M., & Kelty-Stephen, D. G. 2026**, “Constrained interactions and sublinear multifractality in multiscale systems,” *Physical Review E*, 114, 024215 (published 17 August 2026; no preprint). doi:[10.1103/vnl1-v81t](https://doi.org/10.1103/vnl1-v81t) (allied prior art, §6: a mechanism — nonmultiplicative constraints on a cascade’s cross-scale nonlinear correlations — by which a genuine cascade presents with suppressed multifractality. Quoted from the APS abstract; re-read against the published text before submission).
- **Mangalam, M., Likens, A. D., & Kelty-Stephen, D. G. 2025**, “Multifractal nonlinearity as a robust estimator of multiplicative cascade dynamics,” *Physical Review E*, 111, 034126 (arXiv:2312.05653). doi:[10.1103/PhysRevE.111.034126](https://doi.org/10.1103/PhysRevE.111.034126) (the live counterclaim engaged in §6; the 51.31% / 27.07% split is

theirs, over 7,000 generations pooled across additive and multiplicative cascade types and five noise types — see §6 for the exact pool).

- **Mendez, C. G., Mariani, M. C., Beccar-Varela, M. P., Tweneboah, O. K., & Jaroszewicz, S. 2026**, “MF-toolkit: a high-performance Python library for multifractal analysis with automated crossover detection, source identification and application to gravitational-waves data,” arXiv:[2604.16257](#) (§1.3: ships a binomial-cascade generator and an IAAFT module; reports no detection rate). See also arXiv:[2603.04609](#) (finite-size scaling for spurious multifractality in the 2D Ising model) and arXiv:[2606.18901](#) (Jaroszewicz et al., “Persistence without multifractality in tropical Atlantic SST indices,” v2, 7 Aug 2026) — §10’s allied high-side reporting standard.
- **Morel, R., Rochette, G., Leonarduzzi, R., Bouchaud, J.-P., & Mallat, S. 2025**, “Scale dependencies and self-similar models with wavelet scattering spectra,” *Applied and Computational Harmonic Analysis*, 75, 101724 (arXiv:2204.10177). doi:[10.1016/j.acha.2024.101724](#) (their §3.3 self-similarity check compares error bars against estimation variance on a single realisation — no test statistic, no surrogate, no rejection rate; the closest existing move to §7’s omnibus, and not a test).
- **Omi, T., Hirata, Y., & Aihara, K. 2017**, “Hawkes process model with a time-dependent background rate and its application to high-frequency financial data,” *Physical Review E*, 96, 012303. doi:[10.1103/PhysRevE.96.012303](#) (§9: the same data fitted under increasingly flexible backgrounds drops the branching ratio $0.74 \rightarrow 0.41$).
- **Oświęcimka, P., Drożdż, S., Frasca, M., Gębarowski, R., Yoshimura, N., Zunino, L., & Minati, L. 2020**, “Wavelet-based discrimination of isolated singularities masquerading as multifractals in detrended fluctuation analyses,” *Nonlinear Dynamics*, 100, 1689–1704. doi:[10.1007/s11071-020-05581-y](#) (prior art; the published high jaw — isolated singularities artefactually broadening spectra, “a real risk of incorrectly inferring multifractality”).
- **Schreiber, T., & Schmitz, A. 1996**, “Improved Surrogate Data for Nonlinearity Tests,” *Physical Review Letters*, 77, 635–638. doi:[10.1103/PhysRevLett.77.635](#)
- **Wendt, H., & Abry, P. 2007**, “Multifractality Tests Using Bootstrapped Wavelet Leaders,” *IEEE Transactions on Signal Processing*, 55(10), 4811–4820. doi:[10.1109/TSP.2007.896269](#) (prior art; the methodological demand met nineteen years ago for a competing statistic — six bootstrap test declinations with significances, powers and *p*-values assessed by Monte Carlo on synthetic processes).
- **Zhou, W.-X. 2009**, “The components of empirical multifractality in financial returns,” *EPL (Europhysics Letters)*, 88, 28004 (arXiv:0912.4782). doi:[10.1209/0295-5075/88/28004](#)
- **Barunik, J., Aste, T., Di Matteo, T., & Liu, R. 2012**, “Understanding the source of multifractality in financial markets,” *Physica A*, 391, 4234–4251. doi:[10.1016/j.physa.2012.03.037](#)
- **Arsac, L. M. 2023**, “Multifractal Dynamics in Executive Control When Adapting to Concurrent Motor Tasks,” *Entropy*, 25(9), 1364. doi:[10.3390/e25091364](#)

Data availability & acknowledgments

Every quantitative claim above traces to an immutable, versioned results record. Each is deposited as its own citable object; the DOIs below were **published on Zenodo 2026-08-31** and resolve as of that date. The source repository is **private and remains private**, so a record is cited by DOI and never by repository path.

In this repository:

record	what it carries here	DOI
out/v22/	the scattering battery, the 22-source re-audit, and the negative-z diagnosis — pinned seeds, deterministic surrogate ensembles	10.5281/zenodo.22180491 (published 2026-08-31)
out/v7/	the width null audit whose columns v22 quotes unchanged	10.5281/zenodo.22180487 (published 2026-08-31)
out/v21/	the event-time surrogate honesty layer, §9's $\eta \approx 0.956$	10.5281/zenodo.22180489 (published 2026-08-31)
out/v29/	§9's fourth branching-ratio reproduction, on the real GOES stream	10.5281/zenodo.22180497 (published 2026-08-31)
out/v33/	§5.4's mask-transported specificity control	10.5281/zenodo.22180499 (published 2026-08-31)
out/v34/	§6's σ -parameterisation of the power claim	10.5281/zenodo.22180501 (published 2026-08-31)

The analysis package that produces all of them is deposited as a software snapshot: **Fenocosm reproduction snapshot**, [10.5281/zenodo.22180503](#) (published 2026-08-31).

The external in-program records this paper leans on are deposited likewise: **Fenosoma record v3** (the independent same-day scattering replication), [10.5281/zenodo.22180546](#); **Fenopan record v5** (the from-scratch specificity limit and the ant-colony branching-ratio reproduction), [10.5281/zenodo.22180449](#); and **Fenopan record v6** (the labyrinth isolated-singularity limit), [10.5281/zenodo.22180454](#) — each (published 2026-08-31).

The scattering battery and audit figures are committed at [out/v22/p34_scattering_battery.png](#) and [out/v22/p34_scattering_audit.png](#), and regenerate from a fresh clone behind the apparatus gate (`./verify.sh`, which includes the scattering smoke). The independent same-day replication is Fenosoma record v3; the two in-house limits on the scattering omnibus (§7 isolated singularities; §8 from-scratch specificity) and one of §9's branching-ratio reproductions are Fenopan records v6 and v5 (`CONSTITUTION.md` §6.1/§6.2/§10), cited there and traced in Appendix A; §9's fourth reproduction is this repository's [out/v29/](#), the real-stream re-derivation.

The MF DFA/DFA estimators and the MSM-Cox generator are vendored with provenance headers from a CC BY 4.0 analysis package (see [PROVENANCE.md](#)); the scattering transform, the IAAFT and point-process surrogates, and all battery drivers are self-contained modules first written for this repository. Swift/BAT data courtesy of the Swift/BAT Hard X-ray Transient Monitor team (Krimm et al. 2013, ApJS 209, 14). All new work in this paper is CC BY 4.0, © 2026 Evan Tabak Atlas.

Appendix A — Provenance table (every headline number → its sealed

record)

number (as used)	record	file / section
width flags binomial cascade 0.00 (v7 protocol)	v7 (10.5281/zenodo.22180487 ; published 2026-08-31)	out/v7/README.md, control battery
width flags cascade controls 0.125 (v22 protocol)	v22 (10.5281/zenodo.22180491 ; published 2026-08-31)	p34_scattering_battery.json; FINDINGS reconciliation note
scattering flags all three cascade controls 1.00	v22 (10.5281/zenodo.22180491 ; published 2026-08-31)	out/v22/README.md, Table (anchor battery)
specificity AR(1) 0.00, white 0.00 (both statistics, N = 16,384 gap-free)	v22 / v7	anchor battery
mask-transported specificity AR(1) 0.045 (20/448), white 0.036 (16/448) at N ≈ 1,350–7,100 through the v7 pipeline; 2/51 budget, not elevated	v33 (10.5281/zenodo.22180499 ; published 2026-08-31)	out/v33/README.md; p35_maskspec.json
mask-transported cross-scale- C false-flag share 7/896 (flags are S2n-dominant, not C; seams do not fake a cascade)	v33 (10.5281/zenodo.22180499 ; published 2026-08-31)	out/v33/README.md; p35_maskspec.json
z_couple binomial +13.8, msm_logG +3.5, msm_counts +5.1	v22 (10.5281/zenodo.22180491 ; published 2026-08-31)	anchor battery
width flags the deterministic binomial 1/24, the stochastic $\sigma = 1.10$ cascade 21/24, at the same operating point (N = 16,384, n_surr = 50, IAAFT) — the power claim σ -parameterised	v34 (10.5281/zenodo.22180501 ; published 2026-08-31)	out/v34/README.md; p36_sigma.json
width detection rate across the σ -axis 0.042 / 0.375 / 0.750 / 0.875 at $\sigma = 0.30 / 0.50 / 0.70 / 1.10$; omnibus 0.958–1.000 throughout	v34 (10.5281/zenodo.22180501 ; published 2026-08-31)	out/v34/README.md; p36_sigma.json
the binomial's σ -equivalent on that axis, $1.222 = \log_2 p - \log_2(1-p) $ at $p = 0.3$ (computed, and equal to its analytic singularity-support width)	v34	p36_sigma.json, protocol block
the binomial reference's leading 8 seeds reproduce the v22 battery cells (0.125 width, 1.00 scattering) byte-for-byte	v34 (10.5281/zenodo.22180501) / v22 (10.5281/zenodo.22180491); published 2026-08-31	p36_sigma.json cells[].v22_panel; the v22 anchor battery
binomial cascade raw width 0.745, IAAFT z +0.49	v7 (10.5281/zenodo.22180487 ; published 2026-08-31)	out/v7/README.md, control table
power vs length (Table 2)	v22 (10.5281/zenodo.22180491 ; published 2026-08-31)	out/v22/README.md, power curve
per-realisation false-positive rate 2/51 ≈ 0.0392 (ranks 1–2 of 51 clear $p < 0.05$); binomial $p \approx 0.274$ for 1 flag in 8 seeds; $1/24 = 0.042$ indistinguishable from that floor	v22 (10.5281/zenodo.22180491 ; published 2026-08-31)	n_surr = 50, signif_p = 0.05, control_seeds = 8 (protocol block); §5.3, corrected 2026-08-25 and again 2026-08-30
22/22 scattering vs 7/22 width	v22 (10.5281/zenodo.22180491 ; published 2026-08-31)	out/v22/README.md, re-audit; p34_scattering.csv
cross-scale coupling C dominant in 3/22	v22 (10.5281/zenodo.22180491 ; published 2026-08-31)	FINDINGS §F12; p34_scattering.csv

number (as used)	record	file / section
Vela X-1 fold flip $-2.25 \rightarrow +2.56$	v22 (10.5281/zenodo.22180491 ; published 2026-08-31)	out/v22/README.md , negative-z diagnosis
Her X-1 $-5.77 \rightarrow -4.56$, 16% S/N < 1; GX 301-2 $-3.28 \rightarrow -6.48$	v22 (10.5281/zenodo.22180491 ; published 2026-08-31)	negative-z diagnosis
GRS 1915+105 shuffle $z +10.5 \rightarrow$ IAAFT $z +2.1$	v7 (10.5281/zenodo.22180487 ; published 2026-08-31)	out/v7/README.md ; FINDINGS §F5.2
7/22 IAAFT-excess at $p < 0.05$	v7 (10.5281/zenodo.22180487 ; published 2026-08-31)	out/v7/README.md
η median 0.956 vs committed 0.957 , self-excitation-free	v21 (10.5281/zenodo.22180489 ; published 2026-08-31)	out/v21/README.md ; FINDINGS §F11
Fenosoma same-day replication $19/54 \rightarrow 54/54$, τ -dominated	Fenosoma v3 (<i>external; verified against Fenosoma/out/v3/ 2026-08-25 — 54/54 scattering-significant on nsr2db under both editing policies vs the x1 width's 19/54 NN / 21/54 raw, byte-identical IAAFT surrogates regenerated from x1's pinned seeds, excess carried by the time-asymmetry family</i>) — 10.5281/zenodo.22180546 (published 2026-08-31)	Fenosoma record v3, README.md
labyrinth scattering $z = +7.8$, $M(12) = 1,828$ enumerable null	Fenopan v6 (10.5281/zenodo.22180454 ; published 2026-08-31)	Fenopan CONSTITUTION.md §6.2/§10
full-S2 not robustly specific from scratch; IAAFT-of-cascade detects	Fenopan v5 (10.5281/zenodo.22180449 ; published 2026-08-31)	Fenopan record v5; CONSTITUTION.md §6.2
ant-colony η^* higher on self-excitation-free surrogate, 15/16	Fenopan v5 (10.5281/zenodo.22180449 ; published 2026-08-31)	Fenopan CONSTITUTION.md §6.1
t_{MF} split 51.31% positive / 27.07% significant-negative — 7,000 pooled generations = 1,000 processes $\times$ generations 9–15, across additive AND multiplicative cascade types and five noise types , not 7,000 genuine cascades (§6, restated 2026-08-25)	<i>external</i>	Mangalam et al. 2025, PRE 111, 034126 (arXiv:2312.05653)
t_{MF} estimator/null/decision rule: Chhabra–Jensen direct method, 32 IAAFT surrogates, a two-sided threshold at 2.04 (df 31), lengths 2^8 – 2^{14} — not this paper's MF DFA-width + ~ 50 -surrogate rank point (§6)	<i>external</i>	Mangalam et al. 2025 §II.2
Lee & Kelty-Stephen σ -sweep: 200 series at each $\sigma \in [0.1, 1.1]$ step 0.05, two width estimators, narrowing “but more for MF-DFA” (§6)	<i>external</i>	Lee & Kelty-Stephen 2017, Complexity 2017:7015243
sublinear multifractality from nonmultiplicative constraints (allied mechanism, §6)	<i>external</i>	Mangalam & Kelty-Stephen 2026, PRE 114, 024215 — APS abstract; no preprint
isolated singularities artefactually broadening spectra (published high jaw, §1.2/§7)	<i>external</i>	Oświęcimka et al. 2020, Nonlinear Dyn 100:1689

number (as used)	record	file / section
multifractality tests with Monte-Carlo power, wavelet leaders (the demand, 2007; §1.2)	<i>external</i>	Wendt & Abry 2007, IEEE TSP 55(10):4811
<code>t_MF</code> formula	<i>external</i>	Bell et al. 2019, Front. Physiol. 10:998
prior art <code>w_orig < w_surr</code> , bidirectional reading	<i>external</i>	Lee & Kelty-Stephen 2017, Complexity 2017:7015243

Appendix B — Companion work, out of scope for this paper

Two companions are named here for the record and are **not** part of this preprint’s claim:

- **The attribution-power battery.** The measured deficit is on the *existence* test (width vs IAAFT). The three-way *attribution* decomposition (original vs shuffle vs FT vs IAAFT → “mostly fat tails / long-range correlation, not nonlinearity”) is what the largest cross-domain targets run (Zhou 2009, $\approx$ 196 citations; Baruník et al. 2012, $\approx$ 132; the rs-fMRI lineage). Its detection power on the canonical binomial cascade, at those papers’ `N`, is unmeasured; one synthetic session closes it and *calibrates this paper’s blast radius before a referee does*. The cross-domain target roster is maintained in `fenocosm/masses.py:MFDFA_CLAIMS`.
- **The cross-domain re-run census.** OpenAlex counts (fetched 2026-08-24): 314 works matching “multifractal AND surrogate” in title+abstract, 491 with “surrogate OR shuffled”, 3,834 on the MFDFA phrase — a hard *lower* bound (most papers name the surrogate only in Methods). The honestly re-runnable subset (public series + reported protocol) is realistically 50–200; the companion must state its stratum, never “the field.”

Neither companion changes any number in §§1–11; both are deferred so that this paper’s stratum — the width-vs-IAAFT existence test and its `t_MF` distillate — stays exactly measured.