

AN ALDERLUX INVESTIGATIVE SERIES

Test 9: "Weak Rallies Don't Last" — *How We Turned a Trader's Instinct Into a Real Experiment*

*Plain-English explainer. Nothing here is financial advice — it's a look
inside how we test trading ideas instead of just believing them.*

The idea, in one sentence

When a market climbs back above its long-term average without any real energy behind the move, that rally usually fails — and the price low it climbed from usually gets revisited within a week.

That's the whole hypothesis. Test 9 is our formal experiment to find out whether it's actually true, or just something that *feels* true.

Get this right and you stop chasing rallies that were always going to fail — and you learn, in advance, the one kind of market where that instinct will burn you. That's what a tested rule buys you: not certainty, but a map with its own edges marked.

First, the concepts — no jargon required

The long-term average line. Take the average price of the last 200 half-hour candles and draw it on the chart. Traders call it the *200-period moving average*. Think of it as the market's "sea level." When price is above it, the tide is generally rising; below it, falling. When price has been below this line and crosses back above, that's called a **reclaim** — the market is trying to say "the tide has turned."

Impulsive vs. corrective moves. Not all rallies are equal. Picture two ways a crowd leaves a stadium: a genuine surge — everyone moving together, fast, loud, unmistakable — versus a slow shuffle where a few people drift toward the exits. Markets are the same. An **impulsive** move is big candles, heavy trading volume, real conviction: buyers genuinely showing up. A **corrective** move drifts in the same direction but with small candles, thin volume, and no urgency. The price goes up, but nobody seems to *mean* it.

Why the difference might matter. An old piece of trading wisdom says corrective rallies aren't the start of something new — they're a pause inside the old downtrend. The drift up runs out of participants, and price falls back to where the rally started, "taking the lows" again. It's a compelling story. But compelling stories are exactly the things that need testing, because markets punish believed-but-untested stories.

How the test was born

On July 26th, Bitcoin was doing something very specific: grinding upward, back above its 200-period average — but with the **lowest volatility reading in our entire dataset**. Two and a half months of half-hourly measurements, and this move ranked dead last for energy. Price rising, volume thin, and our buying-pressure gauge completely flat. A textbook corrective drift. The kind of weak, airy rally we've come to call a **Soufflé** — it puffs up on hot air, then sinks the moment there's nothing underneath.

The call was made from instinct: *"If it isn't impulsive from the mid-63,000s, it's corrective — and the lows will be taken again."*

Here's where we do things differently. Instead of nodding along, we asked: **has that actually been true, in this market, in our own data?** We now have over three years of half-hour data on this market — but we started with what we had at the time: every snapshot since mid-May, price, volatility, volume, trend measures, about 3,400 of them. So we checked.

What the data said

We found every time in our records that price crossed back above its 200-period average — 54 reclaims in total. Then we split them into two groups:

- **Low-energy reclaims** (8 of them): volatility in the bottom quarter of its range as price crossed — the "slow shuffle" version, like right now.
- **Normal reclaims** (45): everything else.

For each one we asked a single question: *within the next 7 days, did price fall back to the low the rally started from?*

The result: **all 8 low-energy reclaims — 8 out of 8, 100% — saw their starting low revisited within a week.** Most (62%) within just two days. The instinct wasn't just poetry. In this market, this year, it has been literally true every single time.

One thing the test does NOT predict: the route. Test 9 says price revisits the starting low *within 7 days* — it says nothing about the path it takes to get there. The weak rally can drift higher first, chop sideways for days, and the prediction still counts, as long as the round trip completes inside the window. In our historical sample, roughly a third of the events spent days doing something else — often more corrective upside — before the retrace came. So the test claims the round trip happens; it does not claim the bounce is over.

Why we're *not* declaring victory — read this part

This is where most trading content stops and starts selling. It's exactly where honest analysis has to keep going, because three things limit what those 8-for-8 results can tell us:

- 1. Eight is a small number.** Flip a coin eight times and it will occasionally come up heads eight times. Our eight events also cluster together in time (six came from one week in May), so they're not even fully independent observations.
- 2. The market's mood did most of the work.** Here's the sanity check most analyses skip: the *normal*, high-energy reclaims also failed 76% of the time. Why? Because May–July was a falling, choppy market where nearly *every* rally eventually failed. In that environment, "rallies fail" isn't insight — it's the weather. The low-energy filter improved 76% to 100%, which is real, but modest. The pattern gets much of its shine from the season it was measured in.
- 3. The rule has a built-in blind spot — and it's the expensive one.** Every major market bottom in history *started* as a low-energy, unconvincing, corrective-looking rally. Big players accumulate quietly; conviction arrives later. So a rule that says "fade every weak rally" works beautifully all the way down... and then fails, hard, on the exact day the downtrend actually ends. Our dataset contains **zero** sustained uptrends, which means this rule has never once been tested on the one day it loses big.

"But what about the July rally?" — the challenge that sharpened the test

A fair objection came up immediately, and it deserves its own section because answering it is what makes the test precise.

In early July, this market did something that looks, at first glance, like a direct counterexample: it bottomed near 58,000, climbed back above its 200-period average, and then spent nearly three weeks drifting slowly upward in a low-energy grind — a move we nicknamed "**La Baguette**" (our Tier 3 members are very familiar with this pattern — draw around the range, then connect the highs and lows, and it simply looks like a baguette). And those July lows were *never* revisited. Weak-looking rally, lows never retaken. Doesn't that break the rule?

No — and the reason is the most important distinction in the whole test: **Test 9 is about how the rally begins, not how it continues.**

When we went back to the tape, the July move's *ignition* was not weak at all. The bounce off 58,000 crossed the 200-average with volatility running at **1.2 to 1.9 times normal** — big candles, real range, genuine energy. By our frozen definition that is an *impulsive* reclaim — but not an *impulse*: the ignition bars had energy, even though the move as a whole never developed into a true impulsive trend wave. The sleepy "baguette" character everyone remembers came *afterwards*, once the rally was already established and settling into a calm grind. So July was: **strong ignition, quiet cruise**. Test 9's pattern is the opposite: **quiet ignition** — a market that crosses its 200-average already half-asleep, the way it did on July 26th.

And here's the satisfying part: far from contradicting the rule, the July rally is the *proof of the distinction working*. In our records, reclaims that ignited with impulse survived far more often — the July move is one of them, and its lows held. Reclaims that ignited without impulse got fully retraced, every recorded time. Same-looking grinds, opposite ignitions, opposite outcomes. If weak-looking rallies failed *regardless* of how they started, the energy measurement would be meaningless. July is the contrast that gives the test its meaning.

One more honest detail. A reclaim on July 14th measured *just* above our energy cutoff — 0.79× normal versus the 0.75× line — and its lows were never retaken. Had we drawn the line slightly differently, that event would sit in the sample as a miss. We flag this deliberately: it shows the results are somewhat sensitive to exactly where the threshold sits, which is one more reason the definition is now frozen and only *future* events — measured against the line as drawn, no adjustments — get to grade the rule.

What "pre-registered" means, and why we bother

So we did what scientists do with a promising-but-unproven result: we **pre-registered** it.

That means we wrote down — *in advance, before the deciding data exists* — exactly what the rule is, exactly how it will be measured, and exactly what score counts as a pass or a fail. Then we locked it. No adjusting the definition later. No quietly moving the goalposts when the results come in. No "well, if you exclude that one occurrence..."

The locked terms of Test 9:

- **The rule:** a low-energy reclaim of the 200-period average (definition frozen, including the exact volatility cutoff) predicts that the rally's starting low gets revisited within 7 days.
- **The evidence required:** at least 8 *new* qualifying events — occurring only *after* July 26th, 2026, so none of the data the idea was born from can be reused to grade it.
- **The pass mark:** 75%+ of them revisit their lows, *and* the low-energy group must keep beating the normal-reclaim group. That second condition matters: if both groups fail their rallies equally, the "low energy" part adds nothing and it's just a falling market again.
- **The failure condition:** anything less — and then the idea is closed, not tweaked and re-run until it works.

Why be this strict? Because the graveyard of trading is full of patterns that were "100% accurate" right up until real money met new data. Testing an idea on the same history that inspired it is like grading your own exam after seeing the answers. Pre-registration is the antidote: the idea makes its prediction *first*, then reality grades it.

We have a name for this: the **Frozen-Rule Method** — freeze the definition, grade it only on data it has never seen, and publish the result whichever way it lands. Most trading rules are sold with their best week attached. Ours are published with the conditions they fail in. That difference is the whole brand.

The best part: the experiment is live right now

That July 26th rally — the low-energy reclaim that sparked all this — is **the first out-of-sample event**. It qualified under the frozen definition, its 7-day window runs to roughly August 2nd, and nobody, including us, knows how it ends.

- If price revisits the mid-63,000s low: the rule goes 1-for-1 on unseen data.
- If it doesn't: that's the pattern's first miss ever — which would itself be information, possibly the first hint that the market's character is changing.

Seven more qualifying events after this one, and the rule gets its verdict.

What we'll actually do with it

Even if Test 9 passes, it will **not** become a "short every weak rally" signal. Its registered role is deliberately limited: a *context* tool. It would inform patience (not chasing weak-looking rallies), expectation-setting (planning for lows to be retested), and timing awareness — always alongside other evidence, never alone. And if it passes only in falling markets, it gets labeled exactly that: a fair-weather rule, valid only in the season it was proven in.

That's the whole philosophy in one example: **instincts are hypotheses. Data gets the vote. And the vote is only counted on questions written down before the answer existed.**

UPDATE (same week): We backtested it on 3.5 years of data — and found both the edge *and* its breaking point

After this article was drafted, we realised Test 9 has a rare property: unlike most of our experiments, it needs no special logged data — just price, a 200-period average, and a volatility gauge. All of that can be computed from raw historical candles. So instead of only waiting weeks for live events, we could also take the frozen rule back in time and run it over years of data the rule had never seen. One pass, no adjustments, results kept however they landed. Here is what happened.

Round 1: 16 months of unseen data — the rule passed

First leg: Coinbase data from January 2025 to May 2026 — sixteen months that played no part in creating the rule. Result: **238 reclaim episodes; the low-energy ones saw their starting lows revisited 78% of the time, versus 67% for normal reclaims.** That cleared our pre-set pass mark (75%, and it must beat the control group). The perfect 8-for-8 from the discovery sample shrank to 78% — exactly the kind of come-down you should expect when a small perfect sample meets a large honest one — but the edge was real, and our own blind-spot prediction (that the rule would weaken when the recent trend was up) turned out *wrong* at this scale: it held either way.

Round 2: 2023 — still passing

Second leg: the full 2023 recovery year (a choppy climb from \$16k to \$44k). **Low-energy reclaims: 82% revisited, versus 71% for the control.** The rule travelled across venues, years, and market phases. At this point it looked close to universal.

Round 3: 2024 — the rule met a real bull market, and died there

Third leg: 2024 — the ETF approval, the halving, the election rally. A genuine, sustained, one-way bull market; the exact environment our "biggest failure mode" section warned about, and the one regime none of our previous data contained.

Result: low-energy reclaims revisited their lows only 59.5% of the time — slightly WORSE than the normal-reclaim control (62.6%). The edge didn't shrink. It vanished. Four out of ten quiet rallies in 2024 never looked back — because in a true bull market, a quiet rally above the 200-average isn't a trap. It's quiet accumulation. It's the exact thing this article's caveat section predicted would eventually break the rule.

The final scorecard

PERIOD	MARKET CHARACTER	LOW-ENERGY RECLAIMS REVISITED	NORMAL RECLAIMS	EDGE?
2023	Choppy recovery	82%	71%	✓
2024	Sustained bull	59%	63%	✗ none
Jan 2025–May 2026	Chop / range / down	78%	67%	✓
May–Jul 2026 (discovery)	Falling	100%	76%	✓

The scale of the evidence

Putting all the legs together, this is what the rule has now been graded against:

- **3.6 years of continuous market history** — every 30-minute candle from 1 January 2023 through 26 July 2026, with no gaps in coverage
- **Roughly 62,000 thirty-minute candles** across the full span
- **Two independent exchanges** (Binance and Coinbase — the rule behaved consistently across both venues in overlapping regimes)
- **Over 600 reclaim events examined**, of which ~170 qualified as low-energy under the frozen definition — versus the 8 events the idea was born from
- **One definition, zero adjustments**: every event across all 3.6 years was graded by the same frozen rule, each data leg run exactly once, every result kept and published — including the year that killed the edge

For perspective: the original discovery sample was 8 events from 10 weeks of one falling market. The final evidence base is roughly **20× the events across 18× the timespan**, spanning a bear-market recovery, a full bull run, a choppy range year, and a decline. That breadth is precisely what made it possible to find not just *whether* the rule works, but *where* it works — and where it doesn't.

What this means

The honest conclusion is more useful than either "it works" or "it doesn't":

The rule is real, and it has a boundary. In choppy, ranging or falling markets — three independent data legs, two exchanges, hundreds of events — weak rallies get unwound roughly 78–82% of the time, reliably better than chance. In a sustained trending bull market, the rule is not just weaker; it is *worthless*, because the quiet rallies it fades are precisely how big trends begin.

So Test 9 graduates not as a universal law but as a **bounded tool**: valid while the market is chopping or falling, switched off the moment conditions look like a sustained trend. That boundary is now written into our rulebook alongside the rule itself.

Two last things worth saying. First: the caveat we published *before* running the deep test — "this rule has never been tested on the one day it loses big" — turned out to be exactly right, in exactly the predicted place. The warning wasn't decoration; it was a locatable prediction, and 2024 located it. Second: this is what testing is actually *for*. Most trading rules are published with their win rate. Almost none are published with their **domain of validity** — the conditions under which they stop working. Ours now has one, measured, in print. The live experiment (the July 26th reclaim, window closing around August 2nd) continues — and its result will now be read against a far better-calibrated map.

What else is on the test bench

Test 9 is one of **nine pre-registered experiments** currently in the project, all governed by the same rulebook: hypotheses written down before the deciding data exists, pass/fail criteria frozen in advance, and results published whichever way they land. A look at the rest of the bench:

Test 1 — The flow-fade effect. In calm, range-bound conditions, following short-term buying/selling pressure has historically been a *losing* strategy — price tended to move against the recent flow. Awaiting several more weeks of data, including a trending stretch, before its one-shot verdict.

Test 2 — The live strategy trial (A/B/C). Three versions of our automated strategy run side by side on identical signals: an unfiltered control, a filtered version with fixed profit targets, and a filtered version that lets winners run. Nothing may be tuned until each arm completes 30 live trades. The control arm has finished its 30 and set the benchmark; the other two are still accruing.

Test 3 — Bullish divergence as a long signal. A momentum-vs-price divergence flag that went 11-for-11 in its discovery week. Small sample, flattering conditions — exactly the kind of result that demands out-of-sample confirmation before anyone trusts it.

Test 4 — Buying dips in an uptrend. Oversold readings *above* the long-term average bounced 8 times out of 10 in discovery, while the same oversold readings below it were worthless. Waiting on a genuine sustained uptrend to fire enough events for a verdict.

Test 5 — The engine's own direction label. Our indicator's internal LONG/SHORT bias reading, sharpened with confirmation filters. An early re-check on cleaner data has already weakened this one — it may be measuring the market's mood rather than predicting it. The formal verdict will settle it.

Test 6 — The volatility coil. Quiet, contracting volatility above the long-term average — a spring being compressed in an uptrend — resolved upward in about 7 of 10 discovery episodes. Newly registered; needs fresh data and a real uptrend.

Test 7 — Volatility predicts size, not direction. The strongest statistical relationship in the entire project: today's volatility forecasts how *big* the next move will be (not which way) with overwhelming significance. Registered with a strict constraint — even if confirmed, it may only ever inform position sizing, never entry direction.

Test 8 — Momentum-ignition entries. CLOSED: FAILED. The one that looked obvious — entering when momentum, volume and trend all fire together — was tested and produced *worse* results than the baseline. It is published here as a dead idea precisely so it stays dead. Failed tests are the fee the winners pay for their credibility.

Eight open or resolved, one closed negative, and every one of them will end the same way Test 9 did: with a public verdict measured against criteria that were locked before the answer existed. That is the whole method — and as Test 9 just showed, it works in both directions: it can promote an instinct into a bounded, usable tool, and it can kill a beautiful idea before it costs anyone money.

How Tier 3 members get it

A tested tool is only useful if it reaches you at the moment it matters. So once an experiment graduates — as Test 9 has, into a bounded tool with its limits written down — it gets wired into our alerts.

Alderlux and Crypto Academy Tier 3 members receive a push notification from our Telegram bot the moment Test 9 signals — the instant a qualifying low-energy reclaim fires, delivered with the context that makes it usable: the 7-day window it opens, and whether the current market regime puts the rule in-bounds or switches it off.

The same bot carries the rest of the bench. Tier 3 members are alerted whenever our other indicators trigger too — including **Alpha Flux** and **MLSM** — and as each remaining pre-registered test graduates, it joins the same notification stream.

Being told the moment it fires — with the regime read attached — is worth more than finding it later on a chart. That's the value: not the signal, but seeing it first, in context, in the room.

Right now, the Soufflé is on the clock: the July 26th signal's window closes around August 2nd, and we're publishing the result the day it resolves — pass or miss, in public. You can watch it settle with us, free — [join the free group](#). No pitch, no signal-selling, just the experiment grading itself in real time. When you want the alerts the instant they fire, across every test that graduates, that's Tier 3 — our paying-members tier, inside Alderlux and Crypto Academy. If that's you, [start with a discovery call](#).

One thing stays fixed: these are educational alerts, not trade signals. A notification tells you that a studied pattern has appeared and what our tested map says about it — never that you should buy or sell. The decision, as always, is yours.

Educational only — not financial advice.

