

Synthetic Momentum Strategy v1

Strategy Analyzer Standard Report

\$1,020.61
Net P&L

317
Total Trades

65.0%
Win Rate

1.09
Profit Factor

2024-01-02
Start Date

2024-10-20
End Date

Green
MC Signal

2026-07-07
Generated

test-run-000
Run ID

REPORT MAP

Table of Contents

How this report is organized

Report mode: **Standard**. Section numbers are reserved across both report modes, so gaps below (e.g. 8 to 11) mean an optional section isn't available for this run or mode — not missing content. See Methodology & Glossary for details.

0	Decision Brief	Strategy verdict, score, area ratings, radar, key findings, edge concentration
0a	Strategy Health Dashboard	Consolidated verdict, score, badges, confidence, MC/Apex status, warnings
0b	What We Learned	Plain-English synthesis of the report's key findings
0c	Recommendation Engine	Prioritized, evidence-backed recommendations with trade-offs
0d	Rule Test Dashboard	Highest-impact counterfactual rule tests from the stop simulator
0e	Next Research Queue & Validation Plan	Falsifiable follow-up tests with hypothesis, experiment, and pass/fail criteria
0f	Dataset & Assumptions	Symbol, timeframe, date range, commission/slippage assumptions, row accept/reject counts
1	Executive Summary	Key metrics dashboard, performance narrative, MC signal
2	Core Metrics	Full statistics: win rate, PF, Sharpe, Sortino, Calmar
3	Equity Curve & Drawdown	Equity with drawdown overlay chart, top drawdown episodes
4	Split Robustness	Train / Validation / Holdout chronological split, Apex score
5	Session Analysis	P&L and win rate by trading session, session bar chart, hour-of-day breakdown
6	Daily P&L	Day-by-day P&L bars, distribution stats, worst days
7	Rolling State	Rolling profit factor chart, state breakdown
8	Telemetry Tags	GRADE, ACT, WHY, CONF, EDGE tag performance tables
9	Multi-Model Regime	M_CONS, M_AGR, M_ETR, M_VREG model alignment
10	Structure Analysis	BOSDIR, SRSUPD and structure tag performance
11	Monte Carlo / Apex	Scenario matrix, pass-rate chart, survival probability
12	Stop / Exit Simulator	MAE loss-cap and breakeven trigger what-if tables
14a	Worst Days	Top 10 worst trading days
15	Methodology & Glossary	Term definitions and how each metric is computed

Generated by TradingEdgeIQ Strategy Analyzer — tradingedgeiq.com — For informational purposes only. Past performance is not indicative of future results.

SECTION 0

Decision Brief

STRATEGY VERDICT: NEEDS WORK

Key performance metrics are below viable thresholds. Prioritize the Recommendation Engine findings before further deployment.

NET P&L

\$1,020.61

PROFIT FACTOR

1.09

WIN RATE

65.0%

MAX DRAWDOWN

\$2,121.46

AVG P&L / TRADE

\$3.22

STRATEGY SCORE

43/100

Key Findings

- 1. Profit factor of 1.09 is below 1.2 — the strategy may not have a reliable edge on the current data.
- 2. Net P&L / Max DD ratio of 0.5x is poor — max drawdown exceeds net gain, creating capital-efficiency risk.
- 3. Dataset contains 317 trades with a win rate of 65.0% — a large enough sample for reasonably stable conclusions.

Area Ratings

- **Edge Strength** (50/100) **WEAK**
Evidence: Profit factor 1.09 and win rate 65.0% across 317 trades.
Main concern: Profit factor is below the 1.2 reliability floor.
Edge is present but thin; small changes in costs or slippage could erase it.
- **Risk Management** (16/100) **WEAK**

Risk Management score is driven by recovery factor (net P&L ÷ max drawdown) of 0.5x. This sub-score can differ from headline max-drawdown / recovery-factor figures shown elsewhere because it is scaled against a fixed benchmark (recovery factor of 20x = 100/100) rather than reported as a raw ratio.

Evidence: Recovery factor (net P&L ÷ max DD) of 0.5x on max drawdown \$2,121.46.
Main concern: Max drawdown of \$2,121.46 exceeds net P&L — capital efficiency risk.
Drawdown control is the main lever to improve before increasing size.
- **Trade Efficiency** (100/100) **GOOD**

Evidence: Average trade cadence: 1.3 trades/day.
Main concern: No material concern at current sample size.
Trade frequency is in a manageable range for consistent execution.
- **Robustness** (46/100) **WEAK**

Evidence: Composite of profitability, consistency, drawdown control, and trade efficiency sub-scores.
Main concern: No material concern at current sample size.
One or more underlying dimensions is dragging down overall robustness — see the weakest sub-score above.

Baseline vs. Filtered — at a Glance

Metric	Baseline	Filtered	Delta
Trades	420	317	-103
Net P&L	-\$834.07	\$1,020.61	+\$1,854.68
Profit Factor	0.95	1.09	+0.138
Win Rate	59.1%	65.0%	+5.9pp
Avg P&L / Trade	-\$1.99	\$3.22	+\$5.21
Max Drawdown	-\$3,156.39	-\$2,121.46	+\$1,034.93

The filtered set retains 75.5% of baseline trades (317 of 420).

Decision-Summary Radar



Edge Concentration Cards

BEST SESSION New York AM PF 1.91 · Net \$1,572.08	WEAKEST SESSION Asia PF 0.77 · Avg -\$12.22	BEST GRADE A PF 1.99 · Net \$3,960.39
LOWEST PF GRADE C PF 0.52 · Avg -\$18.01 · Ranked by profit factor — see Telemetry Tag Breakdown for full grade comparison.	BEST ACT Enter PF 1.90 · Net \$2,196.11	LOWEST PF ACT Scale PF 0.87 · Avg -\$6.11 · Ranked by profit factor — see Telemetry Tag Breakdown for full grade comparison.

Immediate Action Plan

HIGH PRIORITY Preserve exposure during New York AM Protects the main profit engine before optimizing weaker areas.
HIGH PRIORITY Test reduced size or blocking during Asia May reduce weak-edge exposure and smooth the equity curve.
HIGH PRIORITY Keep A-grade logic as the primary strategy engine Avoids damaging the core edge while tuning weaker states.
How to read this report Start with this Decision Brief and the Recommendation Engine, then use the Session, Telemetry, and Monte Carlo sections to decide which rule changes deserve controlled testing. This Standard report doesn't include raw trade-level tables — regenerate in Full Audit mode, or use the Data Export Pack, if you need the underlying data for auditability.

SECTION 0A

Strategy Health Dashboard

VERDICT: NEEDS WORK · SCORE 43/100 · BADGE: NEEDS WORK

STRATEGY SCORE

43/100

HEALTH BADGE

NEEDS WORK

EDGE STRENGTH

50/100

RISK MANAGEMENT

16/100

TRADE EFFICIENCY

100/100

ROBUSTNESS

46/100

SAMPLE-SIZE CONFIDENCE

High

n=317

DATA-QUALITY CONFIDENCE

Good

MC SURVIVAL STATUS

Survives (78.5% pass rate)

APEX/PROP-STYLE VIABILITY

Not yet viable

Warning Flags

- Profit factor 1.09 is below the 1.2 reliability floor.
- Max drawdown exceeds net P&L (recovery factor < 1.0).
- Edge has fully reversed from Train to Holdout — average P&L flips from positive to negative. (see Section 4 — Split Robustness for details)

Strengths

- Consistency (72/100)
- Trade Efficiency (100/100)

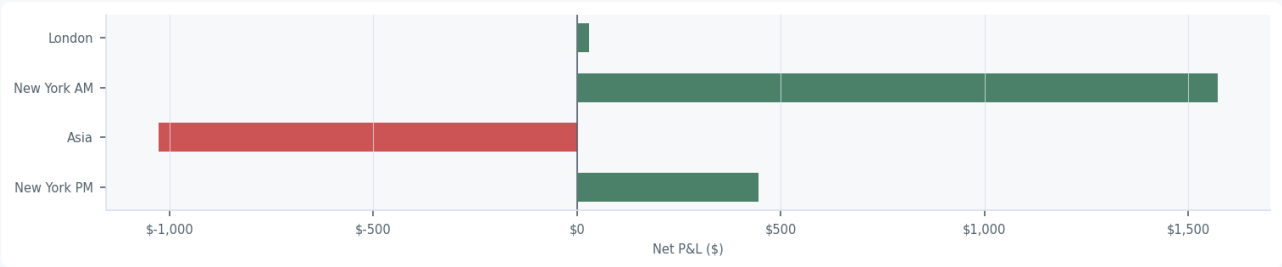
Weaknesses

- Profitability (27/100)
- Drawdown Control (16/100)
- Robustness (46/100)

What This Means

In plain terms: this strategy is currently rated **NEEDS WORK** (43/100). Key metrics are below viable thresholds — treat this as a strategy that needs rule changes, not one ready for live capital. Sample-size confidence is **High**, so the sample is large enough to support reasonably stable conclusions.

Net P&L by Session (preview)



Full breakdown in Section 5 — Session Analysis.

SECTION 0B

What We Learned

This strategy processed **317** trades and produced a net P&L of **\$1,020.61**, with a profit factor of **1.09** and a win rate of **65.0%**. Taken together with a max drawdown of **\$2,121.46**, the headline picture is a strategy that has not demonstrated a reliable edge on this data.

Under Monte Carlo stress-testing, the strategy survived the majority of simulated paths (78.5% pass rate). This is a strong robustness signal.

With 317 trades, the sample size is large enough to support reasonably stable conclusions, though continued monitoring is still recommended as market conditions evolve.

Bottom line: the strategy currently rates as **NEEDS WORK** (43/100 composite score). The priority is revisiting the strategy rules themselves — the current metrics do not support live deployment.

SECTION 0C

Recommendation Engine

How to use this section

These are automatically generated hypotheses for what to preserve, reduce, or test next. Treat them as research priorities, not trading instructions. Validate any rule change with forward testing before using it live.

Top Research Priorities

HIGH · HIGH CONFIDENCE · PRESERVE / PROTECT

Recommendation: Preserve exposure during New York AM

Evidence: PF 1.91, net \$1,572.08, win rate 67.1% across 82 trades.

Expected Impact: Protects the main profit engine before optimizing weaker areas.

Trade-off / Watchout: Do not over-filter; protect this edge while testing weaker sessions in parallel.

HIGH · HIGH CONFIDENCE · SESSION SEGMENTATION

Recommendation: Test reduced size or blocking during Asia

Evidence: Weakest session PF 0.77, avg trade -\$12.22, max DD -\$819.80.

Expected Impact: May reduce weak-edge exposure and smooth the equity curve.

Trade-off / Watchout: Could reduce weak-edge exposure, but may remove recovery/diversification trades — validate out-of-sample.

HIGH · HIGH CONFIDENCE · PRESERVE / PROTECT

Recommendation: Keep A-grade logic as the primary strategy engine

Evidence: A-grade generated \$3,960.39 (388.0% of net P&L) with PF 1.99.

Expected Impact: Avoids damaging the core edge while tuning weaker states.

Trade-off / Watchout: A-grade can still contain tail losses; tune risk controls without over-filtering winners.

HIGH · MEDIUM CONFIDENCE · STOP / PATH RULE TEST

Recommendation: Test a MAE-based loss cap around \$-30

Evidence: Simulation shows net change \$9,469.92, PF 5.18, max DD -\$30.00.

Expected Impact: Potentially improves both P&L and drawdown by cutting early adverse excursions.

Trade-off / Watchout: This is a counterfactual approximation; validate with full execution logic and forward testing.

MEDIUM · MEDIUM CONFIDENCE · SIGNAL-QUALITY SEGMENTATION

Recommendation: Review C-grade / lower-confidence trades for selective filtering or smaller size

Evidence: C-grade PF 0.52, max DD -\$424.22, net -\$1,854.68.

Expected Impact: Potential drawdown reduction if weak C-grade clusters can be isolated.

Trade-off / Watchout: C-grade may still contribute meaningful profit; avoid a blanket block until tested.

Prioritized Recommendations — Full Detail

Priority	Confidence	Experiment	Recommendation	Evidence	Expected Impact	Trade-off / Watchout
High	High	Preserve / protect	Preserve exposure during New York AM	PF 1.91, net \$1,572.08, win rate 67.1% across 82 trades.	Protects the main profit engine before optimizing weaker areas.	Do not over-filter; protect this edge while testing weaker sessions in parallel.
High	High		Test reduced size or blocking during Asia	Weakest session PF 0.77, avg trade -	May reduce weak-edge exposure and	Could reduce weak-edge exposure, but may

Priority	Confidence	Experiment	Recommendation	Evidence	Expected Impact	Trade-off / Watchout
		Session segmentation		\$12.22, max DD - \$819.80.	smooth the equity curve.	remove recovery/diversification trades — validate out-of-sample.
High	High	Preserve / protect	Keep A-grade logic as the primary strategy engine	A-grade generated \$3,960.39 (388.0% of net P&L) with PF 1.99.	Avoids damaging the core edge while tuning weaker states.	A-grade can still contain tail losses; tune risk controls without over-filtering winners.
High	Medium	Stop / path rule test	Test a MAE-based loss cap around \$-30	Simulation shows net change \$9,469.92, PF 5.18, max DD - \$30.00.	Potentially improves both P&L and drawdown by cutting early adverse excursions.	This is a counterfactual approximation; validate with full execution logic and forward testing.
Medium	Medium	Signal-quality segmentation	Review C-grade / lower-confidence trades for selective filtering or smaller size	C-grade PF 0.52, max DD -\$424.22, net -\$1,854.68.	Potential drawdown reduction if weak C-grade clusters can be isolated.	C-grade may still contribute meaningful profit; avoid a blanket block until tested.

SECTION 0D

Rule Test Dashboard

Purpose

This page promotes the highest-impact counterfactual tests from the Stop / Exit Simulator and session breakdown. These results are hypotheses from historical telemetry, not trading instructions — validate in forward testing before promoting any rule change.

BEST MAE LOSS-CAP

\$-30

Δ Net \$9,469.92 · PF 5.18 · DD -\$30.00

PRESERVE WINDOW

New York AM

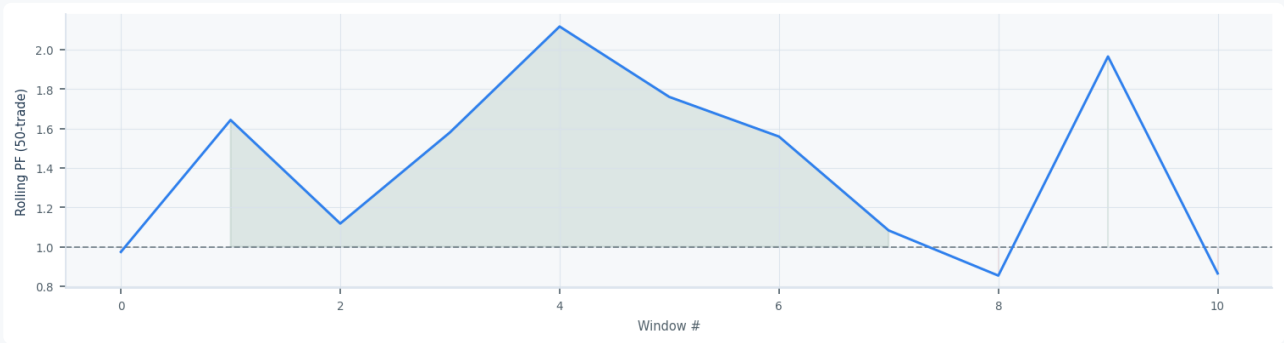
PF 1.91 · Net \$1,572.08

REDUCE / SEGMENT WINDOW

Asia

PF 0.77 · Avg -\$12.22

Rolling Profit Factor (preview)



Full breakdown in Section 7 — Rolling State.

SECTION 0E

Next Research Queue & Validation Plan

How to use this section

Each card is a falsifiable test, not a trading instruction: a hypothesis, the specific experiment that would confirm or reject it, and the pass/fail bar that must be cleared before adopting any resulting rule change live.

Excluding or down-sizing trades during Asia improves risk-adjusted return without materially reducing opportunity.
HIGH PRIORITY**WHY IT MATTERS**

This session currently drags down blended profit factor and may be adding variance without commensurate edge.

TEST

Re-run the backtest with Asia excluded (and separately at half size) and compare PF, net P&L, and max drawdown to baseline.

MODULE

Strategy Analyzer (filtered re-run)

EFFORT

Low — single filter re-run, no new data required.

BENEFIT

Potential PF and drawdown improvement if the session is genuinely low-edge.

RISK

May remove diversification or recovery trades that offset losses elsewhere; validate out-of-sample before adopting permanently.

MIN SAMPLE

25 required (have 84)

PASS/FAIL CRITERIA

Adopt only if filtered PF improves by ≥ 0.15 AND max drawdown does not worsen, using at least 25 held-out trades from Asia.

Sufficient sample to test now.

GRADE=C trades are a net drag and should be filtered or routed to a smaller position size.
MEDIUM PRIORITY**WHY IT MATTERS**

Confirms whether the GRADE tag is a genuinely predictive quality signal or noise, before building automation around it.

TEST

Compare full-sample vs. GRADE=C-excluded metrics; check consistency across at least 3 rolling windows.

MODULE

Strategy Analyzer + Rolling State

EFFORT

Low — tag already captured in telemetry.

BENEFIT

Cleaner rule for what to trade automatically vs. manually review.

RISK

Grade may correlate with session/time-of-day rather than being causal on its own.

MIN SAMPLE

25 required (have 103)

PASS/FAIL CRITERIA

Adopt only if PF improves by ≥ 0.15 with ≥ 50 trades in the excluded grade across at least 2 non-overlapping time windows.

Sufficient sample to test now.

Weak-state windows on this run remain net profitable in isolation, so a governor rule is a monitoring/variance-reduction idea rather than a clear improvement lever. On this run, 50 trades (15.8% of 317 total) fell into a Weak rolling-PF state (rolling window=50, PF below the weak-state threshold), producing PF 1.50 and net \$64.99 in isolation during those windows.

LOW PRIORITY

WHY IT MATTERS

Weak-state trades are net positive on this run — a governor rule would mainly reduce variance, not clearly improve P&L, so this is a lower-priority monitor-and-test-cautiously item rather than a strong recommendation. 50 weak-state trades available — sufficient for Medium confidence, but below the 75-trade threshold required for the strongest confidence tier.

TEST

If pursued, monitor the rolling PF state over more trades and test cautiously with a small sample before committing size changes — do not assume a governor will improve results given weak-state trades are currently net profitable.

MODULE

Strategy Analyzer (Rolling State) + Prop-Firm Simulator

EFFORT

Medium — requires a rule-based re-simulation, not just a filter.

BENEFIT

Primarily reduces variance in an already-profitable weak-state window; upside is secondary to protecting against a future shift to net-negative.

RISK

A lagging indicator (rolling PF) may only detect weak states after most of the damage is already done.

MIN SAMPLE

50 required (have 50)

PASS/FAIL CRITERIA

Adopt only if the governor reduces max drawdown by $\geq 15\%$ while retaining $\geq 80\%$ of baseline net P&L, tested on ≥ 50 trades (50 weak-state trades available — sufficient for Medium confidence, but below the 75-trade threshold required for the strongest confidence tier.)

Sufficient sample to test now.

SECTION 0F

Dataset & Assumptions

Why this matters

Every metric in this report is computed from the data and assumptions listed below. A different symbol, timeframe, or commission/slippage assumption would change the downstream numbers even if the underlying strategy logic is identical.

Field	Value
Symbol / Instrument	MNQ
Timeframe	5m
Date Range	2024-01-02 → 2024-10-20
Trade Count (this run)	317
Commission Assumption	\$0.62
Slippage Assumption	\$0.25
Telemetry Schema Detected	universal_pipe

Rows Accepted / Rejected During Upload

Stage	Count
Raw rows read from file	842
Entry rows recognized	421
Exit rows recognized	421
Duplicate rows dropped	3
Zero-P&L rows dropped	1
Malformed rows dropped	0
Total rows dropped	4

Optional telemetry fields not detected in this dataset: **confidence, regime**. Some optional module tags are missing; see Data Coverage for details.

1

Executive Overview

Performance dashboard, core metrics, and health assessment.

SECTION 1

Executive Summary

NET P&L

\$1,020.61

WIN RATE

65.0%

PROFIT FACTOR

1.09

MAX DRAWDOWN

-\$2,121.46

TOTAL TRADES

317

EXPECTANCY / TRADE

\$3.22

SHARPE RATIO

1.22

CALMAR RATIO

2.13

SORTINO RATIO

1.76

MC SIGNAL

Green

Strategy generated **\$1,020.61** net P&L across **317** trades, with a win rate of **65.0%** and profit factor of **1.09**. Maximum drawdown: **-\$2,121.46**. Expectancy per trade: **\$3.22**. Monte Carlo 10th–90th percentile: **\$408.24 – \$1,632.98** over **1000** paths.

Monte Carlo Signal — Green

Pass probability: **78.5%**. Strategy passes the survival threshold in majority of simulated paths.

SECTION 2

Core Metrics

Health Assessment

Strategy is marginally profitable. Improve average win or reduce average loss to strengthen edge.

Metric	Value
Total Trades	317
Winning Trades	206
Losing Trades	111
Win Rate	65.0%
Net P&L	\$1,020.61
Gross Win	\$13,001.94
Gross Loss	\$11,981.33
Profit Factor	1.09
Expectancy / Trade	\$3.22
Avg Win	—
Avg Loss	—
Win / Loss Ratio	—
Max Win	—
Max Loss	—
Max Drawdown	-\$2,121.46
Recovery Factor	—
Sharpe Ratio	1.220
Sortino Ratio	1.760
Calmar Ratio	2.130
Health Badge	—
Avg Trade Duration	—

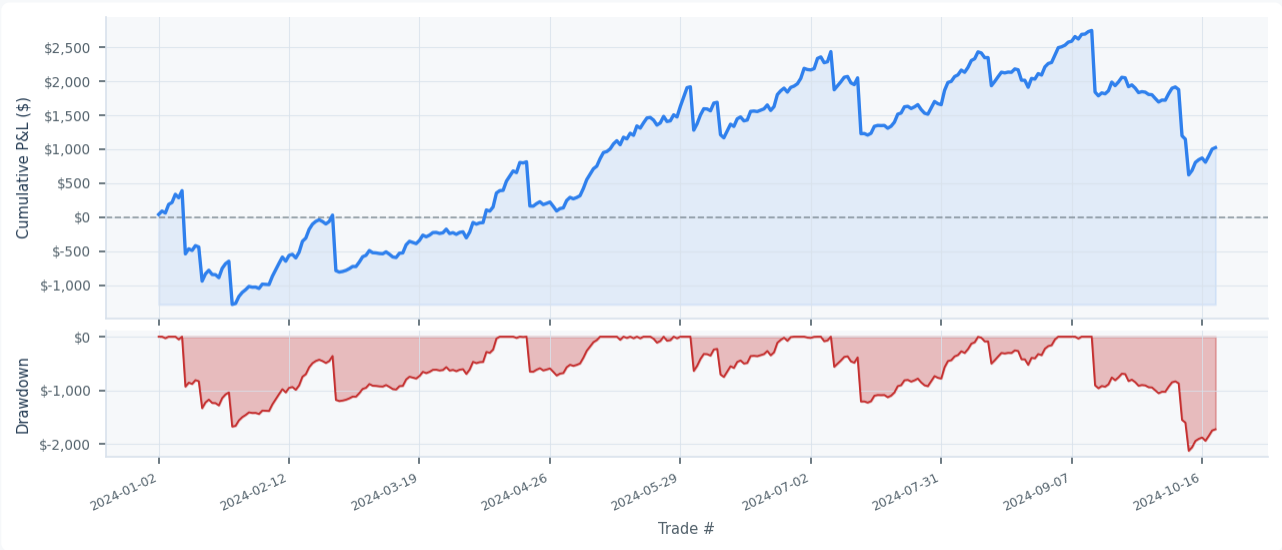
2 Risk & Robustness

Equity curve, drawdown, and chronological split robustness.

SECTION 3

Equity Curve & Drawdown

Equity Curve with Drawdown Overlay



Top panel: cumulative P&L. Bottom panel: underwater drawdown from peak. Deep red zones indicate worst drawdown periods.

Top Drawdown Episodes

Start (trade #)	Trough (trade #)	Recovery (trade #)	Peak Equity	Trough Equity	DD Amount	DD %	Duration
-----------------	------------------	--------------------	-------------	---------------	-----------	------	----------

No material drawdown episodes were detected in this trade history — the equity curve had no peak-to-trough decline that met the minimum episode threshold.

SECTION 4

Chronological Split Robustness

Purpose

Splits the trade history chronologically into Train (70%), Validation (15%), and Holdout (15%). A strategy with genuine edge should show positive average P&L and PF > 1.0 in all three splits. Green rows = positive avg; red rows = negative avg.

Split	Trades	Avg P&L	PF	Win%	Net	Max DD
Train (70%)	221	\$6.31	1.17	67.0%	\$1,394.38	\$928.90
Validation (15%)	48	\$20.53	1.93	62.5%	\$985.58	\$413.18
Holdout (15%)	48	-\$28.32	0.52	58.3%	-\$1,359.35	\$905.29
All	317	\$3.22	1.09	65.0%	\$1,020.61	\$928.90

APEX READINESS SCORE

26.5

Composite: holdout×2 + val×1.5 + PF + log(n) – DD

ALL SPLITS POSITIVE

No

Train, val, and holdout avg > 0

APEX VIABLE

No

n≥75, PF≥1.20, holdout PF≥1.15, DD≤\$1,800

Split Robustness Result

One or more chronological splits has negative average P&L, suggesting the edge may not be stable across all market periods.

Holdout Quality

Weak — holdout average P&L is negative (-\$28.32), the most recent trades are not confirming the earlier edge.

Degradation Warning

Edge has fully reversed from Train to Holdout — average P&L flips from positive to negative.

Apex / Prop-Style Viability Interpretation

Apex/prop-style viability interpretation: This strategy does not yet meet the standard Apex-style viability bar. See the Split Robustness table above for which specific threshold(s) are not met before considering prop-style evaluation.

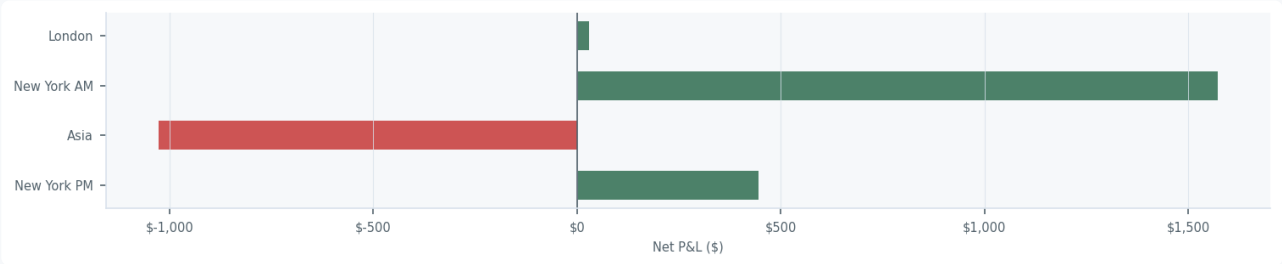
3 Trade Analytics

Session, daily P&L, recent-period, and rolling state diagnostics.

SECTION 5

Session Analysis

Net P&L by Session



Bars left of zero (red) indicate unprofitable sessions. Consider filtering or sizing differently per session.

Session Edge

Highest Net P&L Session: New York AM (\$1,572.08, PF 1.91, n=82). **Lowest Net P&L Session (n≥25):** Asia (-\$1,026.21, n=84). Basis: ranked by net P&L, restricted to sessions with n≥25 trades. Consider whether poor sessions should be filtered or traded smaller.

Session	Trades	Net P&L	Win%	PF	Avg	Max DD
London	74	\$30.09	59.5%	1.01	\$0.41	-\$678.80
New York AM	82	\$1,572.08	67.1%	1.91	\$19.17	-\$928.90
Asia	84	-\$1,026.21	64.3%	0.77	-\$12.22	-\$819.80
New York PM	77	\$444.65	68.8%	1.16	\$5.77	-\$905.29

SECTION 6

Daily P&L

TRADING DAYS

237

% POSITIVE DAYS

66.2%

BEST DAY

\$238.91

WORST DAY

-\$905.29

AVG DAILY P&L

\$4.31

MEDIAN DAILY

\$26.25

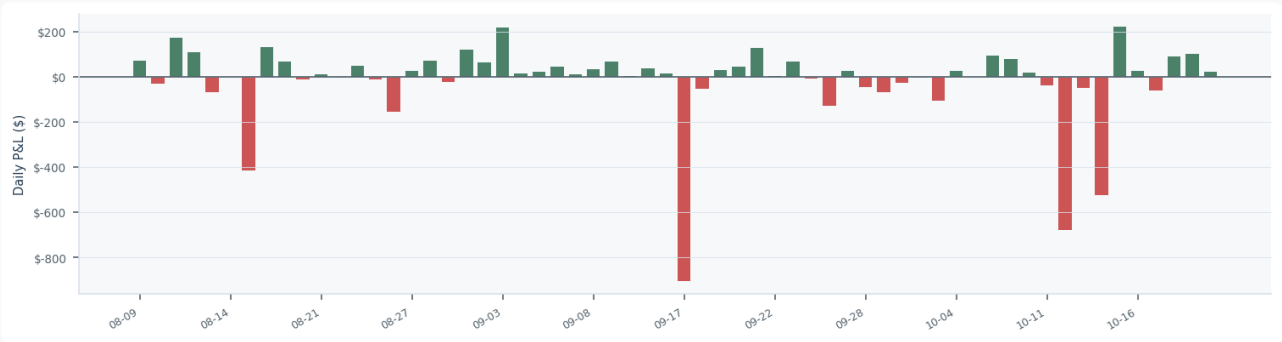
P5 DAILY

-\$413.18

P95 DAILY

\$171.12

Daily P&L (last 60 trading days)



Green bars = profitable days. Red bars = losing days.

Daily Distribution

Positive average daily P&L (\$4.31) but only 1% of days profitable. Check distribution.

Recent 15 Trading Days

Date	Trades	Daily P&L	Cum P&L	W / L
2024-10-04	1	\$27.85	—	1/0
2024-10-05	1	-\$1.18	—	0/1
2024-10-08	1	\$94.82	—	1/0
2024-10-09	1	\$80.22	—	1/0
2024-10-10	1	\$18.89	—	1/0
2024-10-11	1	-\$38.67	—	0/1
2024-10-12	1	-\$678.80	—	0/1
2024-10-13	1	-\$49.72	—	0/1
2024-10-14	1	-\$524.22	—	0/1
2024-10-15	3	\$222.33	—	3/0
2024-10-16	1	\$25.17	—	1/0
2024-10-17	1	-\$61.78	—	0/1
2024-10-18	1	\$91.85	—	1/0
2024-10-19	1	\$100.04	—	1/0
2024-10-20	1	\$22.54	—	1/0

SECTION 6B

Recent-Period Diagnostics

Purpose

Recent trailing-window performance is compared against the strategy's overall profit factor to flag whether the edge is holding up ("Confirms edge") or fading ("Weakens edge") in the most current data, rather than relying solely on the full-history average.

Window	Days	Net P&L	Win%	PF	Max DD	Verdict
Last 30 Days	26	-\$835.44	50.0%	0.52	-\$1,433.23	Weakens edge (sharply)
Last 60 Days	49	-\$1,109.90	61.2%	0.62	-\$2,121.46	Weakens edge
Last 90 Days	72	-\$599.83	62.5%	0.83	-\$2,121.46	Weakens edge

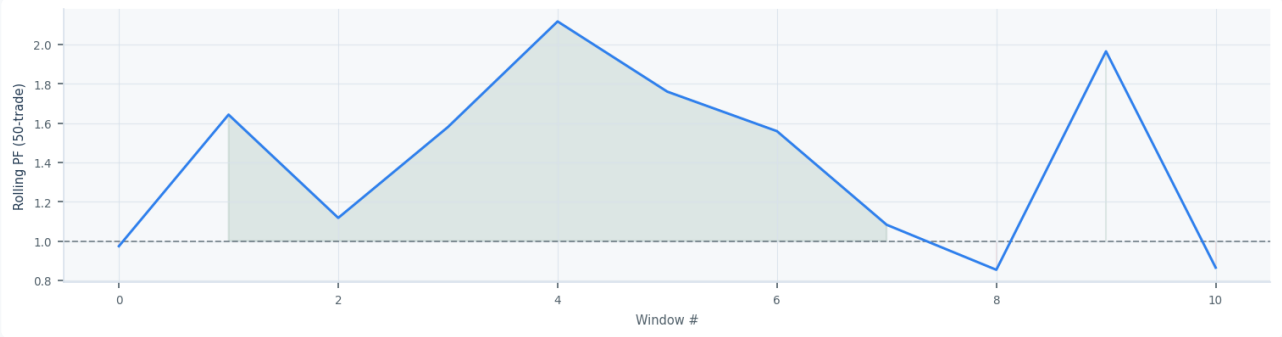
Year-over-year comparison unavailable — dataset spans only a single calendar year.

SECTION 7

Rolling State Diagnostics

Window: 50 trades | Weak PF threshold: 1.2 | Strong PF threshold: 2.0

Rolling Profit Factor



PF above 1.0 (green fill) = profitable window; below 1.0 (red fill) = losing window.

Rolling State Assessment

45% of rolling windows in Weak state. Consider filters to reduce trading during weak periods.

State	Trades	Net P&L	Win%	PF	Avg
Weak	50	\$64.99	—	1.50	—
Strong	50	-\$1,235.74	—	0.56	—

SECTION 7B

Weak-State Diagnostics

Definition

A trading window is classified **Weak** when its rolling profit factor over the most recent 50 trades falls below the weak-PF threshold (1.2). This is a lagging, PF-based definition — it flags periods after the edge has already softened, not before.

TRADES IN WEAK STATE

50

15.8% of 317 total trades

NET P&L (WEAK)

\$64.99

PF (WEAK)

1.50

vs. overall PF

WINDOW SIZE

50

trades per rolling window

Interpretation

Weak-state trades (50, 15.8% of history) remain marginally profitable (PF 1.50). A governor rule would reduce variance but is not clearly required to protect profitability on this evidence alone.

Clustering by Date / Session / Tag / Side

Unavailable — clustering weak-state trades by date/session/tag/side requires a per-trade weak-state flag (rolling PF classification joined back to each trade's session/tag/side), which is not currently included in this payload. Upstream fix: emit a `weak\_state` boolean per trade in the rolling_state computation so this report can group by session/tag/side without re-deriving it.

Governor Candidate

A governor rule (reduce size or pause trading when the rolling PF drops below the weak threshold) is proposed as a specific, falsifiable test in the Next Research Queue & Validation Plan section, gated on the same 50-trade confidence threshold used throughout this report.

4 Telemetry & Structure

Tag breakdowns, multi-model regime, and structure analysis.

SECTION 8

Telemetry Tag Breakdown

How to read tag tables

Identify tag values with high PF and meaningful trade count as candidates for inclusion rules. Negative PF values suggest values to filter or avoid.

GRADE Distribution

GRADE	N	Net P&L	Win%	PF	Avg
A	156	\$3,960.39	68.0%	1.99	\$25.39
C	103	-\$1,854.68	40.8%	0.52	-\$18.01
B	161	-\$2,939.78	62.1%	0.63	-\$18.26

ACT Distribution

ACT	N	Net P&L	Win%	PF	Avg
Hold	114	-\$527.67	60.5%	0.88	-\$4.63
Scale	106	-\$647.83	68.9%	0.87	-\$6.11
Enter	97	\$2,196.11	66.0%	1.90	\$22.64

WHY Distribution

WHY	N	Net P&L	Win%	PF	Avg
News fade	102	-\$246.23	65.7%	0.94	-\$2.41
Breakout	74	-\$364.72	60.8%	0.89	-\$4.93
Trend continuation	76	\$854.15	73.7%	1.28	\$11.24
Mean reversion	65	\$777.41	58.5%	1.54	\$11.96

SECTION 9

Multi-Model Regime Analysis

M_CONS

Value	N	Net P&L	Win%	PF
BX	58	-\$309.64	69.0%	0.90
B	53	\$1,127.34	67.9%	1.79
NB	68	\$21.36	57.4%	1.01
N	73	-\$236.72	65.8%	0.93
NX	65	\$418.27	66.2%	1.23

M_AGR

Value	N	Net P&L	Win%	PF
STRONG	102	\$570.38	65.7%	1.16
MOD	103	-\$1,009.31	57.3%	0.78
SPLIT	112	\$1,459.54	71.4%	1.39

M_ETR

Value	N	Net P&L	Win%	PF
EXTREME	83	\$967.90	60.2%	1.43
STRONG	76	\$1,679.49	75.0%	1.91
MOD	81	-\$1,352.96	58.0%	0.71
WEAK	77	-\$273.82	67.5%	0.92

M_VREG

Value	N	Net P&L	Win%	PF
HIGH_VOL	96	\$1,598.34	64.6%	1.67
NORMAL	111	\$584.23	67.6%	1.14
LOW_VOL	110	-\$1,161.96	62.7%	0.78

m_align

Value	N	Net P&L	Win%	PF
?	317	\$1,020.61	65.0%	1.08

SECTION 10

Structure Analysis

Market Structure Tags

BOSDIR tracks Break-of-Structure direction — bullish or bearish. SRSUPD captures which direction price breaks through SR levels. SRRESL reveals how SR levels resolve after being tested. High PF and positive net P&L in a tag value identifies structural filters worth applying.

BOS Direction (BOSDIR)

Value	Trades	Net P&L	Win%	PF	Avg Trade
Flat	120	\$832.32	66.7%	1.19	\$6.94
Down	100	\$339.89	60.0%	1.09	\$3.40
Up	97	-\$151.60	68.0%	0.96	-\$1.56

SR Update Direction (SRSUPD)

Value	Trades	Net P&L	Win%	PF	Avg Trade
Same	116	\$1,673.50	69.0%	1.43	\$14.43
Higher	97	-\$262.30	61.9%	0.93	-\$2.70
Lower	104	-\$390.59	63.5%	0.91	-\$3.76

SR Resolution Direction (SRRESL)

Value	Trades	Net P&L	Win%	PF	Avg Trade
Resolved	156	\$2,869.02	67.3%	1.75	\$18.39
Pending	161	-\$1,848.41	62.7%	0.77	-\$11.48

5 Monte Carlo & Path Analysis

Apex survival matrix, what-if simulations, and robustness surfaces.

SECTION 11

Monte Carlo / Apex Survival

PASS RATE

78.5%

P10 EQUITY

\$408.24

P50 EQUITY

\$1,020.61

Median outcome

P90 EQUITY

\$1,632.98

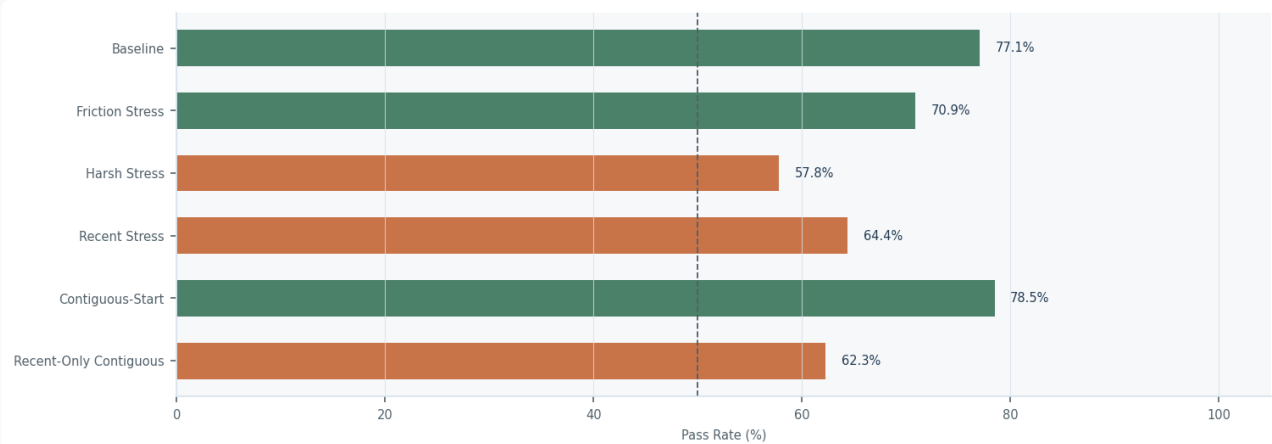
PATHS

1000

Apex Survival Assessment

Monte Carlo baseline pass rate: 78.5% across 1000 simulated paths. Majority of paths survive the profit target threshold — positive MC signal.

Pass Rate by Scenario



Green ≥70%, orange 40–70%, red <40%. Dashed line = 50%.

Apex Survival Battery (6-Scenario)

Baseline MC, Friction MC (added trading costs), Harsh stress MC, Recent-trades-only stress, Contiguous-start tests, and Recent-only contiguous-start tests — run automatically so results don't depend on manually re-running Monte Carlo one variant at a time. Scenarios need at least 25 trades in their source window; scenarios below that threshold are marked "Skipped" with the specific reason (e.g. too few trades overall, or insufficient recent-window data) instead of a fabricated result.

This run has the full Apex Survival Battery (6 scenarios). It is not affected by Standard vs Full Audit report mode — both modes render whichever battery this run's data actually has.

Scenario	Paths/ Starts	P50 Equity	95% DD	99% DD	Pass Rate	DD Breach Rate	Median Trades/ Days	Signal
Baseline	—	\$4,220.82	-\$1,132.37	-\$1,827.27	77.1%	—	—	Green
Friction Stress	—	\$2,529.26	-\$1,338.18	-\$1,855.15	70.9%	—	—	Green
Harsh Stress	—	\$3,458.64	-\$1,019.78	-\$1,549.54	57.8%	—	—	Yellow
Recent Stress	—	\$2,678.69	-\$1,286.86	-\$2,540.14	64.4%	—	—	Yellow
Contiguous-Start	—	\$3,209.01	-\$1,461.20	-\$1,300.87	78.5%	—	—	Green
Recent-Only Contiguous	—	\$2,303.55	-\$946.22	-\$1,771.08	62.3%	—	—	Yellow

SECTION 12

Stop / Exit Simulator

What-If Analysis

These tables simulate applying MAE-based loss caps or breakeven triggers to the existing trade set. Positive Δ Net rows suggest the rule would have improved results. Evaluate the trade-off against reduced trade count and win rate changes.

MAE Loss-Cap What-If

Scenario	Net P&L	Δ Net	Win%	PF	Max DD
\$-30	—	\$9,469.92	—	5.18	-\$30.00
\$-50	—	\$8,504.46	—	3.74	-\$50.00
\$-80	—	\$7,677.76	—	3.02	-\$80.00

6 Evidence & Appendix

Worst days, red-tail diagnostics, raw statistics, and methodology notes.

SECTION 14A

Worst Trading Days

Purpose

The 10 worst trading days reveal the shape of your drawdown risk. Clustering (multiple consecutive bad days) indicates regime sensitivity; isolated outliers may be news events or liquidity gaps. Full Audit mode shows all 25; regenerate in Full Audit for the complete list.

Date	Trades	Daily P&L	Cum P&L
2024-09-17	1	-\$905.29	\$1,836.63
2024-01-13	2	-\$854.72	-\$466.34
2024-07-14	2	-\$818.13	\$1,227.66
2024-02-24	1	-\$815.44	-\$786.68
2024-10-12	1	-\$678.80	\$1,194.40
2024-04-22	2	-\$633.20	\$162.51
2024-06-01	2	-\$624.36	\$1,277.20
2024-01-28	3	-\$584.15	-\$1,268.97
2024-10-14	1	-\$524.22	\$620.46
2024-06-11	2	-\$521.44	\$1,164.92

SECTION 15

Methodology & Glossary

How metrics are computed

All P&L figures are computed directly from the uploaded trade log with no smoothing or survivorship adjustment. Monte Carlo scenarios use trade-level bootstrap resampling (not a parametric return model), so results reflect the empirical distribution of this specific trade set rather than a theoretical one. Percentages are always shown as a fraction scaled to a percent (e.g. 0.62 → 62.0%).

Why the Table of Contents has gaps

Section numbers are reserved consistently across both Standard and Full Audit modes and across runs with different available data — e.g. Section 9 is always Multi-Model Regime and Section 13 is always Robustness Surfaces, whether or not this particular run or mode includes them. This keeps section numbers stable when comparing two reports, but means the Table of Contents can list gaps (e.g. jumping from 8 to 11) when an optional section isn't available for this run, or is omitted in Standard mode — that gap is expected, not missing content.

Glossary

Profit Factor (PF) — Gross winning P&L divided by gross losing P&L. Above 1.0 means the strategy is net profitable; $PF \geq 1.5$ is generally considered a meaningful edge.

Win Rate — Percentage of trades that closed with positive P&L. A high win rate does not by itself imply profitability — it must be read alongside average win/loss size.

Max Drawdown — The largest peak-to-trough decline in cumulative equity observed in the dataset. A key measure of worst-case historical pain, though future drawdowns may exceed it.

Recovery Factor — Net P&L divided by max drawdown. Indicates how many multiples of the worst drawdown the strategy has recovered/earned back over the full period.

Sharpe / Sortino Ratio — Risk-adjusted return measures; Sortino only penalizes downside volatility, while Sharpe penalizes all volatility, so Sortino is typically higher for asymmetric strategies.

Monte Carlo Pass Rate — Percentage of simulated equity paths (via trade-resampling/bootstrap) that meet a defined survival criterion (e.g. profit target reached before a drawdown limit is hit).

Apex Survival Battery — A fixed set of 6 Monte Carlo stress scenarios (baseline, added trading-cost friction, harsh cost/volatility stress, recent-period-only resampling, and two contiguous-block bootstrap tests) used to test robustness beyond a single baseline simulation.

Baseline vs. Filtered — Baseline metrics are computed on the full/unfiltered trade set as recorded; Filtered metrics are computed after applying the strategy's own rule/tag filters. The delta shows the real effect of the filter, not a hypothetical one.

GRADE / ACT / WHY / CONF / EDGE tags — Optional per-trade telemetry tags a strategy can emit describing signal quality (GRADE), the action taken (ACT), the stated rationale (WHY), model confidence (CONF), and a categorical edge label (EDGE). Used to segment performance by decision context rather than just time or symbol.

Rolling Profit Factor / Weak-State — PF computed over a moving trade window (e.g. 50 trades) to reveal whether edge is stable over time or concentrated in specific eras. "Weak state" refers to the window(s) with the lowest rolling PF; "strong state" the highest.

Red-Tail / Outlier Losses — Losses that are large relative to the typical loss distribution for this strategy — often driven by slippage, gap risk, or a missed stop. These trades disproportionately affect max drawdown and tail-risk metrics even when rare.

Sample-Size Guardrail — A rule applied throughout this report: any finding based on fewer than 25 trades is downgraded to a directional note (not a recommendation); fewer than 50 trades caps confidence at Medium. Prevents small-sample noise from being presented as a strong finding.

Expectancy (Avg Trade) — Average net P&L per trade across the dataset ($\text{net P\&L} \div \text{trade count}$). Distinct from Profit Factor: a strategy can have a high PF but low expectancy if it trades rarely, or vice versa.

Calmar Ratio — Annualized return divided by max drawdown. Similar in spirit to Recovery Factor but normalized to an annual basis, making it comparable across strategies with different trade durations or dataset lengths.

DD Breach Rate — In the Apex Survival Battery, the percentage of simulated paths whose drawdown exceeds a defined prop-firm/eval-style drawdown limit at any point — distinct from Pass Rate, which measures reaching a profit target. A path can breach the DD limit even if it eventually recovers, which is why this is tracked separately.

Contiguous-Start Test — A Monte Carlo variant that resamples contiguous (unshuffled) blocks of trades starting from many different points in the historical sequence, rather than fully reshuffling trade order. Tests robustness to "what if I had started trading on a different day" rather than "what if trade order were random."

Holdout Split — The final chronological slice of trades set aside and never used to tune or select the strategy's rules — see Split Robustness. A strategy that performs well on Train/Validation but weak on Holdout is a red flag for overfitting.

MAE (Maximum Adverse Excursion) — The worst unrealized loss a trade reached before it was closed, even if the trade ultimately closed as a winner. Used to test hypothetical tighter stop-loss levels in the Stop / Exit Simulator without needing to re-run the strategy.

MFE (Maximum Favorable Excursion) — The best unrealized profit a trade reached before it was closed, even if it gave back profit before exiting. Used to estimate how much upside a breakeven or trailing-stop rule would have captured or given back.

Apex Readiness / Survival Score — A composite read of the Apex Survival Battery's pass rates and DD breach rates across all its scenarios, used as a directional (not certified) indicator of how a strategy might fare under a prop-style/Apex-style evaluation account. It is a research heuristic, not a guarantee of passing any specific real evaluation.

Report Completeness / Layer Sign-Off

This report is built from 11 layers. Each is independently marked PRESENT, PARTIAL, THIN, UNAVAILABLE, NOT INCLUDED, or ABSENT based on what data was actually available and rendered for THIS specific run in report mode **Standard** — not a blanket status. See the note beside each layer for the specific reason behind its status.

PRESENT 1. Decision — Decision Brief (Section 0), Strategy Health Dashboard (0a), and What We Learned (0b) render the verdict, score, and plain-English synthesis.

NOT INCLUDED 2. Evidence — Detailed Research Evidence (Section 13b) is a Full Audit-only section and is not rendered in Standard mode; Executive Summary, Core Metrics, Session Analysis, and Telemetry Tags (the Standard-mode evidence pages) are still present above.

PRESENT 3. Risk / Robustness — Equity Curve & Drawdown and Chronological Split Robustness (Section 4) cover risk and robustness.

PRESENT 4. MC / Prop-Style Survival — Monte Carlo / Apex Survival (Section 11) fully covers survival probability for this run in Standard mode. Robustness Surfaces (Section 13 — profit, drawdown, profit/DD, MC stress, and Apex pass-rate surfaces) is a Full Audit-only page and is not included here; regenerate this report in Full Audit mode to add it.

PRESENT 5. Rule Testing — Stop / Exit Simulator (Section 12) covers Standard-mode counterfactual rule testing. Red-Tail Rule Lab is Full Audit-only; regenerate this report in Full Audit mode to add it.

PRESENT 6. Research Recommendations — Recommendation Engine (0c) and Next Research Queue & Validation Plan (0e) provide prioritized, evidence-gated recommendations.

PRESENT 7. Validation Plan — Chronological Split Robustness (Train/Validation/Holdout, Section 4) and Monte Carlo scenario testing together fully constitute the validation plan used in this report.

PRESENT 8. Methodology / Glossary — This section (15) defines every term and metric used elsewhere in the report.

NOT INCLUDED 9. Raw Appendix / Export — Appendix: Raw Statistics is a Full Audit-only page and was not rendered for this Standard-mode report — use the Data Export Pack (CSV/JSON) on the Reports page, or regenerate this report in Full Audit mode, for the raw core-metrics dump.

PRESENT 10. Multi-Model Regime — Multi-Model Regime (Section 9) renders M_CONS/M_AGR/M_ETR/M_VREG model-alignment tables for this run.

PRESENT 11. Structure Analysis — Structure Analysis (Section 10) renders BOSDIR/SRSUPD/SRRES structure-tag performance tables for this run.

This report is generated for informational and research purposes only and does not constitute investment advice. Past performance, including all backtested and simulated results shown above, is not indicative of future results.