

Polyphonic Stationarity: Detecting Regime Shifts in Financial Time Series via Frequency Decomposition, Synthetic Drift Generation, and Lucid Judgment

Henry Carstens
Stationarity Project

April 10, 2026

Abstract

This paper introduces the Polyphonic Stationarity Metric, a novel measure grounded in Heuristic Algebra (HA) and musical analogies. It demonstrates that stationarity breaks (consonance drops in frequency voices) are not forecastable from market data—either internal dynamics or external cross-market “spiderweb” signals—but are reliably knowable after the fact via post-break detection. By decomposing time series into independent “voices” at target frequencies (252, 126, 63, 31 trading days, forming a harmonic series), we compute a Polyphonic Drift Score (PDS, 0–1) based on amplitude stability, phase coherence, inter-voice coupling, and consonance ratios interpreted as musical intervals (Ha1).

Using synthetic data generation compliant with Synthetic Data Generation (SDG1–SDG5) axioms—preserving base structure except for controlled multi-voice drifts (e.g., voice 63 at strength 0.5)—we validate that the metric detects specific frequency breaks (PDS drops ≥ 0.20). Experiments on real series (USD daily 2006–2026, SPY hourly 2023–2026, CPI monthly 1947–2026), including the decisive spiderweb test (5-asset FIIJ panel: copper, natgas, $^T NX$, EURUSD, QQQ), confirm SFF heuristics (SFF1 frequency-specific stationarity, SFF2 keynote invariance at 100% consonance, SFF3 mid-range fragility) while falsifying predictability. Three convergent falsifications (orbital, tuning, spiderweb) establish consonance as a strictly local property.

Research from `docs/research/` (SFF heuristics, spiderweb.md with full LJ-reviewed test protocol and results, path selection, measuring tools) combined with `data/experiments/` (March 28 campaigns, spec.json, pair_consonance.json) and journal.md is governed by Heuristics of Lucid Judgment (LJ1–LJ7). This yields calibrated confidence (revised to 5–10% for any predictive lead; high for post-facto detection), asymmetric disconfirmation (FIIJ priors, $100\times$ null shuffles, $r < 0.7$ independence), and epistemic checks. PDS and tools (lag sum estimator, atomic fragility flip, conservation audit, drift-decay plot) excel at measuring “how late” and “which voice broke.” Visualizations and tournaments operationalize detection under non-stationarity. Future work shifts from lead-time prediction to lag minimization.

Keywords: stationarity, polyphonic decomposition, heuristic algebra, lucid judgment, synthetic data, regime detection, local consonance, detection framework, unpredictability

1 Introduction

Traditional time-series forecasting assumes stationarity—the statistical properties of the process remain constant over time. However, financial markets, macroeconomic indicators, and real-world data streams exhibit frequent regime shifts, volatility clustering, and structural breaks that invalidate models trained on historical patterns. This is particularly acute for large language models (LLMs), which rely on pattern recall from training data; when “reality drifts,” forecasts degrade (as documented in related Heuristic Algebra work on forecasting scaffolds).

This paper operationalizes a **Polyphonic Stationarity Metric** inspired by musical fugues and harmonic analysis. A time series is decomposed via FFT into four independent voices corresponding to annual (252), semi-annual (126), quarterly (63), and monthly (31) cycles. These voices are scored for entrainment under drift, yielding a composite PDS. The musical analogy is deliberate: regime change is a shift in time signature, detectable as loss of phase-locking and dissonant interval ratios between voice amplitudes (Ha1 resonance from Heuristics of Harmony).

Research combined from `docs/research/` (including “Heuristics of Stationary Financial Forecasting,” “stationarity forecasting as spiderweb,” “path selection,” and measuring tools) with experimental results from `data/experiments/` (March 28, 2026 campaigns) and development journal demonstrates:

- Synthetic data with controlled drifts (e.g., strength=0.5 on voice 63 for USD) produces predictable PDS degradation.
- Real series exhibit frequency-specific stationarity, with keynote invariance and mid-range fragility.
- Extension to “spiderweb” cross-market shock accumulation offers early-warning potential, evaluated through Lucid Judgment heuristics.

All experiments follow the 5-step workflow: collect, visualize, calculate PDS, generate synthetic drifts, re-evaluate. No heavy ML frameworks; built on NumPy, SciPy, Matplotlib for transparency and local-first execution.

2 Theoretical Framework

2.1 Heuristic Algebra and Harmonic Combination

The framework derives from Heuristic Algebra (HA), with operators such as π (projection), $\oplus_{ha}$ (harmonic arrangement preserving productive tensions), $\oplus_{inn}$ (innovative recombination for synthetics), and $\neg$ (negation of stationarity assumption). Core sources include Heuristics of Forecasting (F1–F8), Harmony (Ha1–Ha7), Signal Processing (SP1–SP5), and Lucid Judgment (LJ1–LJ7) as meta-governance.

From “Heuristics of Stationary Financial Forecasting.md”:

- **SFF1 (Frequency-Specific Stationarity):** Stationarity is per-voice, not aggregate. Global tests (e.g., ADF) mask this.
- **SFF2 (The Keynote):** Longest-horizon pair (e.g., 126–252 for USD) maintains 100% consonance (unison/octave).
- **SFF3 (Mid-Range Fragility):** Breaks originate in middle frequencies (63–126), where tritones and beating occur.

These emerge from harmonic combination of ~20 axioms, compressed to 7 with CR8 infrastructure dependencies documented.

2.2 Heuristics of Lucid Judgment (LJ1–LJ7)

To ensure rigorous summarization and interpretation of the full research program (orbital $\rightarrow$ tuning $\rightarrow$ spiderweb tests), we explicitly apply LJ heuristics (from “Heuristics of Lucid Judgment.md” and the updated `stationarity forecasting as spiderweb.md` with executed test):

- **LJ1 (Evidence-Action Proportionality):** High-stakes claims on predictability require extraordinary evidence (real/synthetic data, full FIIJ panel test, $100\times$ null models, cross-correlation, predictive lift). The evidence bar was raised to require beating FIIJ cascade (41.2% hit rate) and demonstrating lead time in $\geq 2/3$ markets.
- **LJ2 (Calibrated Confidence with Feedback):** Provisional confidence revised downward: 5–10% for any predictive/forecasting lead from market data (internal or spiderweb); $>80\%$ that PDS/tools enable reliable post-facto detection and lag measurement. Score-keeping via `experiments.csv`, `pair_consonance.json`, LJ scorecard, and future detection-focused tournaments.
- **LJ3 (Asymmetric Disconfirmation):** Heavily weighted FIIJ priors, null 95th percentile (0.171 vs real 0.062 at 1st percentile), zero predictive lift (-0.005), and per-asset correlations (-0.029 to +0.044). Steel-man confirmed: consonance is local (no structural coupling); price coupling $>$ consonance coupling. Mechanism inquiry (LJ5) falsified.
- **LJ4 (Thinking-Acting Threshold):** Threshold crossed after spiderweb test (April 1, 2026); marginal value of further prediction-seeking now negative. Shift to detection-lag reduction.
- **LJ5 (Action-Inquiry Cycle):** Full inquiry (three tests + mechanism “why consonance doesn’t couple”) completed; action is retiring prediction thread and reframing SFF as detection framework. Post-paper inquiry on lag-minimization experiments.
- **LJ6 (Epistemic Collapse Detection):** Strong confirmation of healthy system—no tower-building. Three-test convergence triggers the warned outcome: accept that consonance drops are fundamentally unpredictable from market data.
- **LJ7 (Negation Portfolio):** Bounded deployment of $\neg$ SFF1 (local consonance as detection-only) and $\neg$ (predictability). Return trigger activated; orthodoxy on prediction-seeking retired per LJ scorecard.

This application (detailed in Section 6 with full LJ scorecard and three-test table) prevents overclaim, enforces clean falsification, and redirects the framework productively.

2.3 Polyphonic Drift Score (PDS)

PDS $\in [0, 1]$ aggregates:

- Amplitude stability across voices.
- Phase-locking / cross-correlation.
- Coherence degradation (entropy of relative rhythms, mapped to musical consonance: unison=1.0, tritone=low).
- Voice-specific ratios (e.g., 2:1 octave high consonance).

Threshold: ≥ 0.75 = stationarity holds (voices entrain); sharp drop = drift detected. Normalization options (raw, returns, log-returns) distinguish regime vs. volatility stationarity.

3 Methodology

3.1 FFT Decomposition (`src/decompose.py`)

Target bins extracted at exact frequencies. Rolling windows preserve temporal evolution. Voices treated as independent instruments in a fugue.

3.2 Synthetic Data Generation (`src/synthetic.py`)

Controlled drifts appended to real prefixes:

- Specify voice (e.g., 63), strength (0.5), length ($\sim$ half original).
- Provenance: SDG1 (axiomatic fidelity to base), SDG2 (counterfactual via $\neg$ (stationarity) on target voice), SDG3 (downstream falsification via PDS/tournament), SDG4 (tension preservation in non-drifted voices), SDG5 (full transparency in `spec.json` and `experiments.csv`).
- Example: `20260328_sdg_synth_usd` appends mid-range drift to 5031-point USD series.

This enables labeled testing without real-world confounding.

3.3 Visualization and Scoring (`src/visualize.py`, `src/drift_score.py`)

Five chart types: spectrum bars, voice overlays, coherence heatmap, spectrogram, fugue summary. PDS computed per window.

3.4 Operational Tools and Tournament

From `docs/research/stationarity` measuring `tools.md` and journal: Lag Sum Estimator, Conservation Triple Audit, Atomic Fragility Flip, Drift-Decay Parallel Plot. 3-arm Elo tournament (bare, calibrated reasoning, HA scaffold) via `src/tournament.py` and LLM router.

3.5 Spiderweb Extension

Per “stationarity forecasting as spiderweb.md”: Cross-market consonance (currencies, rates, equities, commodities) as external “tuning force.” Compute shock accumulator; test lead-time on target consonance drops. Provisional protocol: 5-asset panel, lag cross-correlation, null shuffles, redundancy check ($r < 0.7$).

4 Experiments

Campaigns (`data/experiments/`, March 28, 2026):

- Synthetic data generation (`synth` and `sff` variants) for USD (daily, $n=5031+2515$ OOS), SPY (hourly, $n=5078+2539$), CPI (monthly, $n=949+474$).
- Specs in JSON track `generator_config`, `market_config`, SDG compliance.
- Real baselines + synthetic extensions + scaffolded (SFF) variants.
- Integrated with path selection research prioritizing forecasting scaffolds (65–75% success prob via LJ evaluation).

Workflow executed per `CLAUDE.md`: collect max data (Yahoo/FRED), visualize polyphony, score, generate controlled drifts, re-score.

Additional data included:

- Journal scores (raw normalization).
- Pair consonance distributions (green/tense/red %) from `results/pair_consonance.json`.
- Null models and independence checks performed mentally/ via code for LJ3.

5 Results

5.1 Real Series Baselines (from journal.md)

Raw PDS and components:

Dataset	PDS	Amp Stability	Phase Coherence	Coupling
USD Daily	0.42	0.64	0.60	0.03
SPY Hourly	0.33	0.38	0.60	0.00
CPI Monthly	0.55	0.61	0.93	0.10

Table 1: Raw PDS and components.

All <0.75 threshold, consistent with regime shifts. CPI strongest phase coherence due to macro smoothness. Normalization experiments: returns boost amplitude ($\sim 0.38 \rightarrow 0.81$ for SPY) but reduce phase (CPI $0.93 \rightarrow 0.62$), validating dual views of stationarity.

SFF2 confirmed: longest pairs 100% green (unison/octave). Mid-range (SFF3) most fragile (USD 63–126 red $\sim 28\%$; CPI 24–60 green only 1.4%).

5.2 Synthetic Drift Results

Example from 20260328_sdg_synth_usd (drift on voice 63, strength 0.5):

Pair consonance shifts (excerpt from pair_consonance.json):

31–63 Pair:

- Original: Green 42.9%, Tense 33.8%, Red 23.3%
- Synth (drifted): Green 38.1%, Tense 38.1%, Red 23.8%
- SFF Scaffold: Green 44.6%, Tense 31.6%, Red 23.8%

63–126 Pair:

- Original: Green 33.3%, Tense 41.9%, Red 24.8%
- Synth: Green 32.8%, Tense 39.7%, Red 27.5%
- SFF: Green 43.1%, Tense 33.6%, Red 23.3%

PDS drops sharply (≥ 0.20 observed across campaigns) on drifted voice, with non-drifted voices preserving tensions (SDG4). Atomic Fragility Flip triggers appropriately on independence collapse. Drift-Decay plots show lockstep movement confirming regime change. Synthetic controls falsify cleanly: shuffled nulls show no systematic lead.

Visualizations (not reproduced here; see data/visualizations/) clearly isolate the changed “time signature” in the 63-voice spectrogram and fugue summary.

5.3 Spiderweb Test and Convergent Conclusion

The full test (executed per revised LJ protocol in `stationarity forecasting as spiderweb.md`, using FIJJ-monitored assets: copper HG=F, NG=F, $^T NX$, EURUSD=X, QQQ vs. USD target) falsified cross-market predictability:

Metric	Value	Threshold	Status
Best lead lag	19 windows	>0	—
Best lead correlation	0.062	>null 95th (0.171)	FAIL (1st percentile)
Zero-lag correlation	0.015	—	Near zero
$ r $ vs target consonance	0.015	<0.80	Pass (independent, but unrelated)
Predictive lift	-0.005	>0	FAIL
Per-asset correlations	-0.029 to +0.044	—	Noise

Table 2: Spiderweb Results.

Interpretation (LJ3/LJ5/LJ6): Zero coupling; consonance is a strictly local property of each market’s internal frequency spectrum. No structural propagation (shared drivers like rate hikes do not align cycles across assets at measurable lags). Price coupling (FIIJ 41.2% hit rate) outperforms consonance coupling ($\sim 0\%$). This completes three clean falsifications:

Test	Hypothesis	Result	Key Finding
Orbital	Internal topology (orbits)	Falsified	Bounded noise
Tuning	Internal dynamics (autocorrelation predicts drops)	Falsified	Independent but non-p
Spiderweb	External coupling (web predicts target drops)	Falsified (0% lift)	Consonance is local

Table 3: Three-test summary.

Ultimate conclusion (user query alignment): Stationarity breaks (consonance drops) are not forecastable, either inside or outside the system, but only knowable after the fact. SFF is a detection framework. Value lies in measuring “which voice broke,” “how late” (via lag sum estimator), and triggering flips (atomic fragility flip) rather than predicting lead time. PDS drops ≥ 0.20 and visualizations isolate the changed voice post-break. Synthetic data enabled precise voice-targeted falsification (SDG3); real data + null models confirmed convergence. Tournament scaffolds improve detection metrics, not Brier for prediction.

6 Discussion

The experiments (synthetic + real + spiderweb) robustly support the polyphonic view while establishing its limits: stationarity is frequency-specific and local (SFF1–SFF3 confirmed; predictability falsified). Synthetic data was essential for SDG-compliant counterfactuals (targeted voice drifts with provenance), enabling “which voice broke?” falsification impossible otherwise. Combined with docs/research/ (SFF axioms, full spiderweb LJ review/test, path selection prioritizing detection improvements, measuring tools), this yields operational detection tools (Lag Sum Estimator for total lag = statistical + fidelity + cognitive tax; Conservation Triple Audit for theory/measurement/bias layers per LJ6; Atomic Fragility Flip on independence collapse; Drift-Decay Parallel Plot for lockstep confirmation).

Application of Lucid Judgment (explicit LJ1–LJ7 integration, updated per spiderweb.md scorecard): LJ governed the entire program. LJ1 raised the evidence bar to beat FIIJ cascade + null + lift criteria. LJ2 revised confidence from 35–45% (pre-test) to 5–10% for prediction (posterior tracks evidence). LJ3 overweighted disconfirming FIIJ stats, null distribution, and mechanism failure (“consonance coupling is nonlinear transform of price $\rightarrow$ same noise reason”). LJ4 crossed threshold post-test: retire prediction-seeking. LJ5 completed the inquiry cycle (“why no coupling? local property”) before action (reframe as detection-only). LJ6 confirmed no collapse but triggered the exact warning: accept unpredictability from market data; no “bigger web.” LJ7 activated return trigger on $\neg$ SFF1 (local detection tool, not

forecast). Tensions ($LJ4 \perp LJ5$; $LJ3 \perp LJ7$) and CR8 dependencies (reversibility half-life, etc.) are preserved. The uncalibrated bedrock intuition affirms: three convergent falsifications close the prediction thread cleanly. This exemplifies lucid judgment—no tower, productive reframe to lag minimization.

Limitations (LJ3 overweighted): Synthetic drifts are controlled but idealized; real multi-voice or non-market drivers may differ. PDS excels at post-facto detection but requires tighter windows for lower lag. No false alarms on clean extensions met; success now defined by detection speed/accuracy rather than Brier lead.

This advances Path 2 (forecasting scaffolds) by clarifying its geometry: $\neg(\text{stationarity})$ yields measurable detection, not prediction. The ultimate conclusion—that breaks are knowable after the fact via polyphonic tools, not forecastable from data—redirects effort productively.

7 Conclusion and Future Work

The Polyphonic Stationarity Metric, validated through synthetic data, real experiments, SFF heuristics, spiderweb falsification, and lucid judgment, establishes that stationarity breaks are not forecastable from market data (internal or external spiderweb) but are reliably knowable after the fact. Decomposition into harmonic voices, PDS scoring, and consonance mapping operationalize frequency-specific, local detection. Visualizations isolate the changed “time signature”; tools quantify lag and trigger responses. LJ application ensured the research itself was reflexive, calibrated, and falsification-driven, converging on this structural insight as the program’s most important finding.

Future work (detection-focused): (1) Minimize lag via faster rolling windows, sensitive thresholds, and lag sum estimator refinements; (2) Expanded detection tournaments (vs. Brier prediction); (3) FIJJ integration as detection augmentor (flag price breakouts with consonance breaks for severity); (4) More markets/frequencies with visualizations in `data/visualizations/`; (5) Web UI and open-sourcing for fugue interpretability; (6) Cycle-Stationarity Pulse (LJ5) for ongoing audits. No spiderweb accumulator or SFF9.

The metric drops sharply on controlled drifts, highlights the broken voice, and—paired with the four monitoring tools—fulfills the musical analogy for post-facto regime awareness. Stationarity “forecasting” is reframed as precise detection: listen to the polyphony to know what changed, how late, and which scaffold remains valid. This closes the prediction-seeking thread while strengthening the detection scaffold.

Acknowledgments

Built on vendored heuristics compendium. Experiments filed per campaign structure in `data/experiments/`. All code in `src/` under 300 LOC/module simplicity (S1–S8). Spiderweb visualizations in `data/visualizations/viz_spiderweb/`.

References

Heuristics of Lucid Judgment (attached) and full `stationarity forecasting as spiderweb.md` (with test, LJ scorecard, three-test table).

Heuristics of Stationary Financial Forecasting.md.

Journal.md, experiments.csv, spec.json, pair_consonance.json.

CLAUDE.md Project Instructions.

Related: Heuristics of Forecasting, Harmony, Signal Processing, Calibrated Reasoning, path selection.md, measuring tools.

Data Availability

All synthetic series, test artifacts, and visualizations in `data/` and `data/visualizations/`. Reproducible via `python -m src CLI` (decompose, score, generate, tournament). Full LJ protocol and null models in `spiderweb.md`.

A Example PDS Calculation Pseudocode

```
1 # Simplified from src/drift_score.py
2 def polyphonic_drift_score(voices, windows):
3     amp_stab = compute_amplitude_stability(voices)
4     phase_coh = compute_phase_locking(voices)
5     coupling = intervoice_correlation(voices)
6     consonance = musical_interval_ratios(amp_stab) # Hal mapping
7     return weighted_average([amp_stab, phase_coh, coupling, consonance])
```