

Reasoning Beyond Prediction: A Neuro-Symbolic Multi-Agent Framework for Explainable Options Trading

Partha P. Adhikari¹, Vineeta Khemchandani², Neetu Sharma³

¹National Institute of Electronics and Information Technology, Delhi

^{2,3}Galgotias University, UP, Noida, India

¹partho@nielit.gov.in, ²vineeta.khemchandani@galgotiasuniversity.edu.in, ³neetu.sharma@galgotiasuniversity.edu.in

Abstract: Price-direction prediction has long been treated as the central task in building automated trading systems. What such systems rarely address is whether a proposed trade makes economic sense, sits within an acceptable risk envelope, or can be justified to a human analyst. This paper presents NeSy-MAT (Neuro-Symbolic Multi-Agent Trader), a council-of-agents architecture that pairs a neural return predictor with lightweight, auditable symbolic checks rather than relying on prediction alone. The verified implementation evaluated here, NeSy-MAT V2, combines a Neural Predictor Agent (an MLP return regressor), two symbolic confirmation agents that read volatility-regime and options-flow signals, and a hard-veto financial ontology that can block a trade outright regardless of what the other agents say. A trade is placed only when the neural predictor and at least one symbolic agent agree, and no ontology rule has fired; there is no iterative re-proposal cycle in the verified system.

We evaluate this implementation on 833,675 raw NSE F&O Bhavcopy contract records for BANKNIFTY weekly options (January 2019 – July 2024) [53], reduced after feature engineering to 1,352 weekly observations. Training uses a single pre-pandemic regime only; Sharpe ratio, maximum drawdown, win rate, and abstain rate are then measured out-of-regime on the COVID-19 shock/recovery period and on the post-2023 high-volume retail-driven period. The result is mixed rather than uniformly favourable: NeSy-MAT V2 trails a plain LSTM baseline on raw Sharpe during the training regime (0.89 vs. 1.32) and during the COVID regime (−0.43 vs. 0.61), and only leads in the post-2023 regime (0.96 vs. 0.35), while abstaining far less often than the LSTM baseline throughout. Bootstrap 95% confidence intervals on all three models' Sharpe ratios are wide and overlap substantially — including overlap with zero — so none of these regime-level differences should be read as statistically significant given the sample sizes available (9–27 trades per regime). Three derivatives practitioners rated a sample of the system's reasoning chains at a moderate 3.3-of-5 median on both actionability and faithfulness. We report these results as they are rather than as we had hoped they would be, include a real logged council decision (24 February 2020) that lost money despite full agent consensus, and are explicit throughout about which claims — an ablation study isolating each component's contribution, and a richer Neo4j/LLM-arbitrated version of this architecture — remain future work rather than what was evaluated.

Keywords: Neuro-Symbolic AI, Multi-Agent Systems, Options Trading, Explainable AI, Rule-Based Symbolic Reasoning, Regime Robustness, LSTM, Bootstrap Confidence Intervals

1. Introduction

The application of AI to financial markets has moved quickly — from rule-based systems to deep-learning models capable of detecting nonlinear patterns in price time series [1][2]. Two problems, however, remain largely unresolved. First, the most accurate models tend to be black boxes that offer no human-interpretable justification for their outputs, which sits uneasily with regulatory regimes such as MiFID II and SEBI's algorithmic-trading guidelines [3]. Second, trading an option is not simply a matter of guessing the direction of the underlying asset: a trader must also weigh volatility-surface dynamics, time-decay effects, event calendars, liquidity constraints, and portfolio-level risk exposure — a multi-dimensional reasoning problem that a single predictive signal cannot capture [4][5].



The first problem is addressed by Neuro-Symbolic (NeSy) paradigm [6] which combines the abilities of neural pattern recognition and symbolic reasoning. The second one is taken care of by the Multi-Agent Systems (MAS) research that demonstrates that complex decisions are better achieved by architectures made up of agents with reductive goals rather than by a unique agent [7][8]. NeSy-MAT has the two seeds as well and extends to Indian equity options markets. One of its core messages is that using structured multi-agent deliberation – instead of adding up predictions or trying to explain afterwards – yields more accurate, better risk-managed and truly explainable trading decisions.

2. Related Work

2.1 Literature Survey and Research Gaps

NeSy-MAT is an integration of six intertwined lines of work completed before. Looking at the side-by-side results of Fischer and Krauss [13] (directional accuracy of 53.2% on an LSTM based sequence model on S&P 500) and Sezer and Ozbayoglu [51] (performance based on chart patterns learned by CNNs from candlestick images on the same S&P 500), it is evident of the ceiling the sequence-only based approaches hit. One can see from the huge and Savine [17] that autodifferentiation of the neural network pricing functions provides Delta and Gamma estimates, which are as stable as those provided by the Monte Carlo simulation but much cheaper to compute – that's what the Risk Agent does to get live Delta and Gamma.

Symbolic relational representations have been shown to deliver 4–12 p.p. absolute increase in stock-prediction accuracy in Financial Knowledge Graphs [26][27][28] while findings of path-based explanations for KG are both useful and convincing, which directly underpins the design of NeSy-MAT's audit trail. The formal statement of the Risk Agent's hard-constraint veto is done by Constrained Policy Optimisation [37] and Andersen et al. [40] provide a statistically valid definition of tail risk as IV rank. Chain-of-thought prompting [9] and the ReAct paradigm [44] have been proven to significantly reduce the factual hallucination in grounded decision tasks, by about a third; and the framework of AutoGen [46] and MetaGPT [47] indicates that output from the multi-agent LLM with different roles is more coherent than the output from a single agent. There is no current framework which merges all four components (option specific deep learning, financial-KG reasoning, constrained-RL risk management, and CoT capable LLM arbitration) in a single statistically validated framework. A complete survey of the 56 papers, included in this work, is provided with complete annotations as supplementary material.

2.2 Why This Gap Persists

The four research streams behind NeSy-MAT's design — deep sequence models for price direction, financial knowledge graphs, constrained-RL risk management, and LLM-based multi-agent reasoning — have largely matured on separate tracks. The rise of LLM-based multi-agent frameworks specifically dates to 2022–2023, evidenced by AutoGen [46], MetaGPT [47], and FinGPT [50] all appearing within that window, well after the deep-learning and knowledge-graph literatures were established. We are not aware of a published framework that combines all four streams — options-specific deep learning, financial-KG reasoning, constrained-RL risk management, and CoT-capable LLM arbitration — for options trading on Indian markets specifically; this is the gap NeSy-MAT targets. An earlier draft of this paper included a bibliometric table quantifying publication growth per stream with year-by-year counts and CAGR figures; we withdrew it because we do not have a documented search methodology or dataset that would let us or a reader reproduce those counts, and we would rather understate this motivation than publish precision we cannot back up.

3. Proposed Architecture: The Council of Agents

NeSy-MAT V2 treats the decision to be taken as a small council, not a single predictor. The architecture is rather conservative: first, a hard-veto layer can outright prevent a trade from happening, and second, if it doesn't, the neural predictor's prediction must agree with at least one independent symbolic signal before trading. Each component's role, tools and output are summarised in Table I and the subsections below contain a description of each

component's internal design. The way this combination works to come up with one result—the order in which checks are activated—and what happens if there is not total agreement, was described separately in Section 4, and is not repeated here.

TABLE I Agent Roles, Tools, and Mandates

Agent	Role	Core Tools	Output
Neural Predictor Agent (NPA)	Forecast forward return from price/vol history	MLP regressor (scikit-learn), 230→128→64→1, trained on Regime A only	Predicted forward return + BUY/no-signal vote
Volatility Regime Agent (VRA)	Judge whether volatility conditions favour the strategy	Percentile thresholds on normalised India VIX and volatility risk premium, fit on Regime A	BUY / abstain vote
Market Microstructure Agent (MMA)	Detect options-flow positioning consistent with a breakout	Percentile thresholds on put–call-ratio momentum and open-interest flow, fit on Regime A	BUY / abstain vote
Financial Ontology	Hard veto on trades with elevated structural risk	Two fixed rules: SEBI-announcement proximity, off-ATM strike (moneyness z-score)	VETO / no-veto

3.1 The Neural Predictor Agent

Here, the Neural Predictor Agent (NPA) is a multi-layer perceptron regressor (scikit-learn MLPRegressor) trained using flattened 10-step lookback sequences with 23 features (230 inputs) from Regime A only, comprising of 2 hidden layers, of 128 and 64 units respectively, using the ReLU activation, the Adam solver, L2 regularisation ($\alpha = 1e-4$) and early stopping on a chronological 80/20 split (patience 15 iterations). Fattening the lookback window instead of relying on a recurrent layer also results in a dependency-light implementation and one that is easily reproducible from a single saved scikit-learning checkpoint, and thus easily compared with the 2-layer LSTM baseline evaluated in parallel with this one (Section 6). If the predicted return of the NPA exceeds a threshold correlation threshold (30 percentile of its Regime A prediction distribution), then the NPA votes BUY and the council decision stops there (Section 4); otherwise, it abstains and the decision of the council stops at the NPA (Section 4).

3.2 The Symbolic Confirmation Agents

To represent the richer Financial Knowledge Graph originally intended for this role (Section 7.3), two simple agents are used—one is a percentile threshold that only activates on Regime A and does not have any learned parameters. The Volatility Regime Agent (VRA) provides signals indicating the volatility regimes of India VIX (VIX_NORM) and the volatility risk premium (VRP, implied minus real volatility) and only suggests voting BUY if VIX_NORM is in the Regime A 25th - 85th percentile and VRP is below Regime A 90th percentile, as this cuts out both too-calm market regimes after which options may be cheap to roll up and too-extreme market regimes in which options are already priced for the move. The Market Microstructure Agent (MMA) is built to give market buy opportunities based on the put–call-ratio momentum (PCR_OI_DELTA) and normalised call/put open-interest flow when both are above their respective Regime A 40th-percentile mark, which is signal of unusual positioning prior to a big swing. No agent can compel a trade; both of these 2 agents must agree to the trade, plus the council rule (section 4) must also approve it.

3.3 The Financial Ontology (Hard-Veto Layer)

The Financial Ontology is a fixed rule set, not a learned model, and it can block a trade regardless of what the NPA, VRA, or MMA vote. Two rules are active, both calibrated on Regime A: a trade is vetoed if a SEBI structural-announcement window is active within 40% of the normalised time-to-expiry, and it is vetoed if the strike’s moneyness is more than 2 standard deviations from the Regime A moneyness distribution. This replaces the Constrained-Policy-Optimisation Risk Agent with Monte Carlo VaR/CVaR estimation originally envisioned for this role (Section 7.3); the verified system’s veto is unconditional and does not compute a position-sizing adjustment — there is no MODIFY path, only VETO or no-VETO.

3.4 The Council Rule

There is no LLM-based Manager Agent in the verified implementation, and no Chain-of-Thought arbitration. The moderating logic is a fixed rule, applied once per decision window with no iteration: if the Financial Ontology vetoes, the outcome is ABSTAIN; otherwise, if the NPA votes BUY and at least one of the VRA or MMA also votes BUY, the outcome is TRADE; otherwise it is ABSTAIN. The natural-language, LLM-arbitrated moderator originally envisioned for this role — including an iterative Propose–Contextualise–Challenge–Resolve loop with MODIFY/HOLD logic — remains a target design (Section 7.3) rather than what produced the results in Section 6.

TABLE II Council Control Flow — Trigger → Transition → Outcome

Trigger	Transition	Outcome
Financial Ontology veto fires (SEBI-window or off-ATM)	Decision short-circuits	ABSTAIN
NPA vote = 0 (predicted return below 30th-percentile threshold)	Decision short-circuits	ABSTAIN
NPA vote = 1, VRA vote + MMA vote = 0	No symbolic corroboration	ABSTAIN
NPA vote = 1, (VRA vote + MMA vote) ≥ 1, no ontology veto	Full council agreement	TRADE (long straddle)

Table II reflects a single-pass evaluation: each decision window is checked once, in the order Ontology → NPA → VRA/MMA, with no re-proposal or MODIFY cycle. It intentionally does not restate what each agent computes internally — see Table I and Sections 3.1–3.4 for that — and is distinct from the walkthrough in Section 4, which explains why each transition happens.

4. Methodology: The Council Decision Rule

NeSy-MAT V2 resolves each decision window with a single evaluation of the rule in Table II, not an iterative debate. Each decision window is one trading day; decisions are aggregated to one position per monthly expiry cycle for evaluation (Section 6), using the last active council decision before expiry. The three checks below run in the order shown, with the Financial Ontology given the first and final word: it can block a trade outright, and neither the NPA nor the symbolic agents can override it.

Check 1 — Ontology Veto. The Financial Ontology (Section 3.3) screens the decision window for a SEBI-announcement window near expiry or an off-ATM strike. A veto here ends the decision immediately, before either the neural or symbolic agents are consulted (Table II).

Check 2 — Neural Signal. If no veto fired, the NPA’s (Section 3.1) predicted return is compared against its Regime A 30th-percentile threshold. Below threshold, the window ends in ABSTAIN.

Check 3 — Symbolic Corroboration. If the NPA votes BUY, the VRA and MMA (Section 3.2) cast independent BUY/abstain votes. At least one of the two must agree with the NPA for the council to TRADE; otherwise the window ends in ABSTAIN.

The full trace of predicted return, individual agent votes, the ontology check, and the realised return is logged for every decision window; the case study in Section 6.5 is drawn directly from this log, not reconstructed after the fact.

5. Technical Implementation

The results in Section 6 come from a concrete, verified implementation. The two deep-learning baselines (LSTM, Transformer) are implemented in PyTorch. NeSy-MAT V2 itself — the system this paper is about — is implemented entirely in scikit-learn and plain NumPy/Pandas: its Neural Predictor Agent is an MLPRegressor, and its symbolic and veto components (Sections 3.2–3.3) are closed-form percentile rules, not learned models. Table III lists the full stack. This is narrower than the LangGraph/Neo4j/LLM-arbitrated architecture originally sketched for this line of work: that remains a target design (Section 7.3), not what produced the numbers in Table VI. The LSTM and Transformer baselines’ hyperparameters were selected via a grid search over depth, hidden width, dropout, and learning rate; the specific configuration values and validation scores are not reported in this draft pending verified experiment logs (Section 7.2).

TABLE III Technology Stack (Verified)

Component	Technology	Purpose
NeSy-MAT V2 — Neural Predictor Agent	scikit-learn MLPRegressor + StandardScaler	Forward-return prediction (Section 3.1)
NeSy-MAT V2 — Symbolic & Veto Agents	Pandas/NumPy percentile rules (no learned parameters)	VRA/MMA corroboration + Ontology veto (Sections 3.2–3.3)
Baseline — LSTM	PyTorch, 2-layer LSTM (hidden=128, dropout=0.2, lookback=10)	Comparison baseline (Table VI)
Baseline — Transformer	PyTorch, 4-layer encoder (d_model=32, nhead=2)	Comparison baseline (Table VI)
Feature Engineering	Pandas + NumPy + SciPy Stats	OHLCV, open interest, India VIX, and event-log feature construction
Data Layer	NSE F&O Bhavcopy [53] (BANKNIFTY), India VIX	Raw contract-record and volatility-index ingestion
Model Persistence	Python pickle / PyTorch state_dict	Saving trained agents and encoders for evaluation
Live-Deployment Track (separate, in progress)	Zerodha Kite Connect, pyotp, pandas-ta	Real-time feed/auth for a future live pilot; not part of the backtest reported here

Reproducibility. The dataset construction pipeline (Table IV) and the regime definitions (Table V) are reported in full. The LSTM/Transformer hyperparameter grid search is not, for the reason noted above: we do not currently have verified logs for the individual configurations and are not willing to publish specific figures we cannot stand behind. The codebase and trained model weights are released publicly (Section 8); once the grid-search logs are re-verified against that release, they will be added here rather than reconstructed from memory.

6. Evaluation Framework

6.1 Dataset: Provenance and Construction

All data are drawn from the NSE F&O Bhavcopy Archive [53], the exchange’s daily end-of-day file for all F&O contracts. The evaluation reported in this paper uses BANKNIFTY weekly index options only, spanning January 2019 to July 2024: 833,675 raw contract-level records (OHLC, volume, open interest, change in open interest,

settlement price), supplemented with India VIX and a curated event log of scheduled macro releases (RBI policy, US FOMC, CPI prints) and SEBI/NSE structural announcements (margin changes, lot-size revisions). No third-party data vendors are used. After feature engineering, the raw records yield 1,352 weekly observations across 43 features. NIFTY50-constituent equity options and full multi-index coverage are not part of the dataset evaluated here; extending beyond BANKNIFTY is future work (Section 7.3).

TABLE IV Dataset Statistics by Year

Year	Raw Contract Records	Regime
2019	83,342	A
2020	159,508	A $\rightarrow$ B (COVID onset, March 2020)
2021	184,036	B
2022	184,215	B $\rightarrow$ C (regime shift, January 2022)
2023	143,074	C
2024 (through July)	79,500	C
Total	833,675	Reduces to 1,352 weekly observations after feature engineering

6.2 Regime-Based Validation Protocol

Rather than a rolling walk-forward scheme, NeSy-MAT V2 is evaluated with a single regime-based split: all model parameters are estimated using Regime A only, and Regimes B and C are held out entirely from training and hyperparameter tuning. This tests out-of-regime generalisation directly — whether a model fitted under calm, pre-pandemic conditions still produces sound decisions once the regime changes — which the financial-ML literature has argued is a harder and more realistic test than in-regime cross-validation. Table V defines the three regimes, their role in the split, and their per-regime record counts.

TABLE V Regime Definitions and Per-Regime Statistics

Regime	Period	Role	Raw Records	Feature Rows	Characterisation
A	2019-01-01 $\rightarrow$ 2020-02-29	Training	103,802	285	Pre-pandemic, low volatility
B	2020-03-01 $\rightarrow$ 2021-12-31	Out-of-regime test	323,084	452	COVID-19 shock and recovery
C	2022-01-01 $\rightarrow$ 2024-07-05	Out-of-regime test	406,789	615	Post-pandemic, high-volume retail phase; several SEBI margin/lot-size interventions

6.3 Evaluation Metrics

For each regime we report the annualised Sharpe ratio (monthly cadence, scaled by $\sqrt{12}$, computed on excess returns over a 6.5% p.a. risk-free rate), maximum drawdown (absolute, in cumulative return units), win rate, number of trades, and abstain rate. $\Delta\text{Sharpe} = \text{Sharpe}(A) - \min(\text{Sharpe}(B), \text{Sharpe}(C))$ summarises worst-case out-of-regime degradation; lower is better. We also report 95% bootstrap confidence intervals on Sharpe (2,000 resamples, seed 42), computed identically for all three models so the intervals are directly comparable. We do not report a Diebold–Mariano

test between models: the return series are not aligned trade-for-trade across models in the way that test requires, and forcing an alignment would overstate what the comparison supports.

TABLE VI Headline Performance by Regime (Verified Results)

Model	Regime	Sharpe	95% CI	Max DD (abs.)	Win Rate	Trades	Abstain Rate
LSTM	A (train)	1.32	[−0.53, 2.69]	1.38	66.7%	9	30.8%
LSTM	B (COVID)	0.61	[−0.96, 1.98]	2.03	69.2%	13	38.1%
LSTM	C (post-2023)	0.35	[−1.03, 1.56]	1.83	61.9%	21	22.2%
Transformer	A (train)	0.81	[−1.88, 1.99]	1.51	55.6%	9	30.8%
Transformer	B (COVID)	0.01	[−1.76, 1.36]	2.89	54.5%	11	47.6%
Transformer	C (post-2023)	−0.07	[−1.60, 1.12]	2.69	55.6%	18	33.3%
NeSy-MAT V2	A (train)	0.89	[−2.00, 2.04]	2.19	46.2%	13	0.0%
NeSy-MAT V2	B (COVID)	−0.43	[−2.58, 0.89]	2.86	47.4%	19	9.5%
NeSy-MAT V2	C (post-2023)	0.96	[−0.30, 2.31]	2.66	57.7%	26	3.7%

ΔSharpe ($A - \min(B,C)$), lower is better: LSTM 0.97, Transformer 0.88, NeSy-MAT V2 1.32. Every model’s 95% CI is wide and overlaps zero in every regime, and the three models’ CIs overlap each other substantially: none of the Sharpe differences in this table, within or across models, should be read as statistically significant — sample sizes (9–27 trades per regime) are simply too small to distinguish these point estimates with any confidence. NeSy-MAT V2 abstains far less often than the LSTM baseline while producing statistically indistinguishable Sharpe ratios (near-zero abstain rate in Regime A versus 30.8% for LSTM). See Section 7 for a candid discussion of what this result does and does not support.

6.4 Baselines and Explanation-Quality Annotation

Two baselines are evaluated on the same regime split: (B1) a standalone LSTM and (B2) a four-layer Transformer encoder, both trained on the same feature set with actions chosen by greedy selection on predicted return. An LSTM-plus-rule-filter baseline and a GPT-4 zero-shot baseline were part of the original evaluation plan but do not currently have verified results and are not reported here. All comparisons assume a transaction cost of 0.05% per leg. Three derivatives practitioners, blind to which system produced each explanation, rated 30 admitted reasoning chains from NeSy-MAT V2 on two 1–5 Likert dimensions: Actionability and Faithfulness. Across the 30 chains and 3 raters, NeSy-MAT V2 scores a median of 3.33 (mean 3.29) on Actionability and a median of 3.33 (mean 3.30) on Faithfulness — a moderate rating, not an exceptional one, and one that leaves clear room for improvement. We have not computed a formal inter-rater reliability statistic (e.g., an intraclass correlation) for these ratings; that is noted as an open item in Section 7.2 rather than left unstated.

6.5 Case Study: 24 February 2020

On 24 February 2020, India VIX rose from 13.70 (20 February) to 17.00 — a 24% jump — while BANKNIFTY spot barely moved. All three of NeSy-MAT V2’s voting components agreed to trade that day: the NPA’s predicted return cleared its threshold, the VRA voted BUY (VIX_NORM was still inside its fitted band — the jump had not yet pushed it into the 85th-percentile panic zone), the MMA voted BUY on options-flow positioning, and the Financial Ontology found no veto condition. The council traded, and the position realised a return of -0.72 against that day’s forward window.

This is a real, logged decision, not an illustrative construction, and it is included precisely because it is not flattering: every component agreed, and the trade still lost. It illustrates a specific, identifiable failure mode — the VRA’s volatility thresholds were fit on Regime A’s calmer historical distribution, and a fast-moving regime change is, almost by construction, exactly the kind of event a percentile threshold calibrated on past data is slow to catch. (India VIX went on to close at 23.24 within four trading days and peaked at 83.61 on 24 March, as the COVID sell-off developed.) A structural-causal extension of this architecture that checks whether the statistical premises behind a threshold still hold under a detected regime shift — rather than relying on a fixed percentile fit once on training data — is discussed as follow-on work in Section 7.3.

7. Discussion

7.1 *What the Results Support*

Three things can be said with the evidence in hand. First, and most important for interpreting everything else in this section: the bootstrap confidence intervals in Table VI are wide enough to overlap zero for every model in every regime, so none of the Sharpe differences reported here — NeSy-MAT V2 against either baseline, or across regimes — should be read as statistically established. This is a direct consequence of small regime-level sample sizes (9–27 trades) and needs to be stated before anything else. Second, within that caveat, the council architecture changes trading behaviour in a measurable way: NeSy-MAT V2’s abstain rate in the training regime is 0.0% against the LSTM’s 30.8%, and it stays lower in both out-of-regime windows, so the architecture is doing something distinct from the baselines, not merely re-deriving their predictions — even though this does not yet show up as a significant Sharpe advantage. Third, the explanation-quality ratings (Section 6.4) sit at a moderate 3.3-of-5 median on both actionability and faithfulness — evidence that the audit trail is intelligible to practitioners, but not evidence of unusually high explanation quality.

7.2 *Limitations*

Several limitations should be stated plainly rather than left for a reader to discover. First, sample sizes are small — 9 to 27 trades per regime — and the bootstrap confidence intervals in Table VI reflect that: this paper cannot claim a statistically significant advantage for NeSy-MAT V2 over either baseline in any regime. A materially larger evaluation window, more underlyings, or a longer Regime C sample would be needed before a significance claim could be supported. Second, there is no controlled ablation study isolating the individual contribution of the NPA, VRA, MMA, or the ontology veto to the results in Table VI; the abstain-rate discipline is observable, but which component drives it is not yet established. Third, the verified system (Sections 3–4) is a substantially simpler realisation of the architecture originally envisioned for this line of work: a scikit-learn MLP and two percentile-threshold rules rather than a deep LSTM/CNN ensemble, a Neo4j financial knowledge graph, and an LLM-arbitrated iterative debate loop. That richer version remains future work (Section 7.3), not a claim this paper makes about what was evaluated. Fourth, the LSTM/Transformer hyperparameter search referenced in Section 5 is not reported in configuration-level detail, for lack of verified logs. Fifth, evaluation is restricted to BANKNIFTY weekly options over three regimes with two verified baselines; the originally planned rule-filtered-LSTM and GPT-4 zero-shot baselines, and the broader Nifty50-constituent dataset, are not part of what is reported here. Sixth, the Likert explanation-quality ratings in Section 6.4 have no accompanying inter-rater reliability statistic. Seventh, transaction costs are held constant at 0.05% per leg regardless of position size.

7.3 *Future Research Directions*

Six directions follow from the limitations above, roughly in priority order. First, run a controlled ablation study that individually disables the NPA, VRA, MMA, and the ontology veto, to attribute Table VI’s abstain-rate discipline and Regime C recovery to specific components rather than stating them as an aggregate observation. Second, expand the evaluation window and underlying set — more regimes, NIFTY alongside BANKNIFTY — specifically to narrow the bootstrap confidence intervals enough to support a significance claim one way or the other. Third, build toward the originally envisioned richer architecture — a deep sequence encoder in place of the flattened MLP, a genuine financial knowledge graph in place of the two percentile-threshold agents, and an LLM-arbitrated moderator with iterative MODIFY logic in place of the fixed council rule — and re-evaluate against the same regimes to test whether the added complexity earns its keep. Fourth, extend the evaluation to the originally planned baselines (rule-filtered LSTM, GPT-4 zero-shot). Fifth, compute a formal inter-rater reliability statistic for the Likert ratings and extend them to a larger sample of reasoning chains. Sixth, the 24 February 2020 case in Section 6.5 shows full council consensus still losing money because a percentile threshold fit on Regime A hadn’t adapted to an emerging regime shift; a structural-causal admissibility check — of the kind explored in our follow-on work on this architecture — that tests whether a threshold’s statistical premises still hold, rather than assuming they do, is a concrete way to address exactly that failure mode.

8. Conclusion

NeSy-MAT V2 reframes options trading as a structured, veto-checked decision problem rather than a forecasting problem, and this paper evaluates a verified implementation of that idea — a scikit-learn neural predictor, two percentile-threshold symbolic agents, and a hard-veto rule set — on 1,352 weekly BANKNIFTY observations across three market regimes. The honest summary is a mixed one: the architecture abstains far more discriminately than a plain LSTM baseline, and it recovers best of the three systems tested in the post-2023 regime, but none of these differences clear a 95% bootstrap confidence interval, so we do not claim statistical significance. Explanation quality, rated by three practitioners on admitted reasoning chains, is moderate rather than exceptional.

We have tried, in this revision, to state plainly what is and is not supported by the evidence: an ablation study attributing specific effects to specific components has not been run; the richer Neo4j/LLM-arbitrated, iterative version of this architecture sketched in earlier drafts is a target design for future work, not what produced these results; and the confidence intervals in Table VI are wide enough that regime-level differences should not be over-read. The 24 February 2020 case in Section 6.5 is included precisely because it is unflattering — every component of the council agreed, and the trade still lost money, because a threshold fit on calmer historical data had not yet caught an emerging regime shift — and it motivates the specific next step (Section 7.3) of testing whether a threshold’s premises still hold rather than assuming they do.

The dataset-construction code and the trained models are released publicly for replication and extension.

REFERENCES

1. Sezer, O. B., Gudelek, M. U., & Ozbayoglu, A. M. (2020). Financial time series forecasting with deep learning: A systematic literature review: 2005–2019. *Applied Soft Computing*, 90, 106181.
2. LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436–444.
3. European Securities and Markets Authority (ESMA). (2018). MiFID II: Guidelines on certain aspects of the MiFID II suitability requirements. ESMA Publication.
4. Hull, J. C. (2022). *Options, Futures, and Other Derivatives* (11th ed.). Pearson Education.
5. Natenberg, S. (2015). *Option Volatility & Pricing: Advanced Trading Strategies and Techniques* (2nd ed.). McGraw-Hill Education.
6. Hitzler, P., et al. (2022). *Neuro-Symbolic Artificial Intelligence: The State of the Art*. IOS Press.
7. Wooldridge, M. (2009). *An Introduction to MultiAgent Systems* (2nd ed.). Wiley.
8. Stone, P., & Veloso, M. (2000). Multiagent systems: A survey from a machine learning perspective. *Autonomous Robots*, 8(3), 345–383.
9. Wei, J., et al. (2022). Chain-of-thought prompting elicits reasoning in large language models. *Advances in Neural Information Processing Systems*, 35, 24824–24837.
10. Hutchinson, J. M., Lo, A. W., & Poggio, T. (1994). A nonparametric approach to pricing and hedging derivative securities via learning networks. *The Journal of Finance*, 49(3), 851–889.

11. Culkin, R., & Das, S. R. (2017). Machine learning in finance: The case of deep learning for option pricing. *Journal of Investment Management*, 15(4), 92–100.
12. Ruf, J., & Wang, W. (2020). Neural networks for option pricing and hedging: A literature review. *Journal of Computational Finance*, 24(1), 1–46.
13. Fischer, T., & Krauss, C. (2018). Deep learning with long short-term memory networks for financial market predictions. *European Journal of Operational Research*, 270(2), 654–669.
14. Zhang, J., & Xiang, Y. (2020). Option pricing with long short-term memory neural networks. *Engineering Letters*, 28(1), 1–9.
15. Ke, A., et al. (2019). Pricing options with LSTM networks: A comparative analysis.
16. Horvath, B., Muguruza, A., & Tomas, M. (2021). Deep learning volatility: A deep neural network perspective on pricing and calibration in (rough) volatility models. *Quantitative Finance*, 21(1), 11–27.
17. Huge, B., & Savine, A. (2020). Differential machine learning. arXiv preprint arXiv:2005.02347.
18. Chen, W., et al. (2019). A novel deep learning framework: CNN-LSTM with attention mechanism for stock price prediction.
19. Yang, Y., et al. (2019). Attention-based option pricing.
20. Lim, B., & Zohren, S. (2021). Time-series forecasting with deep learning: A survey. *Philosophical Transactions of the Royal Society A*, 379(2194), 20200209.
21. Vaswani, A., et al. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.
22. Bao, W., Yue, J., & Rao, Y. (2017). A deep learning framework for financial time series using stacked autoencoders and long-short term memory. *PLOS ONE*, 12(7), e0180944.
23. Patel, J., Shah, S., Thakkar, P., & Kotecha, K. (2015). Predicting stock and stock price index movement using trend deterministic data preparation and machine learning techniques. *Expert Systems with Applications*, 42(1), 259–268.
24. Koshiyama, A., Firoozye, N., & Treleaven, P. (2020). Generative adversarial networks for financial trading strategies fine-tuning and combination. *Quantitative Finance*, 21(5), 797–814.
25. Ji, S., Pan, S., Cambria, E., Martinen, P., & Yu, P. S. (2022). A survey on knowledge graphs: Representation, acquisition, and applications. *IEEE Transactions on Neural Networks and Learning Systems*, 33(2), 494–514.
26. Wang, J., et al. (2024). Financial knowledge graph: A survey. *IEEE Transactions on Knowledge and Data Engineering*, 36(4), 1783–1802.
27. Deng, S., et al. (2019). Knowledge-driven stock trend prediction and explanation via temporal convolutional network. *WWW 2019 Companion Proceedings*, 678–685.
28. Cheng, D., Yang, F., Xiang, S., & Liu, J. (2020). Financial time series forecasting with multi-modality graph neural network. *Pattern Recognition*, 121, 108218.
29. Feng, F., He, X., Wang, X., Luo, C., Liu, Y., & Chua, T. S. (2019). Temporal relational ranking for stock prediction. *ACM Transactions on Information Systems*, 37(2), 1–30.
30. Feng, F., Chen, H., He, X., Ding, J., Sun, M., & Chua, T. S. (2019). Enhancing stock movement prediction with adversarial training. *Proceedings of the 28th IJCAI*, 5843–5849.
31. Kim, R., So, C. H., Jeong, M., Lee, S., Kim, J., & Kang, J. (2019). HATS: A hierarchical graph attention network for stock movement prediction. arXiv preprint arXiv:1908.07999.
32. Li, W., Bao, R., Harimoto, K., Chen, D., Xu, J., & Su, Q. (2020). Modeling the stock relation with graph network for overnight stock movement prediction. *Proceedings of the 29th IJCAI*, 4541–4547.
33. Xu, W., Liu, W., Xu, C., Bian, J., Yin, J., & Liu, T. Y. (2021). REST: Relational event-driven stock trend forecasting. *Proceedings of the Web Conference 2021*, 1–10.
34. Gao, J., et al. (2021). A semantic web approach to financial knowledge extraction. *Proceedings of the AAAI Workshop on Knowledge Graphs*, 2021.
35. Jiang, Z., Xu, D., & Liang, J. (2017). A deep reinforcement learning framework for the financial portfolio management problem. arXiv preprint arXiv:1706.10059.
36. Mihatsch, O., & Neuneier, R. (2002). Risk-sensitive reinforcement learning. *Machine Learning*, 49(2–3), 267–290.
37. Achiam, J., Held, D., Tamar, A., & Abbeel, P. (2017). Constrained policy optimization. *Proceedings of the 34th ICML*, 22–31.
38. Huang, C. Y. (2018). Financial trading as a game: A deep reinforcement learning approach. arXiv preprint arXiv:1807.02787.
39. Bisi, L., Sabbioni, L., Vittori, E., Papini, M., & Restelli, M. (2020). Foreign exchange trading: A risk-averse batch reinforcement learning approach. *Proceedings of the 19th EUMAS*, 2020.
40. Andersen, T. G., Bondarenko, O., & Gonzalez-Perez, M. T. (2015). Exploring return dynamics via corridor implied volatility. *The Review of Financial Studies*, 28(10), 2902–2945.
41. Guan, M., & Liu, X. Y. (2021). Explainable reinforcement learning for portfolio management. arXiv preprint.
42. Hitzler, P., & Sarker, M. K. (2022). *Neuro-Symbolic Artificial Intelligence: The State of the Art*. IOS Press.
43. Brown, T., et al. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877–1901.
44. Yao, S., et al. (2022). ReAct: Synergizing reasoning and acting in language models. arXiv preprint arXiv:2210.03629.

45. Li, G., et al. (2023). CAMEL: Communicative agents for ‘mind’ exploration of large language model society. *Advances in Neural Information Processing Systems*, 36.
46. Wu, Q., et al. (2023). AutoGen: Enabling next-generation LLM applications via multi-agent conversation. *arXiv preprint arXiv:2308.08155*.
47. Hong, S., et al. (2023). MetaGPT: Meta programming for a multi-agent collaborative framework. *arXiv preprint arXiv:2308.00352*.
48. Schick, T., et al. (2023). Toolformer: Language models can teach themselves to use tools. *Advances in Neural Information Processing Systems*, 36, 68539–68551.
49. Park, J. S., et al. (2023). Generative agents: Interactive simulacra of human behavior. *Proceedings of the 36th UIST*, 2023.
50. Yang, H., et al. (2023). FinGPT: Open-source financial large language models. *arXiv preprint arXiv:2306.06031*.
51. Sezer, O. B., & Ozbayoglu, A. M. (2018). Algorithmic financial trading with deep convolutional neural networks: Time series to image conversion approach. *Applied Soft Computing*, 70, 525–538.
52. Chronopoulos, M., Hagspiel, V., & Fleten, S. E. (2021). Value at risk estimation using deep learning. *Journal of Risk*, 23(5), 1–29.
53. National Stock Exchange of India Ltd. (2024). NSE Equity Derivatives — Historical Data (F&O Bhavcopy Archive). Available at: <https://www.nseindia.com/market-data/historical-data>. Accessed: January 2025.
54. Diebold, F. X., & Mariano, R. S. (1995). Comparing predictive accuracy. *Journal of Business & Economic Statistics*, 20(1), 134–144. [Harvey-Leybourne-Newbold finite-sample correction applied.]
55. Ledoit, O., & Wolf, M. (2008). Robust performance hypothesis testing with the Sharpe ratio. *Journal of Empirical Finance*, 15(5), 850–859.
56. Efron, B., & Tibshirani, R. J. (1993). *An Introduction to the Bootstrap*. Chapman & Hall/CRC. [Block bootstrap with block size = 21 trading days to preserve autocorrelation structure.]