

When Do Engineered Features Beat Raw Prices?

A Walk-Forward Benchmark of Classical, Neural, and World-Model Predictors on a Global Equity Dataset

Nilutpaul Sarker Yash^{1†}, Tousif Bashar¹, Md Mahmudul Islam Maruf², Masrur Mouno¹, Sayantan Chakraborty¹, Koushik Howlader³, Ushashi Bhattacharjee¹, Tirtho Roy¹
¹ Iowa State University, USA • ² University of Global Village (UGV), Barishal, Bangladesh • ³ University of Dhaka, Bangladesh
[†] Presenting author

FIGURE 1. CONTROLLED BENCHMARK OVERVIEW — FROM INPUT REPRESENTATION TO EXPERIMENTAL ANSWER

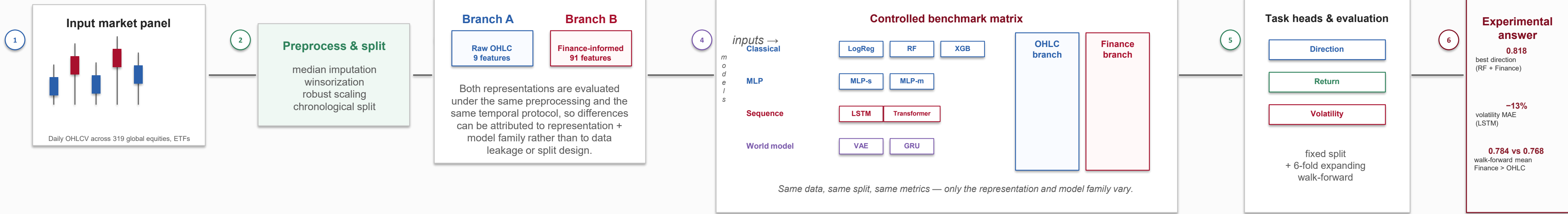


Figure 1. The controlled benchmark compares 9 raw OHLC-derived features against 91 finance-informed features across classical, neural, sequence, and world-model predictors under identical temporal evaluation, then reads off where representation helps, where it fails, and how robust the gains are.

DATASET, TASKS & SCOPE



Prediction tasks

Direction: 3-class {up, flat, down} with neutral band $\tau = 0.005$
Return: next-day log return (secondary target)
Volatility: forward 5-day realized volatility

Preprocessing and controls

Median imputation and winsorization
Robust scaling for volatile distributions
Cross-sectional normalization in the controlled benchmark
Strict chronology to prevent look-ahead leakage

Why this benchmark matters

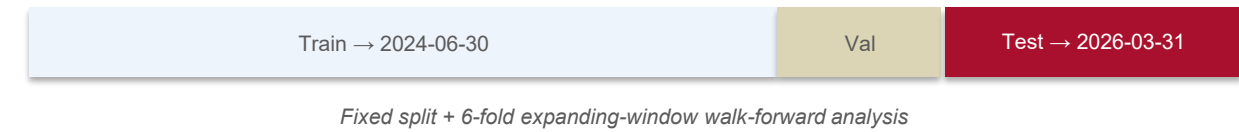
It isolates representation from architecture by keeping the dataset, preprocessing, and evaluation protocol fixed. This lets us ask a cleaner question: are gains coming from better features or just from more complex models?

Key numbers at a glance

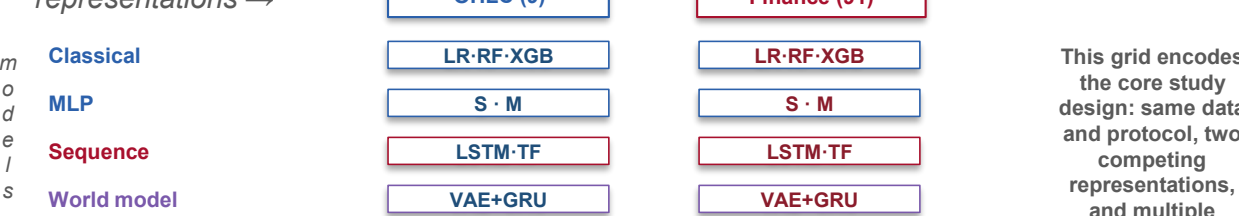
0.818	best directional accuracy
-13%	LSTM volatility MAE reduction
0.784 vs 0.768	walk-forward mean accuracy
-0.039	transformer delta with added features

PROTOCOL & WORLD-MODEL INSET

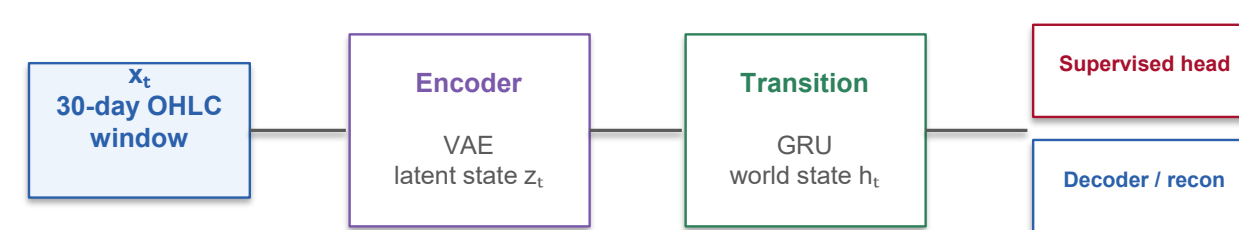
Temporal evaluation protocol



Experimental comparison grid



World-model inset



Composite objective: $L = \lambda_{\text{sup}} L_{\text{sup}} + L_{\text{recon}} + \beta L_{\text{KL}} + \gamma L_{\text{trans}}$
with $\lambda_{\text{sup}} = 5.0, \beta = 0.5, \gamma = 1.0$

Why the world model matters here

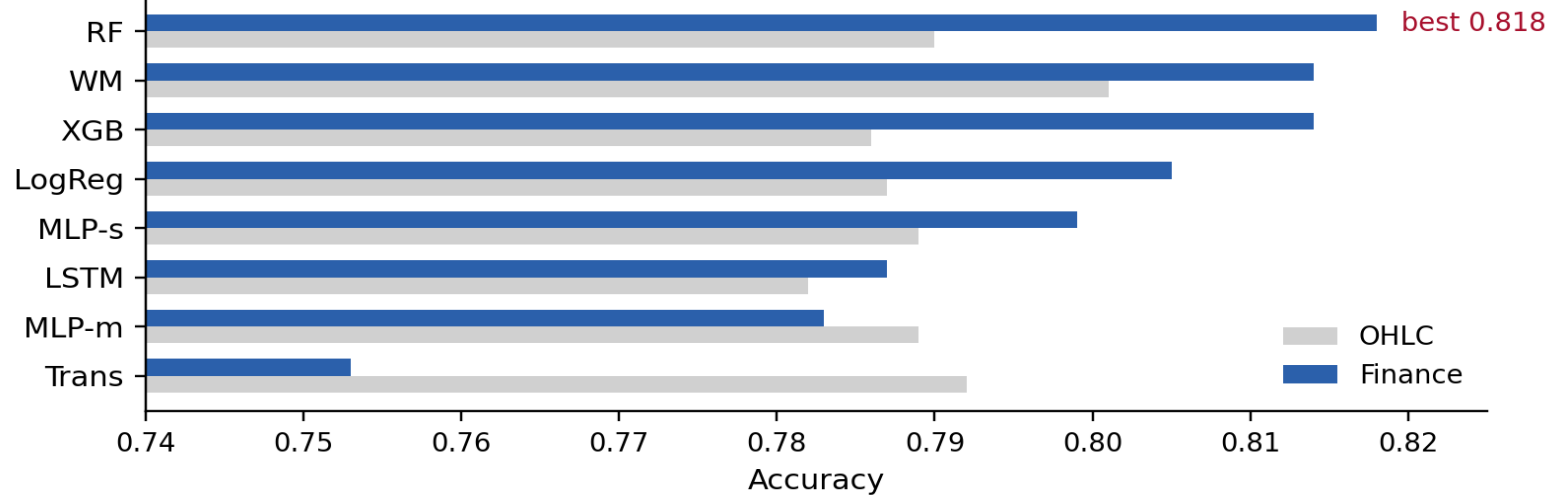
It adds a latent-dynamics comparator rather than only standard tabular and sequence baselines. In the benchmark, it is the strongest OHLC-only classifier. This makes it a natural fit for WM@Booth, where modeling latent state and transition structure is central.

Protocol takeaway

The study is deliberately controlled: the same market panel, the same chronology, and the same evaluation protocol are used throughout. That makes the comparison interpretable. When gains appear, we can attribute them more credibly to representation quality and model family rather than to confounds in splitting or preprocessing.

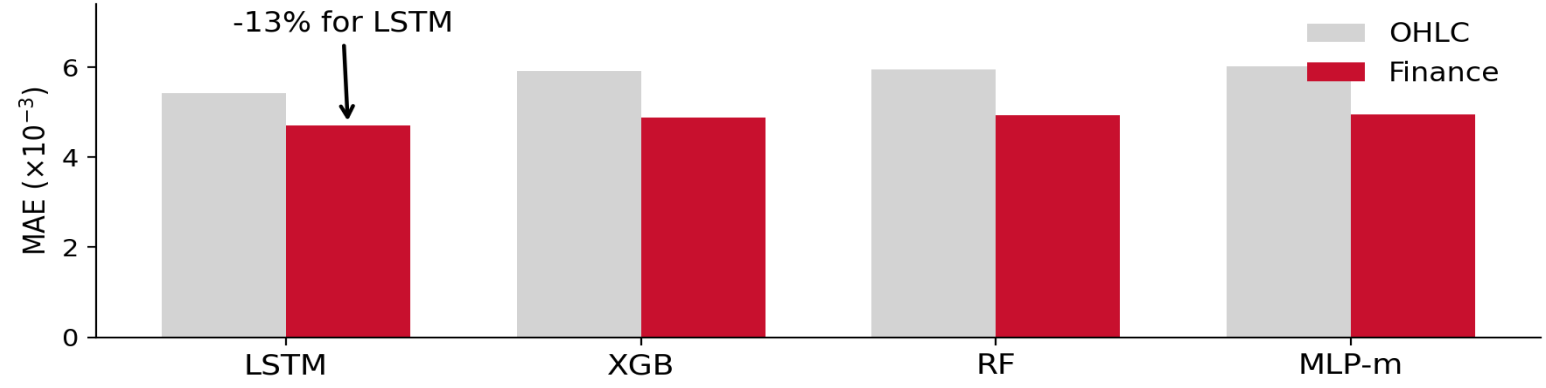
RESULTS & INTERPRETATION

Directional classification performance



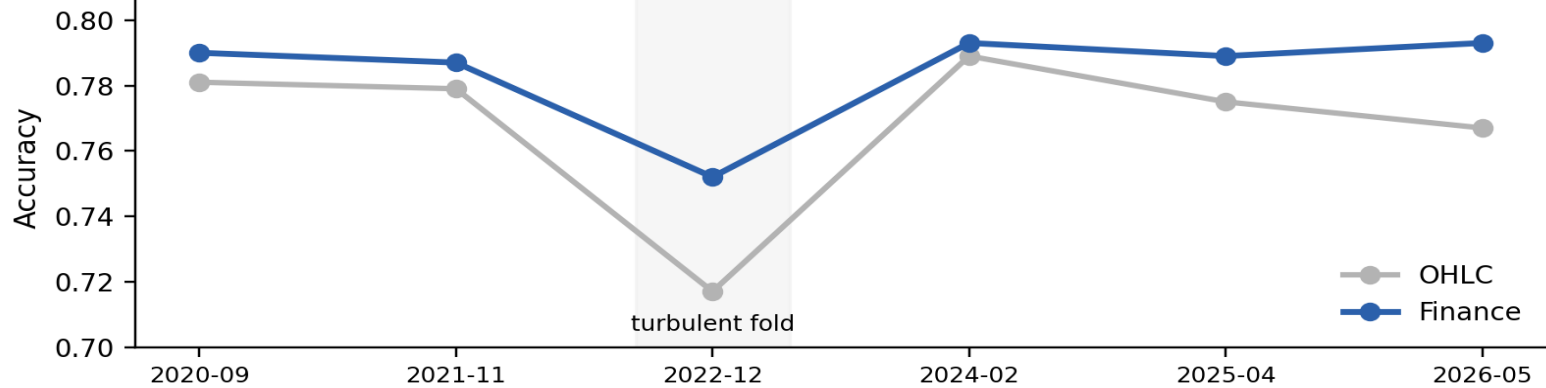
0.818 Best directional accuracy: Random Forest with finance-informed features. The two biggest positive deltas are RF and XGBoost (+0.028 each).

Volatility forecasting error (lower is better)



-13% For volatility, engineered features reduce LSTM MAE by roughly 13%.

Walk-forward robustness across 6 temporal folds



0.784 vs 0.768 Finance-informed features produce a higher and steadier walk-forward mean accuracy.

Interpretation

Feature engineering helps direction and volatility prediction. The gains concentrate in tree-based / simple models rather than in deeper architectures. The compact world model is the strongest OHLC-only classifier. Next-period return regression shows little exploitable signal.

TAKEAWAYS, LIMITATIONS & NEXT STEPS

1. Representation wins

Engineered finance features improve direction and volatility prediction, especially for tabular models.

2. Complexity is not enough

The transformer does not win, while the compact world model is the strongest OHLC-only classifier.

3. Returns remain hard

Raw return regression remains weak; the useful structure lies more in direction and volatility.

Limitations

Single data source, fixed hyperparameters, and no transaction-cost / portfolio-construction evaluation.

Next steps

Add cost-aware backtesting, extend to more asset classes, and tune models individually after the controlled benchmark stage.

Core contribution: a controlled financial ML benchmark showing that representation quality improves both peak performance and temporal robustness more reliably than architectural escalation alone.