

MASTER
APPLIED ECONOMETRICS AND FORECASTING

MASTER'S FINAL WORK
DISSERTATION

**THE PERFORMANCE OF THE TABULAR PRIOR-DATA FITTED
NETWORK IN VOLATILITY FORECASTING: AN EMPIRICAL
COMPARISON WITH HAR AND MACHINE LEARNING MODELS**

GERRIT WILLEM GIJS NIJLAND

SUPERVISION:

JOÃO AFONSO BASTOS

JULY - 2026

GLOSSARY

HAR - Heterogeneous Autoregressive

AR - Autoregressive

ML - Machine learning

RV - Realized variance

KAN - Kolmogorov-Arnold Network

TKAN - Temporal Kolmogorov-Arnold Network

RNN - Recurrent Neural Network

LSTM - Long Short-Term Memory

GRU - Gated Recurrent Unit

TCN - Temporal Convolutional Network

XGB - Extreme Gradient Boosting

LGBM - Light Gradient Boosting Machine

TabPFN - Tabular Prior-data Fitted Network

SCM - Structural Causal Model

MSE - Mean squared error

QLIKE - Quasi-likelihood

MCS - Model Confidence Set

DM - Diebold-Mariano

CW - Clark-West

EnbPI - Ensemble Batch Prediction Intervals

LOO - Leave-one-out

QS - Quantile score

WQS - Weighted quantile score

ADF - Augmented Dickey-Fuller

Keywords: Volatility forecasting, Realized variance, High frequency data, Heterogeneous Autoregressive model, Machine learning, Kolmogorov-Arnold Network, TabPFN, Conformal prediction, Quantile forecasting

ABSTRACT

This paper forecasts realized variance, a nonparametric proxy for volatility constructed from high-frequency data, for the stock of Emerson Electric. Intraday price data spanning approximately 18 years are obtained from the Trade and Quote database hosted by Wharton Research Data Services, and the resulting realized variance series is complemented by a set of exogenous predictors. Forecasts are generated using several variants of the Heterogeneous Autoregressive model alongside a range of machine learning approaches, including the Kolmogorov–Arnold Network and the Tabular Prior-data Fitted Network. For each model, we evaluate conditional mean forecasts, prediction intervals, and quantile forecasts across three horizons: one day ahead, the average week ahead, and the average month ahead. In this single-stock setting, TabPFN achieves the strongest performance across all forecasting tasks at the week- and month-ahead horizons, while no single model dominates at the day-ahead horizon.

Contents

1	Introduction	4
2	Literature review	6
3	Data	10
4	Methodology	13
5	Empirical Results	31
6	Conclusion	49
7	Discussion	51
A	Appendix	58

1 INTRODUCTION

Volatility forecasting is a central topic in financial econometrics, due to its direct relevance for risk management, asset pricing and portfolio allocation. (Poon & Granger 2005, 2003, Andersen et al. 2005) While asset returns are notoriously difficult to predict due to low serial correlation in logarithmic returns, volatility exhibits a higher degree of predictability because it exhibits stylized facts like volatility clustering, long memory persistence, heavy tail return distributions and leverage effects. (Teräsvirta & Zhao 2014, Baillie 1996, Huang & Luo 2024, Bekaert & Wu 2000) This has led to the development of a wide range of econometric and statistical techniques looking to capture the complexity of volatility dynamics, (So et al. 2022) think of the GARCH model as developed by Bollerslev (1986), and the numerous extensions that followed after its original publication. (Teräsvirta 2009) Also among these volatility models is the Heterogeneous Autoregressive (HAR) model of Corsi (2009), which has become a benchmark model due to its ability to approximate long memory patterns using a parsimonious and interpretable linear structure. The model combines non-parametrically constructed realized variance regressors at multiple frequencies, embedded within a parametric AR(3) functional form. After the initial publication of the model, numerous extensions have been proposed to account for nonlinearities, structural breaks and jumps that are not uncommon in financial markets. (Khan 2015, Wen et al. 2016, Andersen et al. 2012)

Despite the development of econometric methods, the increased availability of rich information sets in today's world has stimulated interest in machine learning methods for volatility forecasting. Machine learning models are designed to capture nonlinear interactions and deal with large predictor spaces without requiring strong parametric assumptions. (Gunnarsson et al. 2024) Recent empirical work suggests that machine learning techniques can extract predictive signals beyond those captured by econometric models; traditional econometric models often suffer from overfitting once additional variables are included, whereas machine learning methods are generally better equipped to handle high-dimensional predictor spaces. (Christensen et al. 2023)

However, machine learning models are frequently criticized for being black boxes; their internal mechanisms are often opaque, making it hard to understand how inputs translate into predictions. Unlike traditional econometric models with interpretable parameters, many machine learning algorithms rely on complex nonlinear structures that offer little economic intuition. This lack of transparency can make them harder to explain and implement for practitioners. (Hoepner et al. 2021)

This paper seeks to build on and extend the existing literature on volatility forecasting in several important ways. Recent work has explored the use of machine learning algorithms for forecasting volatility. (Bucci 2020) More specifically, several studies have adopted this framework of comparing the forecasting performance of HAR models with various machine learning approaches. (Rahimikia & Poon 2026, Christensen et al. 2023, Brauneis & Sahiner 2024, Teller et al. 2022, Chun et al. 2025, Taneva-Angelova & Granchev 2025) This paper contributes to that literature by evaluating a broader range of machine learning models and by introducing the relatively new Kolmogorov–Arnold Network and TabPFN models, which, to the best of my knowledge, have not yet been applied in a realized variance forecasting setting.

Lastly, most studies comparing econometric and machine learning models focus mainly on point forecast accuracy. While informative, point forecasts provide limited insight into forecast uncertainty and tails of the forecasting distribution. To address this gap, we also evaluate prediction intervals and extreme outcomes. Specifically, we use Conformal Prediction to construct prediction and quantile regression to forecast tail events.

2 LITERATURE REVIEW

Christensen et al. (2023) compare several variants of the HAR model, including log-HAR, LevHAR, SHAR, and HARQ, with a range of machine-learning methods, including bootstrap aggregation, random forests, gradient boosting, and neural networks with varying depths and numbers of neurons. The analysis is conducted on a dataset of 29 individual stocks contained in the Dow Jones Industrial Average, spanning 17 years, using daily realized variances computed from intraday trading data. They find that machine learning models outperform HAR models in forecasting realized variance in a large regressor space-setting. Moreover, especially random forests and neural networks are preferred as they especially give superior prediction results relative to HAR models, particularly at longer horizons. A similar setup is used by Branco et al. (2024), who forecast realized volatility for ten global equity indices using intraday realized variance data over a period of 21 years, also including exogenous variables, amounting to 27 variables in total. Models include HAR, HAR-X, logHAR, LevHAR and SHAR, among others, as well as regularized versions. With respect to machine learning models, they include bootstrap aggregation, random forests, gradient boosting, fully connected neural networks with different topologies. They find that additional predictors improve forecasts for daily and weekly time horizons, where HAR-X and regularized linear models often have the lowest mean squared error. Also, machine learning models do not systematically outperform linear models. Next, Audrino & Chassot (2025) research the forecasting performance of the HAR model relative to a range of machine learning techniques using a dataset comprising 1,445 individual stocks over a time period of almost 9 years. They include realized volatility and the VIX in the predictor space. Despite extensive hyperparameter tuning for the ML methods, the machine learning models do not outperform the HAR benchmark when the HAR model is carefully fitted. The study finds that the HAR model's linear specification and interpretability as well as a ease of computing make it a practical tool for realized volatility forecasting. The authors conclude that properly specified HAR models provide superior forecasting accuracy when the predictor space is small. Rahimikia & Poon (2026) evaluate the forecasting performance of a broad set of machine learning models for realized volatility using daily data from 2007 to 2022. The authors construct a rich feature set consisting of six HAR variables, 132 limit order book variables and nine news sentiment variables. Across more than seven million model variants, they find that ML models consistently outperform HAR models for roughly 90% of out-of-sample days, particularly when realized volatility is not extremely high. However, during periods of extreme volatility HAR models outperform ML methods, indicating that ML models perform worse in turbulent market regimes. Overall, the authors conclude that

neither ML nor HAR models dominate uniformly; performance depends on model design and volatility regimes. The authors emphasize that ML should not directly replace HAR in practice, and they highlight opportunities for future work, including multi-step forecasting, regime-switching methods and tree-based models. Li et al. (2024) use 5-minute high-frequency data from nine Chinese educational companies (2015–2022) to construct a weighted volatility index and compare the forecasting performance of the HAR model to a set of LSTM neural-network models across daily, weekly, and monthly horizons. Their results show that the HAR framework is effective for capturing short-, medium-, and long-memory components of volatility, but the LSTM models consistently outperform HAR models at all horizons, especially for weekly and monthly forecasts, where nonlinear dynamics are more dominant. The LSTM variants produce substantially lower RMSE, MAE, and MAPE values, with the largest performance improvement observed in the monthly volatility forecasts, indicating that LSTM’s ability to model nonlinear and long-term dependencies gives it a clear advantage over the linear HAR structure. Overall, the study concludes that LSTM delivers superior predictive accuracy on daily, weekly, and monthly volatility, with the performance gap widening as the forecast horizon increases. Li et al. (2025) study realized volatility in the Shanghai Agricultural Stock Index using 5-minute high-frequency data from March 2017 to May 2021, evaluating the comparative forecasting performance of HAR, LSTM, and a hybrid PCA-LSTM model. Across all experiments, the results show that LSTM consistently outperforms the HAR model under identical forecasting conditions, delivering higher predictive accuracy. When the models are expanded to a larger regressor space, the advantage of LSTM becomes even clearer. Most notably, the study finds that the PCA-LSTM model achieves the best overall performance, producing the lowest forecasting errors among all models. The authors conclude that while HAR captures long memory structure effectively, machine learning architectures, especially PCA-LSTM, provide superior accuracy for realized volatility forecasting.

Kolmogorov–Arnold Networks have recently been applied to a range of time-series forecasting problems, among which in financial settings characterized by nonlinear dynamics and evolving regimes. Xu et al. (2024) investigate the use of KANs for volatility forecasting and introduce two variants: Temporal KAN for univariate series and Multivariate Temporal KAN for multivariate environments. Using financial time-series constructed from daily OHLCV data with realized volatility as the target, they show that T-KAN attains predictive performance comparable to or exceeding MLP, RNN, and LSTM benchmarks despite a substantially smaller parameter set, while MT-KAN further improves accuracy by explicitly exploiting cross-series interactions. The authors attribute these re-

sults to the architectural principle of KANs, which model complex temporal relationships through adaptive nonlinear mappings applied at the level of individual inputs, allowing the model to respond flexibly to persistence, regime changes, and nonlinear volatility feedback while remaining interpretable through post-training symbolic analysis. Complementing this evidence, Yamak et al. (2025) provide a comprehensive review of KAN-based models for time-series forecasting across univariate, multivariate, and high-frequency settings, covering benchmark and real-world datasets from electricity load, traffic, weather, financial markets, hydrology, and sensor networks. The review documents that KAN-based models and their extensions, including TKAN, SigKAN, and hybrid variants incorporating attention or convolutional components, frequently match or outperform LSTMs and Transformers while relying on fewer effective parameters and offering direct insight into feature contributions. At the same time, the authors highlight important trade-offs, including increased computational cost during training, scalability challenges in high-dimensional settings, sensitivity to architectural and regularization choices, and a gradual erosion of interpretability as KANs are embedded in more complex hybrid frameworks. Next, Genet & Inzirillo (2024) propose the Temporal Kolmogorov–Arnold Transformer, which replaces recurrent modules in the Temporal Fusion Transformer with Temporal KAN layers tailored to sequential data. The model is evaluated on hourly cryptocurrency trading data over 2020–2022, with the task of multi-step prediction of traded notional volumes using observed market signals and calendar features. The results show that TKAT outperforms GRU- and LSTM-based benchmarks at longer forecasting horizons and that substituting LSTM layers with KAN-based temporal components consistently improves predictive stability within a transformer framework. These gains are attributed to KAN-based temporal processing that captures nonlinear and regime-dependent dynamics more effectively than standard recurrent layers, although the authors note higher training costs and weaker single-step performance due to the overall architectural complexity. Zhang et al. (2025) study the use of Kolmogorov–Arnold Networks in electricity demand forecasting by integrating KAN components into three established time-series models: Temporal Convolutional Networks, Bidirectional LSTMs, and Transformers. Using monthly UK electricity demand data from 2015 to 2024, augmented with economic and meteorological variables such as temperature, precipitation, wind speed, and consumer prices, the authors evaluate the proposed hybrid models against their standard counterparts. The empirical results show that incorporating KAN layers consistently improves forecasting accuracy across all architectures, with notable reductions in RMSE and MAE and higher explanatory power, particularly for medium- to long-horizon forecasts. The performance gains are attributed to the ability of KAN components to refine feature representations and capture complex nonlinear relationships that are insufficiently modeled by standard

fully connected output layers. At the same time, the authors note trade-offs related to increased model complexity and computational cost, especially when KANs are combined with already resource-intensive architectures such as Transformers, highlighting the need to balance accuracy gains against practical deployment considerations. Lastly, Hoo et al. (2025) propose TabPFN-TS, which uses the tabular model TabPFN-v2 for time-series forecasting by converting each time step into a row of features and applying it as a regression model. TabPFN itself is a Transformer trained beforehand on many synthetic tabular datasets to approximate Bayesian predictions: at inference, it takes the observed data as input and directly outputs predictions for new points in a single forward pass, with no need for retraining or optimization on the dataset. (Hollmann et al. 2022) The authors evaluate performance using 28 forecasting tasks across domains like energy, retail, and finance. They find that TabPFN-TS performs best among the compared models in this setting, outperforming time-series foundation models such as TiRex, Chronos-Bolt, Moirai, Toto, and Sundial on the same tasks in terms of MASE and WQL errors.

3 DATA

3.1 Computation of realized variance

We use a dataset containing 18 years of intraday (millisecond timestamps) price data for Emerson Electric, listed on the New York Stock Exchange, covering the period January 3, 2006 to December 29, 2023. This period includes several major macroeconomic events that contributed to increased market-wide volatility, including the Global Financial Crisis (2007–2009), the COVID-19 market crash (2020), and the Russia–Ukraine war (2022–present). The data is retrieved from the NYSE Trade and Quote database, hosted by Wharton Research Data Services. Regarding data filtering, we retain only observations recorded during regular trading hours (09:30–16:00) and remove outliers using algorithms as implemented by Barndorff-Nielsen et al. (2009). Additionally, we exclude the data for December 12, 2006, because Emerson Electric underwent a 2-for-1 stock split on that day, which halved the share price and generated an extremely high daily realized variance that does not reflect any underlying economic event. After following this procedure, we end up with data for $T = 4,528$ trading days. As the daily realized variance will be constructed from the intraday 5-minute logarithmic returns of the stock price, it is useful to understand the assumptions on how the logarithmic price evolves over time. In definition 3.1, those dynamics are explained.

Definition 3.1 *Logarithmic price dynamics*

Consider the stochastic process $X = (\Omega, \mathcal{F}, (\mathcal{F})_{t \geq 0}, (X_t)_{t \geq 0}, \mathbb{P})$ where $(\Omega, \mathcal{F}, \mathbb{P})$ is a probability space, $(\mathcal{F})_{t \geq 0}$ is a filtration, i.e. an increasing family of σ -algebras and $(X_t)_{t \geq 0}$ is the logarithmic price process. The logarithmic price follows the semi-martingale process as specified in 1:

$$dX(t) = \mu(t)dt + \sigma(t)dW(t), \quad \forall t \in [0, T], \quad (1)$$

where $\mu(t)$ is a drift process with finite variation, $\sigma(t)$ is a càdlàg stochastic volatility process and $W(t)$ is a standard Brownian Motion.

To capture volatility, we consider the concept of integrated variance, which is the cumulative variance of a process over a specified time interval. Since the integrated variance over day t is defined as

$$IV_t = \int_{t-1}^t \sigma(s)^2 ds \quad (2)$$

and is unobserved, we require an estimator for it. Under the assumed log-price dynamics in equation 1, the process is a continuous Itô semimartingale without jumps. In

this setting, the quadratic variation of $X(t)$ over $[t-1, t]$ equals the integrated variance, i.e. $QV_t = IV_t$. Barndorff-Nielsen & Shephard (2002) show that realized variance constructed from high-frequency intraday returns satisfies $RV_t \xrightarrow[n \rightarrow \infty]{p} QV_t$. Because $QV_t = IV_t$ under continuity, realized variance provides a consistent estimator of integrated variance, so $RV_t \xrightarrow[n \rightarrow \infty]{p} IV_t$. Moreover, Barndorff-Nielsen & Shephard (2002) showed that as the intraday sampling frequency increases, the error distribution converges to that given in 3.

$$\sqrt{n}(RV_t - IV_t) \xrightarrow{d} \mathcal{N}\left(0, 2 \int_{t-1}^t \sigma_s^4 ds\right) \quad (3)$$

The daily realized variance is constructed from intraday 5-minute logarithmic returns, and the explicit construction of daily realized variance is displayed in 4:

$$RV_t = \sum_{i=1}^n r_{t,i}^2 \quad \text{where} \quad r_{t,i} = \log\left(\frac{P_{t,i}}{P_{t,i-1}}\right), \quad (4)$$

where t is the day and $n = 1/\Delta$ is the number of intraday logarithmic returns, where Δ is the length of intraday sampling interval (5 minutes). Lastly, $r_{t,i}$ is the i -th intraday log-return of day t .

3.2 Exogenous variables

In addition to the regressors constructed from past realized variances, the paper incorporates a set of exogenous variables, namely: the U.S. 3-month Treasury Bill rate (TMTB), the Economic Policy Uncertainty index (EPU), the Chicago Board Options Exchange Volatility Index (VIX), crude oil prices and the Business Conditions index (BC). The data for the first four variables are sourced from the Federal Reserve Bank of St.Louis, and the data for the last variable is sourced from the Federal Reserve Bank Philadelphia.

3.3 Summary statistics

Let $\{RV_t\}_{t=1}^T$ denote the daily realized variance series. In our daily-frequency setting, the monthly HAR regressor $RV_{t-1|t-22}$ is constructed from the preceding 22 trading days. Therefore, the first 22 observations of the daily realized variance series cannot be used for model estimation, as they are required to compute the monthly HAR component. Consequently, this variable first becomes available on the 23rd day, which we set as $t = 1$ in our sample. For the exogenous variables, we handle the few missing values by replacing the missing value with observed value of the day before. In table 1, the summary statistics of all variables are presented.

Table 1: Summary Statistics

Acronym	Mean	Median	Minimum	Maximum	Standard deviation	Skewness	Kurtosis
RV_t	3.604	1.670	0.165	190.694	8.929	10.782	163.061
$RV_{t t-4}$	3.605	1.914	0.281	124.711	7.108	8.564	100.036
$RV_{t t-21}$	3.605	2.171	0.427	68.755	5.951	6.460	50.451
VIX	19.745	17.340	9.140	82.690	9.047	2.408	8.593
EPU	117.920	96.690	3.320	807.660	83.203	2.345	8.662
COP	72.491	71.450	-36.980	145.310	22.365	0.175	-0.384
TMTB	1.302	0.220	-0.050	5.360	1.733	1.227	0.038
BC	-0.287	-0.123	-26.578	9.647	2.220	-6.884	77.098

Because we want to report results for two different information sets, we define the corresponding data sets as $\mathcal{D}_{\text{HAR}}$ and $\mathcal{D}_{\text{HAR-X}}$. Let the lagged realized variance components entering the HAR model be collected in the vector

$$X_t = (RV_t, RV_{t|t-4}, RV_{t|t-21})^\top, \quad (5)$$

where RV_t denotes daily lag of realized variance, $RV_{t|t-4}$ average weekly lag of realized variance, and $RV_{t|t-21}$ average monthly lag of realized variance. In addition, let the set of exogenous variables be defined as

$$Z_t = (EPU_t, VIX_t, BC_t, COP_t, \Delta TMTB_t^1)^\top. \quad (6)$$

Using these components, the two data sets are formally defined as $\mathcal{D}_{\text{HAR}} = \{X_t\}_{t=1}^T$ and $\mathcal{D}_{\text{HAR-X}} = \{(X_t, Z_t)\}_{t=1}^T$, respectively.

¹The variable is differenced to account for nonstationarity, as can be concluded from appendix A.1.

4 METHODOLOGY

4.1 Econometric models

The volatility modelling literature has been strongly influenced by the Heterogeneous Autoregressive model introduced by Corsi (2009), which has become a well-established benchmark in volatility forecasting due to its simplicity, interpretability, and strong empirical performance. By capturing the heterogeneous nature of volatility dynamics through daily, weekly, and monthly components, the HAR framework provides a flexible yet parsimonious way to model the long-memory volatility behavior in financial markets. Building on this foundation, numerous extensions have been proposed to incorporate exogenous predictors, nonlinear dynamics, leverage effects, asymmetries, and structural breaks, thereby broadening the model's applicability in modern volatility forecasting research. In this paper, we will consider a subset of them.

4.1.1 HAR model

The standard heterogeneous autoregressive model of Corsi (2009) captures volatility dynamics by combining daily, weekly, and monthly lagged aggregates, and is specified in 7:

$$RV_t = \beta_0 + \beta_1 RV_{t-1} + \beta_2 RV_{t-1|t-5} + \beta_3 RV_{t-1|t-22} + \varepsilon_t, \quad (7)$$

where the regressors are defined as lagged realized variances: the daily component is RV_{t-1} , the weekly component is $RV_{t-1|t-5} \equiv \frac{1}{5} \sum_{i=1}^5 RV_{t-i}$, and the monthly component is $RV_{t-1|t-22} \equiv \frac{1}{22} \sum_{i=1}^{22} RV_{t-i}$.

4.1.2 HAR-X model

The idea of the HAR-X model, as used by Clements et al. (2024), is an extension of the HAR model in the sense that next to the endogenous realized variance regressors, the regressor space is augmented with the additional exogenous regressors. Hence, the dataset used for estimation and forecast evaluation is $\mathcal{D}_{\text{HAR-X}}$. The HAR-X model functional form is described in 8:

$$RV_t = \beta_0 + \beta_1 RV_{t-1} + \beta_2 RV_{t-1|t-5} + \beta_3 RV_{t-1|t-22} + \beta_4 EPU_{t-1} + \beta_5 VIX_{t-1} + \beta_6 BC_{t-1} + \beta_7 COP_{t-1} + \beta_8 \Delta TMTB_{t-1} + \varepsilon_t. \quad (8)$$

4.1.3 log-HAR model

The log-HAR model (Corsi 2009) models volatility in terms of the logarithm of realized variance rather than its level. Taking logs stabilizes the variance and mitigates the influence of extreme volatility observations that arise from the heavy-tailed nature of realized variance.

$$\log RV_t = \beta_0 + \beta_1 \log RV_{t-1} + \beta_2 \log RV_{t-1|t-5} + \beta_3 \log RV_{t-1|t-22} + \varepsilon_t \quad (9)$$

When converting the log-HAR forecast of $\log RV_t$ back to the level RV_t , it is important to note that one cannot simply apply the exponential function. The log-HAR model produces a prediction $\hat{f}(X_{t-1})$ (or $\hat{f}(X_{t-1}, Z_{t-1})$ in the HAR-X specification) of $\log RV_t$, where $\hat{f}$ denotes the estimated forecast function mapping past observations X_{t-1} to a forecast of $\log RV_t$. Because the exponential transformation is non-linear, the naive back-transformation $\exp(\hat{f})$ is biased by Jensen's inequality. To obtain an unbiased forecast of realized variance, we apply the standard log-normal bias correction, as also used by Christensen et al. (2023):

$$\mathbb{E}\left[\exp\left(\hat{f}(X_{t-1})\right)\right] = \exp\left(\hat{f}(X_{t-1}) + \frac{1}{2} \widehat{\text{Var}}\left(\hat{f}(X_{t-1})\right)\right), \quad (10)$$

and analogously for the HAR-X specification:

$$\mathbb{E}\left[\exp\left(\hat{f}(X_{t-1}, Z_{t-1})\right)\right] = \exp\left(\hat{f}(X_{t-1}, Z_{t-1}) + \frac{1}{2} \widehat{\text{Var}}\left(\hat{f}(X_{t-1}, Z_{t-1})\right)\right). \quad (11)$$

Here, $\widehat{\text{Var}}(\hat{f})$ denotes the estimated variance of the log-HAR residuals computed on the training set. This correction is appropriate because the residuals of the log-HAR regression are approximately Gaussian, i.e., $\log RV_t$ is approximately conditionally normal given its predictors².

4.1.4 HAR-Q model

The HAR-Q model, as introduced by Bollerslev et al. (2016), incorporates realized quarticity as a proxy for intraday return variability, capturing the volatility-of-volatility. By incorporating realized quarticity as a measure of volatility-of-volatility, the HAR-Q model adapts to changes in intraday uncertainty, making it better suited for forecasting realized variance during periods of market turbulence. The HAR-Q model is defined as

²See Appendix A.3 for a visualization of the distribution of the training set residuals.

in 12, where $RQ_t = \frac{n}{3} \sum_{i=1}^n r_{t,i}^4$:

$$RV_t = \beta_0 + \left(\beta_1 + \beta_1^Q RQ_{t-1}^{1/2} \right) RV_{t-1} + \beta_2 RV_{t-1|t-5} + \beta_3 RV_{t-1|t-22} + \varepsilon_t. \quad (12)$$

4.1.5 LevHAR model

The LevHAR (Corsi & Renò 2012) model incorporates leverage terms to capture the asymmetric response of volatility to negative returns. This modification incorporates the well-known leverage effect and might improve forecasting accuracy when downside risk dominates market movements. The LevHAR model is defined as in 13:

$$RV_t = \beta_0 + \beta_1 RV_{t-1} + \beta_2 RV_{t-1|t-5} + \beta_3 RV_{t-1|t-22} + \gamma_1 r_{t-1}^- + \gamma_2 r_{t-1|t-5}^- + \gamma_3 r_{t-1|t-22}^- + \varepsilon_t, \quad (13)$$

where $r_{t-1|t-h} = \frac{1}{h} \sum_{i=1}^h r_{t-i}$ and $r_{t-1|t-h}^- = \min(0, r_{t-1|t-h})$

4.1.6 SHAR model

Another way to model asymmetry is by considering the SHAR (Patton & Sheppard 2015) model, which decomposes realized variance into positive and negative semivariances to explicitly model volatility asymmetry. Since downside and upside returns might have different predictive content for future volatility, including semivariances might yield a more informative HAR structure, as described in 14;

$$RV_t = \beta_0 + \beta_1^- RV_{t-1}^- + \beta_1^+ RV_{t-1}^+ + \beta_2 RV_{t-1|t-5} + \beta_3 RV_{t-1|t-22} + \varepsilon_t, \quad (14)$$

where the positive and negative semivariances are defined as in 15:

$$RV_t^+ = \sum_{i=1}^n r_{t,i}^2 \mathbf{1}_{\{r_{t,i} > 0\}} \quad \text{and} \quad RV_t^- = \sum_{i=1}^n r_{t,i}^2 \mathbf{1}_{\{r_{t,i} < 0\}}. \quad (15)$$

4.2 Machine learning models

We now turn to several machine learning models with network- and tree-based architectures. While our primary focus lies on the (Temporal) Kolmogorov-Arnold Network and TabPFN, and on comparing their forecasting performance to the HAR models introduced earlier, it is also good practice to benchmark KANs against other neural network³

³All machine learning models (except TabPFN and TKAN) are estimated by minimizing a squared-error loss, combined with the architecture-specific regularization term.

and tree-based architectures. In addition, comparing these alternative network models to the HAR benchmarks provides a more complete picture of their relative strengths and weaknesses, particularly because we evaluate forecasting performance across multiple time horizons, where different model structures may perform better or worse relative to one another on different time horizons.

4.2.1 Kolmogorov-Arnold Network

Let $\mathbf{x}_t \in \mathbb{R}^d$ collect the predictors observed at time t , that is the regressors in $\mathcal{D}_{\text{HAR}}$ and, for the extended specification, additionally those in $\mathcal{D}_{\text{HAR-X}}$, and let y_t denote the target. At each forecast origin the estimation window contains the N observations $\{(\mathbf{x}_i, y_i)\}_{i=1}^N$. Bold lowercase symbols denote vectors and bold uppercase symbols matrices or operators. The Kolmogorov-Arnold network of Liu et al. (2024) is derived from the representation theorem in Definition 4.2.

Definition 4.2 *Kolmogorov-Arnold representation theorem*

Let $f : [0, 1]^d \rightarrow \mathbb{R}$ be a multivariate continuous function on a compact domain. Then f can be written as a finite composition of univariate continuous functions and addition: there exist continuous functions $\phi_{q,p} : [0, 1] \rightarrow \mathbb{R}$ and $\Phi_q : \mathbb{R} \rightarrow \mathbb{R}$ such that for all $\mathbf{x} = (x_1, \dots, x_d) \in [0, 1]^d$,

$$f(\mathbf{x}) = \sum_{q=1}^{2d+1} \Phi_q \left(\sum_{p=1}^d \phi_{q,p}(x_p) \right). \quad (16)$$

Theorem 4.2 guarantees a representation with one inner layer of width $2d + 1$. The architecture of Liu et al. (2024) generalizes it to arbitrary depth and width, which is the form used here: layers are indexed by $l = 0, \dots, L$, written as a superscript, with n_l the width of layer l , $n_0 = d$ and $n_L = 1$. Instead of applying a fixed activation function at each node, KANs attach a learnable univariate function to each edge between nodes, parameterized by a B-spline basis.

Definition 4.3 *B-splines*

Let $\{\iota_v\}_{v \in \mathbb{Z}}$ be a non-decreasing knot sequence. The B-spline basis functions $B_v^{(\kappa)} : \mathbb{R} \rightarrow \mathbb{R}$ of degree κ follow the Cox-de Boor recursion

$$B_v^{(0)}(x) = \mathbf{1}_{[\iota_v, \iota_{v+1})}(x), \quad B_v^{(\kappa)}(x) = \frac{x - \iota_v}{\iota_{v+\kappa} - \iota_v} B_v^{(\kappa-1)}(x) + \frac{\iota_{v+\kappa+1} - x}{\iota_{v+\kappa+1} - \iota_{v+1}} B_{v+1}^{(\kappa-1)}(x).$$

Each $B_v^{(\kappa)}$ is a nonnegative piecewise polynomial with local support on $[\iota_v, \iota_{v+\kappa+1}]$ and $\kappa - 1$ continuous derivatives (De Boor 1972).

Let $\phi_{j,p}^{(l)}$ denote the edge function connecting node p in layer l to node j in layer $l + 1$. It is a linear combination of m B-spline basis functions on a bounded interval $[\underline{x}, \bar{x}]$ that covers the range of the standardized inputs,

$$\phi_{j,p}^{(l)}(x) = \sum_{v=1}^m c_{j,p,v}^{(l)} B_v^{(\kappa)}(x), \quad x \in [\underline{x}, \bar{x}], \quad (17)$$

where the coefficients $c_{j,p,v}^{(l)} \in \mathbb{R}$ are learnable and determine the shape of $\phi_{j,p}^{(l)}$. We set $\kappa = 3$, so the edge functions are twice continuously differentiable cubic splines. A node sums the incoming transformed signals, so that a KAN layer computes

$$x_j^{(l+1)} = \sum_{p=1}^{n_l} \phi_{j,p}^{(l)}(x_p^{(l)}), \quad j = 1, \dots, n_{l+1}, \quad (18)$$

with $\mathbf{x}^{(l)} \in \mathbb{R}^{n_l}$ the vector of activations in layer l (Liu et al. 2024). Collecting the edge functions of a layer in the array

$$\Phi^{(l)} = \begin{pmatrix} \phi_{1,1}^{(l)} & \cdots & \phi_{1,n_l}^{(l)} \\ \vdots & \ddots & \vdots \\ \phi_{n_{l+1},1}^{(l)} & \cdots & \phi_{n_{l+1},n_l}^{(l)} \end{pmatrix} \quad \text{gives} \quad \mathbf{x}^{(l+1)} = \Phi^{(l)}(\mathbf{x}^{(l)}), \quad (19)$$

where $\Phi^{(l)} : \mathbb{R}^{n_l} \rightarrow \mathbb{R}^{n_{l+1}}$ is a nonlinear operator whose entries are functions rather than scalars (Somvanshi et al. 2025). Applying it means evaluating each $\phi_{j,p}^{(l)}$ at $x_p^{(l)}$ and summing over p , as in (18); it is therefore not a matrix product, and $\Phi^{(l)}$ is to be distinguished from the scalar outer functions Φ_q in (16). With $\mathbf{x}^{(0)} = \mathbf{x}_t$ and $n_L = 1$, the forecast is the composition of the layers,

$$\hat{y}_t = (\Phi^{(L-1)} \circ \Phi^{(L-2)} \circ \dots \circ \Phi^{(0)})(\mathbf{x}_t). \quad (20)$$

The objective minimized over the estimation window is

$$\mathcal{J}_{\text{KAN}} = \frac{1}{N} \sum_{i=1}^N (y_i - \hat{y}_i)^2 + \lambda_\phi \sum_{l,j,p} \int_{\underline{x}}^{\bar{x}} \left(\frac{\partial^2}{\partial x^2} \phi_{j,p}^{(l)}(x) \right)^2 dx, \quad \lambda_\phi > 0, \quad (21)$$

following Yamak et al. (2025). The second derivative measures how sharply an edge function bends; penalizing its square discourages rapid changes in slope and so limits

how irregular the fitted functions can become.

4.2.2 Temporal Kolmogorov-Arnold Network

The model above maps a single vector $\mathbf{x}_t$ to a forecast. The remaining network models take an ordered input sequence $(\mathbf{x}_{t-S+1}, \dots, \mathbf{x}_t)$ of length S , with sequence positions indexed by $u = t - S + 1, \dots, t$. The building block of a Temporal Kolmogorov-Arnold Network (Genet & Inzirillo 2024, Yamak et al. 2025) is the Recurrent Kolmogorov-Arnold Network layer, which augments (18) with an explicit memory state. Let $x_p^{(l)}(u)$ be the activation of node p in layer l at position u , and $s_p^{(l)}(u)$ the associated memory state. An RKAN layer computes

$$x_j^{(l+1)}(u) = \sum_{p=1}^{n_l} \psi_{j,p}^{(l)}(x_p^{(l)}(u), s_p^{(l)}(u)), \quad j = 1, \dots, n_{l+1}, \quad (22)$$

where $\psi_{j,p}^{(l)} : \mathbb{R}^2 \rightarrow \mathbb{R}$ is a learnable edge function, parameterized as in (17) but taking the memory state as a second argument, so that its behaviour depends on the recent past (Genet & Inzirillo 2024). The memory state of each node evolves as

$$s_p^{(l)}(u) = w_{ss}^{(l)} s_p^{(l)}(u-1) + w_{sx}^{(l)} x_p^{(l)}(u), \quad s_p^{(l)}(t-S) = 0, \quad (23)$$

with learnable scalars $w_{ss}^{(l)}, w_{sx}^{(l)} \in \mathbb{R}$ controlling temporal persistence (Genet & Inzirillo 2024). Equations (22) and (23) define a layer operator

$$\mathbf{x}^{(l+1)}(u) = \Psi^{(l)}(\mathbf{x}^{(l)}(u), \mathbf{s}^{(l)}(u)), \quad \Psi^{(l)} : \mathbb{R}^{n_l} \times \mathbb{R}^{n_l} \rightarrow \mathbb{R}^{n_{l+1}}, \quad (24)$$

with $\mathbf{s}^{(l)}(u) = (s_1^{(l)}(u), \dots, s_{n_l}^{(l)}(u))$. Following Genet & Inzirillo (2024), these layers are still composed across depth as in (20), but the composition is now time-indexed: at each position u the operators are applied in sequence from $\mathbf{x}^{(0)}(u) = \mathbf{x}_u$, with each memory state $\mathbf{s}^{(l)}(u)$ generated from $\mathbf{x}^{(l)}(u)$ through (23) before (24) is applied, while the memory states are propagated recursively along $u = t - S + 1, \dots, t$. In the full TKAN each RKAN transformation is embedded in an LSTM-style cell (Section 4.2.4), whose gates regulate how much of the accumulated cell-state memory is carried forward across positions (Genet & Inzirillo 2024). With $n_L = 1$, the forecast is the activation of the single output node at the final sequence position,

$$\hat{y}_t = x_1^{(L)}(t). \quad (25)$$

4.2.3 Recurrent Neural Network

A recurrent neural network maps the input sequence to a forecast through a recursively updated hidden state (Lipton et al. 2015). Let $\mathbf{h}_u \in \mathbb{R}^{d_h}$ denote this hidden state at position u , a vector of d_h units that carries information forward along the sequence, and let $\tanh(\cdot)$ be the activation applied elementwise. The same weights act at every position, and the state recursion is

$$\mathbf{h}_u = \tanh(\mathbf{W}_x \mathbf{x}_u + \mathbf{W}_h \mathbf{h}_{u-1} + \mathbf{b}_h), \quad u = t - S + 1, \dots, t, \quad (26)$$

where $\mathbf{W}_x \in \mathbb{R}^{d_h \times d}$ is the input weight matrix mapping the d -dimensional input $\mathbf{x}_u$ into the hidden space, $\mathbf{W}_h \in \mathbb{R}^{d_h \times d_h}$ is the recurrent weight matrix acting on the previous state, and $\mathbf{b}_h \in \mathbb{R}^{d_h}$ is a bias vector; all three are learnable. The state is initialised at $\mathbf{h}_{t-S} = \mathbf{0}$, so each sequence starts from a blank memory. By (26), $\mathbf{h}_u$ depends on $\mathbf{x}_u$ and $\mathbf{h}_{u-1}$, and therefore recursively on $(\mathbf{x}_{t-S+1}, \dots, \mathbf{x}_u)$. Since a single value is forecast, the state at the final position is mapped to the target by a linear layer,

$$\hat{y}_t = \mathbf{w}^\top \mathbf{h}_t + b, \quad (27)$$

where $\mathbf{w} \in \mathbb{R}^{d_h}$ is a learnable weight vector, $b \in \mathbb{R}$ a learnable scalar bias, and $\hat{y}_t$ the forecast for position t .

4.2.4 Long Short-Term Memory

A Long Short-Term Memory network (Hochreiter & Schmidhuber 1997) is a recurrent architecture designed to mitigate the vanishing-gradient problem through gated additive state updates. Alongside the hidden state $\mathbf{h}_u \in \mathbb{R}^{d_h}$ it maintains a cell state $\mathbf{c}_u \in \mathbb{R}^{d_h}$, which carries information across positions additively. Let $\sigma(\cdot)$ denote the logistic sigmoid. At each position u , given $\mathbf{x}_u$ and $\mathbf{h}_{u-1}$, the LSTM computes an input, forget and output gate and a candidate cell update,

$$\begin{aligned} \mathbf{i}_u &= \sigma(\mathbf{W}_i \mathbf{x}_u + \mathbf{U}_i \mathbf{h}_{u-1} + \mathbf{b}_i), & \mathbf{f}_u &= \sigma(\mathbf{W}_f \mathbf{x}_u + \mathbf{U}_f \mathbf{h}_{u-1} + \mathbf{b}_f), \\ \mathbf{o}_u &= \sigma(\mathbf{W}_o \mathbf{x}_u + \mathbf{U}_o \mathbf{h}_{u-1} + \mathbf{b}_o), & \tilde{\mathbf{c}}_u &= \tanh(\mathbf{W}_c \mathbf{x}_u + \mathbf{U}_c \mathbf{h}_{u-1} + \mathbf{b}_c), \end{aligned} \quad (28)$$

for $u = t - S + 1, \dots, t$, where the subscripts i, f, o and c label the gate rather than index a sum, $\mathbf{W}_\bullet \in \mathbb{R}^{d_h \times d}$ act on the input, $\mathbf{U}_\bullet \in \mathbb{R}^{d_h \times d_h}$ on the previous hidden state, and $\mathbf{b}_\bullet \in \mathbb{R}^{d_h}$ are bias vectors, all learnable. The two states are updated as

$$\mathbf{c}_u = \mathbf{f}_u \odot \mathbf{c}_{u-1} + \mathbf{i}_u \odot \tilde{\mathbf{c}}_u, \quad \mathbf{h}_u = \mathbf{o}_u \odot \tanh(\mathbf{c}_u), \quad \mathbf{c}_{t-S} = \mathbf{h}_{t-S} = \mathbf{0}, \quad (29)$$

where $\odot$ denotes elementwise multiplication. The forget gate $\mathbf{f}_u$ controls the retention and decay of past information, the input gate $\mathbf{i}_u$ determines how new information is incorporated, and the output gate $\mathbf{o}_u$ governs the exposure of the cell state. Because (29) updates $\mathbf{c}_u$ additively, information can be preserved over long horizons, while the hidden state remains bounded by the $\tanh$ transformation. The forecast is obtained by mapping the final hidden state $\mathbf{h}_t$ to the target through a linear layer, $\hat{y}_t = \mathbf{w}^\top \mathbf{h}_t + b$, with learnable $\mathbf{w} \in \mathbb{R}^{d_h}$ and $b \in \mathbb{R}$.

4.2.5 Gated Recurrent Unit

The Gated Recurrent Unit (Chung et al. 2014) simplifies the LSTM by not maintaining a separate cell state: it keeps only a single hidden state $\mathbf{h}_u \in \mathbb{R}^{d_h}$, on which two gates act directly,

$$\mathbf{z}_u = \sigma(\mathbf{W}_z \mathbf{x}_u + \mathbf{U}_z \mathbf{h}_{u-1} + \mathbf{b}_z), \quad \mathbf{r}_u = \sigma(\mathbf{W}_r \mathbf{x}_u + \mathbf{U}_r \mathbf{h}_{u-1} + \mathbf{b}_r), \quad (30)$$

for $u = t - S + 1, \dots, t$, where z and r label the update and reset gate (not summation indices), $\sigma(\cdot)$ is the logistic sigmoid, and the learnable parameters $\mathbf{W}_\bullet \in \mathbb{R}^{d_h \times d}$, $\mathbf{U}_\bullet \in \mathbb{R}^{d_h \times d_h}$ and $\mathbf{b}_\bullet \in \mathbb{R}^{d_h}$ act on the input, the previous hidden state and as bias respectively. The reset gate determines how much of the previous state enters the candidate state,

$$\tilde{\mathbf{h}}_u = \tanh(\mathbf{W}_h \mathbf{x}_u + \mathbf{U}_h (\mathbf{r}_u \odot \mathbf{h}_{u-1}) + \mathbf{b}_h), \quad (31)$$

while the update gate controls how much of that candidate replaces the previous state,

$$\mathbf{h}_u = (\mathbf{1} - \mathbf{z}_u) \odot \mathbf{h}_{u-1} + \mathbf{z}_u \odot \tilde{\mathbf{h}}_u, \quad \mathbf{h}_{t-S} = \mathbf{0}, \quad (32)$$

where $\odot$ is elementwise multiplication and $\mathbf{1}$ a vector of ones. Because the candidate is added to the retained state rather than overwriting it, information can be carried across long horizons, and unlike the LSTM the state is exposed in full at every position. The forecast is obtained by mapping the final hidden state to the target through a linear layer, $\hat{y}_t = \mathbf{w}^\top \mathbf{h}_t + b$, with learnable $\mathbf{w} \in \mathbb{R}^{d_h}$ and $b \in \mathbb{R}$.

4.2.6 Temporal Convolutional Network

A Temporal Convolutional Network (Lea et al. 2016) models sequential data using one-dimensional convolutions along the temporal dimension rather than recurrent state updates. Time dependence is captured by dilated causal convolutions, which ensure that

the representation at a given position depends only on the current and preceding inputs, so that the time order is preserved. Let $\mathbf{h}_u^{(l)} \in \mathbb{R}^{n_l}$ denote the representation at position u in layer l , with $\mathbf{h}_u^{(0)} = \mathbf{x}_u$ and $\mathbf{h}_u^{(l)} = \mathbf{0}$ for $u < t - S + 1$, so that positions preceding the input sequence contribute nothing to the convolution. Each layer computes

$$\mathbf{h}_u^{(l)} = a \left(\sum_{k=0}^{K-1} \mathbf{W}_k^{(l)} \mathbf{h}_{u-\delta_l k}^{(l-1)} + \mathbf{b}^{(l)} \right), \quad l = 1, \dots, L, \quad u = t - S + 1, \dots, t, \quad (33)$$

where $\mathbf{W}_k^{(l)} \in \mathbb{R}^{n_l \times n_{l-1}}$ are the filter weights, $\mathbf{b}^{(l)} \in \mathbb{R}^{n_l}$ is a bias vector, K is the filter size, k indexes the taps within the filter and δ_l is the dilation factor of layer l , which sets the spacing between the input elements entering the convolution. The elementwise activation $a(\cdot)$ is essential: without it the composition of (33) across layers would collapse to a single linear filter. Because the sum runs only over indices $u - \delta_l k \leq u$, the representation at u is a function of current and past inputs only, and increasing δ_l with depth widens the span of history covered without enlarging K . The forecast is obtained by mapping the final-layer representation at the last position to the target through a linear layer, $\hat{y}_t = \mathbf{w}^\top \mathbf{h}_t^{(L)} + b$, with learnable $\mathbf{w} \in \mathbb{R}^{n_L}$ and $b \in \mathbb{R}$.

4.2.7 Extreme Gradient Boosting

Extreme Gradient Boosting models the prediction as an additive ensemble of regression trees (Chen & Guestrin 2016). Using the estimation window $\{(\mathbf{x}_i, y_i)\}_{i=1}^N$, the prediction after r boosting rounds is

$$\hat{y}_i^{(r)} = \sum_{k=1}^r f_k(\mathbf{x}_i), \quad f_k \in \mathcal{F}, \quad (34)$$

where $\mathcal{F}$ is the space of regression trees and f_k assigns a score to each of its terminal leaves. The final prediction is $\hat{y}_i = \hat{y}_i^{(R)}$, with R the total number of trees. With the squared-error loss, the training objective is

$$\mathcal{J}_{\text{XGB}} = \frac{1}{2} \sum_{i=1}^N (y_i - \hat{y}_i)^2 + \sum_{k=1}^R \left(\gamma J_k + \frac{1}{2} \lambda \sum_{j=1}^{J_k} w_{j,k}^2 \right), \quad (35)$$

where J_k is the number of leaves of tree k , $w_{j,k}$ the prediction assigned to leaf j of tree k , and λ and γ penalize large leaf weights and tree complexity respectively. Trees are added one at a time, each fitted to a second-order Taylor expansion of the loss around the current prediction. Within this subsection let g_i and h_i denote the first- and second-order

derivatives of the loss with respect to $\hat{y}_i^{(r-1)}$, which for the squared-error loss in (35) are

$$g_i = \hat{y}_i^{(r-1)} - y_i, \quad h_i = 1. \quad (36)$$

For a fixed tree structure, the optimal prediction in leaf j is

$$w_j^* = -\frac{\sum_{i \in I_j} g_i}{\sum_{i \in I_j} h_i + \lambda}, \quad (37)$$

with I_j the set of observations falling into that leaf. Splits are selected by the implied reduction in the objective: for a candidate split of a leaf into children I_L and I_R ,

$$\text{Gain} = \frac{1}{2} \left(\frac{(\sum_{i \in I_L} g_i)^2}{\sum_{i \in I_L} h_i + \lambda} + \frac{(\sum_{i \in I_R} g_i)^2}{\sum_{i \in I_R} h_i + \lambda} - \frac{(\sum_{i \in I_L \cup I_R} g_i)^2}{\sum_{i \in I_L \cup I_R} h_i + \lambda} \right) - \gamma, \quad (38)$$

and the split is retained only if this quantity is positive (Chen & Guestrin 2016).

4.2.8 *Light Gradient Boosting Machine*

LightGBM (Ke et al. 2017) is a gradient boosting algorithm designed to improve computational efficiency on large and high-dimensional datasets. Predictions are an additive ensemble of trees as in (34), fitted sequentially to the negative gradients $-g_i$ from (36) using the same second-order approximation and split criterion (38). The differences lie in tree construction and data handling. LightGBM grows trees leaf-wise, splitting at each step the leaf with the largest gain, which produces deeper and asymmetric trees and is therefore usually combined with a cap on the number of leaves J_k . For efficiency it discretizes continuous features into histogram bins and accumulates gradient statistics at the bin level, and it applies Gradient-Based One-Side Sampling, which retains observations with large $|g_i|$ while subsampling those with small gradients, and Exclusive Feature Bundling, which reduces the feature dimension by combining mutually exclusive sparse features (Ke et al. 2017).

4.2.9 *Tabular Prior-data Fitted Network*

TabPFN⁴ is a Prior-Data Fitted Network, abbreviated here as PFN. A prior defines a space $\mathcal{H}$ of data-generating mechanisms relating an input $\mathbf{x}$ to an output y ; a mechanism $\eta \in \mathcal{H}$ has prior probability $p(\eta)$ and induces a distribution $p(D \mid \eta)$ over datasets. Given observed pairs $D_{\text{train}} = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_M, y_M)\}$, the posterior predictive distribution for

⁴In this paper, we use the TabPFN-3 model released by Prior Labs.

a query $\mathbf{x}$ averages over mechanisms, weighting each by how well it explains the data,

$$p(y \mid \mathbf{x}, D_{\text{train}}) \propto \int_{\mathcal{H}} p(y \mid \mathbf{x}, \eta) p(D_{\text{train}} \mid \eta) p(\eta) d\eta. \quad (39)$$

The integral is intractable, so a PFN replaces it by a transformer q_θ fitted once, in advance, to synthetic datasets from the prior: a mechanism is drawn $\eta \sim p(\eta)$, a dataset $D \sim p(D \mid \eta)$ is split into D_{train} and a held-out pair $(\mathbf{x}_{\text{test}}, y_{\text{test}})$, and θ minimizes

$$\mathcal{J}_{\text{PFN}} = \mathbb{E}[-\log q_\theta(y_{\text{test}} \mid \mathbf{x}_{\text{test}}, D_{\text{train}})], \quad (40)$$

the expectation taken over these draws. Minimizing (40) makes q_θ approximate (39) (Hollmann et al. 2022), so no mechanism is ever identified: the network maps observed pairs directly to the predictive distribution Bayesian inference would return. Consequently θ is fixed before the empirical data are seen, and a forecast is a single forward pass over D_{train} and the query point, with no estimation, tuning or model selection at the forecast origin. All structure comes from the prior $p(\eta)$, which in TabPFN mixes structural causal models and Bayesian neural networks drawn with a preference for few nodes and parameters (Hollmann et al. 2022). That prior describes dependence between the columns of a table and encodes nothing about time; moreover D_{train} is passed as a permutation-invariant set, so row order carries no information. Any dependence on the past must therefore enter through the features.

To apply TabPFN to realized-variance forecasting, we follow the TabPFN-TS (Hoo et al. 2025) idea of representing a time series as a table with one row per time step. Because each forecast in our setting is made at origin t using information available at t , we can include lagged realized-variance terms directly as features. For each horizon h , we construct a small conditioning set consisting of the most recent M observations,

$$D_t^{(h)} = \left\{ (\mathbf{x}_u, y_u^{(h)}) \right\}_{u=t-M}^{t-1},$$

where $\mathbf{x}_u$ contains the daily, weekly, and monthly realized-variance averages, and, for the HAR-X specification, the exogenous covariates. TabPFN then performs in-context learning on this set and produces a predictive distribution for the next value. Our forecast is the conditional mean,

$$\hat{y}_t^{(h)} = \mathbb{E}_{q_\theta} \left[y \mid \mathbf{x}_t, D_t^{(h)} \right],$$

so that TabPFN conditions on exactly the same information as the HAR and HAR-

X models, differing only in the functional form used to map past observations to future values.

4.3 Accuracy quantification

4.3.1 Loss functions

Volatility forecast accuracy evaluation is complicated by the fact that the conditional variance is latent and must be proxied. Let r_t denote an asset return with conditional mean zero and conditional variance as in 41:

$$\sigma_t^2 \equiv \mathbb{E}[r_t^2 \mid \mathcal{F}_{t-1}]. \quad (41)$$

Following Patton (2011), realized variance is treated as a conditionally unbiased (but noisy) proxy for the latent conditional variance,

$$RV_t = \sigma_t^2 + \eta_t, \quad \mathbb{E}[\eta_t \mid \mathcal{F}_{t-1}] = 0, \quad (42)$$

where the noise term arises from discrete intraday sampling. Under the diffusion framework described in Definition 3.1, realized variance satisfies:

$$\mathbb{E}[RV_t \mid \mathcal{F}_{t-1}] = \sigma_t^2, \quad \text{MSE}_{t-1}(RV_t) = \mathbb{E}[(RV_t - \sigma_t^2)^2 \mid \mathcal{F}_{t-1}] = \frac{2\sigma_t^4}{n}, \quad (43)$$

Thus, increasing the intraday sampling frequency improves the precision of the volatility proxy at rate n^{-1} . In the empirical analysis, realized variance is constructed using 5-minute returns, corresponding to $n = 78$. Because forecast evaluation minimizes expected loss using realized variance rather than the latent conditional variance, the choice of loss function is crucial. As shown by Patton (2011), a loss function is appropriate for volatility forecast comparison if its minimizer remains the true conditional variance despite measurement error in the proxy. Let $\widehat{RV}_{t|t-1}$ denote a volatility forecast formed at time $t - 1$ for day t . An appropriate loss function must satisfy

$$\widehat{RV}_{t|t-1} = \sigma_t^2, \quad (44)$$

which ensures that forecast rankings are not distorted by noise in realized variance. Both the mean squared error and the quasi-likelihood loss satisfy this condition and therefore yield consistent forecast comparisons even at finite intraday sampling frequencies. In

contrast, alternative loss functions such as the mean absolute error generally imply

$$\widehat{RV}_{t|t-1} = \sigma_t^2 + \mathcal{O}(n^{-1}), \quad (45)$$

so that exact consistency is achieved only asymptotically as the sampling frequency increases. Restricting attention to MSE and QLIKE therefore ensures that volatility forecast evaluation in this study reflects genuine differences in predictive accuracy rather than artefacts of volatility measurement error. Accordingly, forecast accuracy is evaluated using the mean squared error and the quasi-likelihood loss. The MSE is defined as

$$\mathcal{L}^{MSE} = (RV_{t+h} - \hat{RV}_{t+h|t})^2. \quad (46)$$

The QLIKE loss, derived from the Gaussian quasi-likelihood, is defined as

$$\mathcal{L}^{QLIKE} = \frac{RV_{t+h}}{\hat{RV}_{t+h|t}} - \log\left(\frac{RV_{t+h}}{\hat{RV}_{t+h|t}}\right). \quad (47)$$

4.3.2 Model Confidence Set

To compare a potentially large collection of competing forecasting models while accounting for statistical uncertainty, we employ the Model Confidence Set procedure of Hansen et al. (2011). The framework aims to identify the subset of models that cannot be statistically distinguished from the best performing model(s) in terms of expected loss at a given confidence level α . Let $\mathcal{M}_0 = \{1, \dots, m_0\}$ denote the initial set of models. For each model $i \in \mathcal{M}_0$, forecast performance is summarized by a sequence of losses $\mathcal{L}_{i,t} = \mathcal{L}(y_t, \hat{y}_{i,t})$. Relative performance between models i and j is measured by the loss differential:

$$d_{ij,t} = \mathcal{L}_{i,t} - \mathcal{L}_{j,t}, \quad (48)$$

with finite and time-invariant expectation $\mu_{ij} = \mathbb{E}(d_{ij,t})$. Models are ranked according to expected loss, and the set of superior models is defined as

$$\mathcal{M}^* = \{i \in \mathcal{M}_0 : \mu_{ij} \leq 0 \text{ for all } j \in \mathcal{M}_0\}. \quad (49)$$

The MCS procedure constructs a data-dependent confidence set $\widehat{\mathcal{M}}_{1-\alpha}^*$ for $\mathcal{M}^*$ via a sequential testing algorithm. Starting from $\mathcal{M}_0$, the null hypothesis of equal predictive ability

$$H_0^{\mathcal{M}} : \mu_{ij} = 0 \text{ for all } i, j \in \mathcal{M} \quad (50)$$

is tested for a candidate subset $\mathcal{M} \subseteq \mathcal{M}_0$. Inference is based on a bootstrap-based range statistic (Hansen et al. 2003) in 51:

$$T_R = \max_{i,j \in \mathcal{M}} \frac{|\bar{d}_{ij}|}{\sqrt{\widehat{\text{Var}}(\bar{d}_{ij})}}, \quad (51)$$

where $\bar{d}_{ij} = T^{-1} \sum_{t=1}^T d_{ij,t}$ and $\widehat{\text{Var}}(\bar{d}_{ij})$ is estimated using a circular moving-block bootstrap to account for potential dependence in the loss differentials. The block length is set data-dependently to the largest BIC-selected autoregressive order across all loss-differential series, capped at 20 lags, so that the blocks preserve the serial dependence of the most persistent pair.⁵ If the EPA null cannot be rejected, the procedure terminates and the remaining models are retained. If the null hypothesis is rejected, model elimination is based on relative average performance. Define

$$\bar{d}_i = \frac{1}{|\mathcal{M}| - 1} \sum_{j \in \mathcal{M}, j \neq i} \bar{d}_{ij}, \quad (52)$$

the model with the largest $\bar{d}_i$ is eliminated according to $e_M = \arg \max_{i \in \mathcal{M}} \bar{d}_i$ (Zammarchi & Kawasaki 2026) and the equal predictive ability hypothesis is re-tested on the reduced model set. This process is repeated until the null can no longer be rejected, yielding the final Model Confidence Set.

4.3.3 Diebold-Mariano test

To assess whether model i and model j differ in predictive accuracy, we employ the Diebold & Mariano (1995) test based on squared-error loss. For each forecast origin $t \in \mathcal{T}_h$, define the loss differential

$$d_{ij,t} = \mathcal{L}_{i,t} - \mathcal{L}_{j,t}, \quad (53)$$

where $\mathcal{T}_h$ denotes the set of all time points at which h -step-ahead forecasts are evaluated. The sample mean of the loss differentials is

$$\bar{d} = \frac{1}{N} \sum_{t \in \mathcal{T}_h} d_{ij,t}, \quad (54)$$

where N is the number of out-of-sample forecasts. Because h -step-ahead forecast errors overlap, the sequence $\{d_{ij,t}\}$ is serially correlated up to lag $h - 1$. A heteroskedasticity-

⁵We assume that $\mathbb{E}[|d_{ij,t}|] < \infty$ and that $\{d_{ij,t}\}$ is stationary and ergodic.

and autocorrelation-consistent estimator of the long-run variance of the loss differentials is therefore required. Following Diebold & Mariano (1995), we estimate

$$\widehat{\text{Var}}(d) = \hat{\gamma}_d(0) + 2 \sum_{l=1}^{h-1} \left(1 - \frac{l}{h}\right) \hat{\gamma}_d(l), \quad (55)$$

where $\hat{\gamma}_d(l)$ denotes the sample autocovariance of $\{d_{ij,t}\}$ at lag l . The weights $1 - l/h$ correspond to the Bartlett kernel, whose bandwidth we set equal to the forecast horizon, so that the truncation lag is $h - 1$ (0, 4 and 21 lags for $h = 1, 5$ and 22 respectively) and the variance estimate is guaranteed non-negative. The Diebold–Mariano test statistic is then defined as

$$DM = \frac{\bar{d}}{\sqrt{\widehat{\text{Var}}(d)/N}}, \quad (56)$$

and is asymptotically standard normal under the null hypothesis of equal predictive accuracy, $\mathbb{E}[d_{ij,t}] = 0$.

4.3.4 Clark-West test

When comparing nested models (HAR nested within LevHAR, HARQ, or SHAR), the classical Diebold-Mariano test tends to favour the more parsimonious model in finite samples. To correct for this upward bias, we apply the Clark & West (2007) test. In what follows, model i denotes the nested HAR specification, while model j represents an enlarged model that nests HAR. For each forecast origin $t \in \mathcal{T}_h$, where $\mathcal{T}_h$ denotes the set of time points at which h -step-ahead forecasts are evaluated, the Clark-West adjusted loss differential is defined as

$$f_{ij,t} = \mathcal{L}_{i,t} - \mathcal{L}_{j,t} + \left(\widehat{RV}_{t+h|t}^{(i)} - \widehat{RV}_{t+h|t}^{(j)} \right)^2. \quad (57)$$

The sample mean of the adjusted loss differentials is

$$\bar{f} = \frac{1}{N} \sum_{t \in \mathcal{T}_h} f_{ij,t}, \quad (58)$$

where N denotes the number of out-of-sample forecasts. Because h -step-ahead forecast errors overlap, the sequence $\{f_{ij,t}\}$ is serially correlated up to lag $h - 1$. Analogously to the Diebold–Mariano test, we therefore employ a heteroskedasticity- and autocorrelation-

consistent estimator of the long-run variance,

$$\widehat{\text{Var}}(f) = \hat{\gamma}_f(0) + 2 \sum_{l=1}^{h-1} \left(1 - \frac{l}{h}\right) \hat{\gamma}_f(l), \quad (59)$$

where $\hat{\gamma}_f(l)$ is the sample autocovariance of $\{f_{ij,t}\}$ at lag l , and the Bartlett weights and horizon-dependent truncation are the same as in (55). The Clark-West test statistic is then defined as

$$CW = \frac{\bar{f}}{\sqrt{\widehat{\text{Var}}(f)/N}}, \quad (60)$$

and is asymptotically standard normal under the null hypothesis that the larger model provides no improvement in predictive accuracy.

4.4 Uncertainty quantification

Conformal prediction provides a distribution-free framework for constructing prediction sets with finite-sample validity. (Shafer & Vovk 2008) The key idea is to assign conformity scores that measure how well observations agree with the data and to construct prediction sets by comparing new samples to these scores. Under exchangeability, the resulting prediction interval $C_\alpha(X_{n+1})$ satisfies 61, independently of the underlying data distribution.

$$\mathbb{P}(Y_{n+1} \in C_\alpha(X_{n+1})) \geq 1 - \alpha \quad (61)$$

In time series settings, exchangeability is violated, motivating extensions such as Ensemble Batch Prediction Intervals (EnbPI) as proposed by Xu & Xie (2021). Consider the model

$$y_t = f(x_t) + \varepsilon_t, \quad (62)$$

where the error process $\{\varepsilon_t\}$ may exhibit temporal dependence. EnbPI constructs residual-based conformity scores in a manner analogous to conformal prediction using leave-one-out (LOO) ensemble predictors. Given bootstrap estimators $\{\hat{f}_b\}_{b=1}^B$, the LOO predictor is

$$\hat{f}_{-i}(x_i) = \mathcal{G}(\{\hat{f}_b(x_i) : i \notin S_b\}), \quad (63)$$

where S_b denotes the set of indices sampled (with replacement) to form the b -th bootstrap dataset used to fit $\hat{f}_b$. The corresponding residuals $\hat{\varepsilon}_i = |y_i - \hat{f}_{-i}(x_i)|$ serve as conformity scores. Prediction intervals at time t are constructed as

$$C_t^\alpha = \left[\hat{f}_{-t}(x_t) \pm q_{1-\alpha}(\{\hat{\varepsilon}_i\}_{i=t-T}^{t-1}) \right], \quad (64)$$

where $q_{1-\alpha}(\cdot)$ denotes the empirical $(1 - \alpha)$ -quantile of the most recent residuals. Hence, intervals are centered at the LOO prediction and their width is determined by past residuals, updated in a rolling window to adapt to temporal changes. EnbPI does not guarantee exact finite-sample marginal coverage. Instead, under mild assumptions on the error process and the approximation quality of the predictors, it achieves approximately valid marginal coverage, as is displayed in 65 (Xu & Xie 2021):

$$\mathbb{P}(Y_t \in C_t^\alpha) \geq 1 - \alpha - e_T, \quad \text{with } e_T \rightarrow 0. \quad (65)$$

4.5 Quantile-based forecast evaluation

Point forecasts summarize the conditional mean of realized variance but provide no information about forecast accuracy in different regions of the conditional distribution. Since practitioners are particularly concerned with increased uncertainty in their risk management, we complement point forecast evaluation with quantile-based methods that explicitly assess predictive performance in the right tail of the distribution.

Quantile forecasts are obtained by minimizing the loss as shown in 66 (Koenker & Bassett Jr 1978), which penalizes under-prediction more heavily than over-prediction for $\tau > 0.5$.

$$\rho_\tau = (RV_{t+h} - \hat{Q}_{t+h|t}(\tau))(\tau - \mathbf{1}_{\{RV_{t+h} < \hat{Q}_{t+h|t}(\tau)\}}) \quad (66)$$

Quantile forecast accuracy is evaluated using the quantile score in 67, as introduced by Gneiting & Raftery (2007),

$$QS_{t+h|t}(\tau) = (RV_{t+h} - \hat{Q}_{t+h|t}(\tau))(\tau - \mathbf{1}_{\{RV_{t+h} < \hat{Q}_{t+h|t}(\tau)\}}), \quad (67)$$

where lower values indicate better predictive accuracy and the score provides a local evaluation at a fixed quantile level. To test equal quantile forecast accuracy relative to a benchmark model between model i and j , we employ the following test statistic, as mentioned by Meligkotsidou et al. (2019):

$$t_{i,j} = \frac{\overline{QS}_i(\tau) - \overline{QS}_j(\tau)}{\sqrt{\widehat{\text{Var}}(QS_i(\tau) - QS_j(\tau)) / N}}, \quad (68)$$

where $\overline{QS}_i(\tau)$ denotes the average quantile score at level τ , and $\widehat{\text{Var}}(QS_i(\tau) - QS_j(\tau))$ is a HAC-consistent estimator of the long-run variance of the quantile score differential. Under the null of equal expected quantile accuracy, $t(\tau)$ is asymptotically standard normal. To augment our single quantile assessment with insight into forecasting performance

over a specific region of the distribution, we also compute the weighted quantile score as in 69:

$$WQS_{t+h|t} = \int_0^1 QS_{t+h|t}(\tau) \omega(\tau) d\tau, \quad (69)$$

which aggregates quantile scores using a weighting function $\omega(\tau)$ that corresponds to a specific region of interest of the distribution. As said, we focus on the right tail of the distribution and therefore set $\omega(\tau) = \tau^2$, as suggested by Meligkotsidou et al. (2019). The differences in WQS are evaluated by using WQS in place of QS in 68.

5 EMPIRICAL RESULTS

In reporting our results, we make a distinction between the results generated using the HAR dataset $\mathcal{D}_{\text{HAR}}$ and results generated using the HAR-X dataset $\mathcal{D}_{\text{HAR-X}}$, to see if certain models perform better or worse when additional predictors are included. Additionally, we want to assess forecasting power for different time horizons, namely one day-ahead, average-week ahead and average-month ahead. We do so by reporting Model Confidence Sets and Diebold-Mariano test results; the two are complementary. Because our model set implies a large number of pairwise comparisons, we rely on the MCS algorithm to limit the risk of us drawing flawed conclusions from repeated pairwise testing; it delivers the subset of models that cannot be rejected as inferior to the best-performing model(s) at confidence level $1 - \alpha$. The Diebold-Mariano test is then used to gain further insight into the relative performance of the reduced number of models retained in the MCS. We evaluate forecast performance using both We evaluate forecast performance using both the mean squared error and the quasi-likelihood loss function. MSE provides a symmetric measure of absolute forecast accuracy, while QLIKE is an asymmetric loss that penalizes volatility underestimation more strongly. This asymmetry is particularly relevant for practitioners, as underestimating volatility is typically more costly than overestimation for risk management. Reporting results under both loss functions therefore allows us to assess accuracy from an absolute level perspective as well as from a risk-sensitive perspective that places greater emphasis on under prediction. Next to accuracy, uncertainty will be evaluated by reporting prediction intervals.

As part of the forecasting procedure, the data are first split into a training set containing the initial 80% of the sample and a test set containing the remaining 20%. For both the econometric models and the machine learning models, forecasts are generated using a rolling window approach with a fixed window length of $W = 3,622$ observations, so that each forecast origin re-estimates the model on the most recent W observations, discarding the oldest observation as one new observation enters the window. After examining the stationarity⁶ of the variables, we find that, at the 5% significance level, the 3-month Treasury Bill rate is non-stationary. To address this issue the variable is transformed by taking its first difference before being included in the HAR-type models. For the machine learning models, all input variables are standardized before training; while such scaling is not required for tree-based algorithms, it is nevertheless applied as it does not impair their performance. For the HAR-type models, when exogenous variables are included, all input variables are standardized prior to model estimation.

⁶The results of the Augmented Dickey-Fuller tests are provided in Appendix A.1.

For the machine learning models, we use rolling window cross-validation on the training set to tune hyperparameters⁷. Hyperparameter tuning is performed separately for each forecasting horizon, to account for the fact that different horizons may benefit from different model architectures and degrees of regularization.

Finally, it is important to clarify the interpretation of the forecasting horizons. The one-day-ahead forecast is denoted by RV_{t+1} . The week-ahead forecast refers to the average realized variance over the subsequent trading week and is denoted by $RV_{t+5|t+1}$. Similarly, the month-ahead forecast corresponds to the average realized variance over the following trading month and is denoted by $RV_{t+22|t+1}$. Here, $RV_{t+h|t+1} \equiv \frac{1}{h} \sum_{i=1}^h RV_{t+i}$.

5.1 Model Confidence Set results

In our application, we use (part of) the procedure as described by Hansen et al. (2003), specifically as how it is outlined in 'Appendix B' in their paper. One can interpret the results as follows: the MCS is a set-valued confidence region that contains the best forecasting model with probability at least $1 - \alpha$. Models included in the MCS cannot be statistically rejected as inferior given the data, but inclusion does not imply equal performance or dominance. (Hansen et al. 2003) We report results for $\mathcal{L}^{QLIKE}$ and $\mathcal{L}^{MSE}$, which are displayed in table 2 and table 3, respectively.

Table 2: Models included in the Model Confidence Set across estimation sets and forecast horizons, $\alpha = 0.05$, $B = 10,000$, loss function: $\mathcal{L}^{QLIKE}$

Estimated on $\mathcal{D}_{HAR}$			Estimated on $\mathcal{D}_{HAR-X}$		
RV_{t+1}	$RV_{t+5 t+1}$	$RV_{t+22 t+1}$	RV_{t+1}	$RV_{t+5 t+1}$	$RV_{t+22 t+1}$
HAR	TabPFN	TabPFN	Log-HARX	TabPFN	TabPFN
Log-HAR			TKANX		
HAR-Q			LGBMX		
SHAR			XGBX		
LevHAR			RNNX		
TKAN					
LGBM					
XGB					
RNN					
TabPFN					

For the day-ahead, we observe that all the HAR models cannot be ruled out as statistically inferior when models are estimated on $\mathcal{D}_{HAR}$, under both loss functions. Under $\mathcal{L}^{MSE}$, only the LSTM gets ruled out. Under $\mathcal{L}^{QLIKE}$, the MCS eliminates the KAN, LSTM, GRU and TCN. When we visually inspecting the forecasts in figures 1 and 2, it can be observed that none of the models is really able to capture the magnitude of the

⁷The hyperparameter tuning procedure is available in Appendix A.4.

Table 3: Models included in the Model Confidence Set across estimation sets and forecast horizons, $\alpha = 0.05$, $B = 10,000$, loss function: $\mathcal{L}^{MSE}$

Estimated on $\mathcal{D}_{HAR}$			Estimated on $\mathcal{D}_{HAR-X}$		
RV_{t+1}	$RV_{t+5 t+1}$	$RV_{t+22 t+1}$	RV_{t+1}	$RV_{t+5 t+1}$	$RV_{t+22 t+1}$
HAR	TabPFN	TabPFN	Log-HARX	TabPFNX	TabPFNX
Log-HAR			TKANX		
HAR-Q			LGBMX		
SHAR			RNNX		
LevHAR			GRUX		
KAN			TabPFNX		
TKAN					
LGBM					
XGB					
RNN					
GRU					
TCN					
TabPFN					

largest variance spikes. Also, quite a lot of models forecast with some lagged response, so they forecast a spike (although of too small magnitude), when the actual surge in volatility already happened.

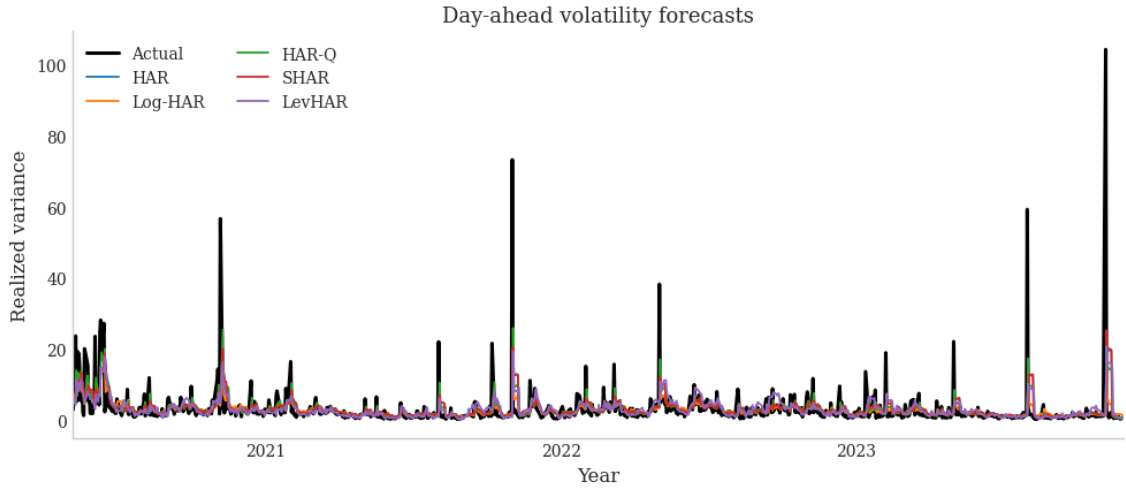


Figure 1: Forecasts of HAR-type models estimated on $\mathcal{D}_{HAR}$

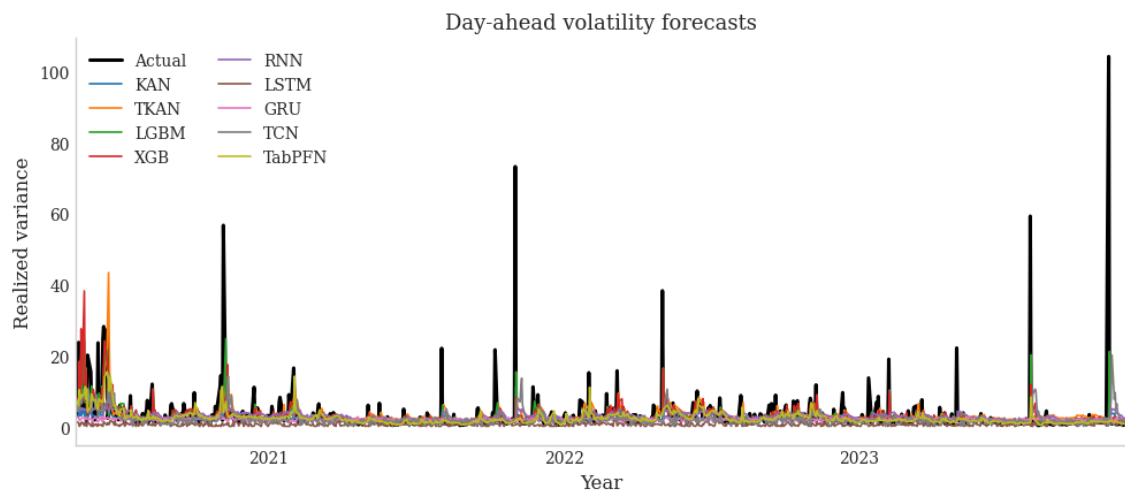


Figure 2: Forecasts of machine learning models estimated on $\mathcal{D}_{HAR}$

Next, when models are estimated on $\mathcal{D}_{HAR-X}$, the only HAR-type model retained in the MCS for the day-ahead horizon under both loss functions is Log-HARX. Moreover, substantially fewer machine learning models are included compared to the case where models are estimated on $\mathcal{D}_{HAR}$, suggesting that the excluded HAR and ML models may not be sufficiently capable of handling the larger predictor set. When we consider the visualizations of the forecasts of models estimated on $\mathcal{D}_{HAR-X}$ in figures 3 and 4, again it is observed neither the HAR nor the machine learning models are able to capture the magnitude of the largest realized variance spikes. This suggests that incorporating additional macroeconomic variables does not substantially improve the models' forecasting accuracy.

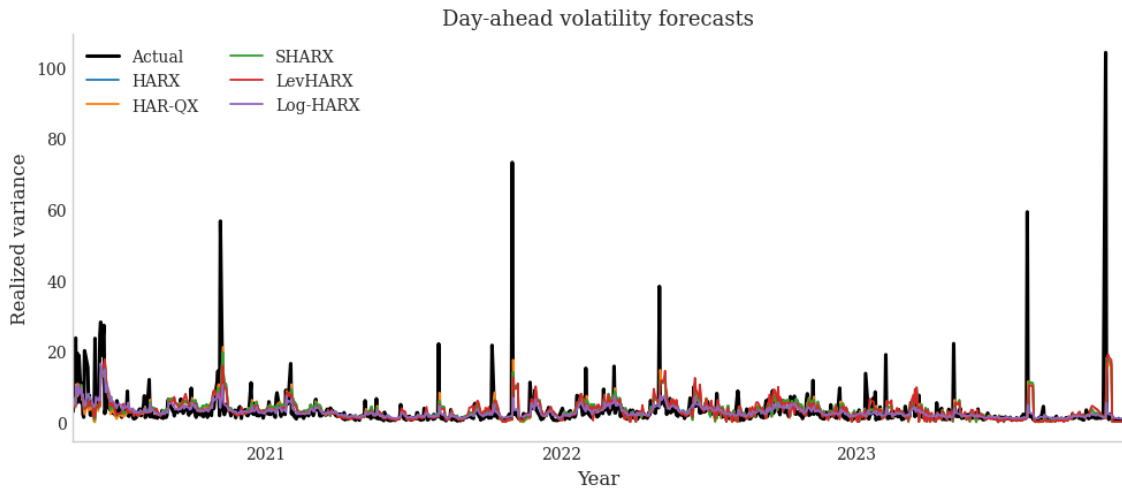


Figure 3: Forecasts of HAR-type models estimated on $\mathcal{D}_{HAR-X}$

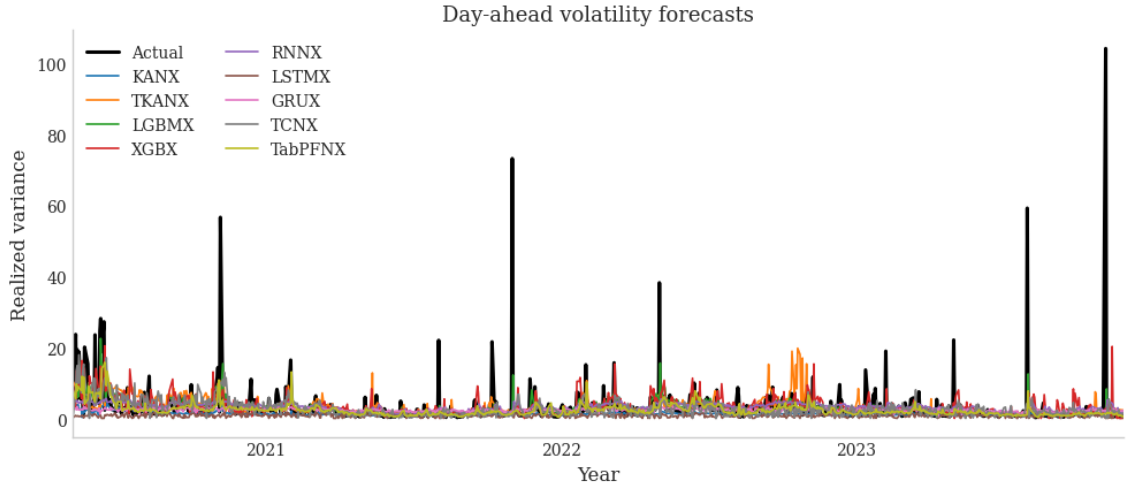


Figure 4: Forecasts of machine learning models estimated on $\mathcal{D}_{HAR-X}$

The results for the week-ahead and month-ahead forecasting horizons strongly favor TabPFN. As shown in Tables 2 and 3, TabPFN is the only model that remains in the Model Confidence Set across both loss functions and both information sets. This indicates that all competing models are statistically rejected as inferior at the 5% significance level, while TabPFN cannot be rejected as belonging to the set of best-performing models. These findings provide strong evidence that TabPFN achieves highly competitive predictive performance relative to all other models at medium and longer forecasting horizons. This conclusion is further supported by the full pairwise Diebold–Mariano test results reported in Appendix A.6. In particular, for models estimated on $\mathcal{D}_{HAR}$, TabPFN consistently exhibits superior predictive ability in pairwise comparisons against every alternative model, under both loss functions and for both the average week-ahead and average month-ahead horizons. The superiority of TabPFN is especially conveniently observable when considering the mean squared error, where it delivers substantially more accurate forecasts than all competing models. The results for the models estimated on $\mathcal{D}_{HAR}$ are largely consistent with those obtained for $\mathcal{D}_{HAR-X}$, with one exception: the comparison between SHARX and TabPFN is statistically insignificant. A visual inspection suggests that this is driven by instability in the SHARX forecasts. In particular, SHARX exhibits severe overprediction toward the end of the test sample, leading to large spikes in the QLIKE loss. This inflates the variance of the loss differentials, reducing the power of the Diebold–Mariano test and resulting in insignificance. Despite this, the relative mean squared error remains low, indicating that the lack of significance is primarily due to irregular behavior of the SHARX forecasts rather than comparable predictive performance.

Increasing the significance level from $\alpha = 0.05$ to $\alpha = 0.10$ makes the Model Confidence Set procedure less conservative and thus more aggressive in eliminating models;

the corresponding results are reported in Appendix A.5. Notably, the LevHAR model is then eliminated under $\mathcal{L}^{QLIKE}$, which, together with Figure 5, suggests that it may underpredict relative to other models, as the quasi-likelihood loss penalizes this strongly; an effect that becomes statistically significant at the 10% level. In addition, TabPFN is also eliminated when α increases, while under $\mathcal{L}^{MSE}$, TCN is further excluded compared to the 5% MCS. For models estimated on $\mathcal{D}_{HAR-X}$, Log-HARX is the only one that survives under $\mathcal{L}^{MSE}$, indicating that a simple log transform of realized variance in combination with its parsimonious structure dominates other more complicated linear and machine learning models, as well as the simpler HAR model. The conclusions for the average week-ahead and month-ahead horizons remain unchanged.

5.2 Pairwise Diebold-Mariano test results

While the Model Confidence Set identifies models that cannot be jointly rejected as inferior, it does not imply equality of predictive accuracy for the models in the MCS. To further characterize relative performance within the confidence set, we report pairwise Diebold–Mariano (and Clark–West) tests.⁸ We only display the loss matrices for the day-ahead forecasts here, as the average week-ahead and average month-ahead sets only contain TabPFN. In figures 5 and 6, the results are displayed. Each cell reports the relative loss between two models, defined as in 70:

$$\frac{\text{MSE}_j}{\text{MSE}_i} \quad \text{or} \quad \frac{\text{QLIKE}_j}{\text{QLIKE}_i}, \quad (70)$$

where i corresponds to row the model in row i and j corresponds to the model in column j . Alongside the relative losses, statistical significance is reported for 10%, 5%, and 1% significance levels. The DM test evaluates

$$H_0 : \mathbb{E}[d_{ij,t}] = 0 \quad \text{vs.} \quad H_1 : \mathbb{E}[d_{ij,t}] \neq 0, \quad (71)$$

where a relative loss smaller (larger) than 1, when accompanied by statistically significant test results, indicates that the model listed in column j has superior (inferior) forecasting performance compared to the model listed in row i , and we indicate this by colouring the cell green (red).

Because several HAR-type models are nested versions of one another, we also report the Clark–West test alongside the Diebold–Mariano test. In nested model settings, the

⁸Full pairwise relative MSE and test results (for both predictor spaces and loss functions) for all evaluated models are provided in appendix A.6.

standard Diebold–Mariano test is biased toward the smaller model, while the Clark–West corrects for this bias. The one sided Clark–West test compares the forecasting accuracy of the HAR(-X) to the HARQ(-X), SHAR(-X) and LevHAR(-X), where the following hypothesis are formulated:

$$H_0 : \mathbb{E}[f_{ij,t}] \leq 0 \quad \text{vs.} \quad H_1 : \mathbb{E}[f_{ij,t}] > 0. \quad (72)$$

Rejection of the null hypothesis in favor of the alternative indicates that the enlarged model delivers a statistically significant improvement in forecasting accuracy relative to the nested HAR benchmark, after correcting for estimation bias arising from model nesting. When interpreting figures 5 and 6, it is important to note that, for nested model comparisons, significance is assigned as follows: the Clark–West test result is reported first, and if this test is insignificant, the Diebold–Mariano test result is used instead. Also note that the models not mentioned in any of the figures are already considered inferior to the models that are mentioned under that specific loss function and predictor set.

When focusing on the models estimated on $\mathcal{D}_{HAR}$ in Figure 5, Log-HAR and TabPFN emerge as the strongest performers under $\mathcal{L}^{MSE}$, outperforming most alternatives in pairwise tests. This suggests that, under a symmetric loss function, these models most frequently deliver the most accurate point forecasts in direct comparisons. In contrast, under $\mathcal{L}^{QLIKE}$, XGB dominates the majority of pairwise comparisons, followed by Log-HAR, indicating they performs better most often when underprediction is penalized more heavily and are therefore particularly effective when avoiding underestimation is crucial. The models with the most inferior pairwise performances under $\mathcal{L}^{MSE}$ seem to be HAR, followed by HAR-Q, GRU and TCN, supporting the argument that even when one only wants to consider HAR-type models, it is still worthwhile to consider alternative specifications from the standard HAR for short horizon forecasting. Under $\mathcal{L}^{QLIKE}$, LevHAR appears to perform worst in most pairwise comparisons among the models in the MCS, followed by TabPFN.

	HAR	Log-HAR	HAR-Q	SHAR	LevHAR	KAN	TKAN	LGBM	XGB	RNN	GRU	TCN	TabPFN
HAR		0.917	1.003	1.006	0.991	0.959	0.961	0.957	0.961	0.923	0.976	1.003	0.910
Log-HAR	1.090		1.093	1.096	1.081	1.045	1.048	1.043	1.047	1.006	1.063	1.093	0.991
HAR-Q	0.997	0.915		1.003	0.989	0.956	0.959	0.954	0.958	0.921	0.973	1.000	0.907
SHAR	0.994	0.912	0.997		0.986	0.953	0.956	0.951	0.955	0.918	0.970	0.997	0.904
LevHAR	1.009	0.925	1.011	1.014		0.967	0.969	0.965	0.969	0.931	0.984	1.011	0.917
KAN	1.043	0.957	1.046	1.049	1.034		1.003	0.998	1.002	0.963	1.018	1.046	0.949
TKAN	1.041	0.955	1.043	1.047	1.032	0.997		0.996	1.000	0.960	1.015	1.043	0.946
LGBM	1.045	0.959	1.048	1.051	1.036	1.002	1.004		1.004	0.964	1.020	1.048	0.950
XGB	1.041	0.955	1.043	1.047	1.032	0.998	1.000	0.996		0.961	1.015	1.044	0.947
RNN	1.084	0.994	1.086	1.090	1.074	1.039	1.041	1.037	1.041		1.057	1.086	0.986
GRU	1.025	0.940	1.028	1.031	1.016	0.982	0.985	0.981	0.985	0.946		1.028	0.932
TCN	0.997	0.915	1.000	1.003	0.989	0.956	0.958	0.954	0.958	0.920	0.973		0.907
TabPFN	1.099	1.009	1.102	1.106	1.090	1.054	1.057	1.052	1.056	1.015	1.073	1.102	

(a) Loss function: $\mathcal{L}^{MSE}$

	HAR	Log-HAR	HAR-Q	SHAR	LevHAR	TKAN	LGBM	XGB	RNN	TabPFN
HAR		0.991	1.008	0.999	1.072	0.988	0.992	0.985	1.024	1.037
Log-HAR	1.009		1.017	1.008	1.081	0.996	1.001	0.993	1.033	1.046
HAR-Q	0.992	0.983		0.991	1.063	0.979	0.984	0.977	1.015	1.029
SHAR	1.001	0.993	1.010		1.073	0.989	0.994	0.986	1.025	1.038
LevHAR	0.933	0.925	0.941	0.932		0.921	0.926	0.919	0.955	0.968
TKAN	1.012	1.004	1.021	1.011	1.085		1.005	0.997	1.037	1.050
LGBM	1.008	0.999	1.016	1.007	1.080	0.995		0.992	1.032	1.045
XGB	1.015	1.007	1.024	1.014	1.088	1.003	1.008		1.039	1.053
RNN	0.977	0.968	0.985	0.976	1.047	0.965	0.969	0.962		1.013
TabPFN	0.964	0.956	0.972	0.963	1.033	0.952	0.957	0.950	0.987	

(b) Loss function: $\mathcal{L}^{QLIKE}$

Figure 5: Relative loss matrices of models estimated on $\mathcal{D}_{HAR}$ and retained in MCS at the day-ahead horizon; cell colours indicate statistical significance of pairwise forecast comparisons (light, medium, and dark shading corresponding to the 10%, 5%, and 1% levels, respectively, with green (red) denoting superior (inferior) performance of the column model relative to the row model).

	Log-HARX	TKANX	LGBMX	RNNX	GRUX	TabPFN
Log-HARX		1.114	1.026	1.043	1.081	1.038
TKANX	0.897		0.921	0.936	0.970	0.931
LGBMX	0.975	1.086		1.017	1.054	1.012
RNNX	0.959	1.068	0.984		1.037	0.995
GRUX	0.925	1.030	0.949	0.965		0.960
TabPFN	0.964	1.074	0.989	1.005	1.042	

(a) Loss function: $\mathcal{L}^{MSE}$

	Log-HARX	TKANX	LGBMX	XGBX	RNNX
Log-HARX		1.029	1.010	1.053	1.030
TKANX	0.972		0.982	1.023	1.001
LGBMX	0.990	1.018		1.042	1.019
XGBX	0.950	0.977	0.960		0.978
RNNX	0.971	0.999	0.981	1.023	

(b) Loss function: $\mathcal{L}^{QLIKE}$

Figure 6: Relative loss matrices of models estimated on $\mathcal{D}_{HAR-X}$ and retained in MCS at the day-ahead horizon; cell colours indicate statistical significance of pairwise forecast comparisons (light, medium, and dark shading corresponding to the 10%, 5%, and 1% levels, respectively, with green (red) denoting superior (inferior) performance of the column model relative to the row model).

Next, when considering the models retained in the MCS when estimated on the extended predictor set in figure 6, it can be observed that Log-HARX is statistically superior to all models retained in the MCS in the respective pairwise comparisons under $\mathcal{L}^{MSE}$, which is noteworthy given its relatively simple specification compared to extensively hyperparameter-tuned machine learning models. Under $\mathcal{L}^{QLIKE}$, all pairwise comparisons are insignificant.

5.3 Prediction Intervals

In this implementation, we use what is denoted as 'Algorithm 1' in Xu & Xie (2021). In our application, we use the mean as the aggregation function $\mathcal{G}$ and we use online updating by setting $s = 1$. Prediction intervals are generated using Ensemble Batch Prediction Intervals for all models except TabPFN. This is because TabPFN directly outputs an approximation of the posterior predictive distribution, from which prediction intervals can be obtained by extracting the corresponding quantiles.

To assess each model's ability to quantify uncertainty, we consider two metrics. Equation 73 defines the empirical coverage (EC), which measures the proportion of observations for which the true value y_t lies within the prediction interval $[L_t, U_t]$:

$$EC = \frac{1}{N} \sum_{t=1}^N \mathbf{1}\{L_t \leq y_t \leq U_t\} \quad (73)$$

In addition, we evaluate the sharpness of the prediction intervals using the median relative interval width, defined as in 74:

$$MRIW = \text{median}_{t=1, \dots, N} \left(\frac{U_t - L_t}{y_t} \right) \quad (74)$$

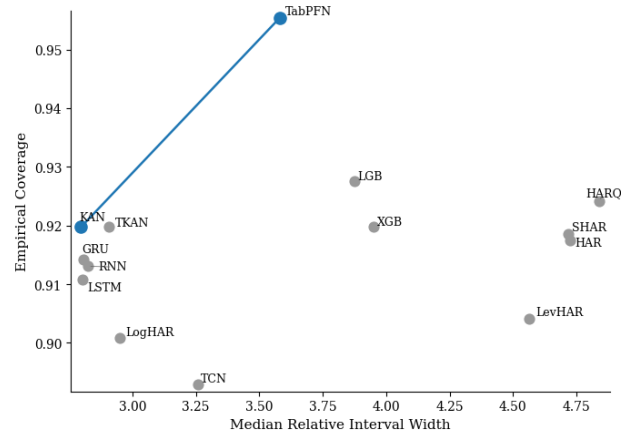
which captures the median width of the prediction intervals relative to the actual values.

Although EnbPI does not guarantee exact marginal coverage in finite samples, it provides approximate coverage guarantees under mild assumptions. In practice, these guarantees may not hold exactly, particularly in finite samples or under model misspecification. Therefore, empirical coverage serves as a practical measure of calibration, indicating how well prediction intervals capture the actual values. When combined with the median relative interval width, which measures the sharpness of the intervals, this allows for a comprehensive assessment of uncertainty quantification by evaluating the trade-off between reliability and informativeness. In Figures 7 and 8, Pareto plots based on 95%

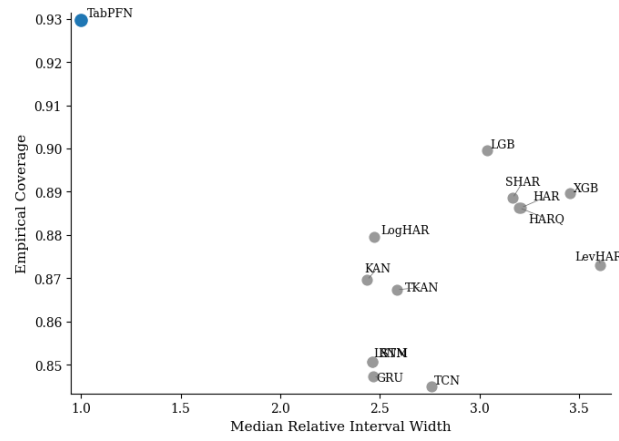
prediction intervals are presented with empirical coverage and median relative interval width on the axes, enabling the identification of Pareto-efficient models.

For the day-ahead horizon, both KAN and TabPFN lie on the Pareto frontier for models estimated on $\mathcal{D}_{HAR}$ and $\mathcal{D}_{HAR-X}$, indicating that neither model is strictly dominated in terms of the trade-off between coverage and interval width. Specifically, KAN achieves the lowest median relative interval width under both information sets, producing the narrowest prediction intervals, whereas TabPFN attains the highest empirical coverage. Another noteworthy result is that the HAR models (except for Log-HAR) seem to produce larger median relative interval widths than the machine learning models, under both information sets. A possible explanation for this finding is the limited flexibility of HAR-type models, which rely on a linear structure. This restricts their ability to capture nonlinearities and more complex dynamics (e.g. regime shifts) in the data, leading to larger residuals. Since EnbPI constructs prediction intervals based on the empirical distribution of these residuals, this translates directly into wider intervals. In contrast, machine learning models can often better approximate the underlying data-generating process, leading to smaller residual dispersion and, consequently, sharper prediction intervals.

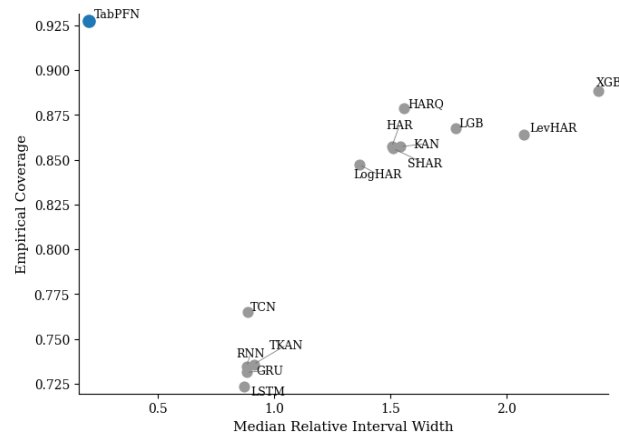
For the average week-ahead and average month-ahead horizons, this trade-off disappears, as TabPFN clearly dominates the competing models. In these settings, it simultaneously achieves the highest empirical coverage and the smallest median relative interval width, and does so with a substantial margin, indicating both highly reliable and informative prediction intervals.



(a) Day-ahead horizon

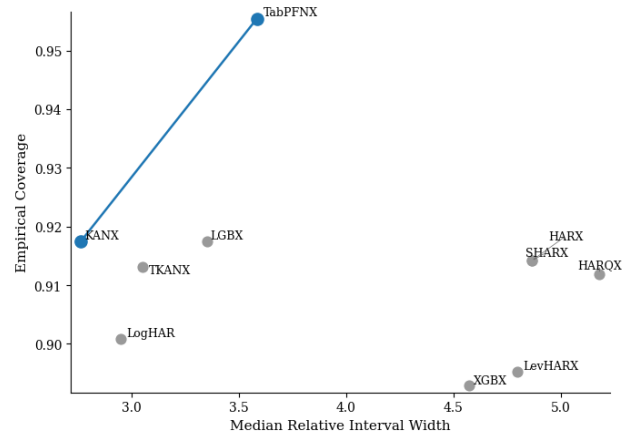


(b) Average week-ahead horizon

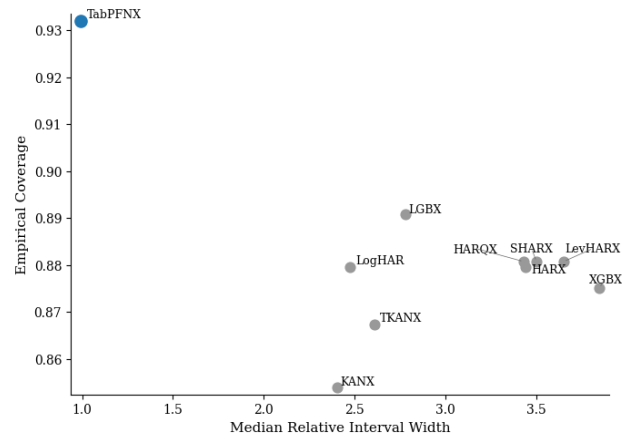


(c) Average month-ahead horizon

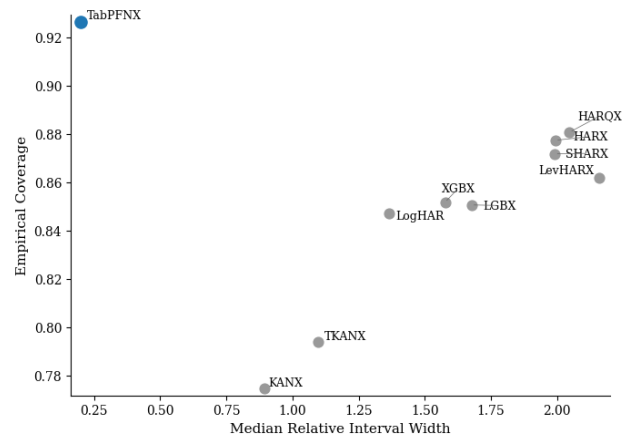
Figure 7: Pareto plot of interval width and empirical coverage, for models estimated on $\mathcal{D}_{HAR}$



(a) Day-ahead horizon



(b) Average week-ahead horizon



(c) Average month-ahead horizon

Figure 8: Pareto plot of interval width and empirical coverage, for models estimated on $\mathcal{D}_{HAR-X}$

5.4 Upper-tail quantile forecasting accuracy

Now, we consider quantile forecast⁹ accuracy, specifically the right tail. This is done by report relative Weighted Quantile Scores¹⁰, calculated as in 75:

$$\frac{WQS_j}{WQS_i}, \quad (75)$$

where i corresponds to row the model in row i and j corresponds to the model in column j . Alongside the relative losses, statistical significance is indicated by light, normal and dark colour codes, corresponding to the 10%, 5%, and 1% significance levels. These colours reflect p-values from pairwise equal quantile forecast accuracy tests, the hypotheses for these tests are shown in 76:

$$H_0 : \mathbb{E}[WQS_i - WQS_j] = 0 \quad \text{vs.} \quad H_1 : \mathbb{E}[WQS_i - WQS_j] \neq 0, \quad (76)$$

where a relative loss smaller (larger) than 1, when accompanied by statistically significant test results, indicates that the model listed in column j has superior (inferior) quantile forecasting performance compared to the model listed in row i , and we indicate this by colouring the cell green (red).

The results are shown in Figures 9 and 10. For day-ahead quantile forecasts based on models estimated with $\mathcal{D}_{HAR}$, the LSTM performs worst, being inferior in all pairwise comparisons. GRU, KAN, RNN, TCN, and TKAN are also outperformed in most cases, indicating relatively weak performance. In contrast, LGBM, LevHAR, Log-HAR, SHAR, TabPFN, and XGB exhibit either statistically insignificant differences or significant superiority over competitors. A similar pattern holds for $\mathcal{D}_{HAR-X}$, with one key difference: TabPFNX is now outperformed by LGBMX and Log-HARX, which achieve stronger predictive accuracy. Overall, no single model consistently dominates at the day-ahead horizon.

In contrast, for the average week-ahead and month-ahead horizons, TabPFN uniformly dominates all other models across both predictor sets. This indicates that TabPFN exhibits superior quantile forecasting performance relative to all other models.

To confirm the conclusions drawn from the Weighted Quantile Score pairwise tests, the

⁹Quantile forecasts for TabPFN are obtained by extracting them from its posterior predictive distribution rather than by directly optimizing the quantile loss function in 66.

¹⁰In this application, the integral is approximated discretely as $WQS_{t+h|t} = \sum_{\tau \in \{0.90, 0.95, 0.99\}} \omega(\tau) QS_{t+h|t}(\tau)$.

Quantile Score pairwise tests at $\tau = 0.95$ in Appendices 15 and 16 and at $\tau = 0.99$ in Appendices 17 and 18 yield largely consistent results, aside from a limited number of individual pairwise differences in significance. The overall pattern remains unchanged: at the day-ahead horizon, LSTM is dominated in all pairwise comparisons under both information sets, while GRU, KAN, RNN, TCN, and TKAN are outperformed in most comparisons. The conclusions for the remaining models are consistent with those based on the Weighted Quantile Score, while at the average week-ahead and month-ahead horizons, TabPFN continues to exhibit superior quantile forecasting accuracy, despite some pairwise comparisons now being significant at $\alpha = 0.05$ or $\alpha = 0.10$ when considering models estimated on $\mathcal{D}_{HAR-X}$ instead of at $\alpha = 0.01$ everywhere.

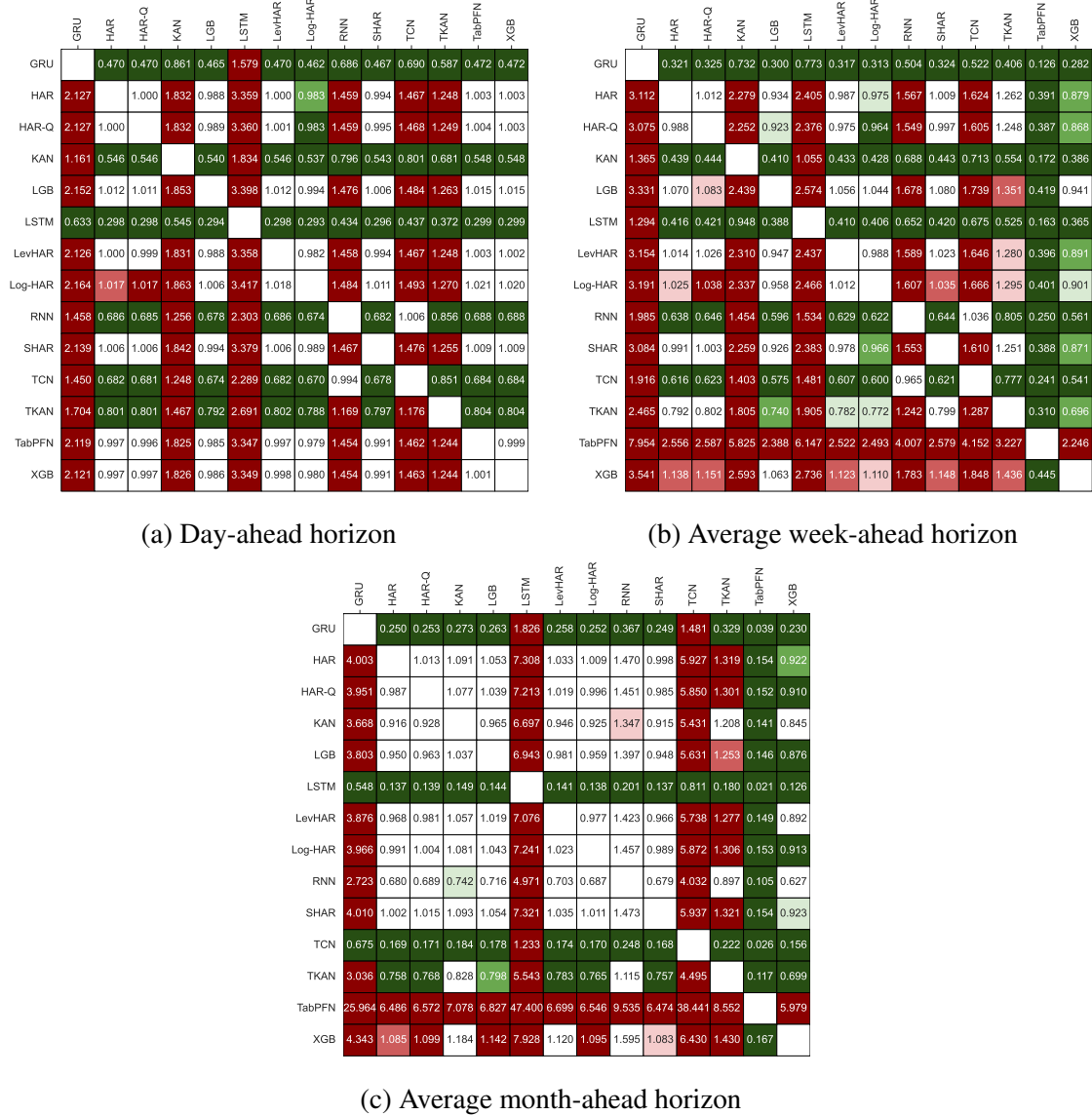


Figure 9: Pairwise equal predictive accuracy test results with Weighted Quantile Scores, estimated on $\mathcal{D}_{HAR}$; cell colours indicate statistical significance of pairwise forecast comparisons (light, medium, and dark shading corresponding to the 10%, 5%, and 1% levels, respectively, with green (red) denoting superior (inferior) performance of the column model relative to the row model).

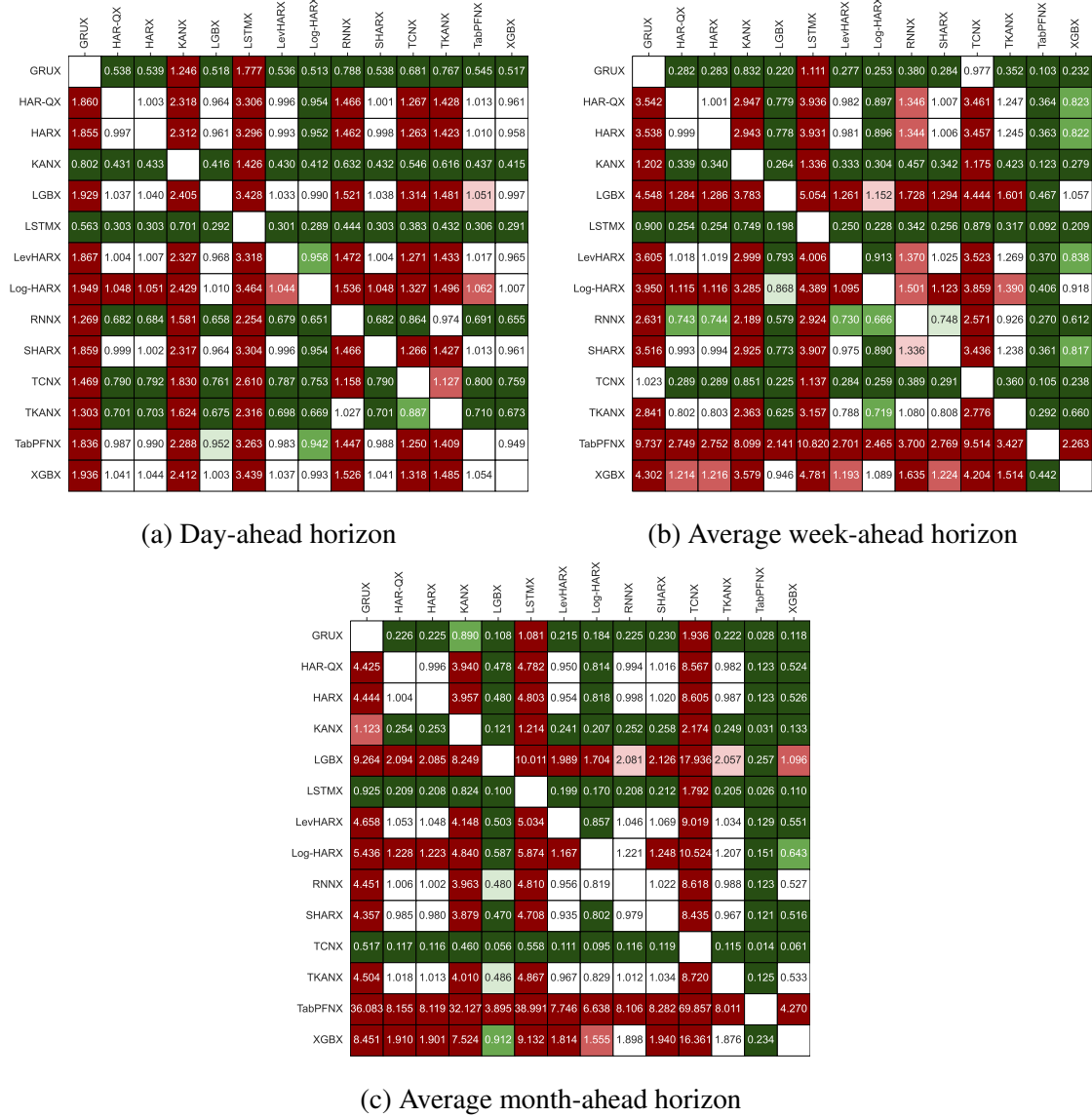


Figure 10: Pairwise equal predictive accuracy test results with Weighted Quantile Scores, estimated on $\mathcal{D}_{HAR-X}$; cell colours indicate statistical significance of pairwise forecast comparisons (light, medium, and dark shading corresponding to the 10%, 5%, and 1% levels, respectively, with green (red) denoting superior (inferior) performance of the column model relative to the row model).

5.5 *Interpreting TabPFN’s performance across horizons*

Two components of the synthetic prior used in TabPFN-3 are especially relevant for realized variance forecasting. First, the prior includes a temporal component: SCMs are extended into discrete-time Dynamic Structural Causal Models, allowing the model to learn from synthetic tasks where rows arrive in time order and where future values depend on past ones (Grinsztajn et al. 2026). This mirrors the structure of realized variance, which is itself a time-indexed sequence with persistent dynamics. Because TabPFN is pretrained on synthetic temporal tasks, it enters our forecasting problem with an inductive bias toward reasoning over ordered observations rather than treating them as i.i.d. samples. This aligns with our rolling-window setup, where each forecast uses only past realized variance to predict a future value.

Second, the prior includes explicit out-of-distribution prediction tasks. These tasks expose the model to synthetic situations where the test data differ from the training data, and where correct predictions require extrapolation rather than interpolation. Grinsztajn et al. (2026) highlight that this out-of-distribution prior enables TabPFN to extrapolate beyond the range of recently observed values. Realized variance frequently exhibits sudden spikes or regime shifts that push the series outside the range covered by the current estimation window. Models trained only on historical data must estimate parameters from these limited windows, whereas TabPFN can rely on inference strategies learned during pretraining to remain stable under such shifts.

These two prior components matter most at the average week- and month-ahead horizons. At these horizons, the forecast target is a moving average of several future days, so consecutive targets overlap heavily and contain relatively little independent information. HAR and other ML models must infer parameters from these weak signals. TabPFN, by contrast, can fall back on its pretrained temporal reasoning and out-of-distribution-aware inference, which are designed to generalize robustly from small, overlapping sets of observations and to handle situations where the target lies outside the range of recent values.

At the one-day-ahead horizon, realized variance is close to linearly persistent, and simple autoregressive or HAR-type models already capture most of the predictive structure. In this regime, the advantages of TabPFN’s temporal and out-of-distribution priors are less pronounced, which aligns with the absence of a single clearly dominant model at the shortest horizon.

6 CONCLUSION

For the one-day-ahead horizon, the results show no clear dominance of one single model across all evaluation criteria. Instead, forecasting performance depends on the chosen loss function and predictor set. When models are estimated on $\mathcal{D}_{HAR}$, the Model Confidence Set results indicate that multiple HAR-type specifications cannot be statistically rejected as inferior under both MSE and QLIKE. However, pairwise comparisons reveal that specific models do stand out conditionally: Log-HAR and TabPFN perform best under MSE in the sense that they are superior in most pairwise comparisons, while XGB and Log-HAR are most effective under QLIKE, indicating stronger robustness against volatility underestimation. For models estimated on the extended information set $\mathcal{D}_{HAR-X}$, performance becomes more selective. Log-HARX emerges as the only HAR-type model consistently retained in the MCS, and it is statistically superior to all competing models under MSE in pairwise comparisons. This highlights that a parsimonious specification with a log transformation can outperform more complex models when additional predictors are included. At the same time, many machine learning models are excluded, suggesting difficulties in effectively exploiting the larger information set. Across both predictor sets, an important limitation persists: none of the models accurately captures the magnitude of extreme volatility spikes. From an uncertainty perspective, no single model dominates the trade-off between calibration and sharpness. TabPFN achieves the highest empirical coverage, while KAN produces the narrowest intervals. For quantile forecasting, no model consistently dominates across all pairwise comparisons; instead, performance should be interpreted in terms of the frequency with which models dominate or are dominated. Under this interpretation, linear and tree-based methods, together with TabPFN, perform relatively well, whereas network-based architectures are dominated in a substantial share of comparisons.

In contrast to the short horizon, the results for the average week-ahead and month-ahead forecasts are remarkably clear and consistent across all evaluation dimensions. The Model Confidence Set results show that TabPFN is the only model retained across both loss functions and both predictor sets. This implies that all competing models are statistically rejected as inferior at the 5% level. These findings are strongly reinforced by the pairwise Diebold–Mariano tests, which demonstrate that TabPFN consistently outperforms all alternative models in terms of predictive accuracy. This dominance extends beyond point forecasts. In terms of uncertainty quantification, TabPFN clearly outperforms all other models by simultaneously achieving the highest empirical coverage and the lowest median relative interval width. Unlike the day-ahead case, there is no trade-off between

reliability and sharpness: TabPFN provides both highly accurate and highly informative prediction intervals. Similarly, for upper-tail forecasting, TabPFN uniformly dominates all competing models across both predictor sets and horizons. The Weighted Quantile Score results indicate that it delivers superior quantile forecasts for the upper tail. Importantly, these findings hold regardless of whether additional macroeconomic predictors are included. This suggests that TabPFN is able to effectively extract relevant information from both smaller and larger predictor sets, in contrast to some models whose performance deteriorates when the dimensionality increases. Taken together, the results provide consistent evidence that, for this series, TabPFN produced the most accurate and best-calibrated week- and month-ahead forecasts under every criterion considered. Whether this ranking extends to other assets, sample periods, or market conditions cannot be established from a single-asset design and remains a topic for further research. The consistency of the result across both loss functions, both predictor sets, and both medium-term horizons nevertheless makes a broader cross-sectional evaluation a worthwhile next step for further research. Although we do not explicitly evaluate computational cost or implementation complexity, it is worth noting that TabPFN is relatively convenient to use compared to many alternative machine learning models. In particular, it does not require explicit hyperparameter tuning by the user, which substantially simplifies the modeling pipeline relative to methods such as gradient boosting or neural networks. In addition, TabPFN produces an approximation to the posterior predictive distribution, which allows prediction intervals to be obtained directly by extracting the relevant quantiles. This avoids the need for auxiliary procedures to construct prediction intervals.

7 DISCUSSION

Several limitations of the empirical analysis should be taken into account when interpreting the results. First, the modeling framework does not explicitly account for the presence of jumps in the log-price process. While realized variance inherently captures both the continuous and jump components of price variation, ignoring this decomposition may affect how volatility dynamics are modeled. In particular, jump-driven movements tend to be abrupt and non-persistent, whereas most of the models considered are designed to capture smoother, persistent dynamics. As a result, the inability of the models to accurately capture the magnitude of extreme volatility spikes may, at least in part, be attributed to this omission rather than purely to model inadequacy. Second, the construction of the test set may have influenced the evaluation of forecasting performance, especially with respect to extreme events. The test period begins immediately after the realized variance spike associated with the COVID-19 market crash, meaning that the most extreme volatility realization is not included in the out-of-sample evaluation. At the same time, the empirical results show that none of the models is able to accurately predict large volatility spikes, regardless of the predictor set or modeling approach. While this suggests a general limitation in forecasting extreme realized variance, the exclusion of such an extreme event from the test set may have affected the relative performance rankings, particularly for models that are better suited to capturing nonlinearities or tail behavior. Third, in the quantile forecasting analysis, hyperparameter configurations were not optimized specifically for quantile loss functions, but instead taken from the point forecasting setup to keep computational cost manageable. This may have disadvantaged machine learning models in particular, as their performance is typically sensitive to the choice of objective function. Therefore, the relatively weaker performance of several machine learning models in the right-tail accuracy evaluation may partly reflect suboptimal tuning rather than solely limitations of the models themselves. Fourth, we did not explicitly model regimes of high and low volatility. Financial markets often exhibit regime-dependent behavior, where volatility dynamics differ between calm and turbulent periods, which was not taken into account in this research. Fifth, the way the models arrive at their forecasts differs fundamentally, which complicates a direct comparison. The HAR-type and machine learning models are re-estimated on the rolling training window at each forecast origin, with the machine learning models additionally relying on hyperparameters selected through a grid search over a pre-specified set of candidate values, so that an unfavourable or insufficiently wide grid may have led to configurations that understate their true forecasting ability. TabPFN, by contrast, is never fitted to the training data at all: its parameters are fixed from synthetic pre-training and it produces forecasts purely by conditioning on the observations supplied

in context, meaning that the relative performance we report partly reflects a difference in how the models learn rather than in forecasting ability alone.

REFERENCES

- Andersen, T. G., Bollerslev, T., Christoffersen, P. & Diebold, F. X. (2005), ‘Volatility forecasting’.
- Andersen, T. G., Dobrev, D. & Schaumburg, E. (2012), ‘Jump-robust volatility estimation using nearest neighbor truncation’, *Journal of Econometrics* **169**(1), 75–93.
- Audrino, F. & Chassot, J. (2025), ‘Hard to beat: The overlooked impact of rolling windows in the era of machine learning’, *International Journal of Forecasting* .
- Baillie, R. T. (1996), ‘Long memory processes and fractional integration in econometrics’, *Journal of Econometrics* **73**(1), 5–59.
- Barndorff-Nielsen, O. E., Hansen, P. R., Lunde, A. & Shephard, N. (2009), ‘Realized kernels in practice: Trades and quotes’.
- Barndorff-Nielsen, O. E. & Shephard, N. (2002), ‘Econometric analysis of realized volatility and its use in estimating stochastic volatility models’, *Journal of the Royal Statistical Society Series B: Statistical Methodology* **64**(2), 253–280.
- Bekaert, G. & Wu, G. (2000), ‘Asymmetric volatility and risk in equity markets’, *The Review of Financial Studies* **13**(1), 1–42.
- Bollerslev, T. (1986), ‘Generalized autoregressive conditional heteroskedasticity’, *Journal of Econometrics* **31**(3), 307–327.
- Bollerslev, T., Patton, A. J. & Quaedvlieg, R. (2016), ‘Exploiting the errors: A simple approach for improved volatility forecasting’, *Journal of Econometrics* **192**(1), 1–18.
- Branco, R. R., Rubesam, A. & Zevallos, M. (2024), ‘Forecasting realized volatility: Does anything beat linear models?’, *Journal of Empirical Finance* **78**, 101524.
- Brauneis, A. & Sahiner, M. (2024), ‘Crypto volatility forecasting: Mounting a har, sentiment, and machine learning horserace’, *Asia-Pacific Financial Markets* pp. 1–33.
- Bucci, A. (2020), ‘Realized volatility forecasting with neural networks’, *Journal of Financial Econometrics* **18**(3), 502–531.
- Chen, T. & Guestrin, C. (2016), Xgboost: A scalable tree boosting system, in ‘Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining’, pp. 785–794.

- Christensen, K., Siggaard, M. & Veliyev, B. (2023), ‘A machine learning approach to volatility forecasting’, *Journal of Financial Econometrics* **21**(5), 1680–1727.
- Chun, D., Cho, H. & Ryu, D. (2025), ‘Volatility forecasting and volatility-timing strategies: A machine learning approach’, *Research in International Business and Finance* **75**, 102723.
- Chung, J., Gulcehre, C., Cho, K. & Bengio, Y. (2014), ‘Empirical evaluation of gated recurrent neural networks on sequence modeling’, *arXiv preprint arXiv:1412.3555*.
- Clark, T. E. & West, K. D. (2007), ‘Approximately normal tests for equal predictive accuracy in nested models’, *Journal of Econometrics* **138**(1), 291–311.
- Clements, A., Preve, D. & Tee, C. (2024), ‘Harvesting the har-x volatility model’, *Available at SSRN 4733597*.
- Corsi, F. (2009), ‘A simple approximate long-memory model of realized volatility’, *Journal of Financial Econometrics* **7**(2), 174–196.
- Corsi, F. & Renò, R. (2012), ‘Discrete-time volatility forecasting with persistent leverage effect and the link with continuous-time volatility modeling’, *Journal of Business & Economic Statistics* **30**(3), 368–380.
- De Boor, C. (1972), ‘On calculating with b-splines’, *Journal of Approximation Theory* **6**(1), 50–62.
- Diebold, F. X. & Mariano, R. S. (1995), ‘Comparing predictive accuracy’, *Journal of Business and Economic Statistics* **13**(3), 253–263.
- Genet, R. & Inzirillo, H. (2024), ‘A temporal kolmogorov-arnold transformer for time series forecasting’, *arXiv preprint arXiv:2406.02486*.
- Gneiting, T. & Raftery, A. E. (2007), ‘Strictly proper scoring rules, prediction, and estimation’, *Journal of the American Statistical Association* **102**(477), 359–378.
- Grinsztajn, L., Flöge, K., Key, O., Birkel, F., Jund, P., Roof, B., Manium, M., Hoo, S. B., Bühler, M., Garg, A. et al. (2026), ‘TabPFN-3: Technical report’, *arXiv preprint arXiv:2605.13986*.
- Gunnarsson, E. S., Isern, H. R., Kaloudis, A., Rissstad, M., Vigdel, B. & Westgaard, S. (2024), ‘Prediction of realized volatility and implied volatility indices using ai and machine learning: A review’, *International Review of Financial Analysis* **93**, 103221.

- Hansen, P. R., Lunde, A. & Nason, J. M. (2003), ‘Choosing the best volatility models: the model confidence set approach’, *Oxford Bulletin of Economics and Statistics* **65**, 839–861.
- Hansen, P. R., Lunde, A. & Nason, J. M. (2011), ‘The model confidence set’, *Econometrica* **79**(2), 453–497.
- Hochreiter, S. & Schmidhuber, J. (1997), ‘Long short-term memory’, *Neural Computation* **9**(8), 1735–1780.
- Hoepner, A. G., McMillan, D., Vivian, A. & Wese Simen, C. (2021), ‘Significance, relevance and explainability in the machine learning age: an econometrics and financial data science perspective’, *The European Journal of Finance* **27**(1-2), 1–7.
- Hollmann, N., Müller, S., Eggensperger, K. & Hutter, F. (2022), ‘Tabpfn: A transformer that solves small tabular classification problems in a second’, *arXiv preprint arXiv:2207.01848*.
- Hoo, S. B., Müller, S., Salinas, D. & Hutter, F. (2025), ‘From tables to time: Extending tabpfn-v2 to time series forecasting’, *arXiv preprint arXiv:2501.02945*.
- Huang, Y. & Luo, Y. (2024), ‘Forecasting conditional volatility based on hybrid garch-type models with long memory, regime switching, leverage effect and heavy-tail: Further evidence from equity market’, *The North American Journal of Economics and Finance* **72**, 102148.
- Ke, G., Meng, Q., Finley, T., Wang, T., Chen, W., Ma, W., Ye, Q. & Liu, T.-Y. (2017), ‘Lightgbm: A highly efficient gradient boosting decision tree’, *Advances in Neural Information Processing Systems* **30**.
- Khan, M. Y. (2015), *Advances in applied nonlinear time series modeling*, PhD thesis, lmu.
- Koenker, R. & Bassett Jr, G. (1978), ‘Regression quantiles’, *Econometrica* pp. 33–50.
- Lea, C., Vidal, R., Reiter, A. & Hager, G. D. (2016), Temporal convolutional networks: A unified approach to action segmentation, in ‘European Conference on Computer Vision’, Springer, pp. 47–54.
- Li, H., Huang, X., Luo, F., Zhou, D., Cao, A. & Guo, L. (2025), ‘Revolutionizing agricultural stock volatility forecasting: a comparative study of machine learning and har-rv models’, *Journal of Applied Economics* **28**(1), 2454081.

- Li, X., Li, D., Cheng, Y. & Li, W. (2024), ‘Forecasting the volatility of educational firms based on har model and lstm models considering sentiment and educational policy’, *Heliyon* **10**(19).
- Lipton, Z. C., Berkowitz, J. & Elkan, C. (2015), ‘A critical review of recurrent neural networks for sequence learning’, *arXiv preprint arXiv:1506.00019*.
- Liu, Z., Wang, Y., Vaidya, S., Ruehle, F., Halverson, J., Soljačić, M., Hou, T. Y. & Tegmark, M. (2024), ‘Kan: Kolmogorov-arnold networks’, *arXiv preprint arXiv:2404.19756*.
- Meligkotsidou, L., Panopoulou, E., Vrontos, I. D. & Vrontos, S. D. (2019), ‘Quantile forecast combinations in realised volatility prediction’, *Journal of the Operational Research Society* **70**(10), 1720–1733.
- Patton, A. J. (2011), ‘Volatility forecast comparison using imperfect volatility proxies’, *Journal of Econometrics* **160**(1), 246–256.
- Patton, A. J. & Sheppard, K. (2015), ‘Good volatility, bad volatility: Signed jumps and the persistence of volatility’, *Review of Economics and Statistics* **97**(3), 683–697.
- Poon, S.-H. & Granger, C. (2005), ‘Practical issues in forecasting volatility’, *Financial Analysts Journal* **61**(1), 45–56.
- Poon, S.-H. & Granger, C. W. J. (2003), ‘Forecasting volatility in financial markets: A review’, *Journal of Economic Literature* **41**(2), 478–539.
- Rahimikia, E. & Poon, S.-H. (2026), ‘Machine learning for realised volatility forecasting’, *Journal of Empirical Finance* p. 101739.
- Shafer, G. & Vovk, V. (2008), ‘A tutorial on conformal prediction.’, *Journal of Machine Learning Research* **9**(3).
- So, M. K., Chu, A. M., Lo, C. C. & Ip, C. Y. (2022), ‘Volatility and dynamic dependence modeling: Review, applications, and financial risk management’, *Wiley Interdisciplinary Reviews: Computational Statistics* **14**(5), e1567.
- Somvanshi, S., Javed, S. A., Islam, M. M., Pandit, D. & Das, S. (2025), ‘A survey on kolmogorov-arnold network’, *ACM Computing Surveys* **58**(2), 1–35.
- Taneva-Angelova, G. & Granchev, D. (2025), ‘Deep learning and transformer architectures for volatility forecasting: Evidence from us equity indices’, *Journal of Risk and Financial Management* **18**(12), 685.

- Teller, A., Pigorsch, U. & Pigorsch, C. (2022), ‘Short-to long-term realized volatility forecasting using extreme gradient boosting’, *Available at SSRN 4267541* .
- Teräsvirta, T. (2009), An introduction to univariate garch models, *in* ‘Handbook of Financial Time Series’, Springer, pp. 17–42.
- Teräsvirta, T. & Zhao, Z. (2014), Stylized facts of return series, robust estimates and three popular models of volatility, *in* ‘Perspectives on Econometrics and Applied Economics’, Routledge, pp. 67–94.
- Wen, F., Gong, X. & Cai, S. (2016), ‘Forecasting the volatility of crude oil futures using har-type models with structural breaks’, *Energy Economics* **59**, 400–413.
- Xu, C. & Xie, Y. (2021), Conformal prediction interval for dynamic time-series, *in* ‘International Conference on Machine Learning’, PMLR, pp. 11559–11569.
- Xu, K., Chen, L. & Wang, S. (2024), ‘Kolmogorov-arnold networks for time series: Bridging predictive power and interpretability’, *arXiv preprint arXiv:2406.02496* .
- Yamak, P. T., Li, Y., Zhang, T. & Pathan, M. S. (2025), ‘Kolmogorov-arnold networks for time series forecasting: a comprehensive review’, *Cluster Computing* **28**(14), 929.
- Zammarchi, G. & Kawasaki, Y. (2026), ‘Assessing the quality of forecasting quantitatively using model confidence set: from naive models to neural networks’, *Quality & Quantity* pp. 1–25.
- Zhang, Y., Cui, L. & Yan, W. (2025), ‘Integrating kolmogorov–arnold networks with time series prediction framework in electricity demand forecasting’, *Energies* **18**(6), 1365.

A APPENDIX

A.1 Stationarity test results

The Augmented Dickey-Fuller test results are reported in table 4. The regression to derive the test statistic is

$$\Delta y_t = \alpha + \beta t + \gamma y_{t-1} + \sum_{i=1}^p \delta_i \Delta y_{t-i} + \varepsilon_t$$

where γ is the parameter of interest. Then, the ADF test statistic is

$$ADF = \frac{\hat{\gamma}}{SE(\hat{\gamma})}.$$

The hypotheses that are being considered in this test are: $H_0 : \gamma = 0$ vs. $H_1 : \gamma < 0$. Here, if the null hypothesis indicates the presence of a unit root (hence nonstationarity), whereas the alternative hypothesis indicates absence of a unit root. As a consequence, the 3-Month Treasury Bill rate variable appears nonstationary if we consider a 5% significance level. It can be observed that the first difference, $\Delta TMTB$, is stationary.

Table 4: Augmented Dickey-Fuller test results

variable	test statistic	p-value
RV	-7.913	0.000
RV_{t-1}	-7.915	0.000
$RV_{t-1 t-5}$	-7.139	0.000
$RV_{t-1 t-22}$	-7.227	0.000
VIX	-5.249	0.000
EPU	-5.412	0.000
COP	-2.997	0.035
TB3M	-0.901	0.788
BC	-8.419	0.000
$\Delta TMTB$	-9.967	0.000

A.2 Multicollinearity

Table 5: Correlation Matrix of regressors in $\mathcal{D}_{HAR-X}$

	RV_{t-1}	$RV_{t t-4}$	$RV_{t t-21}$	VIX_t	EPU_t	COP_t	$TMTB_t$	BC_t
RV_{t-1}	1.000	0.787	0.582	0.572	0.312	-0.132	-0.066	-0.367
$RV_{t t-4}$	0.787	1.000	0.786	0.711	0.401	-0.174	-0.085	-0.499
$RV_{t t-21}$	0.582	0.786	1.000	0.749	0.485	-0.234	-0.107	-0.669
VIX_t	0.572	0.711	0.749	1.000	0.497	-0.084	-0.210	-0.361
EPU_t	0.312	0.401	0.485	0.497	1.000	-0.126	-0.207	-0.281
COP_t	-0.132	-0.174	-0.234	-0.084	-0.126	1.000	-0.006	0.126
$TMTB_t$	-0.066	-0.085	-0.107	-0.210	-0.207	-0.006	1.000	0.019
BC_t	-0.367	-0.499	-0.669	-0.361	-0.281	0.126	0.019	1.000

A.3 Distribution of residuals of log-realized variance

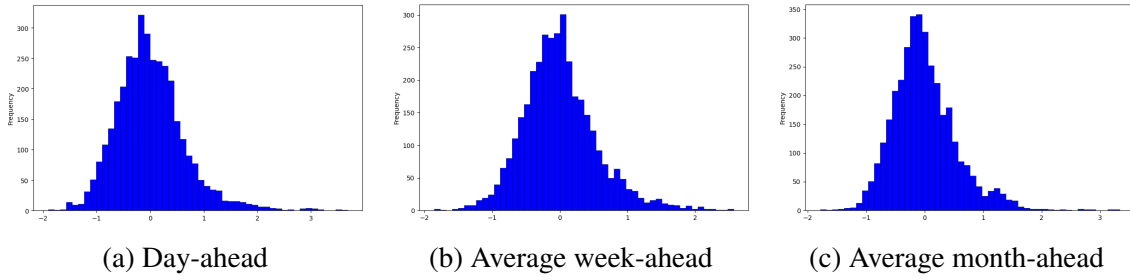


Figure 11: Histograms of the training set residuals of the log-HAR model at all forecast horizons.

A.4 Hyperparameter tuning

Table 6: Hyperparameter tuning for network-based models

Shared hyperparameters	KAN/TKAN	RNN	LSTM	GRU	TCN
$B \in \{32, 64, 128, 256, 350\}$	$W \in \{2, 3, 4, 5, 6, 7, 8, 9, 10\}$ $SG \in \{3, 5, 7, 10\}$	$H \in \{5, 8, 13\}$ $L \in \{1, 2\}$	$H \in \{2, 3, 5\}$ $L \in \{1, 2\}$	$H \in \{3, 4, 6\}$ $L \in \{1, 2\}$	$H \in \{3, 6\}$ $L = 1$
$S \in \{60, 120, 200, 300\}$					
$\lambda \in \{10^{-5}, 10^{-4}, 10^{-3}, 10^{-2}, 10^{-1}, 1, 10\}$ (except TKAN)					
LR = 10^{-3} (default Adam)					
D = 0.2					

Notes. Abbreviations: W = width; SG = spline grid; H = hidden units; L = number of layers; B = batch size; S = training steps; LR = learning rate; D = dropout rate. During structural hyperparameter selection (prior to regularization tuning), the regularization parameter λ was fixed at 10^{-4} . Also, batch size $B = 128$ and training steps $S = 50$ were used prior to optimization and regularization tuning.

Table 7: Hyperparameter tuning for tree-based models

LightGBM	XGBoost
Leaves $\in \{15, 30, 60, 100\}$	Max depth $\in \{3, 5, 7\}$
Depth $\in \{3, 5, 7\}$	Depth $\in \{3, 5, 7\}$
$\lambda \in \{10^{-5}, 10^{-4}, 10^{-3}, 10^{-2}, 10^{-1}, 1, 10\}$	$\lambda \in \{10^{-5}, 10^{-4}, 10^{-3}, 10^{-2}, 10^{-1}, 1, 10\}$
Learning rate = default	Learning rate = default
Min. leaf samples = 50	Min. child weight = 5
Number of trees = 1000	Number of trees = 1000

Notes. LR = learning rate. During structural hyperparameter selection (prior to regularization tuning), the regularization parameter λ was fixed at 10^{-4} .

A.5 Model Confidence Set under alternative significance level

Table 8: Models included in the Model Confidence Set across estimation sets and forecast horizons, $\alpha = 0.10$, $B = 10,000$, loss function: $\mathcal{L}^{QLIKE}$

Estimated on $\mathcal{D}_{HAR}$			Estimated on $\mathcal{D}_{HAR-X}$		
RV_{t+1}	$RV_{t+5 t+1}$	$RV_{t+22 t+1}$	RV_{t+1}	$RV_{t+5 t+1}$	$RV_{t+22 t+1}$
HAR	TabPFN	TabPFN	Log-HARX	TabPFN	TabPFN
Log-HAR			TKANX		
HAR-Q			LGBMX		
SHAR			XGBX		
TKAN			RNNX		
LGBM					
XGB					
RNN					

Table 9: Models included in the Model Confidence Set across estimation sets and forecast horizons, $\alpha = 0.10$, $B = 10,000$, loss function: $\mathcal{L}^{MSE}$

Estimated on $\mathcal{D}_{HAR}$			Estimated on $\mathcal{D}_{HAR-X}$		
RV_{t+1}	$RV_{t+5 t+1}$	$RV_{t+22 t+1}$	RV_{t+1}	$RV_{t+5 t+1}$	$RV_{t+22 t+1}$
HAR	TabPFN	TabPFN	Log-HARX	TabPFN	TabPFN
Log-HAR					
HAR-Q					
SHAR					
LevHAR					
KAN					
TKAN					
LGBM					
XGB					
RNN					
GRU					
TabPFN					

A.6 Full pairwise Diebold-Mariano test results

A.6.1 Day-ahead forecasts

	HAR	Log-HAR	HAR-Q	SHAR	LevHAR	KAN	TKAN	LGBM	XGB	RNN	LSTM	GRU	TCN	TabPFN
HAR		0.917	1.003	1.006	0.991	0.959	0.961	0.957	0.961	0.923	1.110	0.976	1.003	0.910
Log-HAR	1.090		1.093	1.096	1.081	1.045	1.048	1.043	1.047	1.006	1.210	1.063	1.093	0.991
HAR-Q	0.997	0.915		1.003	0.989	0.956	0.959	0.954	0.950	0.921	1.108	0.973	1.000	0.907
SHAR	0.994	0.912	0.997		0.986	0.953	0.956	0.951	0.955	0.918	1.104	0.970	0.997	0.904
LevHAR	1.009	0.925	1.011	1.014		0.967	0.969	0.965	0.969	0.931	1.120	0.984	1.011	0.917
KAN	1.043	0.957	1.046	1.049	1.034		1.003	0.998	1.002	0.963	1.158	1.018	1.046	0.949
TKAN	1.041	0.955	1.043	1.047	1.032	0.997		0.996	1.000	0.960	1.155	1.015	1.043	0.946
LGBM	1.045	0.959	1.048	1.051	1.036	1.002	1.004		1.004	0.964	1.160	1.020	1.048	0.950
XGB	1.041	0.955	1.043	1.047	1.032	0.998	1.000	0.996		0.961	1.156	1.015	1.044	0.947
RNN	1.084	0.994	1.086	1.090	1.074	1.039	1.041	1.037	1.041		1.203	1.057	1.086	0.986
LSTM	0.901	0.826	0.903	0.906	0.893	0.863	0.865	0.862	0.865	0.831		0.879	0.903	0.819
GRU	1.025	0.940	1.028	1.031	1.016	0.982	0.985	0.981	0.985	0.946	1.138		1.028	0.932
TCN	0.997	0.915	1.000	1.003	0.989	0.956	0.958	0.954	0.950	0.920	1.107	0.973		0.907
TabPFN	1.099	1.009	1.102	1.106	1.090	1.054	1.057	1.052	1.056	1.015	1.221	1.073	1.102	

(a) $\mathcal{D}_{HAR}, \mathcal{L}^{MSE}$

	HAR	Log-HAR	HAR-Q	SHAR	LevHAR	KAN	TKAN	LGBM	XGB	RNN	LSTM	GRU	TCN	TabPFN
HAR		0.991	1.008	0.999	1.072	1.067	0.988	0.992	0.985	1.024	4.192	1.123	1.453	1.037
Log-HAR	1.009		1.017	1.008	1.081	1.076	0.996	1.001	0.993	1.033	4.229	1.133	1.465	1.046
HAR-Q	0.992	0.983		0.991	1.063	1.058	0.970	0.984	0.977	1.015	4.158	1.114	1.441	1.029
SHAR	1.001	0.993	1.010		1.073	1.068	0.989	0.994	0.986	1.025	4.197	1.124	1.454	1.038
LevHAR	0.933	0.925	0.941	0.932		0.995	0.921	0.926	0.918	0.955	3.911	1.048	1.355	0.968
KAN	0.938	0.929	0.945	0.937	1.005		0.926	0.930	0.923	0.960	3.931	1.053	1.362	0.973
TKAN	1.012	1.004	1.021	1.011	1.085	1.080		1.005	0.997	1.037	4.245	1.137	1.471	1.050
LGBM	1.008	0.999	1.016	1.007	1.080	1.075	0.995		0.992	1.032	4.225	1.132	1.464	1.045
XGB	1.015	1.007	1.024	1.014	1.088	1.083	1.003	1.008		1.039	4.257	1.140	1.475	1.053
RNN	0.977	0.968	0.985	0.976	1.047	1.042	0.965	0.969	0.962		4.095	1.097	1.419	1.013
LSTM	0.239	0.236	0.241	0.238	0.256	0.254	0.236	0.237	0.235	0.244		0.268	0.346	0.247
GRU	0.891	0.883	0.898	0.890	0.955	0.950	0.880	0.884	0.877	0.912	3.734		1.294	0.924
TCN	0.688	0.682	0.694	0.688	0.738	0.734	0.680	0.683	0.678	0.705	2.886	0.773		0.714
TabPFN	0.964	0.956	0.972	0.963	1.033	1.028	0.952	0.957	0.950	0.987	4.042	1.083	1.400	

(b) $\mathcal{D}_{HAR}, \mathcal{L}^{QLIKE}$

	HARX	Log-HARX	HAR-QX	SHARX	LevHARX	KANX	TKANX	LGBMX	XGBX	RNNX	LSTMX	GRUX	TCNX	TabPFN
HARX		0.880	1.011	1.006	1.007	1.001	0.981	0.903	0.965	0.918	1.091	0.952	0.953	0.914
Log-HARX	1.136		1.148	1.143	1.144	1.137	1.114	1.026	1.096	1.043	1.237	1.081	1.082	1.038
HAR-QX	0.985	0.871		0.995	0.996	0.990	0.970	0.893	0.954	0.908	1.077	0.942	0.942	0.904
SHARX	0.994	0.875	1.005		1.001	0.995	0.975	0.898	0.959	0.913	1.082	0.946	0.947	0.908
LevHARX	0.993	0.874	1.004	0.999		0.994	0.974	0.897	0.958	0.912	1.081	0.945	0.946	0.907
KANX	0.990	0.879	1.010	1.005	1.006		0.980	0.902	0.964	0.917	1.088	0.951	0.951	0.913
TKANX	1.019	0.897	1.031	1.026	1.027	1.021		0.921	0.984	0.936	1.110	0.970	0.971	0.931
LGBMX	1.107	0.975	1.119	1.114	1.115	1.108	1.086		1.068	1.017	1.206	1.054	1.054	1.012
XGBX	1.036	0.912	1.048	1.043	1.044	1.038	1.017	0.936		0.952	1.129	0.987	0.987	0.947
RNNX	1.089	0.959	1.101	1.096	1.097	1.090	1.068	0.984	1.051		1.186	1.037	1.037	0.995
LSTMX	0.918	0.808	0.928	0.924	0.925	0.919	0.901	0.829	0.886	0.843		0.874	0.875	0.839
GRUX	1.050	0.925	1.062	1.057	1.058	1.052	1.030	0.949	1.014	0.965	1.144		1.001	0.960
TCNX	1.050	0.924	1.061	1.056	1.057	1.051	1.030	0.948	1.013	0.964	1.143	0.999		0.959
TabPFN	1.094	0.964	1.106	1.101	1.102	1.096	1.074	0.989	1.056	1.005	1.192	1.042	1.042	

(c) $\mathcal{D}_{HAR-X}, \mathcal{L}^{MSE}$

	HARX	Log-HARX	HAR-QX	SHARX	LevHARX	KANX	TKANX	LGBMX	XGBX	RNNX	LSTMX	GRUX	TCNX	TabPFN
HARX		0.747	4.693	1.093	1.060	0.993	0.768	0.755	0.786	0.769	2.544	0.803	1.116	0.787
Log-HARX	1.339		6.283	1.463	1.419	1.329	1.029	1.010	1.053	1.030	3.406	1.074	1.495	1.054
HAR-QX	0.213	0.159		0.233	0.226	0.212	0.164	0.161	0.168	0.164	0.542	0.171	0.238	0.168
SHARX	0.915	0.683	4.294		0.970	0.908	0.703	0.691	0.720	0.704	2.328	0.734	1.021	0.720
LevHARX	0.944	0.705	4.428	1.031		0.937	0.725	0.712	0.742	0.726	2.400	0.757	1.053	0.743
KANX	1.007	0.752	4.727	1.101	1.067		0.774	0.760	0.792	0.775	2.562	0.808	1.124	0.793
TKANX	1.301	0.972	6.107	1.422	1.379	1.282		0.982	1.023	1.001	3.310	1.044	1.453	1.025
LGBMX	1.325	0.990	6.219	1.448	1.404	1.316	1.018		1.042	1.019	3.371	1.063	1.479	1.043
XGBX	1.272	0.950	5.968	1.390	1.348	1.263	0.977	0.960		0.978	3.235	1.021	1.420	1.001
RNNX	1.300	0.971	6.103	1.421	1.378	1.291	0.999	0.981	1.023		3.308	1.044	1.452	1.024
LSTMX	0.393	0.294	1.845	0.430	0.417	0.390	0.302	0.297	0.309	0.302		0.315	0.439	0.310
GRUX	1.248	0.931	5.848	1.362	1.321	1.237	0.958	0.940	0.980	0.958	3.170		1.391	0.981
TCNX	0.898	0.669	4.204	0.979	0.949	0.889	0.688	0.676	0.704	0.689	2.279	0.719		0.705
TabPFN	1.270	0.949	5.960	1.388	1.346	1.261	0.976	0.953	0.998	0.977	3.231	1.019	1.418	

(d) $\mathcal{D}_{HAR-X}, \mathcal{L}^{QLIKE}$

Figure 12: Day-ahead relative loss matrices; cell colours indicate statistical significance of pairwise forecast comparisons (light, medium, dark = 10%, 5%, 1%; green/red = column better/worse than row).

A.6.2 Week-ahead forecasts

	HAR	Log-HAR	HAR-Q	SHAR	LevHAR	KAN	TKAN	LGBM	XGB	RNN	LSTM	GRU	TCN	TabPFN
HAR		0.736	0.995	1.038	0.988	0.836	1.275	0.825	1.087	0.738	0.807	0.939	0.920	0.310
Log-HAR	1.359		1.352	1.410	1.343	1.135	1.732	1.121	1.476	0.999	1.097	1.275	1.250	0.421
HAR-Q	1.005	0.739		1.043	0.993	0.839	1.280	0.829	1.092	0.738	0.811	0.943	0.924	0.311
SHAR	0.963	0.709	0.959		0.952	0.805	1.228	0.795	1.047	0.708	0.778	0.904	0.887	0.299
LevHAR	1.012	0.745	1.007	1.050		0.845	1.290	0.835	1.100	0.744	0.817	0.950	0.931	0.314
KAN	1.197	0.881	1.191	1.242	1.183		1.525	0.988	1.301	0.880	0.966	1.123	1.101	0.371
TKAN	0.785	0.577	0.781	0.814	0.775	0.656		0.645	0.853	0.577	0.633	0.736	0.722	0.243
LGBM	1.212	0.892	1.206	1.257	1.197	1.012	1.544		1.317	0.891	0.978	1.137	1.115	0.375
XGB	0.920	0.677	0.916	0.955	0.909	0.769	1.173	0.760		0.676	0.743	0.864	0.847	0.285
RNN	1.360	1.001	1.354	1.412	1.345	1.137	1.734	1.123	1.479		1.098	1.277	1.252	0.422
LSTM	1.239	0.912	1.233	1.286	1.224	1.035	1.579	1.023	1.346	0.911		1.163	1.140	0.384
GRU	1.065	0.784	1.061	1.106	1.053	0.890	1.358	0.879	1.158	0.783	0.860		0.980	0.330
TCN	1.087	0.800	1.082	1.128	1.074	0.908	1.385	0.897	1.181	0.799	0.877	1.020		0.337
TabPFN	3.227	2.376	3.213	3.350	3.190	2.697	4.114	2.664	3.507	2.372	2.605	3.029	2.970	

(a) $\mathcal{D}_{HAR}, \mathcal{L}^{MSE}$

	HAR	Log-HAR	HAR-Q	SHAR	LevHAR	KAN	TKAN	LGBM	XGB	RNN	LSTM	GRU	TCN	TabPFN
HAR		0.995	1.001	1.001	1.051	1.063	1.033	0.989	0.982	1.013	1.039	1.279	1.442	0.862
Log-HAR	1.005		1.006	1.006	1.057	1.069	1.038	0.994	0.987	1.019	1.045	1.286	1.449	0.867
HAR-Q	0.999	0.994		1.000	1.050	1.062	1.032	0.988	0.981	1.012	1.038	1.278	1.440	0.861
SHAR	0.999	0.994	1.000		1.050	1.062	1.032	0.988	0.981	1.012	1.038	1.278	1.440	0.861
LevHAR	0.951	0.946	0.952	0.952		1.011	0.982	0.941	0.934	0.964	0.989	1.217	1.371	0.820
KAN	0.941	0.936	0.942	0.942	0.989		0.972	0.930	0.924	0.953	0.978	1.203	1.356	0.811
TKAN	0.968	0.963	0.969	0.969	1.018	1.029		0.958	0.951	0.981	1.006	1.239	1.396	0.835
LGBM	1.011	1.006	1.012	1.012	1.063	1.075	1.044		0.993	1.025	1.051	1.294	1.458	0.872
XGB	1.019	1.013	1.019	1.020	1.071	1.083	1.052	1.007		1.032	1.059	1.303	1.468	0.878
RNN	0.987	0.982	0.988	0.988	1.039	1.049	1.019	0.976	0.968		1.026	1.263	1.423	0.851
LSTM	0.962	0.957	0.963	0.963	1.011	1.023	0.994	0.951	0.945	0.975		1.231	1.387	0.830
GRU	0.782	0.778	0.782	0.782	0.822	0.831	0.807	0.773	0.767	0.792	0.812		1.127	0.674
TCN	0.694	0.690	0.694	0.694	0.729	0.737	0.716	0.686	0.681	0.703	0.721	0.887		0.598
TabPFN	1.160	1.154	1.161	1.161	1.219	1.233	1.198	1.147	1.139	1.175	1.205	1.484	1.672	

(b) $\mathcal{D}_{HAR}, \mathcal{L}^{QLIKE}$

	HARX	Log-HARX	HAR-QX	SHARX	LevHARX	KANX	TKANX	LGBMX	XGBX	RNNX	LSTMX	GRUX	TCNX	TabPFN
HARX		0.650	0.990	1.024	1.011	0.986	0.782	0.741	0.636	0.742	1.170	1.093	1.128	0.307
Log-HARX	1.538		1.523	1.575	1.555	1.517	1.202	1.140	0.979	1.141	1.800	1.682	1.735	0.473
HAR-QX	1.010	0.656		1.034	1.021	0.996	0.789	0.748	0.643	0.740	1.181	1.104	1.139	0.310
SHARX	0.976	0.635	0.967		0.987	0.963	0.763	0.724	0.621	0.724	1.142	1.068	1.101	0.300
LevHARX	0.989	0.643	0.979	1.013		0.975	0.773	0.733	0.630	0.734	1.157	1.081	1.115	0.304
KANX	1.014	0.659	1.004	1.039	1.025		0.793	0.752	0.646	0.752	1.187	1.109	1.144	0.312
TKANX	1.279	0.832	1.267	1.310	1.293	1.261		0.949	0.814	0.949	1.497	1.398	1.443	0.393
LGBMX	1.349	0.877	1.338	1.382	1.364	1.330	1.055		0.859	1.001	1.579	1.475	1.522	0.415
XGBX	1.571	1.021	1.556	1.609	1.588	1.549	1.228	1.164		1.165	1.838	1.718	1.772	0.483
RNNX	1.348	0.876	1.335	1.381	1.363	1.329	1.054	0.999	0.858		1.577	1.474	1.520	0.414
LSTMX	0.855	0.556	0.846	0.875	0.864	0.843	0.668	0.634	0.544	0.634		0.935	0.964	0.263
GRUX	0.915	0.595	0.906	0.937	0.925	0.902	0.715	0.678	0.582	0.678	1.070		1.031	0.281
TCNX	0.887	0.576	0.878	0.908	0.897	0.874	0.693	0.657	0.564	0.658	1.037	0.970		0.273
TabPFN	3.254	2.115	3.222	3.332	3.289	3.208	2.543	2.411	2.071	2.413	3.807	3.558	3.669	

(c) $\mathcal{D}_{HAR-X}, \mathcal{L}^{MSE}$

	HARX	Log-HARX	HAR-QX	SHARX	LevHARX	KANX	TKANX	LGBMX	XGBX	RNNX	LSTMX	GRUX	TCNX	TabPFN
HARX		0.620	0.765	0.756	0.768	0.862	0.632	0.659	0.600	0.635	1.513	1.306	1.925	0.545
Log-HARX	1.613		1.234	1.220	1.239	1.391	1.019	1.064	0.968	1.024	2.441	2.107	3.105	0.879
HAR-QX	1.308	0.811		0.989	1.004	1.128	0.826	0.862	0.784	0.830	1.978	1.708	2.516	0.712
SHARX	1.322	0.820	1.011		1.016	1.140	0.836	0.872	0.793	0.839	2.000	1.727	2.545	0.720
LevHARX	1.302	0.807	0.996	0.985		1.123	0.823	0.859	0.781	0.826	1.970	1.700	2.506	0.709
KANX	1.160	0.719	0.887	0.877	0.891		0.733	0.765	0.695	0.736	1.754	1.514	2.232	0.632
TKANX	1.582	0.981	1.210	1.197	1.215	1.365		1.043	0.945	1.004	2.394	2.067	3.045	0.862
LGBMX	1.517	0.940	1.160	1.147	1.165	1.308	0.958		0.910	0.963	2.294	1.981	2.919	0.826
XGBX	1.667	1.034	1.275	1.261	1.280	1.438	1.054	1.099		1.058	2.522	2.177	3.209	0.908
RNNX	1.575	0.977	1.205	1.191	1.210	1.359	0.996	1.039	0.945		2.383	2.057	3.032	0.858
LSTMX	0.661	0.410	0.506	0.500	0.508	0.570	0.418	0.436	0.396	0.420		0.863	1.272	0.360
GRUX	0.766	0.475	0.586	0.579	0.588	0.660	0.484	0.505	0.459	0.486	1.158		1.474	0.417
TCNX	0.520	0.322	0.397	0.393	0.399	0.448	0.328	0.343	0.312	0.330	0.786	0.679		0.283
TabPFN	1.835	1.138	1.404	1.388	1.410	1.583	1.160	1.210	1.101	1.165	2.777	2.397	3.532	

(d) $\mathcal{D}_{HAR-X}, \mathcal{L}^{QLIKE}$

Figure 13: Week-ahead relative loss matrices; cell colours indicate statistical significance of pairwise forecast comparisons (light, medium, dark = 10%, 5%, 1%; green/red = column better/worse than row).

A.6.3 Month-ahead forecasts

	HAR	Log-HAR	HAR-Q	SHAR	LevHAR	KAN	TKAN	LGBM	XGB	RNN	LSTM	GRU	TCN	TabPFN
HAR		0.790	1.063	1.033	1.075	1.226	3.216	1.700	2.154	0.698	1.825	0.965	1.664	0.047
Log-HAR	1.265		1.345	1.308	1.361	1.552	4.069	2.152	2.726	0.883	2.310	1.221	2.106	0.060
HAR-Q	0.941	0.743		0.972	1.012	1.153	3.025	1.599	2.026	0.656	1.717	0.908	1.565	0.045
SHAR	0.968	0.765	1.029		1.040	1.186	3.111	1.645	2.084	0.675	1.766	0.934	1.610	0.046
LevHAR	0.930	0.735	0.989	0.961		1.140	2.990	1.581	2.003	0.649	1.698	0.897	1.548	0.044
KAN	0.816	0.644	0.867	0.843	0.877		2.622	1.387	1.757	0.569	1.489	0.787	1.357	0.039
TKAN	0.311	0.246	0.331	0.321	0.334	0.381		0.529	0.670	0.217	0.568	0.300	0.518	0.015
LGBM	0.588	0.465	0.625	0.609	0.632	0.721	1.891		1.267	0.410	1.074	0.568	0.979	0.028
XGB	0.464	0.367	0.494	0.480	0.499	0.569	1.493	0.789		0.324	0.848	0.448	0.773	0.022
RNN	1.433	1.133	1.524	1.481	1.541	1.758	4.609	2.437	3.087		2.616	1.383	2.385	0.068
LSTM	0.548	0.433	0.582	0.566	0.589	0.672	1.762	0.931	1.180	0.382		0.529	0.912	0.026
GRU	1.036	0.819	1.102	1.071	1.114	1.271	3.332	1.762	2.232	0.723	1.892		1.725	0.049
TCN	0.601	0.475	0.639	0.621	0.646	0.737	1.932	1.022	1.294	0.419	1.097	0.580		0.028
TabPFN	21.06416	66.222	4.1521	7.9122	6.7325	85.457	79.695	84.945	41.314	71.088	48.820	34.535	0.088	

(a) $\mathcal{D}_{HAR}, \mathcal{L}^{MSE}$

	HAR	Log-HAR	HAR-Q	SHAR	LevHAR	KAN	TKAN	LGBM	XGB	RNN	LSTM	GRU	TCN	TabPFN
HAR		1.007	1.002	1.000	1.019	1.025	1.158	1.017	1.057	1.008	2.330	1.092	2.351	0.903
Log-HAR	0.993		0.995	0.992	1.012	1.018	1.151	1.009	1.049	0.999	2.312	1.084	2.333	0.896
HAR-Q	0.998	1.005		0.998	1.017	1.023	1.157	1.015	1.055	1.004	2.325	1.090	2.346	0.901
SHAR	1.000	1.008	1.002		1.020	1.025	1.160	1.017	1.057	1.006	2.330	1.092	2.351	0.903
LevHAR	0.981	0.988	0.983	0.981		1.006	1.137	0.997	1.037	0.987	2.285	1.071	2.306	0.886
KAN	0.975	0.983	0.977	0.975	0.994		1.131	0.992	1.031	0.981	2.272	1.065	2.293	0.881
TKAN	0.863	0.869	0.864	0.862	0.879	0.884		0.877	0.912	0.868	2.009	0.942	2.027	0.779
LGBM	0.984	0.991	0.986	0.983	1.003	1.009	1.140		1.040	0.990	2.291	1.074	2.312	0.888
XGB	0.946	0.953	0.948	0.946	0.964	0.970	1.097	0.962		0.952	2.203	1.033	2.223	0.854
RNN	0.994	1.001	0.996	0.994	1.013	1.019	1.152	1.011	1.051		2.316	1.086	2.336	0.897
LSTM	0.429	0.432	0.430	0.429	0.438	0.440	0.498	0.436	0.454	0.432		0.469	1.009	0.388
GRU	0.916	0.922	0.917	0.915	0.933	0.939	1.061	0.931	0.968	0.921	2.133		2.152	0.827
TCN	0.425	0.429	0.426	0.425	0.434	0.436	0.493	0.433	0.450	0.428	0.991	0.465		0.384
TabPFN	1.108	1.116	1.110	1.107	1.129	1.136	1.284	1.126	1.171	1.114	2.580	1.210	2.603	

(b) $\mathcal{D}_{HAR}, \mathcal{L}^{QLIKE}$

	HARX	Log-HARX	HAR-QX	SHARX	LevHARX	KANX	TKANX	LGBMX	XGBX	RNNX	LSTMX	GRUX	TCNX	TabPFN
HARX		0.516	1.017	1.007	1.018	1.145	0.885	0.481	0.306	0.607	1.239	1.006	1.932	0.039
Log-HARX	1.938		1.971	1.953	1.973	2.220	1.715	0.933	0.593	1.177	2.401	1.950	3.745	0.075
HAR-QX	0.983	0.507		0.991	1.001	1.126	0.870	0.473	0.301	0.597	1.218	0.989	1.900	0.038
SHARX	0.993	0.512	1.010		1.011	1.137	0.879	0.478	0.304	0.603	1.230	0.999	1.918	0.038
LevHARX	0.982	0.507	0.999	0.990		1.125	0.869	0.473	0.300	0.597	1.217	0.988	1.898	0.038
KANX	0.873	0.450	0.888	0.880	0.889		0.773	0.420	0.267	0.530	1.081	0.878	1.687	0.034
TKANX	1.130	0.583	1.149	1.138	1.150	1.294		0.544	0.346	0.686	1.400	1.137	2.183	0.044
LGBMX	2.078	1.072	2.113	2.093	2.115	2.379	1.838		0.635	1.262	2.573	2.090	4.014	0.080
XGBX	3.270	1.687	3.325	3.294	3.329	3.745	2.894	1.574		1.986	4.050	3.290	6.318	0.127
RNNX	1.646	0.849	1.674	1.658	1.676	1.885	1.457	0.792	0.503		2.039	1.656	3.181	0.064
LSTMX	0.807	0.417	0.821	0.813	0.822	0.925	0.714	0.389	0.247	0.490		0.812	1.560	0.031
GRUX	0.994	0.513	1.011	1.001	1.012	1.138	0.880	0.478	0.304	0.604	1.231		1.920	0.038
TCNX	0.518	0.267	0.526	0.521	0.527	0.593	0.458	0.249	0.158	0.314	0.641	0.521		0.020
TabPFN	25.8423	33.226	27.86	0.326	30.429	59.722	86.812	4.387	9.0215	69.882	0.0725	99.949	9.926	

(c) $\mathcal{D}_{HAR-X}, \mathcal{L}^{MSE}$

	HARX	Log-HARX	HAR-QX	SHARX	LevHARX	KANX	TKANX	LGBMX	XGBX	RNNX	LSTMX	GRUX	TCNX	TabPFN
HARX		0.874	1.032	1.043	1.131	1.312	1.164	0.892	0.842	0.909	1.619	1.157	3.794	0.799
Log-HARX	1.144		1.181	1.651	1.294	1.501	1.331	1.020	0.963	1.040	1.851	1.323	4.340	0.914
HAR-QX	0.969	0.847		1.398	1.096	1.271	1.128	0.864	0.815	0.881	1.568	1.121	3.676	0.774
SHARX	0.693	0.606	0.715		0.784	0.909	0.807	0.618	0.583	0.630	1.122	0.802	2.629	0.553
LevHARX	0.884	0.773	0.912	1.276		1.160	1.029	0.788	0.744	0.804	1.431	1.023	3.354	0.706
KANX	0.762	0.668	0.787	1.100	0.862		0.887	0.680	0.641	0.693	1.234	0.882	2.892	0.609
TKANX	0.859	0.751	0.887	1.240	0.972	1.127		0.766	0.723	0.781	1.391	0.994	3.260	0.686
LGBMX	1.121	0.980	1.157	1.618	1.268	1.471	1.305		0.944	1.015	1.815	1.297	4.254	0.896
XGBX	1.188	1.039	1.227	1.715	1.344	1.559	1.383	1.060		1.080	1.923	1.375	4.509	0.949
RNNX	1.100	0.962	1.135	1.587	1.244	1.443	1.280	0.981	0.926		1.780	1.272	4.173	0.878
LSTMX	0.618	0.540	0.638	0.892	0.689	0.810	0.719	0.551	0.520	0.562		0.715	2.344	0.493
GRUX	0.864	0.756	0.892	1.247	0.978	1.134	1.006	0.771	0.727	0.786	1.399		3.280	0.690
TCNX	0.264	0.230	0.272	0.380	0.298	0.346	0.307	0.235	0.222	0.240	0.427	0.305		0.211
TabPFN	1.252	1.095	1.292	1.807	1.416	1.642	1.457	1.117	1.054	1.138	2.026	1.448	4.750	

(d) $\mathcal{D}_{HAR-X}, \mathcal{L}^{QLIKE}$

Figure 14: Month-ahead relative loss matrices; cell colours indicate statistical significance of pairwise forecast comparisons (light, medium, dark = 10%, 5%, 1%; green/red = column better/worse than row).

A.7 Quantile Score pairwise tests

A.7.1 $\tau = 0.95$

	GRU	HAR	HAR-Q	KAN	LGB	LSTM	LevHAR	Log-HAR	RNN	SHAR	TCN	TKAN	TabPFN	XGB
GRU		0.510	0.508	0.860	0.507	1.581	0.511	0.501	0.687	0.509	0.613	0.567	0.504	0.520
HAR	1.959		0.996	1.686	0.993	3.097	1.002	0.982	1.345	0.997	1.201	1.111	0.987	1.019
HAR-Q	1.967	1.004		1.692	0.996	3.109	1.006	0.986	1.350	1.001	1.205	1.115	0.991	1.023
KAN	1.162	0.593	0.591		0.589	1.838	0.594	0.583	0.798	0.592	0.712	0.659	0.585	0.605
LGB	1.974	1.008	1.004	1.698		3.121	1.010	0.989	1.355	1.005	1.210	1.119	0.994	1.027
LSTM	0.633	0.323	0.322	0.544	0.320		0.323	0.317	0.434	0.322	0.388	0.359	0.319	0.329
LevHAR	1.955	0.998	0.994	1.682	0.991	3.091		0.980	1.343	0.995	1.198	1.108	0.985	1.017
Log-HAR	1.995	1.018	1.015	1.717	1.011	3.154	1.020		1.370	1.015	1.223	1.131	1.005	1.038
RNN	1.456	0.743	0.740	1.253	0.738	2.302	0.745	0.730		0.741	0.892	0.826	0.734	0.758
SHAR	1.965	1.003	0.999	1.690	0.995	3.106	1.005	0.985	1.349		1.204	1.114	0.990	1.022
TCN	1.632	0.833	0.830	1.404	0.827	2.580	0.835	0.818	1.121	0.831		0.925	0.822	0.849
TKAN	1.764	0.900	0.897	1.518	0.894	2.789	0.902	0.884	1.211	0.898	1.081		0.889	0.918
TabPFN	1.986	1.013	1.010	1.708	1.006	3.139	1.015	0.995	1.363	1.010	1.217	1.125		1.033
XGB	1.923	0.981	0.977	1.654	0.974	3.039	0.983	0.964	1.320	0.978	1.178	1.090	0.968	

	GRU	HAR	HAR-Q	KAN	LGB	LSTM	LevHAR	Log-HAR	RNN	SHAR	TCN	TKAN	TabPFN	XGB
GRU		0.352	0.359	0.733	0.346	0.775	0.351	0.344	0.508	0.355	1.185	0.405	0.130	0.330
HAR	2.841		1.021	2.083	0.982	2.203	0.998	0.978	1.443	1.009	3.368	1.151	0.369	0.938
HAR-Q	2.783	0.980		2.041	0.962	2.158	0.978	0.958	1.413	0.989	3.299	1.127	0.362	0.919
KAN	1.364	0.480	0.490		0.471	1.057	0.479	0.470	0.693	0.485	1.617	0.552	0.177	0.450
LGB	2.694	1.019	1.040	2.122		2.244	1.017	0.996	1.470	1.028	3.431	1.172	0.376	0.955
LSTM	1.290	0.454	0.463	0.946	0.446		0.453	0.444	0.655	0.458	1.529	0.522	0.168	0.426
LevHAR	2.847	1.002	1.023	2.087	0.983	2.207		0.980	1.446	1.011	3.374	1.153	0.370	0.939
Log-HAR	2.905	1.022	1.044	2.130	1.004	2.252	1.020		1.475	1.032	3.443	1.176	0.377	0.959
RNN	1.969	0.693	0.707	1.444	0.680	1.526	0.692	0.678		0.700	2.334	0.797	0.256	0.650
SHAR	2.815	0.991	1.011	2.064	0.972	2.182	0.989	0.969	1.429		3.336	1.140	0.366	0.929
TCN	0.844	0.297	0.303	0.619	0.291	0.654	0.296	0.290	0.428	0.300		0.342	0.110	0.278
TKAN	2.469	0.869	0.887	1.811	0.853	1.914	0.867	0.850	1.254	0.877	2.927		0.321	0.815
TabPFN	7.699	2.710	2.766	5.646	2.660	5.969	2.705	2.651	3.910	2.735	9.126	3.118		2.541
XGB	3.030	1.066	1.089	2.222	1.047	2.349	1.064	1.043	1.539	1.076	3.592	1.227	0.394	

	GRU	HAR	HAR-Q	KAN	LGB	LSTM	LevHAR	Log-HAR	RNN	SHAR	TCN	TKAN	TabPFN	XGB
GRU		0.268	0.272	0.277	0.285	1.808	0.278	0.279	0.367	0.274	2.712	0.383	0.041	0.246
HAR	3.735		1.014	1.034	1.065	6.752	1.038	1.041	1.373	1.025	10.130	1.432	0.153	0.919
HAR-Q	3.683	0.986		1.020	1.050	6.657	1.024	1.026	1.353	1.010	9.987	1.412	0.151	0.906
KAN	3.612	0.967	0.981		1.029	6.528	1.004	1.006	1.327	0.991	9.794	1.385	0.148	0.889
LGB	3.508	0.939	0.953	0.971		6.342	0.975	0.977	1.289	0.962	9.515	1.345	0.144	0.863
LSTM	0.553	0.148	0.150	0.153	0.158		0.154	0.154	0.203	0.152	1.500	0.212	0.023	0.136
LevHAR	3.597	0.963	0.977	0.996	1.025	6.503		1.002	1.322	0.987	9.756	1.380	0.147	0.885
Log-HAR	3.589	0.961	0.975	0.994	1.023	6.488	0.998		1.319	0.984	9.734	1.376	0.147	0.883
RNN	2.721	0.728	0.739	0.753	0.776	4.919	0.756	0.758		0.746	7.380	1.044	0.111	0.670
SHAR	3.646	0.976	0.990	1.010	1.039	6.590	1.014	1.016	1.340		9.888	1.398	0.149	0.897
TCN	0.369	0.099	0.100	0.102	0.105	0.667	0.103	0.103	0.136	0.101		0.141	0.015	0.091
TKAN	2.608	0.698	0.708	0.722	0.743	4.714	0.725	0.727	0.958	0.715	7.072		0.107	0.642
TabPFN	24.422	6.538	6.631	6.762	6.961	44.143	6.789	6.804	8.974	6.698	66.229	9.365		6.010
XGB	4.064	1.088	1.103	1.125	1.158	7.345	1.130	1.132	1.493	1.115	11.020	1.558	0.166	

(a) Day-ahead horizon

(b) Average week-ahead horizon

(c) Average month-ahead horizon

Figure 15: Pairwise equal predictive accuracy test results with Quantile Scores ($\tau = 0.95$), estimated on $\mathcal{D}_{HAR}$; cell colours indicate statistical significance of pairwise forecast comparisons (light, medium, and dark shading corresponding to the 10%, 5%, and 1% levels, respectively, with green (red) denoting superior (inferior) performance of the column model relative to the row model).

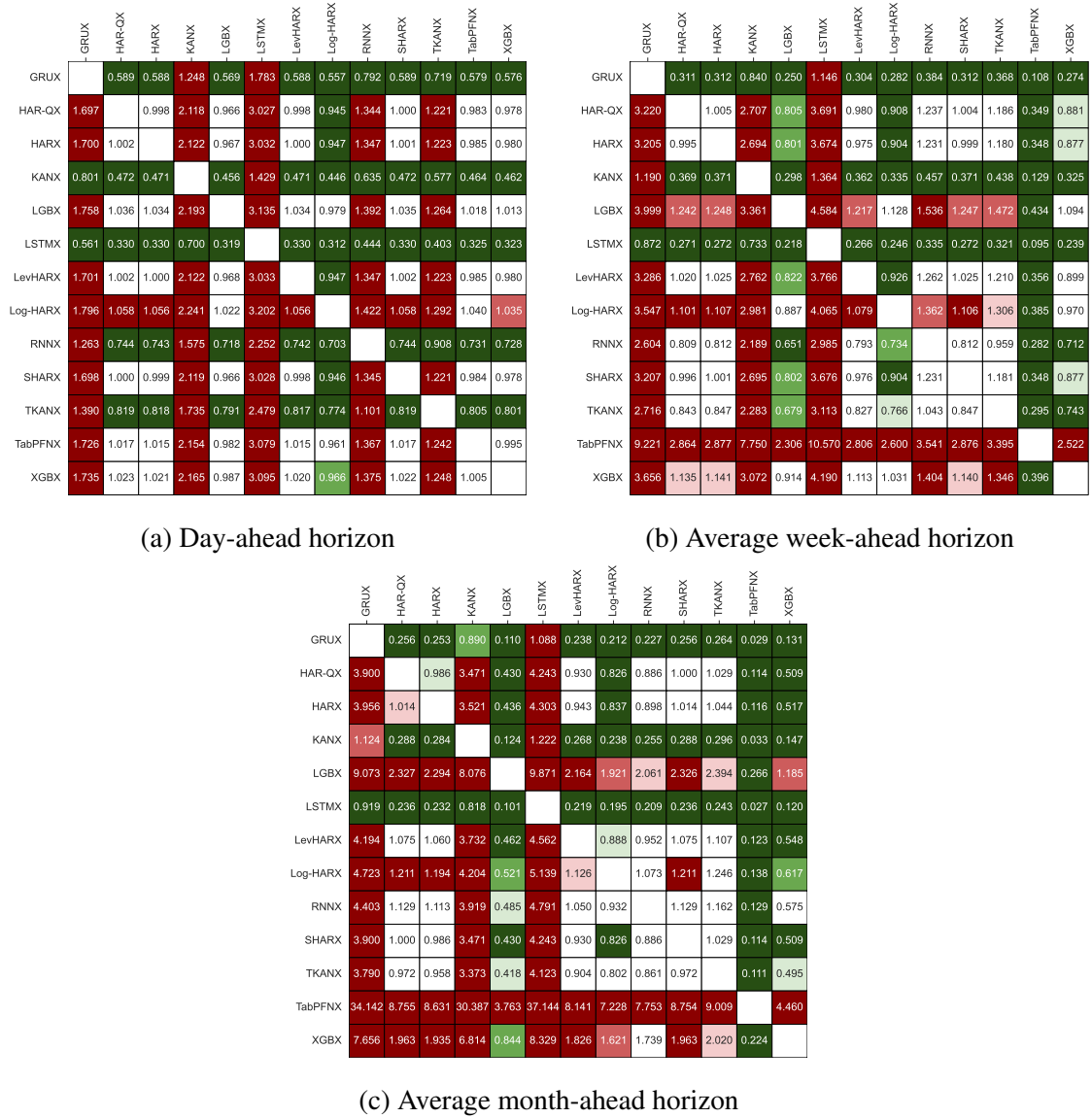


Figure 16: Pairwise equal predictive accuracy test results with Quantile Scores ($\tau = 0.95$), estimated on $\mathcal{D}_{HAR-X}$; cell colours indicate statistical significance of pairwise forecast comparisons (light, medium, and dark shading corresponding to the 10%, 5%, and 1% levels, respectively, with green (red) denoting superior (inferior) performance of the column model relative to the row model).

A.7.2 $\tau = 0.99$

	GRU	HAR	HAR-Q	KAN	LGB	LSTM	LevHAR	Log-HAR	RNN	SHAR	TCN	TKAN	TabPFN	XGB
GRU		0.310	0.310	0.856	0.305	1.609	0.310	0.310	0.634	0.305	0.582	0.535	0.333	0.308
HAR	3.227		0.999	2.764	0.983	5.193	1.000	0.999	2.045	0.984	1.878	1.725	1.075	0.993
HAR-Q	3.230	1.001		2.766	0.984	5.198	1.001	1.000	2.047	0.985	1.880	1.727	1.076	0.994
KAN	1.168	0.362	0.362		0.356	1.879	0.362	0.362	0.740	0.356	0.680	0.624	0.389	0.359
LGB	3.281	1.017	1.016	2.810		5.280	1.017	1.016	2.079	1.000	1.910	1.754	1.093	1.010
LSTM	0.621	0.193	0.192	0.532	0.189		0.193	0.192	0.394	0.189	0.362	0.332	0.207	0.191
LevHAR	3.227	1.000	0.999	2.764	0.984	5.194		0.999	2.045	0.984	1.878	1.725	1.075	0.993
Log-HAR	3.229	1.001	1.000	2.766	0.984	5.196	1.001		2.046	0.985	1.879	1.726	1.076	0.994
RNN	1.578	0.489	0.489	1.352	0.481	2.539	0.489	0.489		0.481	0.919	0.844	0.526	0.486
SHAR	3.280	1.016	1.015	2.809	1.000	5.277	1.016	1.016	2.078		1.909	1.753	1.092	1.009
TCN	1.718	0.532	0.532	1.471	0.524	2.765	0.532	0.532	1.089	0.524		0.918	0.572	0.529
TKAN	1.871	0.580	0.579	1.602	0.570	3.010	0.580	0.579	1.185	0.570	1.089		0.623	0.576
TabPFN	3.002	0.930	0.929	2.571	0.915	4.831	0.930	0.930	1.902	0.915	1.747	1.605		0.924
XGB	3.249	1.007	1.006	2.782	0.990	5.228	1.007	1.006	2.059	0.991	1.891	1.737	1.082	

(a) Day-ahead horizon

	GRU	HAR	HAR-Q	KAN	LGB	LSTM	LevHAR	Log-HAR	RNN	SHAR	TCN	TKAN	TabPFN	XGB
GRU		0.176	0.175	0.717	0.124	0.754	0.165	0.168	0.441	0.180	0.429	0.318	0.082	0.117
HAR	5.692		0.997	4.083	0.706	4.290	0.938	0.958	2.511	1.022	2.440	1.809	0.467	0.667
HAR-Q	5.706	1.003		4.093	0.707	4.301	0.941	0.960	2.518	1.025	2.446	1.813	0.468	0.669
KAN	1.394	0.245	0.244		0.173	1.051	0.230	0.235	0.615	0.250	0.598	0.443	0.114	0.163
LGB	8.066	1.417	1.414	5.786		6.080	1.330	1.357	3.559	1.449	3.458	2.563	0.662	0.945
LSTM	1.327	0.233	0.233	0.952	0.164		0.219	0.223	0.585	0.238	0.569	0.422	0.109	0.155
LevHAR	6.065	1.066	1.063	4.350	0.752	4.571		1.020	2.676	1.089	2.600	1.927	0.498	0.711
Log-HAR	5.944	1.044	1.042	4.263	0.737	4.480	0.980		2.623	1.067	2.548	1.889	0.488	0.696
RNN	2.266	0.398	0.397	1.626	0.281	1.708	0.374	0.381		0.407	0.972	0.720	0.186	0.266
SHAR	5.568	0.978	0.976	3.994	0.690	4.197	0.918	0.937	2.457		2.387	1.769	0.457	0.652
TCN	2.333	0.410	0.409	1.673	0.289	1.758	0.385	0.392	1.029	0.419		0.741	0.191	0.273
TKAN	3.147	0.553	0.552	2.257	0.390	2.372	0.519	0.529	1.389	0.568	1.349		0.258	0.369
TabPFN	12.188	2.141	2.136	8.742	1.511	9.186	2.010	2.051	5.378	2.189	5.225	3.873		1.428
XGB	8.535	1.500	1.496	6.122	1.058	6.433	1.407	1.436	3.766	1.533	3.659	2.712	0.700	

(b) Average week-ahead horizon

	GRU	HAR	HAR-Q	KAN	LGB	LSTM	LevHAR	Log-HAR	RNN	SHAR	TCN	TKAN	TabPFN	XGB
GRU		0.110	0.106	0.185	0.105	1.869	0.128	0.111	0.300	0.104	1.480	0.122	0.026	0.093
HAR	9.119		0.963	1.688	0.956	17.041	1.170	1.009	2.736	0.948	13.493	1.114	0.234	0.844
HAR-Q	9.470	1.038		1.753	0.993	17.696	1.215	1.048	2.841	0.984	14.011	1.157	0.243	0.876
KAN	5.403	0.592	0.571		0.567	10.097	0.693	0.598	1.621	0.561	7.994	0.660	0.139	0.500
LGB	9.537	1.046	1.007	1.765		17.821	1.223	1.055	2.861	0.991	14.111	1.165	0.245	0.882
LSTM	0.535	0.059	0.057	0.099	0.056		0.069	0.059	0.161	0.056	0.792	0.065	0.014	0.050
LevHAR	7.795	0.855	0.823	1.443	0.817	14.567		0.863	2.339	0.810	11.534	0.952	0.200	0.721
Log-HAR	9.038	0.991	0.954	1.673	0.948	16.889	1.159		2.711	0.939	13.373	1.104	0.232	0.836
RNN	3.333	0.366	0.352	0.617	0.350	6.229	0.428	0.369		0.346	4.932	0.407	0.086	0.308
SHAR	9.624	1.055	1.016	1.781	1.009	17.984	1.235	1.065	2.887		14.240	1.176	0.247	0.890
TCN	0.676	0.074	0.071	0.125	0.071	1.263	0.087	0.075	0.203	0.070		0.083	0.017	0.063
TKAN	8.186	0.898	0.864	1.515	0.858	15.297	1.050	0.906	2.456	0.851	12.112		0.210	0.757
TabPFN	38.895	4.265	4.107	7.199	4.079	72.684	4.990	4.304	11.668	4.042	57.551	4.752		3.598
XGB	10.809	1.185	1.141	2.001	1.133	20.199	1.387	1.196	3.243	1.123	15.993	1.320	0.278	

(c) Average month-ahead horizon

Figure 17: Pairwise equal predictive accuracy test results with Quantile Scores ($\tau = 0.99$), estimated on $\mathcal{D}_{HAR}$; cell colours indicate statistical significance of pairwise forecast comparisons (light, medium, and dark shading corresponding to the 10%, 5%, and 1% levels, respectively, with green (red) denoting superior (inferior) performance of the column model relative to the row model).

	GRUX	HAR-QX	HARX	KANX	LGBX	LSTMX	LevHARX	Log-HARX	RNNX	SHARX	TCNX	TKANX	TabPFNX	XGBX
GRUX		0.359	0.361	1.279	0.356	1.835	0.352	0.352	0.738	0.357	0.556	0.815	0.395	0.338
HAR-QX	2.787		1.005	3.565	0.991	5.114	0.981	0.980	2.057	0.994	1.549	2.270	1.101	0.942
HARX	2.773	0.995		3.548	0.986	5.090	0.976	0.975	2.047	0.989	1.541	2.259	1.096	0.938
KANX	0.782	0.281	0.282		0.278	1.435	0.275	0.275	0.577	0.279	0.434	0.637	0.309	0.264
LGBX	2.811	1.009	1.014	3.596		5.159	0.989	0.989	2.075	1.003	1.562	2.290	1.111	0.951
LSTMX	0.545	0.196	0.196	0.697	0.194		0.192	0.192	0.402	0.194	0.303	0.444	0.215	0.184
LevHARX	2.842	1.020	1.025	3.635	1.011	5.216		0.999	2.097	1.014	1.579	2.315	1.123	0.961
Log-HARX	2.843	1.020	1.025	3.637	1.011	5.218	1.001		2.098	1.014	1.580	2.316	1.123	0.961
RNNX	1.355	0.486	0.489	1.733	0.482	2.487	0.477	0.477		0.483	0.753	1.104	0.535	0.458
SHARX	2.804	1.006	1.011	3.587	0.997	5.146	0.987	0.986	2.069		1.558	2.284	1.108	0.948
TCNX	1.799	0.646	0.649	2.302	0.640	3.302	0.633	0.633	1.328	0.642		1.466	0.711	0.608
TKANX	1.228	0.441	0.443	1.570	0.437	2.253	0.432	0.432	0.906	0.438	0.682		0.485	0.415
TabPFNX	2.531	0.908	0.913	3.238	0.900	4.646	0.891	0.890	1.868	0.903	1.407	2.062		0.856
XGBX	2.957	1.061	1.066	3.783	1.052	5.427	1.041	1.040	2.183	1.055	1.644	2.409	1.168	

	GRUX	HAR-QX	HARX	KANX	LGBX	LSTMX	LevHARX	Log-HARX	RNNX	SHARX	TCNX	TKANX	TabPFNX	XGBX
GRUX		0.174	0.174	0.824	0.094	1.085	0.168	0.142	0.327	0.179	0.958	0.284	0.070	0.099
HAR-QX	5.749		0.999	4.735	0.542	6.238	0.968	0.815	1.880	1.026	5.506	1.520	0.401	0.569
HARX	5.754	1.001		4.739	0.543	6.243	0.969	0.816	1.882	1.027	5.511	1.521	0.402	0.570
KANX	1.214	0.211	0.211		0.114	1.317	0.204	0.172	0.397	0.217	1.163	0.321	0.085	0.120
LGBX	10.605	1.845	1.843	8.736		11.507	1.786	1.504	3.468	1.893	10.157	2.804	0.740	1.050
LSTMX	0.922	0.160	0.160	0.759	0.087		0.155	0.131	0.301	0.165	0.883	0.244	0.064	0.091
LevHARX	5.938	1.033	1.032	4.891	0.580	6.443		0.842	1.942	1.060	5.687	1.570	0.414	0.588
Log-HARX	7.050	1.226	1.225	5.808	0.685	7.650	1.187		2.306	1.259	6.753	1.864	0.492	0.698
RNNX	3.058	0.532	0.531	2.519	0.288	3.318	0.515	0.434		0.546	2.929	0.809	0.213	0.303
SHARX	5.601	0.974	0.973	4.614	0.528	6.077	0.943	0.794	1.832		5.365	1.481	0.391	0.555
TCNX	1.044	0.182	0.181	0.860	0.098	1.133	0.176	0.148	0.341	0.186		0.276	0.073	0.103
TKANX	3.782	0.658	0.657	3.115	0.357	4.103	0.637	0.536	1.237	0.675	3.622		0.264	0.374
TabPFNX	14.327	2.492	2.490	11.802	1.351	15.546	2.413	2.032	4.686	2.558	13.723	3.789		1.419
XGBX	10.099	1.757	1.755	8.319	0.952	10.958	1.701	1.432	3.303	1.803	9.673	2.670	0.705	

(a) Day-ahead horizon

(b) Average week-ahead horizon

	GRUX	HAR-QX	HARX	KANX	LGBX	LSTMX	LevHARX	Log-HARX	RNNX	SHARX	TCNX	TKANX	TabPFNX	XGBX
GRUX		0.123	0.128	0.894	0.039	1.084	0.109	0.067	0.173	0.134	1.947	0.107	0.018	0.048
HAR-QX	8.137		1.039	7.273	0.318	8.822	0.887	0.545	1.405	1.089	15.842	0.871	0.150	0.390
HARX	7.835	0.963		7.003	0.306	8.494	0.854	0.525	1.353	1.048	15.254	0.839	0.144	0.375
KANX	1.119	0.137	0.143		0.044	1.213	0.122	0.075	0.193	0.150	2.178	0.120	0.021	0.054
LGBX	25.568	3.142	3.263	22.853		27.719	2.786	1.714	4.414	3.421	49.776	2.737	0.470	1.225
LSTMX	0.922	0.113	0.118	0.824	0.036		0.101	0.062	0.159	0.123	1.796	0.099	0.017	0.044
LevHARX	9.177	1.128	1.171	8.202	0.359	9.949		0.615	1.584	1.228	17.865	0.982	0.169	0.440
Log-HARX	14.921	1.834	1.904	13.336	0.584	16.176	1.626		2.576	1.996	29.048	1.597	0.274	0.715
RNNX	5.792	0.712	0.739	5.177	0.227	5.280	0.631	0.388		0.775	11.276	0.620	0.107	0.277
SHARX	7.474	0.918	0.954	6.681	0.292	8.103	0.814	0.501	1.290		14.551	0.800	0.137	0.358
TCNX	0.514	0.063	0.066	0.459	0.020	0.557	0.056	0.034	0.089	0.069		0.055	0.009	0.025
TKANX	9.342	1.148	1.192	8.350	0.365	10.128	1.018	0.626	1.613	1.250	18.188		0.172	0.447
TabPFNX	54.376	6.682	6.940	48.602	2.127	58.951	5.925	3.644	9.388	7.275	105.856	5.820		2.605
XGBX	20.877	2.566	2.665	18.661	0.817	22.634	2.275	1.399	3.604	2.793	40.644	2.235	0.384	

(c) Average month-ahead horizon

Figure 18: Pairwise equal predictive accuracy test results with Quantile Scores ($\tau = 0.99$), estimated on $\mathcal{D}_{HAR-X}$; cell colours indicate statistical significance of pairwise forecast comparisons (light, medium, and dark shading corresponding to the 10%, 5%, and 1% levels, respectively, with green (red) denoting superior (inferior) performance of the column model relative to the row model).