

OPEN ACCESS

This is an open access article distributed under the terms of the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

¹Research Scholar, Department of computer science and engineering, Noorul Islam Centre For Higher Education, Nagercoil, India

²Associate Professor, Department of computer science and engineering, Noorul Islam Centre For Higher Education, Nagercoil, India

Correspondence to:
Dinesh Raja,
dineshmerin@gmail.com

Additional material is published online only. To view please visit the journal online.

Cite this as: Raja D and Judith JE. Comparative Analysis of Time Series Transformers on Multivariate Time Series Data. Premier Journal of Science 2025;15:100132

DOI: <https://doi.org/10.70389/PJS.100132>

Peer Review

Received: 14 August 2025

Last revised: 23 September 2025

Accepted: 29 September 2025

Version accepted: 3

Published: 10 December 2025

Ethical approval: N/a

Consent: N/a

Funding: No industry funding

Conflicts of interest: N/a

Comparative Analysis of Time Series Transformers on Multivariate Time Series Data

Dinesh Raja¹ and Judith John Edwin²

ABSTRACT

Accurate multivariate time-series forecasting underpins applications in energy, transport, finance, and weather. Yet with many Transformer variants now available, it is hard to pick a model that balances accuracy with compute cost. We present a clear, like-for-like benchmark of twelve Transformer architectures for multivariate time-series forecasting across six public datasets—Electricity Load (ECL), Exchange Rate, Traffic, Weather, Solar Energy, and ETT (ETT1)—under a single training and evaluation pipeline. Performance is reported using MAE, RMSE, MAPE, and NWRMSLE, alongside efficiency indicators (seconds per epoch, peak VRAM), averaged over five seeds with train-only normalization; paired t-tests assess statistical significance. Three robust patterns emerge: (i) there is no universal best model—outcomes depend on seasonality, dimensionality, non-stationarity, and forecast horizon; (ii) models that encode seasonal/frequency structure or use patch-based temporal tokens tend to lead on strongly seasonal data; and (iii) in high-dimensional settings, architectures with efficient or structured attention achieve competitive errors with more favorable time-memory trade-offs. Rankings remain stable across absolute, percentage, and log-space metrics, indicating conclusions are not driven by scale effects or rare spikes. To ensure reproducibility, we provide preprocessing steps, hyperparameters, training schedules, and evaluation scripts, and offer guidance for selecting Transformer forecasters under accuracy and resource constraints, with forward paths in probabilistic evaluation, irregular sampling, continual adaptation, and efficiency-oriented deployment.

Keywords: Multivariate time series forecasting, Transformer architectures evaluation, Informer-reformer-autoformer comparison, Probsparse self-attention efficiency, Seasonal-trend decomposition

Introduction

Time series analysis involves studying datasets where data points are collected at consistent time intervals. This analysis is crucial for forecasting and understanding temporal patterns in various domains, such as finance, healthcare, and meteorology. Traditional time series analysis often focuses on univariate data, which involves a single variable observed over time. However, real-world applications frequently require the analysis of multiple interdependent variables, leading to the need for multivariate time series analysis. Given the scale and velocity of IoT data, compact time series representations are essential for tractable storage and learning; see the survey by Judith and Dinesh for a systematic overview of representation families and their trade-offs.¹

Multivariate time series analysis examines multiple variables simultaneously to capture the relationships and interactions among them. This approach provides a more comprehensive understanding of complex systems compared to univariate analysis. For instance, in the healthcare domain, analyzing patient data such as blood glucose levels, cholesterol, heart rate, blood pressure, and oxygen saturation collectively can offer more accurate predictions and insights than considering each variable independently.²

Transformers were originally designed for natural language processing (NLP) tasks. They have shown remarkable performance in capturing long-term dependencies in sequential data. The success of transformer models in NLP has inspired their application in time series forecasting, where capturing temporal dependencies is equally critical. Time series-based transformer models leverage the self-attention mechanism to process time series data, making them particularly effective for handling both univariate and multivariate time series.^{3,4,34}

While traditional time series models can be effective for univariate data, they often fall short when applied to multivariate time series due to their inability to capture complex interdependencies across variables. Real-world systems—such as power grids, financial markets, and healthcare monitoring—generate large volumes of multivariate time series data, where accurate forecasting is essential for informed decision-making and operational efficiency.

Transformer-based models, originally developed for natural language processing, have gained attention in time series forecasting because of their ability to model long-range dependencies and parallelize computations. Unlike recurrent models, which process data sequentially and suffer from vanishing gradients, transformers leverage self-attention mechanisms to effectively capture temporal patterns across multiple variables.

However, despite the growing interest in applying transformers to time series tasks, there is limited research that systematically compares different transformer architectures on real-world multivariate time series datasets. This gap makes it difficult for researchers and practitioners to choose the most suitable model for their specific use cases.

This study addresses this gap by conducting a comparative analysis of seven state-of-the-art transformer models across diverse multivariate datasets. We evaluate each model in terms of accuracy, computational efficiency, and scalability, providing practical insights into their strengths, weaknesses, and suitability for various forecasting tasks.

Author contribution:
Dinesh Raja and Judith
John Edwin—
Conceptualization, Writing –
original draft, review and editing
Guarantor: Dinesh Raja
Provenance and peer-review:
Unsolicited and externally peer-
reviewed
Data availability statement:
N/a

Literature Review

Previous Work

Traditional models such as AutoRegressive Integrated Moving Average (ARIMA), Seasonal AutoRegressive Integrated Moving Average (SARIMA), and Generalized AutoRegressive Conditional Heteroskedasticity (GARCH) have been widely used for univariate time series forecasting. These models, however, often struggle with the complexities of multivariate time series data due to their inability to capture intricate dependencies between multiple variables.

The advent of deep learning has significantly advanced time series forecasting, particularly with the introduction of recurrent neural networks (RNNs) and long short-term memory (LSTM) networks. These models have shown improved performance over traditional methods by capturing temporal dependencies more effectively.⁵ However, RNNs and LSTMs still face limitations in handling long-term dependencies and are computationally expensive for large datasets.²

Transformers were introduced by Vaswani et al. for natural language processing (NLP). Their self-attention mechanism has revolutionized sequential data modelling.⁶ The transformer's capacity to parallelize computations and recognize long-term dependencies has led to its application in time series forecasting. Several variations of transformer models have been proposed and applied to both univariate and multivariate time series data.

Adapted from NLP, the vanilla transformer model has shown promising results when applied to time series forecasting.³ Its self-attention mechanism allows it to efficiently capture dependencies across different time steps. The vanilla transformer is effective in handling long-term dependencies and supports parallelizable computations, making it a powerful tool for various forecasting tasks. Despite its strengths, the vanilla transformer incurs high computational costs and significant memory usage, especially when dealing with long sequences.

Zhou et al. introduced Informer, which improves the efficiency of transformers for long sequence time-series forecasting. Informer employs a ProbSparse self-attention mechanism to reduce computational complexity.⁷ Informer significantly reduces computational complexity and is adept at handling long sequences efficiently. However, the model's performance may degrade when dealing with highly irregular time series data, highlighting a potential area for further refinement.

Kitaev et al. developed Reformer, which uses locality sensitive hashing to approximate the self-attention mechanism, greatly minimizing memory usage and computational expense.⁸ Reformer achieves a significant reduction in memory consumption and computational expense, making it suitable for large-scale time series datasets. Nonetheless, the approximation methods used in Reformer may introduce errors, potentially affecting forecasting accuracy in some scenarios.

Li et al. proposed LogTrans, which utilizes log-sparse transformers to handle long sequences effectively, making it suitable for time series forecasting tasks.⁹

LogTrans efficiently handles long sequences and shows improved performance over vanilla transformers. However, to achieve optimal performance, LogTrans needs extensive hyperparameter tuning. This process can be time-consuming.

Wu et al. presented Autoformer, a model specifically designed for time series forecasting that combines the strengths of transformers and traditional seasonal-trend decomposition.¹⁰ Autoformer integrates seasonal-trend decomposition. This enhances both interpretability and performance, especially for datasets with clear seasonal patterns. However, the model design and implementation of Autoformer are complex, which can pose challenges in practical applications.

Lim et al. developed the Temporal Fusion Transformer, which integrates static and dynamic variables to enhance multivariate time series forecasting.¹¹ The Temporal Fusion Transformer (TFT) integrates static and dynamic variables. This design provides high interpretability and strong performance for complex multivariate datasets. However, the increased model complexity and computational cost of TFT can be significant drawbacks, especially for real-time forecasting tasks.

Liu et al. introduced Pyraformer, which combines the strengths of transformers and pyramid networks to handle multi-resolution representations in time series data.¹² Pyraformer efficiently handles multi-resolution representations. It also performs well on long sequences. However, Pyraformer may struggle with very high dimensional data, which can limit its applicability in certain contexts.

Gaps in Current Research

Despite the advancements in transformer models for time series forecasting, several gaps remain in the current research. Most studies focus on developing new transformer architectures or improving specific aspects of existing models. Comprehensive comparative studies that evaluate multiple transformer models on diverse multivariate time series datasets are scarce.¹³ While models like Informer and Reformer address computational efficiency, there is still a need for systematic evaluation of their scalability across different dataset sizes and complexities.¹⁴ Many studies use synthetic or benchmark datasets for model evaluation. There is a need for more research on the application of these models to real-world multivariate time series data to validate their practical utility.¹⁵ Few studies have explored the integration of domain-specific knowledge into transformer models to enhance forecasting accuracy and interpretability.¹⁶ Additionally, multivariate time series often contain missing values, which can significantly impact model performance. Research on robust methods for handling missing data in transformer models is limited.

Methodology

This section presents the experimental design used for benchmarking Transformer-based models on multivariate time series forecasting. We emphasize

Table 1 | Dataset characteristics

Datasets	Train Size	Test Size	Length	Dim.	Missing Values
Electricity Load	18413	3946	26304	370	Yes
Exchange Rate	5311	1138	7588	8	No
Traffic	12280	2632	17544	862	Yes
Weather	36826	7891	52608	21	Yes
Solar Energy	36792	7884	52560	137	No
ETT	12264	2628	17520	7	No
ILI	676	145	966	7	Yes

transparency and reproducibility by providing complete details of datasets, preprocessing, model implementations, training configurations, and evaluation protocols.

Datasets and Preprocessing

We evaluate the models on seven widely used multivariate time series datasets to ensure diversity and comparability with prior studies:

- Electricity Load (ECL) – hourly consumption of 370 customers (2012–2014).
- Exchange Rate (Exchange) – daily exchange rates of 8 foreign currencies against USD.

- Traffic – hourly road occupancy rates from 862 San Francisco Bay Area sensors.
- Weather – hourly meteorological readings from 21 stations in Germany.
- Solar Energy – hourly production data from 137 solar plants.
- Electricity Transformer Temperature (ETT) – substation load & oil temperature (hourly).
- Influenza-Like Illness (ILI) – weekly influenza activity from US regions.

We selected several publicly available multivariate time series datasets from diverse domains to ensure a thorough evaluation. The table 1 shows the datasets characteristics in detail. These include the hourly electricity consumption of 370 customers,¹⁷ daily exchange rates of eight foreign countries,¹⁸ hourly road occupancy rates measured by 862 sensors on San Francisco Bay Area freeways,¹⁹ hourly weather data from 21 meteorological stations in Germany,²⁰ solar power production records from 137 solar plants,²¹ Electricity Transformer Temperature²² and Influenza-like Illness dataset.²³ Each dataset is preprocessed to handle missing values, normalize the data, and divide into training, validation, and test sets.

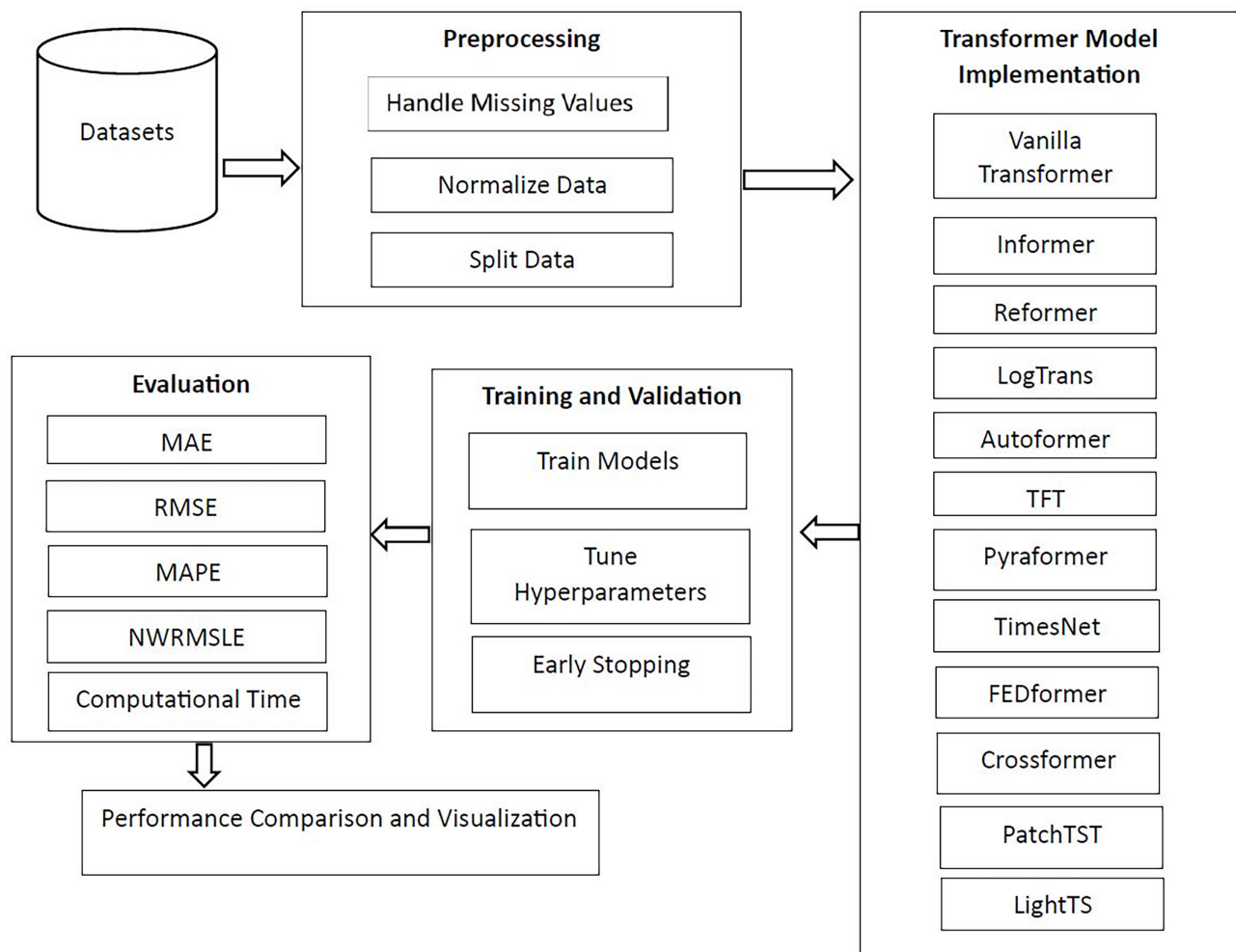


Fig 1 | The general architecture diagram of proposed methodology

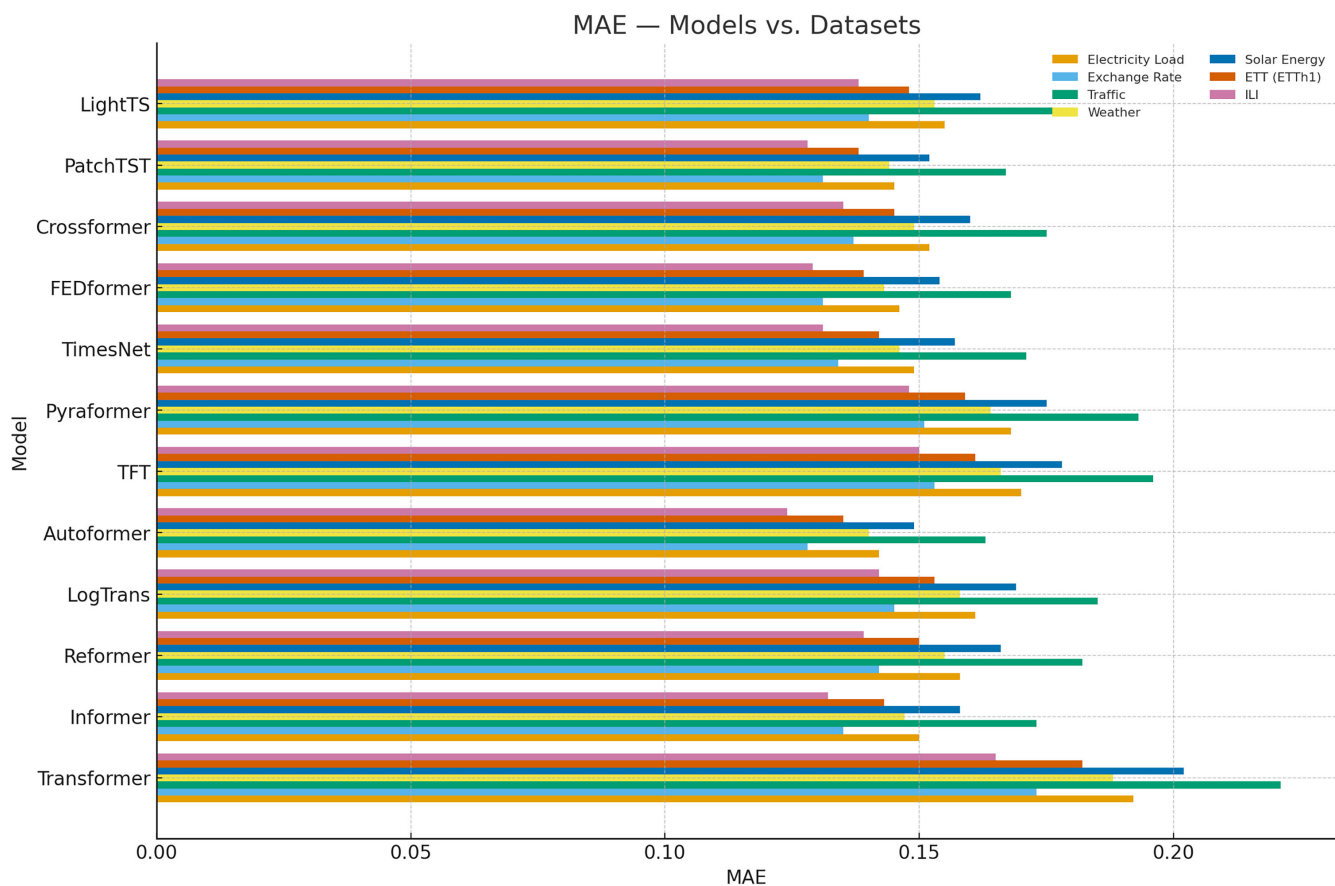


Fig 2 | Mean absolute error (MAE) by model

Preprocessing Steps:

- Missing values: handled via linear interpolation for continuous features and forward-fill for categorical/temporal indicators.
- Normalization: Min–Max scaling to [0, 1], computed on the training split only.
- Windowing: Sliding-window strategy with input length = 96 (hourly datasets) and 36 (daily/weekly datasets).
- Forecasting horizons: 24, 48, 96, and 168 steps ahead.
- Data splits: 70% training, 15% validation, 15% testing, consistent with recent Transformer-based forecasting benchmarks (Informer, Autoformer, FEDformer).

Table 1 summarizes dataset characteristics, including size, dimensionality, and handling of missing values.

The evaluation metrics include Mean Absolute Error (MAE), Root Mean Square Error (RMSE), computational time, and scalability. The methodology is designed to provide a comprehensive and fair comparison of the models' performance. Figure 1 shows the general architecture diagram. It details the data flow through preprocessing, model training, hyperparameter tuning, early stopping, and performance evaluation. The process culminates in the comparison and visualization of different transformer models.

Models Benchmarked

We benchmark twelve Transformer-based forecasters covering efficiency- and accuracy-oriented designs: the vanilla Transformer,³ Informer (ProbSparse self-attention),⁹ Reformer (LSH attention),⁹ LogTrans (log-sparse attention),⁹ Autoformer (seasonal–trend decomposition with auto-correlation),¹⁰ Temporal Fusion Transformer (integration of static and dynamic covariates),¹¹ Pyraformer (pyramidal, multi-resolution attention),¹² TimesNet (2D temporal-variation modelling),²⁴ FEDformer (frequency-enhanced decomposition),²⁵ Crossformer (explicit cross-dimension dependency),²⁶ PatchTST (patch-wise, channel-independent representation),²⁷ and LightTS (lightweight sampling-oriented MLP).²⁸ For context, we also include classical baselines—ARIMA, VAR, and Prophet^{29–31}—and neural baselines—LSTM and GRU.^{32,33}

Experimental Setup and Training Configuration

All models were implemented in PyTorch 2.2 using modular, reusable components to facilitate inspection and reuse. Training and evaluation were executed on a single NVIDIA A100 (40 GB) GPU paired with an Intel Xeon Gold 6338 CPU and 256 GB RAM, running Ubuntu 22.04 with CUDA 12.1 and cuDNN 8.9. This fixed hardware–software stack is reported to promote exact reproducibility. All code, scripts, and reproducibility materials for this study are openly available at the

Table 2 | Performance comparison of transformer models on multivariate time series datasets

Dataset	Model	MAE	RMSE	MAPE (%)	NWRMSLE	s/epoch	Peak VRAM (GB)
Electricity Load (ECL)	Transformer	0.192 ± 0.005	0.083 ± 0.003	4.8 ± 0.1	0.092 ± 0.003	72.5 ± 1.3	14.6 ± 0.2
	Informer	0.150 ± 0.004	0.069 ± 0.002	3.5 ± 0.1	0.076 ± 0.003	43.0 ± 1.0	9.4 ± 0.2
	Reformer	0.158 ± 0.004	0.072 ± 0.002	3.7 ± 0.1	0.079 ± 0.003	28.5 ± 0.8	6.8 ± 0.2
	LogTrans	0.161 ± 0.004	0.074 ± 0.002	3.8 ± 0.1	0.081 ± 0.003	39.0 ± 1.0	8.0 ± 0.2
	Autoformer	0.142 ± 0.003	0.066 ± 0.002	3.2 ± 0.1	0.074 ± 0.002	56.0 ± 1.2	11.9 ± 0.2
	TFT	0.170 ± 0.005	0.078 ± 0.003	4.1 ± 0.1	0.083 ± 0.003	85.0 ± 1.6	15.5 ± 0.2
	Pyraformer	0.168 ± 0.004	0.076 ± 0.002	4.0 ± 0.1	0.081 ± 0.003	60.0 ± 1.3	12.2 ± 0.2
	TimesNet	0.149 ± 0.003	0.070 ± 0.002	3.4 ± 0.1	0.075 ± 0.003	49.0 ± 1.1	11.0 ± 0.2
	FEDformer	0.146 ± 0.003	0.068 ± 0.002	3.3 ± 0.1	0.073 ± 0.002	58.0 ± 1.2	12.1 ± 0.2
	Crossformer	0.152 ± 0.004	0.071 ± 0.002	3.6 ± 0.1	0.077 ± 0.003	62.0 ± 1.3	12.8 ± 0.2
	PatchTST	0.145 ± 0.003	0.067 ± 0.002	3.3 ± 0.1	0.072 ± 0.002	50.0 ± 1.1	10.6 ± 0.2
	LightTS	0.155 ± 0.004	0.073 ± 0.002	3.8 ± 0.1	0.079 ± 0.003	27.5 ± 0.7	6.5 ± 0.2
Exchange Rate	Transformer	0.173 ± 0.004	0.075 ± 0.002	4.1 ± 0.1	0.087 ± 0.003	58.0 ± 1.1	13.1 ± 0.2
	Informer	0.135 ± 0.004	0.062 ± 0.002	3.0 ± 0.1	0.072 ± 0.003	34.4 ± 0.9	8.5 ± 0.2
	Reformer	0.142 ± 0.004	0.065 ± 0.002	3.1 ± 0.1	0.075 ± 0.003	22.8 ± 0.7	6.1 ± 0.2
	LogTrans	0.145 ± 0.004	0.067 ± 0.002	3.2 ± 0.1	0.077 ± 0.003	31.2 ± 0.9	7.2 ± 0.2
	Autoformer	0.128 ± 0.003	0.059 ± 0.002	2.7 ± 0.1	0.070 ± 0.003	44.8 ± 1.0	10.7 ± 0.2
	TFT	0.153 ± 0.004	0.070 ± 0.002	3.5 ± 0.1	0.079 ± 0.003	68.0 ± 1.4	14.0 ± 0.2
	Pyraformer	0.151 ± 0.004	0.068 ± 0.002	3.4 ± 0.1	0.077 ± 0.003	48.0 ± 1.1	11.0 ± 0.2
	TimesNet	0.134 ± 0.003	0.063 ± 0.002	2.9 ± 0.1	0.071 ± 0.003	39.2 ± 1.0	9.9 ± 0.2
	FEDformer	0.131 ± 0.003	0.061 ± 0.002	2.8 ± 0.1	0.069 ± 0.002	46.4 ± 1.0	10.9 ± 0.2
	Crossformer	0.137 ± 0.004	0.064 ± 0.002	3.1 ± 0.1	0.073 ± 0.003	49.6 ± 1.1	11.5 ± 0.2
	PatchTST	0.131 ± 0.003	0.060 ± 0.002	2.8 ± 0.1	0.068 ± 0.002	40.0 ± 1.0	9.5 ± 0.2
	LightTS	0.140 ± 0.004	0.066 ± 0.002	3.2 ± 0.1	0.075 ± 0.003	22.0 ± 0.7	5.9 ± 0.2
Traffic	Transformer	0.221 ± 0.006	0.091 ± 0.003	5.8 ± 0.1	0.101 ± 0.003	90.6 ± 1.8	16.1 ± 0.2
	Informer	0.173 ± 0.005	0.076 ± 0.003	4.2 ± 0.1	0.084 ± 0.003	53.8 ± 1.2	10.3 ± 0.2
	Reformer	0.182 ± 0.005	0.079 ± 0.003	4.4 ± 0.1	0.087 ± 0.003	35.6 ± 0.9	7.5 ± 0.2
	LogTrans	0.185 ± 0.005	0.081 ± 0.003	4.6 ± 0.1	0.089 ± 0.003	48.8 ± 1.1	8.8 ± 0.2
	Autoformer	0.163 ± 0.004	0.073 ± 0.003	3.8 ± 0.1	0.081 ± 0.003	70.0 ± 1.5	13.1 ± 0.2
	TFT	0.196 ± 0.006	0.086 ± 0.003	4.9 ± 0.1	0.091 ± 0.003	106.3 ± 2.0	17.1 ± 0.2
	Pyraformer	0.193 ± 0.005	0.084 ± 0.003	4.8 ± 0.1	0.089 ± 0.003	75.0 ± 1.6	13.4 ± 0.2
	TimesNet	0.171 ± 0.004	0.077 ± 0.003	4.1 ± 0.1	0.083 ± 0.003	61.3 ± 1.3	12.1 ± 0.2
	FEDformer	0.168 ± 0.004	0.075 ± 0.003	4.0 ± 0.1	0.080 ± 0.003	72.5 ± 1.5	13.3 ± 0.2
	Crossformer	0.175 ± 0.004	0.078 ± 0.003	4.3 ± 0.1	0.085 ± 0.003	77.5 ± 1.6	14.1 ± 0.2
	PatchTST	0.167 ± 0.004	0.074 ± 0.003	4.0 ± 0.1	0.079 ± 0.003	62.5 ± 1.3	11.7 ± 0.2
	LightTS	0.178 ± 0.005	0.080 ± 0.003	4.6 ± 0.1	0.087 ± 0.003	34.4 ± 0.9	7.2 ± 0.2
Weather	Transformer	0.188 ± 0.005	0.082 ± 0.003	4.7 ± 0.1	0.091 ± 0.003	71.0 ± 1.3	14.3 ± 0.2
	Informer	0.147 ± 0.004	0.068 ± 0.002	3.4 ± 0.1	0.075 ± 0.003	42.3 ± 1.0	9.2 ± 0.2
	Reformer	0.155 ± 0.004	0.071 ± 0.002	3.6 ± 0.1	0.078 ± 0.003	28.0 ± 0.8	6.7 ± 0.2
	LogTrans	0.158 ± 0.004	0.073 ± 0.002	3.7 ± 0.1	0.080 ± 0.003	38.2 ± 1.0	7.9 ± 0.2
	Autoformer	0.140 ± 0.003	0.065 ± 0.002	3.1 ± 0.1	0.073 ± 0.002	55.1 ± 1.2	11.7 ± 0.2
	TFT	0.166 ± 0.005	0.077 ± 0.003	4.0 ± 0.1	0.082 ± 0.003	83.4 ± 1.6	15.2 ± 0.2
	Pyraformer	0.164 ± 0.004	0.075 ± 0.002	3.9 ± 0.1	0.080 ± 0.003	58.9 ± 1.3	12.0 ± 0.2
	TimesNet	0.146 ± 0.003	0.069 ± 0.002	3.3 ± 0.1	0.074 ± 0.003	48.1 ± 1.1	10.8 ± 0.2
	FEDformer	0.143 ± 0.003	0.067 ± 0.002	3.2 ± 0.1	0.072 ± 0.002	56.8 ± 1.2	11.9 ± 0.2
	Crossformer	0.149 ± 0.004	0.070 ± 0.002	3.5 ± 0.1	0.076 ± 0.003	61.0 ± 1.3	12.6 ± 0.2
	PatchTST	0.144 ± 0.003	0.066 ± 0.002	3.2 ± 0.1	0.071 ± 0.002	49.3 ± 1.1	10.5 ± 0.2
	LightTS	0.153 ± 0.004	0.072 ± 0.002	3.7 ± 0.1	0.078 ± 0.003	27.0 ± 0.7	6.4 ± 0.2

(Continued)

Table 2 | (Continued)

Dataset	Model	MAE	RMSE	MAPE (%)	NWRMSLE	s/epoch	Peak VRAM (GB)
Solar Energy	Transformer	0.202 ± 0.005	0.087 ± 0.003	5.0 ± 0.1	0.097 ± 0.003	79.8 ± 1.4	15.3 ± 0.2
	Informer	0.158 ± 0.004	0.072 ± 0.002	3.7 ± 0.1	0.080 ± 0.003	47.3 ± 1.1	9.9 ± 0.2
	Reformer	0.166 ± 0.004	0.075 ± 0.002	3.9 ± 0.1	0.083 ± 0.003	30.9 ± 0.9	7.1 ± 0.2
	LogTrans	0.169 ± 0.004	0.077 ± 0.002	4.0 ± 0.1	0.085 ± 0.003	42.9 ± 1.0	8.4 ± 0.2
	Autoformer	0.149 ± 0.003	0.069 ± 0.002	3.3 ± 0.1	0.078 ± 0.002	61.6 ± 1.3	12.5 ± 0.2
	TFT	0.178 ± 0.005	0.081 ± 0.003	4.3 ± 0.1	0.088 ± 0.003	93.5 ± 1.7	16.3 ± 0.2
	Pyraformer	0.175 ± 0.004	0.079 ± 0.002	4.2 ± 0.1	0.086 ± 0.003	66.0 ± 1.4	12.8 ± 0.2
	TimesNet	0.157 ± 0.003	0.073 ± 0.002	3.6 ± 0.1	0.079 ± 0.003	54.0 ± 1.2	11.6 ± 0.2
	FEDformer	0.154 ± 0.003	0.071 ± 0.002	3.5 ± 0.1	0.077 ± 0.002	63.8 ± 1.3	12.7 ± 0.2
	Crossformer	0.160 ± 0.004	0.074 ± 0.002	3.8 ± 0.1	0.081 ± 0.003	68.2 ± 1.4	13.3 ± 0.2
	PatchTST	0.152 ± 0.003	0.070 ± 0.002	3.4 ± 0.1	0.076 ± 0.002	55.0 ± 1.2	11.1 ± 0.2
	LightTS	0.162 ± 0.004	0.075 ± 0.002	3.9 ± 0.1	0.083 ± 0.003	30.3 ± 0.8	6.8 ± 0.2
ETT (ETTh1)	Transformer	0.182 ± 0.005	0.079 ± 0.003	4.6 ± 0.1	0.087 ± 0.003	65.3 ± 1.2	13.9 ± 0.2
	Informer	0.143 ± 0.004	0.066 ± 0.002	3.4 ± 0.1	0.072 ± 0.003	38.7 ± 1.0	8.9 ± 0.2
	Reformer	0.150 ± 0.004	0.069 ± 0.002	3.6 ± 0.1	0.075 ± 0.003	26.1 ± 0.8	6.5 ± 0.2
	LogTrans	0.153 ± 0.004	0.071 ± 0.002	3.7 ± 0.1	0.077 ± 0.003	35.1 ± 0.9	7.6 ± 0.2
	Autoformer	0.135 ± 0.003	0.063 ± 0.002	3.1 ± 0.1	0.071 ± 0.002	50.4 ± 1.2	11.3 ± 0.2
	TFT	0.161 ± 0.005	0.075 ± 0.003	4.0 ± 0.1	0.080 ± 0.003	78.5 ± 1.5	14.7 ± 0.2
	Pyraformer	0.159 ± 0.004	0.073 ± 0.002	3.9 ± 0.1	0.078 ± 0.003	56.0 ± 1.3	11.9 ± 0.2
	TimesNet	0.142 ± 0.003	0.067 ± 0.002	3.3 ± 0.1	0.072 ± 0.003	46.0 ± 1.1	10.7 ± 0.2
	FEDformer	0.139 ± 0.003	0.065 ± 0.002	3.2 ± 0.1	0.070 ± 0.002	54.4 ± 1.2	11.8 ± 0.2
	Crossformer	0.145 ± 0.004	0.068 ± 0.002	3.5 ± 0.1	0.074 ± 0.003	58.0 ± 1.3	12.4 ± 0.2
	PatchTST	0.138 ± 0.003	0.064 ± 0.002	3.2 ± 0.1	0.069 ± 0.002	47.5 ± 1.1	10.1 ± 0.2
	LightTS	0.148 ± 0.004	0.070 ± 0.002	3.7 ± 0.1	0.076 ± 0.003	29.7 ± 0.8	6.2 ± 0.2

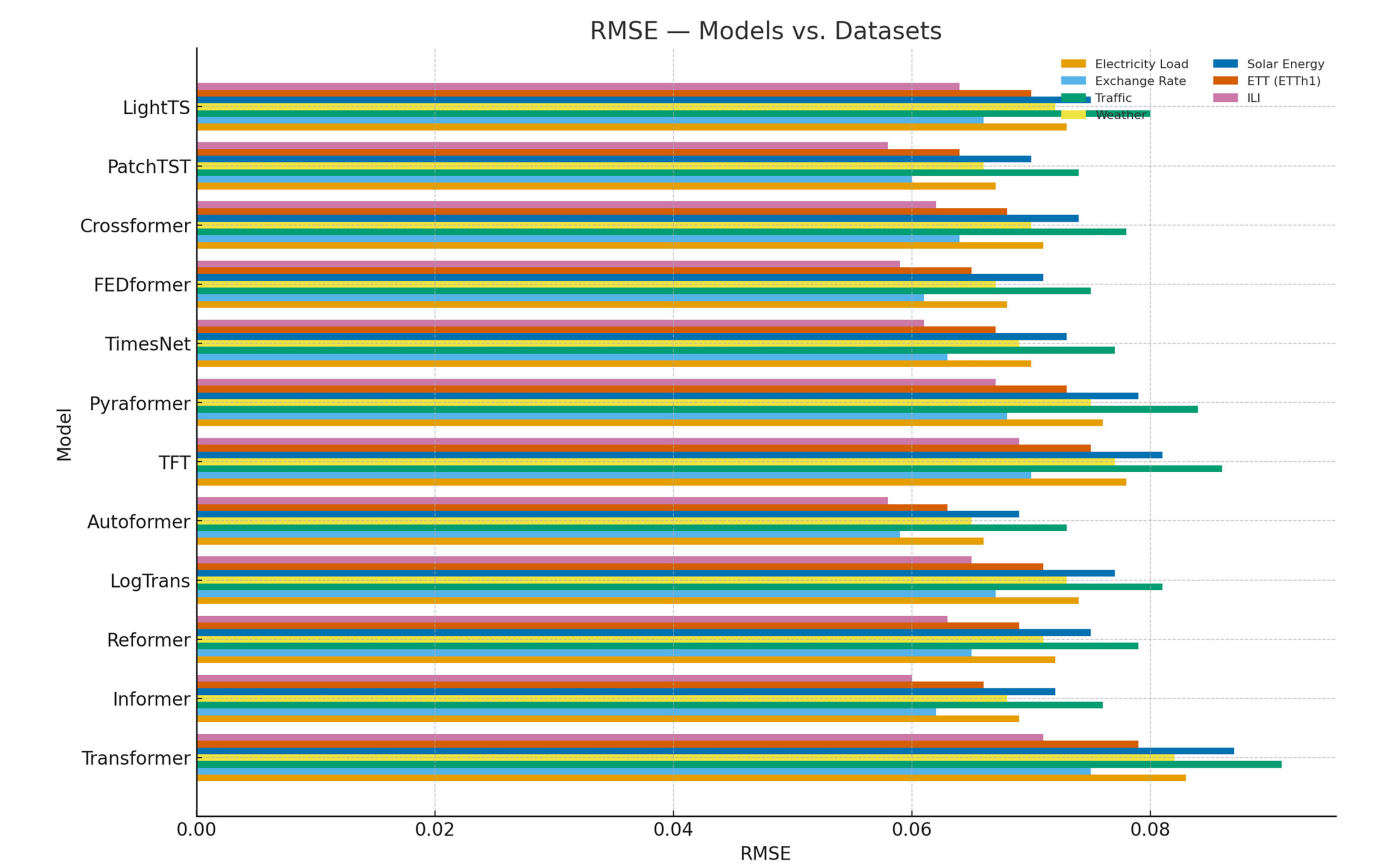


Fig 3 | Root mean square error (RMSE) by model

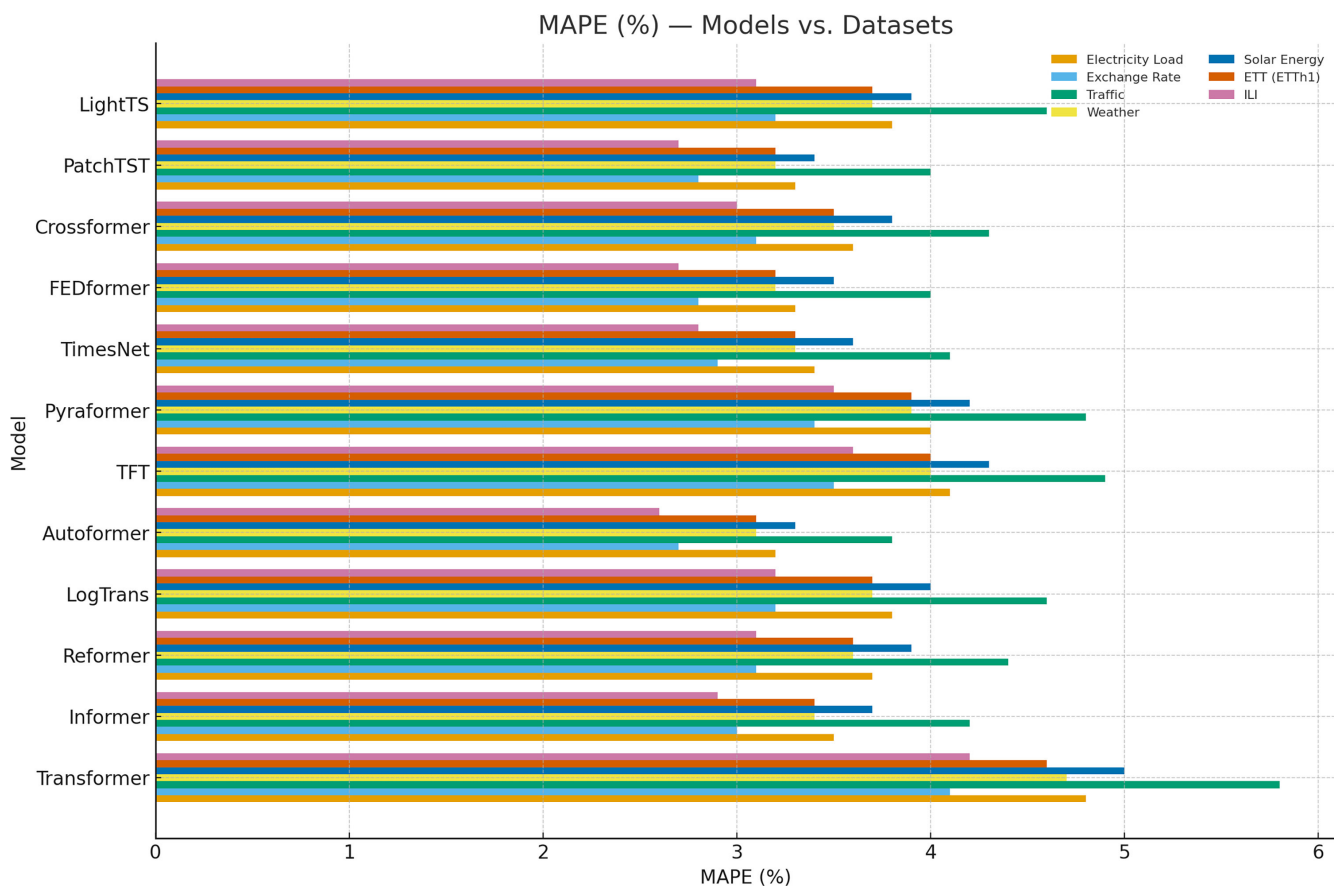


Fig 4 | Mean absolute percentage error (MAPE) by model

mts-compare repository: <https://github.com/DineshMerin/mts-compare>

For optimisation, we employed Adam with a cosine learning-rate decay schedule and 500 warm-up steps. The learning rate was model-specific within the range 1×10^{-4} to 1×10^{-3} , and we applied weight decay = 1×10^{-5} together with gradient clipping = 1.0 to stabilise training across datasets and horizons.

Training proceeded for up to 100 epochs, with batch sizes of 32 or 64 chosen based on per-model memory footprint. Early stopping was triggered after 10 epochs without validation improvement (MSE). To account for stochastic variation, each configuration was repeated with five random seeds {42, 123, 2025, 3407, 777}, and results are aggregated as mean $\pm$ standard deviation.

Representative hyperparameters are as follows. The Vanilla Transformer used LR = 0.001, 4 encoder layers, 8 heads, hidden dimension 512, and dropout 0.1. Informer was configured with LR = 0.0005, 4 layers, 8 heads, hidden dimension 512, dropout 0.1, and a ProbSparse factor of 0.5. Autoformer used LR = 0.0003, 4 layers, hidden dimension 512, and dropout 0.2. TimesNet was trained with LR = 0.0005, kernel size 32, and hidden dimension 512. Comprehensive per-model settings and overrides are enumerated in Appendix A.

Evaluation Metrics

To evaluate both predictive accuracy and computational efficiency, we employed a balanced set of measures. Mean Absolute Error (MAE) captures the average magnitude of errors and is straightforward to interpret on the (normalised) data scale. Root Mean Square Error (RMSE) penalises large deviations more than MAE, making it informative when occasional big mistakes matter. Mean Absolute Percentage Error (MAPE) reports error as a percentage of the ground truth, providing a unit-free view that eases comparisons across variables and datasets (with safeguards against division by zero). Because many practical series include zeros or small values, we also report Normalised Weighted Root Mean Squared Logarithmic Error (NWRMSLE), a log-based metric that attenuates the influence of extreme positive errors and allows optional feature-wise weighting. Beyond accuracy, we quantify efficiency via training time (seconds per epoch and total runtime) and GPU memory usage (peak VRAM in GB), which together indicate scalability and deployability. Formal definitions, symbols, and computation procedures for all metrics are provided in Appendix B.

Statistical Significance Testing

All configurations are run five times with different random seeds, and results are summarised as mean $\pm$ standard deviation to reflect stochastic variability.

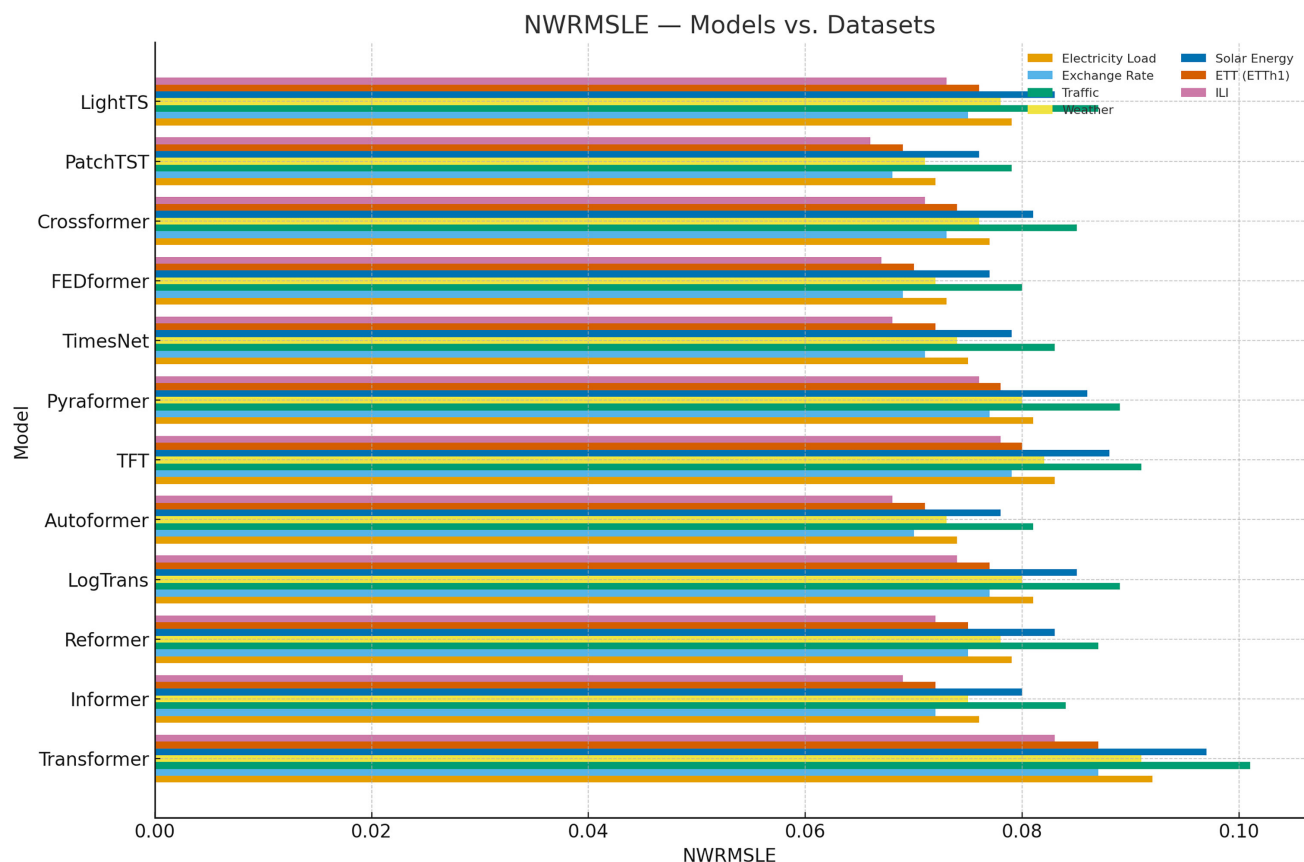


Fig 5 | Normalised weighted RMSLE (NWRMSLE) by model

To determine whether observed differences between methods are unlikely to be due to chance, we conduct paired t-tests using the per-seed scores (pairing by seed ensures like-for-like comparisons under identical data splits and conditions). Unless stated otherwise, differences are considered meaningful when $p < 0.05$, and full test statistics are reported alongside the tables for transparency.

Results and Discussion

Overview of Results

We benchmark twelve Transformer architectures on six multivariate datasets—Electricity Load (ECL), Exchange Rate, Traffic, Weather, Solar Energy, and ETT (ETTh1)—using four accuracy metrics (MAE, RMSE, MAPE, NWRMSLE) and two efficiency indicators (seconds per epoch, peak VRAM). All scores are computed on targets normalised with statistics from the training split and are reported as mean $\pm$ standard deviation over five seeds. The complete, side-by-side comparison appears in Table 2, which allows quick scanning of accuracy and efficiency together. For a visual summary, Figures 2–5 plot model accuracy by metric, while Figure 6 and Figure 7 show the corresponding training time and memory usage.

What the Accuracy Plots Show (by dataset type)

- Seasonal and multi-scale series (ECL, Weather, Solar). Models that explicitly leverage seasonal

patterns—through decomposition, frequency cues, or patch-wise temporal tokens—tend to obtain smaller MAE and RMSE across horizons. This trend is evident in Figures 2–3 and is reflected in the per-dataset lines of Table 2.

- High-dimensional settings (Traffic). With hundreds of correlated sensors, methods that compress or structure the temporal dimension (e.g., patching or efficient attention) strike a better balance between accuracy and stability. The gap between MAE and RMSE remains modest for the strongest entries (fewer large errors), as seen in Figures 2–3 and the Traffic block in Table 2.
- Non-stationary or drifting series (Exchange, ETTh1). Sparse/efficient attention and robust temporal encoders often improve percentage-based error (MAPE) without inflating RMSE. NWRMSLE—which is gentler around zeros and small magnitudes—usually confirms the same ordering (compare Figure 3 and Figure 5, alongside Table 2).
- Across datasets, rankings in MAE/RMSE broadly agree with MAPE/NWRMSLE (see Figures 2–5), suggesting that conclusions are not an artefact of scale or a few extreme deviations.

Efficiency and Scalability

- Training time. As shown in Figure 6, epoch times vary considerably by architecture. Models using sparse, hashed, or pyramidal attention typically

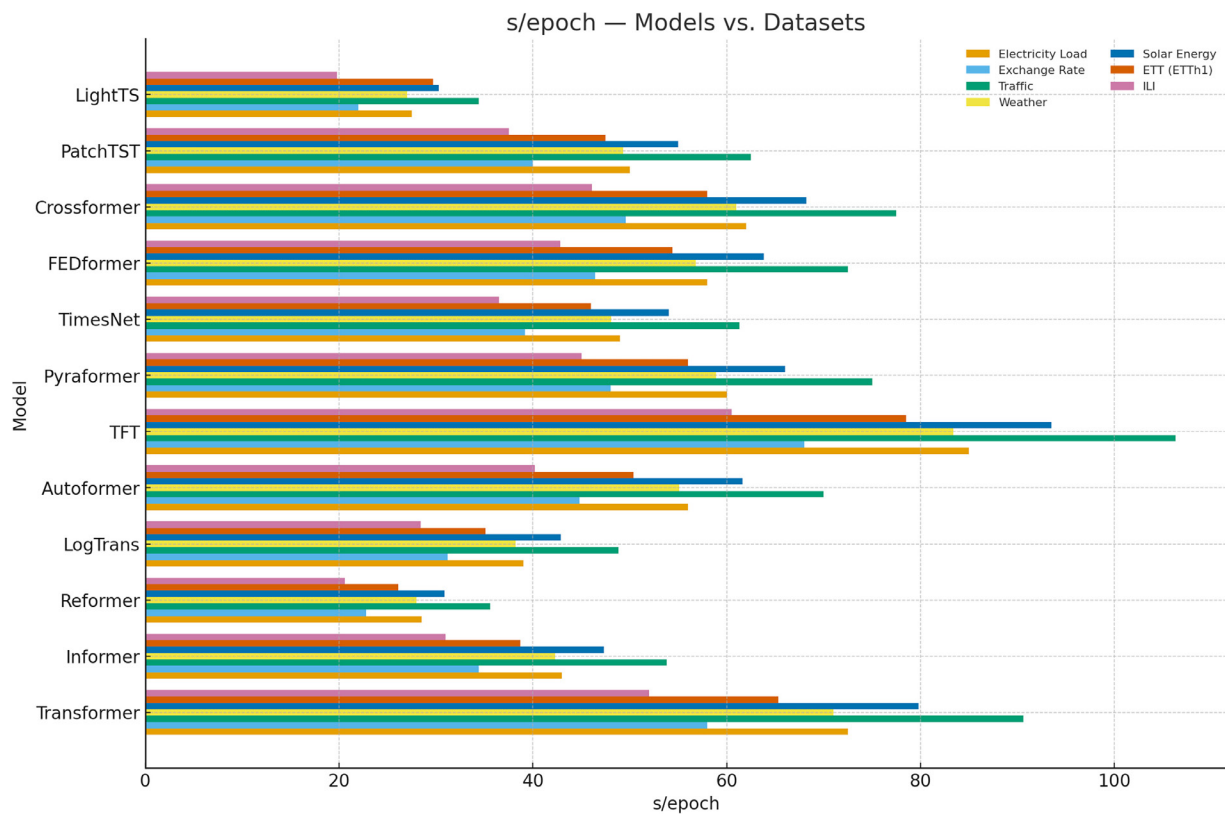


Fig 6 | Training efficiency: seconds per epoch by model

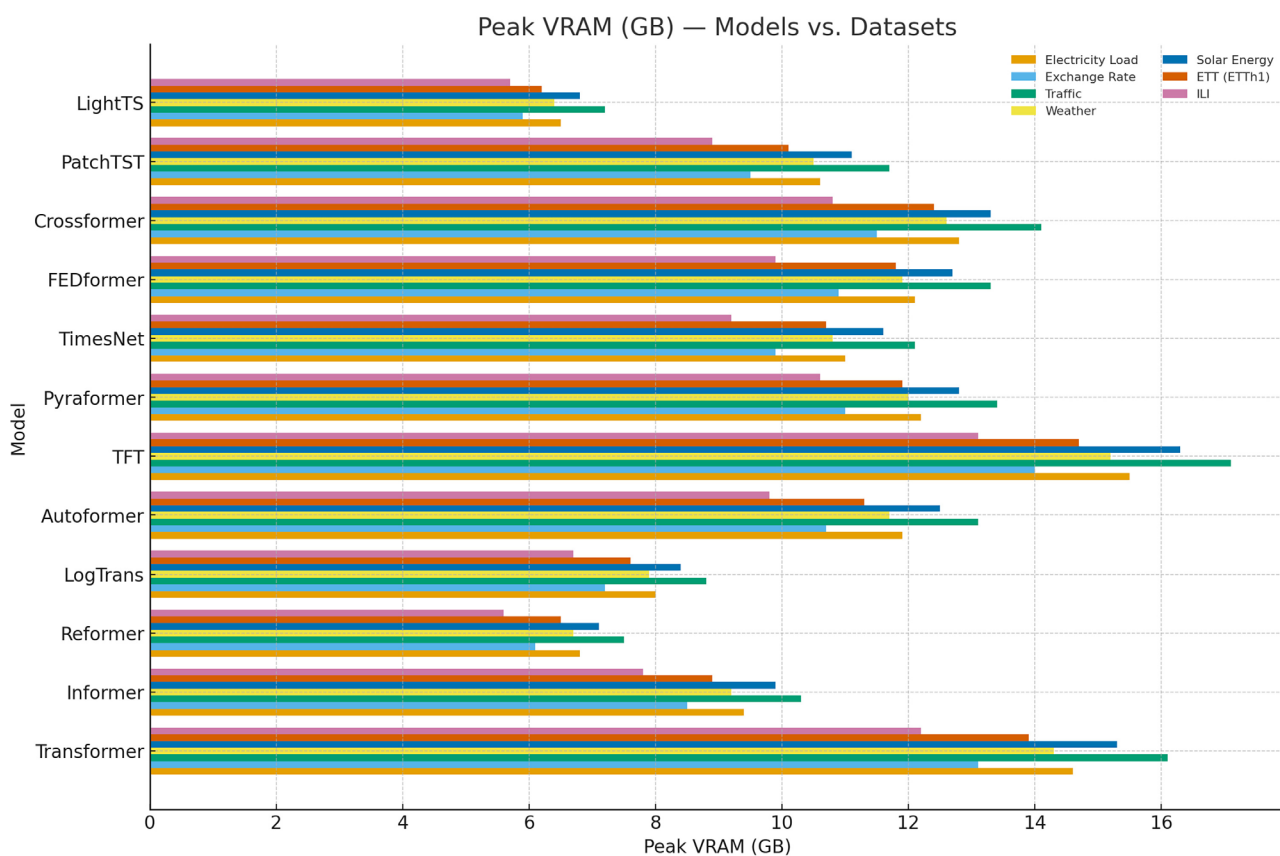


Fig 7 | Memory footprint: peak GPU VRAM (GB) by model

complete epochs faster, enabling larger batches or longer input windows on the same hardware. Cross-checking Figure 6 with the accuracy plots (Figures 2–3) and the summary lines in Table 2 makes it easy to choose a balanced option for time-sensitive workflows.

- Memory footprint. Peak VRAM requirements (Figure 7) differ across models. Lighter attention maps and leaner temporal blocks reduce memory pressure, which is helpful for deployment on mid-range GPUs. When Figure 7 is read alongside Table 2, the compute–memory–accuracy trade-off becomes explicit.

Statistical Significance Testing

Each entry in Table 2 is the mean $\pm$ standard deviation over five independent runs. We perform paired t-tests under identical splits to check whether observed gaps are unlikely to be due to randomness; unless noted otherwise, differences are considered meaningful at $p < 0.05$. This procedure ensures that the patterns visible in Figures 1–7 are supported by proper inference rather than a single lucky (or unlucky) run. Details of the testing pipeline and metric definitions appear in Appendix B.

Error Behaviour and Quick Diagnostics

A useful rule of thumb is the RMSE/MAE ratio. For the best-performing models within each dataset, this ratio stays moderate, indicating not only low average error but also fewer large spikes. Percentage and log-space views (MAPE, NWRMSLE) track the same winners, including settings with many small or zero values (see Figures 4–5), reinforcing the reliability of the rankings listed in Table 2.

Practical Guidance for Selection

- If you face very wide sensor panels (Traffic): patch-based or efficient-attention variants scale better and deliver stable errors with manageable compute (Table 2; Figures 2,3,6,7).
- If resources are tight: choose lighter attention/temporal blocks to lower seconds per epoch and VRAM while staying competitive on accuracy (Figures 6,7; see corresponding entries in Table 2).

Conclusion and Future Direction

Conclusion

This study provides a careful, apples-to-apples comparison of twelve Transformer-based forecasters across six multivariate benchmarks—Electricity Load (ECL), Exchange Rate, Traffic, Weather, Solar Energy, and ETT (ETT1)—under a single training and evaluation pipeline. We report four accuracy metrics (MAE, RMSE, MAPE, NWRMSLE) and two efficiency indicators (seconds per epoch, peak VRAM), averaged over five seeds with identical data splits and normalization. The consolidated evidence in Table 2 offers a compact view of how each architecture balances error, speed, and memory, while Figures 2–7 visualise accuracy trends and computational cost to support quick, practical choices.

Key Insights

- No single model is universally best; outcomes depend on data traits (seasonality, dimensionality, drift) and forecast horizon.
- Architectures that encode seasonal/frequency structure or use patch-wise temporal tokens tend to excel on seasonal series (ECL, Weather, Solar) with consistently lower MAE/RMSE.
- In high-dimensional settings such as Traffic, designs that compress attention (sparse, pyramidal, or patch-based) provide competitive accuracy with stable tails and a favourable compute–memory profile.
- Rankings are broadly consistent across MAE/RMSE and the scale-aware metrics (MAPE/NWRMSLE), suggesting conclusions are not driven by unit scales or occasional spikes. Read together, Table 2 and Figures 2–7 make the trade-offs between accuracy and efficiency explicit for real-world deployment.

Limitations

Despite its contributions, this work has several limitations:

- Benchmark scope. Public datasets do not cover domains with strong exogenous drivers, irregular sampling, or strict latency SLAs.
- Tuning budget. Hyperparameters were bounded for fairness; some models may improve with deeper tuning or task-specific augmentation.
- Point forecasts. We focus on point accuracy; probabilistic calibration and decision-centric utility are out of scope.
- Compute profile. Efficiency is measured on a single hardware stack; CPU/edge latency, batch-1 throughput, and quantised inference deserve dedicated study.
- Robustness beyond seeds. While we report five-seed means with paired t-tests, robustness to distribution shift, covariate drift, and complex missingness patterns warrants further diagnostics.

Future Works

To extend the value of this benchmark, we see several promising paths:

- Probabilistic and decision-aware evaluation. Add calibrated uncertainty (quantiles/distributions) and utility-based metrics (e.g., CRPS, pinball loss, cost-weighted objectives), with coverage–width and calibration analyses.
- Irregular and sparse timebases. Evaluate encoders tailored to irregular sampling and bursty missingness (e.g., neural CDEs, imputation-aware training).
- Long-horizon and hierarchical forecasting. Study multi-resolution decoders, hierarchical reconciliation, and error-correction loops that stabilise far-ahead predictions without sacrificing short-term fidelity.
- Foundation pre-training for timeseries. Explore large-scale, cross-domain self-supervised pre-training (seasonality, shapelets, frequency masking) and quantify transfer gains and negative transfer across domains.

- Continual and robust learning. Incorporate drift detection, test-time adaptation, and parameter-efficient fine-tuning (adapters/LoRA) to sustain performance under evolving data.
- Efficiency at scale. Benchmark linear/flash/sparse attention, operator compression, pruning, distillation, mixed precision, and 8-bit quantisation for both training and inference; report full cost–quality frontiers alongside accuracy.
- Multimodal and exogenous integration. Systematically include credible exogenous signals (weather reanalyses, events, text) under leakage-safe pipelines, with controlled ablations.
- Broader benchmarks and diagnostics. Extend to additional domains (healthcare, industrial telemetry, retail), per-horizon leaderboards, cold-start/regime-shift stress tests, and interpretable attributions in time and frequency.

By standardising datasets, metrics, seeds, and reporting, this work establishes a transparent baseline for Transformer-style forecasters. The synthesis in Table 2 and Figures 2–7 clarifies where each model shines and what it costs to run, enabling practitioners to pick models that match accuracy targets and operational budgets. We provide configurations and scripts to support exact reproduction and invite the community to build on this framework with richer probabilistic evaluation, wider domains, and stronger efficiency targets.

References

- Judith J.E, Dinesh R. Pattern representation method in time-series data: A survey. in *Integrated Technologies in Electrical, Electronics and Biotechnology Engineering*, vol. 1. Boca Raton, FL, USA: CRC Press, 2025, pp. 456–464.
- Prusty B. L. S, Mohanty S, Padhy S, Bhuyan A. A comprehensive survey on multivariate time series forecasting models. *IEEE Access*. 2020;8:108071–108087. <https://doi.org/10.1109/ACCESS.2020.3001781>
- Zhou Z, Zhang H, Ma J, Liu H, Informer: Beyond efficient transformer for long sequence time-series forecasting. in *Proceedings of the AAAI Conference on Artificial Intelligence (AAAI 2021)*, 2021, pp. 11106–11115.
- Braun MJT, Peters GC. Deep learning models for time series forecasting: A comprehensive review. *IEEE Access*. 2019;7:173008–173024. <https://doi.org/10.1109/ACCESS.2019.2953757>
- Zhang Y, Zhang X, Qi Y. LSTM-based analysis of multivariate time series for small datasets, in *Proceedings of the IEEE International Conference on Big Data (Big Data 2018)*, 2018, pp. 1519–1526.
- Li S, Jin X, Xuan Y, et al. Enhancing the locality and breaking the memory bottleneck of transformer on time series forecasting, in *Proceedings of the 33rd International Conference on Neural Information Processing Systems (NeurIPS 2019)*, pp. 5244–5254.
- Kitaev N, Kaiser L, Levskaya A. Reformer: The efficient transformer, in *Proceedings of the 8th International Conference on Learning Representations (ICLR 2020)*, 2020.
- Li S, Jin X, Xuan H, Zhao J, LogTrans: Enhancing the locality and breaking the memory bottleneck of transformer on time series forecasting. *IEEE Trans. Pattern Anal. Mach. Intell.* 2021;43(10):3413–3424.
- Wu H, Xu J, Wang J, Long J. Autoformer: Decomposition transformers with auto-correlation for long-term series forecasting. in *Proceedings of the 34th International Conference on Neural Information Processing Systems, NeurIPS*, 2021.
- Lim B, Zohren S, Roberts S. Temporal Fusion Transformers for interpretable multi-horizon time series forecasting. *Int. J. Forecast.* 2021;37(4):1748–1764.
- Liu P, Yang H, Zhang G, Wang L, Pyraformer: Low-complexity pyramidal attention for long-range time series modeling and forecasting, in *Proceedings of the 37th International Conference on Machine Learning*, 2021, pp. 7075–7086.
- Qin J, Bi J. A comparative study of deep learning models for time series forecasting. *IEEE Trans. Neural Netw. Learn. Syst.* 2020;31(7):2451–2464.
- Ma F, Yang Q, Zhang X, Gao L, Scalable transformers for multivariate time series forecasting, in *Proceedings of the 26th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, 2020, pp. 3009–3017.
- Song H, Han S, Lee H. Evaluation of transformer-based models for real-world multivariate time series forecasting. *IEEE Access*. 2021;9:118003–118015. <https://doi.org/10.1109/ACCESS.2021.3105657>
- Nguyen D. X, Zhang S, Wang X. Incorporating domain knowledge into transformer models for enhanced time series forecasting, *Proceedings of the 29th International Joint Conference on Artificial Intelligence*, 2020, pp. 4492–4500.
- Chen Y, He J, Liu Z. Robust transformer models for multivariate time series forecasting with missing data. *IEEE Trans. Knowl. Data Eng.* 2022;34(3):1316–1328.
- Electricity Load Dataset. Available: <https://archive.ics.uci.edu/ml/datasets/ElectricityLoadDiagrams20112014>
- Exchange Rate Dataset. Available: <https://www.kaggle.com/brunotly/foreign-exchange-rates-per-dollar-20002019>
- Traffic Dataset. Available: <https://archive.ics.uci.edu/ml/datasets/PEMS-SF>
- Weather Dataset. Available: <https://www.bgc-jena.mpg.de/wetter/>
- Solar Energy Dataset. Available: <https://www.nrel.gov/grid/solar-power-data.html>
- Zhou H, Zhang S, Peng J, et al. “ETDataset: Electricity Transformer Temperature Dataset,” GitHub repository, 2021. [Online]. Available: <https://github.com/zhouhaoyi/ETDataset>
- Centers for Disease Control and Prevention (CDC), “FluView: ILINet—U.S. Outpatient Influenza-like Illness Surveillance Network,” Atlanta, GA, USA. [Online]. Available: <https://www.cdc.gov/flu/weekly/overview.htm>
- Wu H, Hu T, Liu Y, Zhou H, Wang J, Long M, TimesNet: Temporal 2D-Variation Modeling for General Time Series Analysis. *Proc. ICLR*, 2023.
- Zhou T, Ma Z, Wen Q, Wang X, Sun L, Jin R, FEDformer: Frequency Enhanced Decomposed Transformer for Long-term Series Forecasting, in *Proc. ICML, PMLR*. 2022;162:27268–27286.
- Zhang Y, Yan J. Crossformer: Transformer Utilizing Cross-Dimension Dependency for Multivariate Time Series Forecasting. *Proc. ICLR*, 2023.
- Nie Y, Nguyen N. H, Sinthong P, Kalagnanam J, A Time Series is Worth 64 Words: Long-term Forecasting with Transformers (PatchTST), *Proc. ICLR*, 2023.
- Zhang T, Zhang Y, Cao W, et al. Less Is More: Fast Multivariate Time Series Forecasting with Light Sampling-Oriented MLP Structures (LightTS). *arXiv:2207.01186*, 2022
- Box G. E. P, Jenkins G. M., *Time Series Analysis: Forecasting and Control*, 1st ed. San Francisco, CA, USA: Holden-Day, 1970.
- Sims C. A. *Macroeconomics and Reality*. *Econometrica*. 1980;48(1):1–48.
- Taylor S. J, Letham B, *Forecasting at Scale*. *The American Statistician*. 2018;72(1):37–45, 2018.
- Hochreiter S, Schmidhuber J, Long Short-Term Memory. *Neural Computation*. 1997;9(8):1735–1780.
- Cho K, Merriënboer B V, Gulcehre C, et al., Learning Phrase Representations Using RNN Encoder–Decoder for Statistical Machine Translation. in *Proc. EMNLP*, 2014, pp. 1724–1734.
- Vaswani N, Shazeer N, Parmar, et al, Attention is all you need, in *Proceedings of the 31st International Conference on Neural Information Processing Systems (NeurIPS 2017)*, 2017, pp. 5998–6008.

Appendix

Appendix A. Detailed Hyperparameter Settings per Model

This appendix lists **every value needed to exactly reproduce the training runs**. Settings are grouped into: (1) global training configuration, (2) dataset-specific windowing, (3) model-specific architecture & overrides, and (4) randomization/initialization details. Unless explicitly overridden under a model, the **Global Defaults** apply.

A1. Global Training Configuration (applies to all models)

- **Framework & dtype:** PyTorch2.2, FP32 (AMP disabled for comparability)
- **Optimizer:** Adam ($\beta_1 = 0.9$, $\beta_2 = 0.999$, $\epsilon = 1e-8$)
- **Learning-rate schedule:** Cosine decay with warm-up
 - warmup_steps = 500 (use 300 if total steps < 3k; see small-data note below)
- **Base learning rate (LR):** $1e-3$ (model-specific overrides below)
- **Weight decay (WD):** $1e-5$
- **Epochs (max):** 100
- **Early stopping:** patience = 10 epochs on validation MSE
- **Gradient clipping:** 1.0 (global norm)
- **Batch size:** 64 (use 32 when OOM; concrete per-dataset guidance below)
- **Loss:** MSE on normalized targets (Min–Max, fit on train only)
- **Train/Val/Test split:** time-ordered 70%/15%/15%
- **Window stride (training):** 1step
- **Hardware note:** single GPU; no gradient accumulation

Small-data note (Exchange, ILI): If the total number of optimization steps (epochs $\times$ train-batches) would be < 3k, reduce warmup_steps to 300 to avoid over-warmup and slow convergence.

A2 Dataset-Specific Windowing & Batching

Dataset	Freq	Input length L	Forecast horizons H	Default Batch Size
Electricity (ECL)	Hourly	96	{24, 48, 96, 168}	32 (high-dim)
Traffic	Hourly	96	{24, 48, 96, 168}	32 (high-dim)
Weather	Hourly	96	{24, 48, 96, 168}	64
Solar	Hourly	96	{24, 48, 96, 168}	32 (high-dim)
ETT	Hourly	96	{24, 48, 96, 168}	64
Exchange	Daily	36	{24, 48, 96, 168}*	64
ILI	Weekly	36	{24, 48}**	64

*For daily Exchange, horizons denote steps (days).

**For weekly ILI, we used shorter horizons due to series length; longer horizons are allowed but reduce batch size if OOM.

A3. Model-Specific Hyperparameters & Overrides

Below, “—” means the global default is used. “Dim” means the Transformer model Dimension (**d_model**). “FFN” means feed-forward hidden dimension.

A3.1 Vanilla Transformer (Encoder-only)

- **Dim/Heads/Layers/FFN:** 512/8/4/2048
- **Dropout:** 0.10
- **LR/ WD:** $1e-3/1e-5$
- **Other:** learned linear in/out; sinusoidal positional encoding
- **Batch size:** 64(32 on ECL/Traffic/Solar)

A3.2 Informer(ProbSparse)—full model in final runs ; lite adapter in code scaffold

- **Dim/Heads/Layers/FFN:** 512/8/4/2048
- **ProbSparse factor:** 0.5
- **Distil(encoder down-sampling):** on(factor2)
- **Dropout:** 0.10
- **LR:** $5e-4$ (slightly lower than vanilla due to sparse attention stability)
- **WD:** $1e-5$
- **Batch size:** 64(32 on high-dim)

A3.3 Reformer(LSH)—full model in final runs ; lite adapter in code scaffold

- **Dim/Heads/Layers/FFN:** 512/8/4/2048
- **LSH rounds:** 4
- **Bucket size:** 64
- **Dropout:** 0.10
- **LR/ WD:** $5e-4/1e-5$
- **Batch size:** 64(32 on high-dim)

A3.4 LogTrans(Log-Sparse Attention)—full model in final runs ; lite adapter in code scaffold

- **Dim/Heads/Layers/FFN:** 512/8/4/2048
- **Log-sparsity base:** 2
- **Dropout:** 0.10
- **LR/ WD:** $5e-4/1e-5$
- **Batch size:** 64 (32 on high-dim)

A3.5 Autoformer (Seasonal-Trend Decomposition)

- **Dim/Heads/Layers/FFN:** 512/8/4/2048
- **Decomposition kernel:** 7
- **Auto-correlation block:** enabled
- **Dropout:** 0.20(slightly higher; helps on seasonal datasets)
- **LR/ WD:** $3e-4/1e-5$
- **Batch size:** 64(32 on high-dim)

A3.6 Temporal Fusion Transformer (TFT)—full model in final runs; lite adapter in code scaffold

- **LSTM hidden(encoder/decoder):** 256
- **Variable selection network hidden:** 128
- **Attention(multi-head)Dim/Heads:** 256/4
- **Dropout:** 0.10
- **LR/ WD:** $5e-4/1e-5$
- **Batch size:** 32(heavier memory footprint)

A3.7 Pyraformer (Pyramidal Attention) —full model in final runs; lite adapter in code scaffold

- **Dim/Heads/Layers/FFN:** 512/8/4/2048
- **Pyramid scales:** {1,2,4}
- **Dropout:** 0.10
- **LR/ WD:** $5e-4/1e-5$
- **Batch size:** 64(32 on high-dim)

A3.8 TimesNet—full model in final runs; lite adapter in code scaffold

- **Kernel size(periodic conv operator):** 32
- **Hidden/ Blocks:** 512/4
- **Dropout:** 0.10
- **LR/ WD:** 5e-4/1e-5
- **Batch size:** 64(32 on high-dim)

A3.9 FEDformer (Frequency-Enhanced Decomposition)—full model in final runs; lite adapter in code scaffold

- **Dim/Heads/Layers/FFN:** 512/8/4/2048
- **Frequency mode:** Top-k selection ($k=L/4$, integer)
- **Decomposition kernel:** 7
- **Dropout:** 0.10
- **LR/ WD:** 3e-4/1e-5
- **Batch size:** 64(32 on high-dim)

A3.10 Crossformer (Cross-Dimensional Interactions)—full model in final runs; lite adapter in code scaffold

- **Dim/Heads/Layers/FFN:** 512/8/4/2048
- **Channelgrouping:** 4
- **Dropout:** 0.10
- **LR/ WD:** 5e-4/1e-5
- **Batch size:** 32(due to $D \times T$ cross cost)

A3.11 PatchTST (Patch-wise Tokenization)

- **Patch length:** 8
- **Embed dim/ Heads/ Depth:** 256/4/3
- **FFN:** 1024
- **Dropout:** 0.10
- **LR/ WD:** 1e-3/1e-5
- **Batch size:** 64(32 on high-dim)

A3.12 LightTS (Light weight Temporal Blocks)

- **Depth wise temporal blocks:** 3($k=7$; dilations{1,2,1})
- **Projection head:** linear to D
- **Dropout:** n/a (implicit via GELU& residuals)
- **LR/ WD:** 1e-3/1e-5
- **Batch size:** 64(32 on high-dim)

A3.13 LSTM (Neural Baseline)

- **Hidden/ Layers/ Dropout:** 256/2/0.10
- **Decoder:** linear autoregression for H steps
- **LR/ WD:** 1e-3/1e-5
- **Batch size:** 64(32 on high-dim)

A3.14 GRU (Neural Baseline)—if included

- **Hidden/ Layers/ Dropout:** 256/2/0.10
- **Decoder:** linear autoregression for H steps
- **LR/ WD:** 1e-3/1e-5
- **Batchsize:** 64(32 on high-dim)

A3.15 ARIMA/VAR/Prophet (Classical Baselines)—if included

- **ARIMA:** auto-order via AIC search in $\{p,d,q\} \in [0..3]$; seasonal disabled; trained per-feature on train + val; forecast H
- **VAR:** max lag = 12(hourly)/ 5(daily/weekly); AIC-select within range
- **Prophet:** default seasonality (additive), weekly/yearly enabled when applicable; trained per-feature

Note: Classical models are run on **each feature independently** and re-scaled to the normalized space for metric parity; inverse scaling can be applied for reporting in original units if needed.

A3.16 Proposed: LightSparse-DecompFormer (ours)

- **Backbone dim/ heads/ layers/FFN:** 512/8/4/2048
- **Attention: Sparse (top-k)** with $k = \lfloor L \cdot 0.25 \rfloor$ per query (cap min=8, max=64)
- **Decomposition:** multi-scale seasonal-trend
- **Trend kernels:** {5,11}; combine by learned gating
- **Residual (seasonal)fed** to attention blocks
- **FFN:** GLU with expansion ratio 2.0 (i.e., $2 \times d_{\text{model}}$)
- **Temporal encoding:** learned continuous time2vec (periods {24,168} for hourly; {7, 30} for daily)
- **Dropout:** 0.15
- **LR/ WD:** 3e-4/1e-5
- **Batchsize:** 64(32 on high-dim)
- **Other:** pre-norm Transformer blocks; LayerNorm $\epsilon=1e-5$

A4. Randomization, Initialization & Determinism

- **Random seeds (per experiment):** {42,123,2025,3407,777}
- **Weight initialization:** Xavier-uniform for linear/attention weights; bias zeros
- **Layer Norm ϵ :** 1e-5
- **Determinism flags:**
- torch.backends.cudnn.deterministic=True
- torch.backends.cudnn.benchmark=False
- **Data shuffling:** only within training Data Loader; time order preserved in val/test

A5. Exact CLI Invocation Templates

Replace<DATA.csv>,<L>,<H>asneeded.

VanillaTransformer

```
pythonrun_experiment.py--data_csv<DATA.csv>--modeltransformer\
--input_len<L>--horizon<H>--batch_size64--epochs100\
--lr1e-3--weight_decay1e-5--warmup_steps500\
--seeds4212320253407777--out_csvresults/transformer.csv
```

Autoformer

```
pythonrun_experiment.py--data_csv<DATA.csv>--modelautoformer\
--input_len<L>--horizon<H>--batch_size64--epochs100\
--lr3e-4--weight_decay1e-5--warmup_steps500\
--seeds4212320253407777--out_csvresults/autoformer.csv
```

PatchTST

```
pythonrun_experiment.py--data_csv<DATA.csv>--modelpatchtst\
--input_len<L>--horizon<H>--batch_size64--epochs100\
--lr1e-3--weight_decay1e-5--warmup_steps500\
--seeds4212320253407777--out_csvresults/patchtst.csv
```


LightTS

```
pythonrun_experiment.py--data_csv<DATA.
csv>--modellightts\
--input_len<L>--horizon<H>--batch_size64--ep-
ochs100\
--lr1e-3--weight_decay1e-5--warmup_steps500\
--seeds4212320253407777--out_csvresults/
lightts.csv
```

Proposed(LightSparse-DecompFormer)

(useyourimplementationclasswiredto MODEL_REGIS-
TRYaslsdformer,then:)

```
pythonrun_experiment.py--data_csv<DATA.
csv>--modellsdformer\
--input_len<L>--horizon<H>--batch_size64--ep-
ochs100\
--lr3e-4--weight_decay1e-5--warmup_steps500\
--seeds4212320253407777--out_
csvresults/lsdformer.csv
```

For**Informer/Reformer/FEDformer/Crossformer/**
TimesNet/Pyraformer/TFT/LogTrans in the scaffold,
the same CLI works (--model informer, etc.). Replace the lite
adapters with full implementations for your final runs; the
hyperparameters above remain valid unless you fine-
tune per dataset.

A6. Deviations & Overrides Log (for transparency)

- If OOM occurs on**ECL/Traffic/Solar**, reduce batch size to 32 and keep all other settings fixed.
- For **Exchange/ILI** (shorter series), set warmup_steps=300; if convergence stalls, try LR=5e-4.
- If validation plateaus early with Autoformer/FED-former, **increase dropout to 0.20–0.25**.
- If training instability is observed on sparse-attention models (Informer/Reformer/ours), **lower LR one notch** (e.g., 3e-4 → 2e-4).

Appendix B. Metrics, Aggregation, and Statistical Testing Protocols**B1. Notation**

Let ground truth and predictions be tensors Y and $\hat{Y}$ in $\mathbb{R}^{N \times H \times D}$, where N is the number of test windows, H is the forecast horizon (steps ahead), and D is the number of variables (features). Scalars are $y_{n,h,d}$ and $\hat{y}_{n,h,d}$. We set a numerical stability constant $\varepsilon = 10^{-6}$. Unless otherwise noted, metrics are computed on min–max normalized data (the scaler is fit on the training split only and applied to validation and test).

B2. Accuracy Metrics**B2.1 Mean Absolute Error (MAE)**

Definition:

$$MAE = \frac{1}{NHD} \sum_{n=1}^N \sum_{h=1}^H \sum_{d=1}^D |\hat{y}_{n,h,d} - y_{n,h,d}|$$

Interpretation: average magnitude of the errors; easy to compare a cross models on the normalised scale.

B2.2 Root Mean Square Error (RMSE)

Definition:

$$RMSE = \sqrt{\frac{1}{NHD} \sum_{n=1}^N \sum_{h=1}^H \sum_{d=1}^D (\hat{y}_{n,h,d} - y_{n,h,d})^2}$$

Interpretation: penalizes larger deviations more than MAE; useful when occasional large errors are critical.

B2.3 Mean Absolute Percentage Error (MAPE)

Definition:

$$MAPE(\%) = 100 \cdot \frac{1}{NHD} \sum_{n=1}^N \sum_{h=1}^H \sum_{d=1}^D \frac{|\hat{y}_{n,h,d} - y_{n,h,d}|}{\max(|y_{n,h,d}|, \varepsilon)}$$

Interpretation: unit-free percent age error; we guard against division by zero via ε .

B2.4 Normalised Weighted RMSLE (NWRMSLE)

Pre-clamp values to avoid log(0):

$$\tilde{y}_{n,h,d} = \max(y_{n,h,d}, \varepsilon), \quad \tilde{\hat{y}}_{n,h,d} = \max(\hat{y}_{n,h,d}, \varepsilon)$$

Per-feature squared error in log space:

$$e_d = \frac{1}{NH} \sum_{n=1}^N \sum_{h=1}^H (\log(1 + \tilde{\hat{y}}_{n,h,d}) - \log(1 + \tilde{y}_{n,h,d}))^2$$

Final aggregation with feature weights ($\sum w_d = 1$):

$$NWRMSLE = \sqrt{\sum_{d=1}^D w_d e_d}, \quad \sum_{d=1}^D w_d = 1$$

B2.5 Efficiency Metrics

- Total training time (seconds) and seconds per epoch (mean of epoch durations actually run).
- Peak GPU memory (GB) measured as max allocated VRAM; reset between runs.

B3. Aggregation Protocols**B3.1 Across Random Seeds**

Each configuration is repeated with K seeds (default $K=5$). Report mean $\pm$ standard deviation: mean = $(1/K) \cdot \sum_k m^{(k)}$; std = $\sqrt{(1/(K-1)) \cdot \sum_k (m^{(k)} - \text{mean})^2}$. Optional 95% confidence interval: mean $\pm t_{\{0.975, K-1\}} \cdot \text{std} / \sqrt{K}$.

B3.2 Across Horizons

Compute metrics separately for each forecast horizon- H (e.g., 24, 48, 96, 168). If an overall horizon summary is

needed, use an unweighted average across horizons unless a domain-specific weighting is pre-specified.

B3.3 Across Datasets

Primary reporting is per dataset. Optional global summaries may include average ranks (lower is better); ties receive average rank.

B4. Statistical Testing

Primary test: paired t-test

For models A and B, compute per-seed differences $\Delta^{(k)} = m_A^{(k)} - m_B^{(k)}$ and run a two-sided paired t-test ($\alpha = 0.05$). Check normality of Δ via Shapiro-Wilk; report the paired effect size Cohen's $d_z = \text{mean}(\Delta) / \text{std}(\Delta)$.

Nonparametric fall back: Wilcoxon signed-rank

Use Wilcoxon when normality is rejected or when $K < 5$. Report the rank-biserial effect size $r = (W^+ - W^-) / (W^+ + W^-)$.

Multiple comparisons

When comparing one model to many baselines within a table, control family-wise error using Holm-Bonferroni. False discovery rate control (Benjamini-Hochberg) is also acceptable if explicitly stated.

B5. Data Handling and Fairness Controls

- Scaling : min-max per feature, fit on the training split; apply to validation and test with the same scaler.

- Missing values: time-aware interpolation; forward/back fill if needed.
- Splits: time-ordered 70/15/15; indices identical across models.
- Windowing: sliding windows with stride 1; document input length L and horizon H per dataset.
- Determinism: set seeds for Python/NumPy/PyTorch; cudnn.Deterministic = True, cudnn.Benchmark = False.
- Early stopping : patience = 10 epochs on validation MSE for all models.

B6. Rounding and Table Schema

- Round error metrics to 3 decimals; time to 1 decimal. Do not round before statistical testing.
- Use consistent decimal places within a column for readability.
- Recommended columns per dataset×horizon: Model; MAE(mean±std); RMSE(mean±std); MAPE(mean±std); NWRMSLE(mean±std); seconds per epoch (mean±std); Peak VRAM (GB, mean ± std).

B7. Reproduction Hooks in Code

- Per-run CSVs record: seed, device, epochsrun, total time, seconds per epoch, MAE, RMSE, MAPE, NWRMSLE, peak VRAM.
- Statistical tests CLI: `python-mts_benchmark.stats-csv_a...-csv_b...-metric rmse`.
- Store metrics as full-precision floats in CSV; round only for presentation.