
SWIFTTS: A SWIFT SELECTION FRAMEWORK FOR TIME SERIES PRE-TRAINED MODELS VIA MULTI-TASK META-LEARNING

Tengxue Zhang¹, Biao Ouyang¹, Yang Shu^{1*}, Xinyang Chen^{2*}, Chenjuan Guo¹, Bin Yang¹

¹East China Normal University, ²Harbin Institute of Technology (Shenzhen)

{txzhang, bouyang}@stu.ecnu.edu.cn

{yshu, cjguo, byang}@dase.ecnu.edu.cn

chenxinyang@hit.edu.cn

ABSTRACT

Pre-trained models exhibit strong generalization to various downstream tasks. However, given the numerous models available in the model hub, identifying the most suitable one by individually fine-tuning is time-consuming. In this paper, we propose **SwiftTS**, a swift selection framework for time series pre-trained models. To avoid expensive forward propagation through all candidates, SwiftTS adopts a learning-guided approach that leverages historical dataset-model performance pairs across diverse horizons to predict model performance on unseen datasets. It employs a lightweight dual-encoder architecture that embeds time series and candidate models with rich characteristics, computing patchwise compatibility scores between data and model embeddings for efficient selection. To further enhance the generalization across datasets and horizons, we introduce a horizon-adaptive expert composition module that dynamically adjusts expert weights, and the transferable cross-task learning with cross-dataset and cross-horizon task sampling to enhance out-of-distribution (OOD) robustness. Extensive experiments on 14 downstream datasets and 8 pre-trained models demonstrate that SwiftTS achieves state-of-the-art performance in time series pre-trained model selection.

1 INTRODUCTION

Time series forecasting (Wu et al., 2023; Nie et al., 2023; Chen et al., 2024) is a fundamental task with broad applications in finance, weather prediction, and energy management. Inspired by the success of pre-trained models in natural language processing (Hurst et al., 2024; Yang et al., 2025) and computer vision (Dosovitskiy et al., 2021), numerous time series foundation models have been developed (Das et al., 2024). Pre-trained on large and diverse datasets, these models acquire transferable knowledge that can be adapted to downstream tasks through fine-tuning, eliminating the need for training from scratch (Kumar et al., 2022; Jia et al., 2022; Dettmers et al., 2023).

However, no single pre-trained model excels across all tasks (Li et al., 2025), making model selection for time series forecasting challenging. While fine-tuning all candidates provides ground-truth performance, it is often computationally infeasible for large model pools. Therefore, developing efficient methods to identify the optimal pre-trained model is crucial for real-world deployment. Existing approaches, primarily designed for image models, are feature-analytic methods that analyze features extracted from the target dataset using pre-trained models: some assess feature-task alignment via statistical metrics (Nguyen et al., 2020), while others investigate intrinsic properties of the feature space (Pándy et al., 2022; Wang et al., 2023). Only a few are learning-based (Zhang et al., 2023), learning similarity functions between data and model representations for model selection. Despite recent advances, several challenges remain unresolved for time series pre-trained models:

Challenge 1: Oversight of model heterogeneity and time series data characteristics. Current time series pre-trained models are typically heterogeneous in both architecture and training objectives, unlike the standardized feature extractors common in vision models. This diversity hinders

*Corresponding authors

unified feature extraction and limits the applicability of many existing methods. Moreover, extracting features also requires costly forward passes through each model, leading to substantial computational overhead as the model hub or datasets scale. Some learning-based methods attempt to mitigate this cost via a shared feature extractor, but often sacrifice performance. In addition, time series data exhibit temporal dependencies and sequential patterns that are critical for accurate forecasting but largely ignored in current approaches. Valuable prior knowledge and intuitive insights, such as “models generally perform better when the downstream domain aligns with the pre-training domain,” are also rarely incorporated into current model selection criteria.

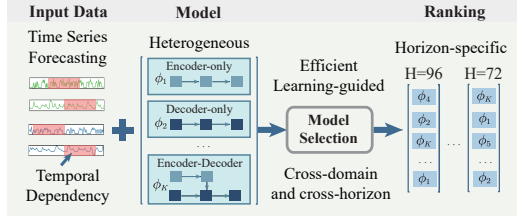


Figure 1: SwiftTS employs an efficient learning-guided selection framework for time series forecasting, enabling horizon-specific selection and improved cross-domain generalization via multi-task meta-learning.

tions may degrade on longer horizons. Ignoring such variability can result in misleading performance rankings that fail to reflect the dynamic, horizon-dependent nature of time series forecasting. Current methods often lack mechanisms to effectively select models tailored to specific forecasting horizons, as in Figure 1, thereby limiting their generalization across varying forecasting ranges.

To address **Challenge 1**, we propose an efficient learning-guided selection framework that avoids inconsistent feature extraction and reduces forward propagation costs. We collect the performance of various dataset-model pairs across different horizons as training data to learn how models perform on unseen datasets. The framework is tailored to time series characteristics and incorporates prior knowledge to enhance selection. The data encoder segments the input series into patches and generates patch-level embeddings that capture local temporal patterns and preserve sequential dependencies. The model encoder incorporates meta-information, topological structure, and functionality to construct a comprehensive representation of candidate models. By employing a lightweight dual-encoder architecture, our framework independently learns informative embeddings for both downstream datasets and candidate models. Finally, patch-level cross-attention assesses the fine-grained compatibility score between them, enabling accurate and efficient model selection.

To address **Challenge 2**, we propose a generalizable multi-task meta-learning strategy. To equip our framework with the multi-task flexibility to accommodate varying horizons, we incorporate a horizon-adaptive expert composition. This design dynamically assigns adaptive weights to multiple experts based on the target forecasting horizon, enabling horizon-specific ranking predictions. To further enhance generalization across both domains and horizons, we propose transferable cross-task learning to improve robustness in out-of-distribution (OOD) scenarios. We introduce a meta-learning paradigm with two task sampling strategies: (1) cross-dataset sampling, where tasks are drawn from different datasets to encourage inter-domain generalization, and (2) cross-horizon sampling, where tasks are constructed using varying forecasting horizons to improve horizon-level adaptability. By meta-learning from tasks across diverse datasets and horizons, our framework learns transferable knowledge that captures both task-specific characteristics and shared cross-task patterns. This ultimately improves its performance in real-world applications.

To the best of our knowledge, this is the first model selection method for time series pre-trained models. The contributions are summarized as follows:

- We propose a swift learning-guided framework that leverages a dual-encoder to embed datasets and models, computing patchwise compatibility scores for model selection.
- We introduce a multi-task meta-learning strategy with a horizon-adaptive expert composition to enhance generalization across datasets and forecasting horizons.

- Extensive experiments on benchmarks comprising 14 real-world downstream datasets and 8 pre-trained models show that our method achieves state-of-the-art performance in pre-trained model selection for time series forecasting.

2 RELATED WORK

Time series Pre-trained Model. Existing models can be broadly categorized into three main architectures: (1) *Encoder-only models*: MOIRAI (Woo et al., 2024) flattens multivariate sequences into unified sequences for Transformer pre-training, Moment (Goswami et al., 2024) employs masked reconstruction to train a versatile Transformer, and UniTS (Gao et al., 2024) introduces task tokenization and dynamic self-attention across temporal and variable dimensions. (2) *Decoder-only models*: TimesFM Das et al. (2024) and Timer Liu et al. (2024) adopt GPT-style designs for next-token prediction, achieving strong zero-shot performance. (3) *Encoder-decoder models*: TTM Ekambaram et al. (2024) leverages MLP-Mixer blocks with multi-resolution sampling to capture cross-channel patterns. ROSE (Wang et al., 2024) enhances generalization through decomposed frequency learning and time series register components. Chronos (Ansari et al., 2024) adapts the T5 (Raffel et al., 2020) language foundation model to time series by discretizing data via binning and scaling. As the number and variety of time series pre-trained models continue to grow, efficiently and accurately selecting the most suitable model from a diverse model hub remains a challenge.

Pre-trained Model Selection. Pre-trained model selection aims to quickly identify the best model for downstream tasks from a model hub. Existing methods fall into two broad categories: (1) *Feature-analytic methods* analyze features extracted by pre-trained models from the target dataset. Early approaches, such as NCE (Tran et al., 2019) and LEEP (Nguyen et al., 2020), leverage statistical metrics but depend on pre-trained classifiers, limiting their applicability to self-supervised models. LogME (You et al., 2021) overcomes this by estimating the maximum label marginalized likelihood. RankME (Garrido et al., 2023) posits that models with higher feature matrix ranks exhibit superior transferability. Other methods focus on class separability during the fine-tuning process. GBC (Pándy et al., 2022) measures the degree of overlap between pairwise target classes based on extracted features. SFDA (Shao et al., 2022) enhances class separability by projecting features into a Fisher space. Etran (Gholami et al., 2023) introduces an energy-based transferability metric, while DISCO (Zhang et al., 2025) proposes a framework for evaluating pre-trained models based on the distribution of spectral components. (2) *Learning-based methods* aim to predict model transferability through a learning framework. Model Spider (Zhang et al., 2023) learns model representations and a similarity function by aligning them with downstream task representations, enabling model selection via the learned similarity. Despite the diversity of existing methods, they often rely on costly feature extraction and generalize poorly across domains and forecasting horizons. To address these issues, we propose SwiftTS, a swift model selection framework via multi-task meta-learning. It infuses prior knowledge and adopts a learning-guided paradigm, avoiding expensive feature analysis used in prior work while improving OOD robustness.

3 PROBLEM FORMULATION

Given a model hub $Z = \{\phi_k\}_{k=1}^K$ of K time series pre-trained models and a target dataset $D = \{x_i, y_i\}_{i=1}^N$ with N samples, the goal is to select the model that achieves the best performance on the time series forecasting task. Brute-force fine-tuning of all models yields the ground-truth performances $\{r_k\}_{k=1}^K$ for the model hub, but at prohibitive computational cost. To avoid this, model selection methods estimate transferability without fine-tuning by assigning each model ϕ_k an assessment score $\hat{r}_k$, where a larger $\hat{r}_k$ indicates stronger expected performance:

$$\hat{r}_k = f(\phi_k, D, H) \quad (1)$$

The f is a scoring function that measures the compatibility between the model ϕ_k and the target dataset D under the forecasting horizon H . Ideally, the predicted scores $\{\hat{r}_k\}_{k=1}^K$ should strongly correlate with the fine-tuning results $\{r_k\}_{k=1}^K$, enabling the selection of the most transferable model.

4 METHODS

We propose an efficient framework for pre-trained model selection in time series forecasting via multi-task meta-learning (Figure 2). The **learning-guided selection framework** adopts a dual-encoder architecture: a temporal-aware data encoder captures sequential patterns by segmenting time series into patches to generate data embeddings, while a knowledge-infused model encoder incorporates prior knowledge about models to construct rich model embeddings. Then, we employ patch-level cross-attention to evaluate fine-grained compatibility scores between them. To facilitate multi-task forecasting and enhance generalization, we further adopt a **generalizable multi-task meta-learning** strategy: a horizon-adaptive expert composition module adaptively assigns weights to experts based on the target horizon, and the transferable cross-task learning with cross-dataset and cross-horizon task sampling improves robustness under OOD scenarios. The framework is trained on a meta-dataset of N samples, $\mathcal{D}_{\text{meta}} = \{D^i, Z, H^i, \mathbf{r}^i\}_{i=1}^N$, where the i -th sample includes a downstream dataset D^i , a shared model hub Z , the horizon H^i and the corresponding ranking scores $\mathbf{r}^i$ of the model hub. The ranking scores reflect the relative performance of models in Z on dataset D^i under horizon H^i . This meta-dataset allows the framework to learn how models perform on various datasets, enabling accurate performance prediction and model selection on unseen datasets.

4.1 LEARNING-GUIDED SELECTION FRAMEWORK

Temporal-Aware Data Encoder. In time series modeling, capturing temporal dependencies and sequential patterns is essential for accurate forecasting. To effectively model these characteristics, we divide the time series $X \in \mathbb{R}^{L \times C}$ with L time steps across C variates into patches of size S , yielding $P = \lfloor L/S \rfloor$ patches. Each patch $X_p \in \mathbb{R}^{S \times C}$ is linearly projected into a d -dimensional embedding $X'_p \in \mathbb{R}^{1 \times d}$, forming patch embeddings $E_{\text{patch}} \in \mathbb{R}^{P \times d}$. To preserve temporal order, the positional encodings $E_{\text{pos}} \in \mathbb{R}^{P \times d}$ following (Vaswani et al., 2017) are added: $E_{\text{inp}} = E_{\text{patch}} + E_{\text{pos}}$.

The resulting embeddings E_{inp} are fed into a self-attention module to capture long-range dependencies, where $W_Q^{sa} \in \mathbb{R}^{d \times d}$, $W_K^{sa} \in \mathbb{R}^{d \times d}$, and $W_V^{sa} \in \mathbb{R}^{d \times d}$ are the learnable projection matrices:

$$E_{\text{sa}} = \text{SA}(E_{\text{inp}}) = \text{softmax} \left(E_{\text{inp}} W_Q^{sa} (E_{\text{inp}} W_K^{sa})^T / \sqrt{d_k} \right) E_{\text{inp}} W_V^{sa} \quad (2)$$

Downstream datasets are often large, making full-dataset encoding costly and limiting representation diversity. We address this by repeatedly sampling subsets of B time series and aggregating their attention outputs E_{sa} into $E_{\text{sub}} \in \mathbb{R}^{B \times P \times d}$. Mean pooling along the batch dimension yields a compact data embedding $E_d \in \mathbb{R}^{P \times d}$ that captures shared temporal patterns within the subset while remaining robust to sample variability. The resulting E_d then serves as an expressive summary of downstream tasks to compute compatibility scores with model embeddings for model selection.

Knowledge-Infused Model Encoder. To characterize a pre-trained model ϕ_k , we infuse three key components: meta-information, topological structure, and functionality. The **meta-information** embedding $v_a^k \in \mathbb{R}^{1 \times d_a}$ encodes prior knowledge from the pre-training to guide model selection. We consider: (1) Model architecture: Categorized into *encoder-only*, *decoder-only*, and *encoder-decoder*, reflecting structural design and training characteristics. (2) Model capacity: Estimated by parameter count, indicating ability to capture complex patterns. (3) Model complexity: Measured in Giga Multiply-Accumulate operations (GMACs). Generally, higher complexity allows the model to capture richer patterns. (4) Model dimension: The hidden dimension size across inputs, states, and outputs. Larger dimensions enable the model to learn more expressive and detailed information. (5) Pre-trained domain: Models pre-trained on similar domains typically transfer better.

The **topological structure** provides a detailed view of a model’s architecture and inductive biases. Structural information, such as layer depth, data flow, and connectivity reveals how models process inputs, extract features, and make predictions. We first represent the architecture of the model ϕ_k as a directed acyclic graph (DAG), denoted by $G_k = (V_k, E_k, A_k)$, where V_k and E_k denote vertices and edges, and A_k represents the node attributes. Once the DAG is constructed, we employ graph2vec (Narayanan et al., 2017; Rozemberczki et al., 2020), an unsupervised graph embedding method inspired by doc2vec (Le & Mikolov, 2014) to obtain the topological embedding $v_t^k \in \mathbb{R}^{1 \times d_t}$.

The **functionality** reflects how pre-trained parameters encapsulate the biases and knowledge acquired during pre-training. Since directly embedding millions of parameters is infeasible, we adopt

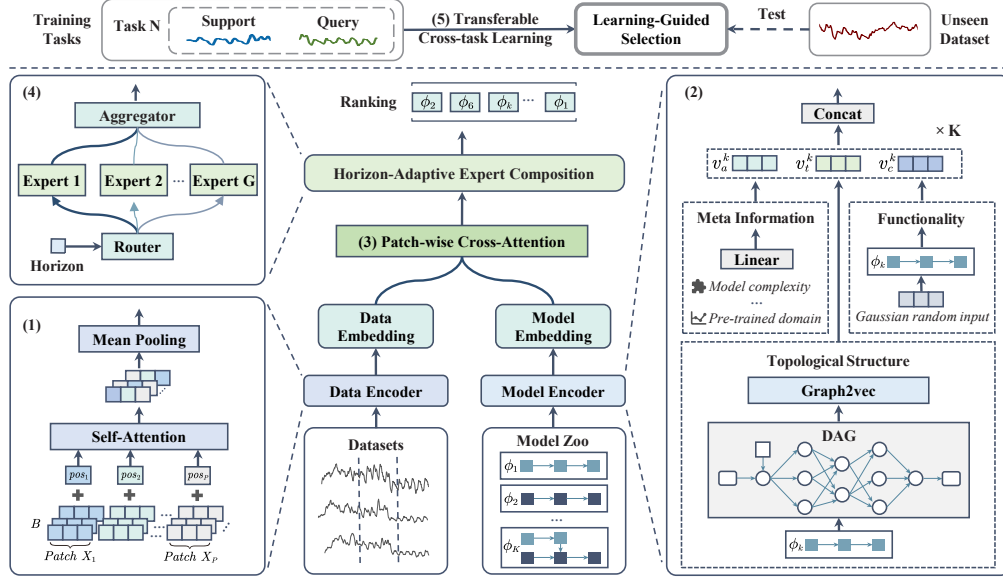


Figure 2: The framework of SwiftTS, consisting of (1) a temporal-aware data encoder, (2) a knowledge-infused model encoder, (3) patchwise cross-attention, (4) a horizon-adaptive expert composition module, and (5) the transferable cross-task learning.

a functional embedding inspired by model distillation, which characterizes a model through its input-output behavior. The intuition is simple: models with different architectures or parameters implement distinct functions and thus are expected to produce distinguishable outputs on identical inputs. Specifically, we feed a fixed set of Gaussian random noise inputs $\epsilon \sim \mathcal{N}(0, I)$ into each model ϕ_k and record the outputs $v_c^k = \phi_k(\epsilon)$ as its functional embedding, where $v_c^k \in \mathbb{R}^{1 \times d_c}$.

Finally, we integrate the meta-information embeddings $v_a \in \mathbb{R}^{K \times d_a}$, the topological embeddings $v_t \in \mathbb{R}^{K \times d_t}$, and the functional embeddings $v_c \in \mathbb{R}^{K \times d_c}$ of the model hub by concatenation, then project the result through a linear transformation $W_m \in \mathbb{R}^{d \times (d_a + d_t + d_c)}$ and a nonlinear activation σ to generate the final model embedding $E_m \in \mathbb{R}^{K \times d}$:

$$E_m = \sigma([v_a, v_t, v_c]W_m^T) \quad (3)$$

Patchwise Compatibility Score. To facilitate a fine-grained and context-aware comparison between downstream datasets and pre-trained models, we compute a compatibility score using patchwise cross-attention (CA). Unlike global similarity measures, patchwise cross-attention captures localized correspondences by assessing each data patch’s contribution to overall compatibility. Specifically, the model embedding E_m as the query, and the data embedding E_d as the key and value:

$$E_{ca} = CA(E_m, E_d) = \text{softmax} \left(E_m W_Q^{ca} (E_d W_K^{ca})^T / \sqrt{d_k} \right) E_d W_V^{ca} \quad (4)$$

where $W_Q^{ca} \in \mathbb{R}^{d \times d}$, $W_K^{ca} \in \mathbb{R}^{d \times d}$, $W_V^{ca} \in \mathbb{R}^{d \times d}$ are projection matrices. This mechanism enables the model to focus on semantically meaningful regions in the data that are most relevant to the characteristics of the model. Finally, a multi-layer perceptron (MLP) produces the ranking prediction, where $\hat{r} \in \mathbb{R}^K$ denotes the predicted ranking scores for K candidate pre-trained models:

$$\hat{r} = \text{MLP}(E_{ca}) \quad (5)$$

Learn-to-select Optimization. During training, we adopt a joint objective combining ranking regularization and prediction accuracy. The ranking loss enforces correct relative orderings among pre-trained models, while the prediction loss (MSE) ensures precise performance estimation:

$$\mathcal{L}_{\text{total}} = \underbrace{-\sum_{k=1}^K p_k(\hat{r}) \log q_k(r)}_{\text{ranking loss}} + \lambda \cdot \underbrace{\sum_{k=1}^K \|r_k - \hat{r}_k\|_2^2}_{\text{prediction loss}} \quad (6)$$

where $\hat{r}$ and r denote the predicted and ground-truth ranking scores, $p_k(\hat{r})$ and $q_k(r)$ are their softmax-normalized forms of the k -th model, and $\hat{r}_k, r_k$ are the corresponding individual scores.

4.2 GENERALIZABLE MULTI-TASK META-LEARNING

Horizon-Adaptive Expert Composition. Time series models often perform inconsistently across forecasting horizons, leading to varying rankings. To equip our framework with the multi-task flexibility to accommodate varying horizons, we propose a horizon-adaptive expert composition module that dynamically integrates specialized experts for different horizons. A lightweight router network assigns softmax-normalized weights to G experts based on the target horizon H :

$$\mathbf{w} = \text{softmax}(\text{Router}(H; \theta_s)) \quad (7)$$

where θ_s denotes the parameters of the router, and $\mathbf{w} \in \mathbb{R}^G$ the expert weights. Each expert, implemented as an MLP, processes the cross-attention output E_{sa} to generate the final prediction through a linear combination of the expert outputs to replace Equation (5):

$$\hat{\mathbf{r}} = \sum_{g=1}^G w_g \cdot \text{MLP}_g(E_{ca}) \quad (8)$$

This design flexibly adapts to diverse horizons without the need for retraining, enhancing both parameter sharing across tasks and improving computational efficiency.

Transferable Cross-Task Learning. Existing methods often struggle to generalize to datasets that deviate from the pre-training distribution, posing a critical issue in time series forecasting, where performance is also sensitive to the forecasting horizon. To achieve robust model ranking and selection, we target two OOD scenarios: (1) generalizing model rankings from seen to unseen datasets, and (2) transferring performance prediction across different horizons.

To enable transferable cross-task learning, we incorporate meta-learning into our framework. Given the constructed meta-dataset $\mathcal{D}_{\text{meta}} = \{D^i, Z, H^i, \mathbf{r}^i\}_{i=1}^N$, we sample a set of diverse tasks $\mathcal{T} = \{\mathcal{T}_1, \mathcal{T}_2, \dots, \mathcal{T}_n\}$, where each task is divided into a support set and a query set. During training, as shown in Figure 3, the support set is used to simulate fast adaptation to new conditions without updating the parameters, referred to as the inner-loop update. Subsequently, the query set evaluates the model’s performance after adaptation and involves actual parameter updates, constituting the outer-loop update. To explicitly promote cross-domain and cross-horizon generalization, we introduce two task sampling strategies: (1) **cross-dataset sampling**, where tasks and their support and query sets are drawn from different datasets to promote generalization across domains; and (2) **cross-horizon sampling**, where tasks and their support and query sets are constructed from varying forecasting horizons to enhance adaptability at the horizon level. These strategies help the model learn shared and domain-specific patterns, boosting performance across diverse forecasting scenarios.

To formalize the cross-task learning process, we treat each sampled task as an independent learning episode during training. In each episode, the model first adapts to a *support set*, simulating rapid adjustment to a new domain or forecasting horizon. The adapted parameters are then evaluated on a *query set*, and the meta-gradients are computed based on the performance of this query set. These gradients are subsequently used to update the initial model parameters in a direction that improves generalization across tasks. Formally, let θ represent the initial parameters of the model. For a given task $\mathcal{T}_i$, we compute the task-specific parameters θ'_i via a few gradient steps on the support set:

$$\theta'_i = \theta - \alpha \nabla_{\theta} \mathcal{L}_{\text{supp}}(\mathcal{T}_i; \theta) \quad (9)$$

where α is the inner-loop learning rate and $\mathcal{L}_{\text{supp}}(\mathcal{T}_i; \theta)$ denotes the loss function evaluated on the support set of task $\mathcal{T}_i$, as defined in Equation (6). The adapted parameters θ'_i are then evaluated

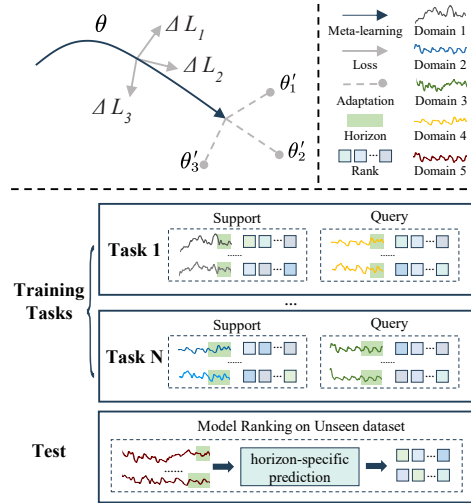


Figure 3: Parameter update process in cross-task learning (top-left), and sampling strategies for cross-horizon and cross-dataset tasks (bottom).

on the query set to compute the meta-objective that measures how well the model generalizes after adaptation: $\mathcal{L}_{\text{query}}(\mathcal{T}_i; \theta'_i)$. This loss is also computed as defined in Equation (6). Finally, the model parameters θ are updated using the meta-gradient across all tasks:

$$\theta \leftarrow \theta - \gamma \nabla_{\theta} \sum_{\mathcal{T}_i} \mathcal{L}_{\text{query}}(\mathcal{T}_i; \theta'_i), \quad (10)$$

where γ is the outer-loop learning rate. By repeatedly performing the two-step optimization process, which consists of inner-loop adaptation and outer-loop generalization, the model learns parameters that enable rapid adaptation to new domains or forecasting horizons with minimal data and computational resources. Please refer to Algorithm 1 for the overall algorithmic process.

5 EXPERIMENTS

5.1 EXPERIMENTAL DESIGN

Datasets. We evaluate SwiftTS on 14 public time series forecasting datasets across diverse domains, including electricity (ETTh1/ETTh2 (Zhou et al., 2021), ETTm1/ETTM2 (Zhou et al., 2021), Electricity (Trindade, 2015)), energy (Solar (Lai et al., 2018), Wind (Li et al., 2022)), traffic (PEMS08 (Song et al., 2020), Traffic (Wu et al., 2021)), environment (Weather (Wu et al., 2021), AQShunyi (Zhang et al., 2017)), natural (ZafNoo (Poyatos et al., 2021), CzeLan (Poyatos et al., 2021)), and economics (Exchange (Lai et al., 2018)). Dataset statistics are detailed in Appendix A.1. To ensure reproducibility, we follow standard dataset splits: training and validation sets are used for model selection, while the test set is reserved for evaluating ground-truth fine-tuning performance.

Pre-trained Models. To ensure robust evaluation, we select eight state-of-the-art time series pre-trained models from diverse architectures and training paradigms: (1) Encoder-only: MOIRAI (Woo et al., 2024), UniTS (Gao et al., 2024), and Moment (Goswami et al., 2024); (2) Decoder-only: TimesFM (Das et al., 2024) and Timer (Liu et al., 2024); (3) Encoder-decoder: TTM (Ekambaram et al., 2024), ROSE (Wang et al., 2024), and Chronos (Ansari et al., 2024). We collect ground-truth performance results from the TSFM-Bench benchmark (Li et al., 2025) for common forecasting horizons of $\{96, 192, 336, 720\}$, under both zero-shot and full-shot settings.

Baselines and Metrics. We compare various model selection methods under three paradigms: (1) Feature-analytic methods: RankME (Garrido et al., 2023), LogME (You et al., 2021), Regression (Gholami et al., 2023), Etran (Gholami et al., 2023), DISCO (Zhang et al., 2025). (2) Learning-based method: Model Spider (Zhang et al., 2023). (3) Brute-force method: Zero-shot performance. Further details are provided in Appendix A.2. For evaluation, we use weighted Kendall’s τ_{ω} to measure the correlation between estimated scores $\{\hat{r}_k\}_{k=1}^K$ and fine-tuned results $\{r_k\}_{k=1}^K$, following previous work (Shao et al., 2022; Gholami et al., 2023; Li et al., 2023; Zhang et al., 2025). A larger τ_{ω} indicates stronger alignment between estimated and true rankings. Details are in Appendix A.3.

Implementation. To assess generalization on unseen tasks and prevent data leakage, we adopt a strict splitting protocol: in each training run, we randomly select 3 out of the 14 benchmark datasets held out for testing, while the remaining 11 are split into training and validation sets using an 8:1 ratio. All experiments are conducted on an NVIDIA GeForce RTX 3090 GPU with batch size 16 for 80 epochs, using $G = 4$ experts. Optimization is performed with Adam($\beta_1 = 0.9$, $\beta_2 = 0.999$). In the meta-learning process, the inner-loop and outer-loop learning rates are $\alpha = 0.001$ and $\gamma = 0.005$. The loss trade-off coefficient is $\lambda = 0.7$, with sensitivity analysis in Section A.5.

5.2 EXPERIMENTAL RESULTS

Main Results. Table 1 compares SwiftTS with various baselines across 14 datasets and four horizons of $\{96, 192, 336, 720\}$. The results demonstrate that SwiftTS consistently outperforms the baselines in average τ_{ω} across all horizons, highlighting its effectiveness in selecting high-performing pre-trained models. Feature-analytic methods often suffer from inconsistent features derived from pre-trained models with diverse architectures and paradigms. Model Spider alleviates this issue by learning a similarity function between datasets and models, but it overlooks sequential dependencies in time series and prior knowledge of the models. In contrast, SwiftTS employs a dual architecture consisting of a temporal-aware data encoder and a knowledge-infused encoder, which

Table 1: Method comparison of the weighted Kendall’s τ_ω across 14 datasets and their average. The best and second-best results are in bold and underlined. Our method achieves the best overall τ_ω .

	Horizon	RankME	LogME	Regression	Etran	DISCO	Model spider	zero-shot	SwiftTS
ETTh1	H=96	0.292	0.257	0.257	0.287	0.213	0.309	0.091	0.464
	H=192	0.182	0.014	0.104	0.147	0.059	0.386	-0.058	0.479
	H=336	0.161	0.137	0.120	0.158	-0.036	<u>0.321</u>	-0.126	0.436
	H=720	0.086	0.196	0.108	0.147	0.019	0.334	0.241	0.543
ETTh2	H=96	0.116	0.089	0.078	0.256	0.156	0.398	0.211	0.436
	H=192	0.101	-0.150	-0.059	0.333	-0.165	0.451	0.430	0.330
	H=336	-0.011	-0.134	-0.083	<u>0.269</u>	-0.154	0.389	0.256	0.219
	H=720	0.053	-0.014	-0.069	0.106	0.188	0.345	0.474	0.351
ETTm1	H=96	0.005	0.126	0.295	0.691	-0.407	0.734	0.063	0.501
	H=192	0.005	0.079	0.079	<u>0.552</u>	-0.406	0.567	-0.040	0.501
	H=336	-0.529	-0.048	-0.048	0.099	0.191	0.590	-0.276	0.652
	H=720	-0.479	-0.027	-0.027	0.173	0.166	<u>0.266</u>	-0.041	0.269
ETTm2	H=96	0.140	0.080	0.334	0.913	0.613	0.667	0.285	0.847
	H=192	0.128	-0.020	0.164	0.568	0.520	0.696	0.319	1.000
	H=336	-0.118	-0.021	0.159	<u>0.700</u>	0.473	0.698	0.727	0.663
	H=720	-0.068	0.050	0.163	0.485	0.165	0.356	0.078	0.534
Electricity	H=96	-0.037	0.317	0.317	0.559	-0.026	0.311	0.162	0.361
	H=192	0.043	0.306	0.306	0.685	-0.585	0.549	0.093	0.240
	H=336	-0.330	-0.419	-0.419	-0.009	0.091	<u>0.358</u>	-0.279	0.520
	H=720	-0.403	-0.344	-0.344	0.027	0.148	0.356	0.029	0.394
Traffic	H=96	-0.138	0.059	-0.124	-0.096	0.148	0.049	-0.312	0.247
	H=192	-0.444	-0.432	-0.432	-0.338	<u>0.225</u>	-0.005	-0.198	0.351
	H=336	-0.569	-0.568	-0.568	-0.692	0.643	0.030	-0.026	0.066
	H=720	-0.833	-0.745	-0.565	-0.394	0.223	0.420	0.222	0.702
Solar	H=96	-0.214	0.120	0.222	0.359	-0.569	0.452	0.308	0.222
	H=192	-0.017	0.473	0.430	0.246	-0.022	0.344	0.110	0.119
	H=336	-0.003	<u>0.553</u>	0.553	-0.359	-0.296	0.250	-0.395	0.277
	H=720	-0.062	<u>0.527</u>	<u>0.527</u>	0.055	-0.703	0.051	-0.027	0.512
Weather	H=96	0.015	-0.381	-0.034	0.499	0.273	0.543	0.128	0.285
	H=192	-0.125	-0.023	-0.118	0.466	0.024	0.257	0.341	0.129
	H=336	-0.171	0.102	-0.066	0.539	0.404	0.340	0.321	0.108
	H=720	-0.065	-0.057	0.241	0.745	0.375	0.251	0.702	0.316
Exchange	H=96	-0.032	-0.492	-0.343	0.193	0.112	0.059	-0.284	0.251
	H=192	-0.040	-0.597	-0.414	<u>0.246</u>	0.152	0.154	-0.273	0.252
	H=336	-0.112	-0.536	-0.306	0.148	0.250	0.230	-0.143	0.233
	H=720	-0.210	-0.617	-0.444	-0.024	0.536	0.322	0.126	-0.386
ZafNoo	H=96	-0.117	0.058	-0.148	0.113	-0.153	0.511	-0.384	0.656
	H=192	-0.285	-0.235	-0.274	-0.044	-0.035	<u>0.436</u>	-0.017	0.786
	H=336	-0.285	0.040	-0.133	-0.073	-0.200	0.580	0.023	0.732
	H=720	-0.224	-0.106	-0.256	0.129	-0.235	0.552	0.164	0.668
CzeLan	H=96	0.103	0.071	0.012	0.632	-0.343	-0.121	0.171	0.575
	H=192	-0.037	-0.376	-0.258	0.337	<u>0.362</u>	0.217	0.171	0.527
	H=336	-0.069	-0.090	-0.155	-0.109	-0.154	-0.009	<u>0.301</u>	0.847
	H=720	-0.125	-0.214	-0.074	-0.125	-0.136	0.349	0.282	0.839
AQShunyi	H=96	-0.371	-0.283	-0.270	-0.117	0.107	0.414	-0.349	0.939
	H=192	-0.407	<u>0.328</u>	<u>0.328</u>	-0.045	-0.224	0.126	0.309	0.734
	H=336	-0.377	-0.184	-0.255	-0.277	0.276	0.084	0.420	0.723
	H=720	-0.332	-0.140	-0.209	-0.253	0.119	0.335	0.675	0.701
Wind	H=96	0.211	0.142	0.062	0.417	-0.196	0.244	-0.106	0.395
	H=192	0.211	0.258	0.231	0.136	0.482	0.355	-0.155	0.395
	H=336	0.045	0.545	0.319	0.338	-0.532	0.281	0.349	0.262
	H=720	-0.040	0.474	0.443	0.126	-0.377	0.202	0.133	0.162
PEMS08	H=96	0.140	0.118	0.118	-0.253	0.110	-0.103	0.445	0.401
	H=192	0.038	-0.003	-0.003	-0.321	-0.068	-0.318	<u>0.420</u>	0.505
	H=336	-0.445	-0.296	-0.296	-0.791	-0.029	-0.025	0.440	0.016
	H=720	-0.637	-0.244	-0.244	-0.178	0.070	-0.345	0.452	0.442
avg	H=96	0.008	0.020	0.056	0.318	0.003	<u>0.319</u>	0.031	0.470
	H=192	-0.046	-0.027	0.006	0.212	0.023	<u>0.301</u>	0.104	0.453
	H=336	-0.201	-0.066	-0.084	-0.004	0.066	<u>0.294</u>	0.114	0.411
	H=720	-0.238	-0.090	-0.054	0.073	0.040	<u>0.271</u>	0.251	0.432
Num.Top-1		0	3	2	8	4	6	5	28

together enhance selection performance. Additionally, SwiftTS features a horizon-adaptive expert composition module, allowing it to handle multiple horizons simultaneously within a unified framework. By comparison, all baselines require recomputation or retraining for different horizons. This efficiency makes SwiftTS well-suited for real-world applications demanding flexible multi-horizon forecasting. Moreover, while other methods exhibit varying degrees of negative correlation in different datasets and horizons, SwiftTS adopts transferable cross-task learning to maintain predominantly positive correlations across 14 datasets and different horizons, showing its OOD robustness.

Top- k Performance. We report the top- k selection probability $\text{Pr}(\text{top-}k)$ as used in (Zhang et al., 2025; Gholami et al., 2023). This metric evaluates the likelihood that the best-performing model appears within the top- k of the estimated ranking. Results in Table 2 demonstrate that SwiftTS achieves the best top- k performance, validating the model selection effectiveness of our method.

Table 2: Methods comparison of $\Pr(\text{top-}k)$ and average τ_ω across horizons on 14 datasets.

	Pr(top1)	Pr(top2)	Pr(top3)	τ_ω
RankME	0.000	0.000	0.196	-0.119
LogME	0.071	0.196	0.268	-0.041
Regression	0.036	0.125	0.286	-0.019
Etran	<u>0.304</u>	0.393	0.536	0.150
DISCO	0.232	0.375	0.536	0.033
Model Spider	<u>0.304</u>	<u>0.482</u>	0.571	<u>0.296</u>
zero shot	0.286	0.464	<u>0.589</u>	0.125
SwiftTS	0.339	0.500	0.607	0.442

Table 3: Ablation studies for model embedding: τ_ω and average across horizons are listed below.

v_a	v_t	v_c	96	192	336	720	avg
✓			0.341	0.283	0.331	0.401	0.339
	✓		0.225	0.270	0.318	0.227	0.267
		✓	0.365	0.401	0.317	0.397	0.370
✓	✓		0.361	0.383	0.315	<u>0.417</u>	0.369
	✓	✓	<u>0.427</u>	<u>0.430</u>	0.328	0.391	0.394
✓		✓	<u>0.380</u>	0.422	0.437	0.403	<u>0.411</u>
✓	✓	✓	0.470	0.453	<u>0.411</u>	0.432	0.442

5.3 FURTHER ANALYSIS

Embedding ablation. The model encoder integrates three embeddings to represent a neural network: meta-information embedding (v_a), topological embedding (v_t), and functional embeddings (v_c). To assess their individual contributions, we perform ablation studies by removing one or two embedding type at a time and measuring the resulting performance. As shown in Table 3, meta-information embedding v_a and the functional embedding v_c contribute the most to performance. The best results are obtained when all three embeddings are utilized together, demonstrating the complementary nature and necessity of each type of embedding.

Cross-task learning ablation. To validate the effectiveness of the proposed transferable cross-task learning, we conduct an ablation experiment as illustrated in Figure 4b. Specifically, we evaluate the model’s average performance on 14 downstream datasets across four horizons. A clear upward trend is observed when comparing the results from the w/o cross-task learning setting to those with the cross-task learning enabled. This consistent improvement across horizons demonstrates that our cross-task learning effectively enhances the model’s predictive capability and robustness.

IID vs. OOD. In Figure 4a, we compare the average τ_ω performance under IID and OOD settings across four datasets and four forecasting horizons. Compared to OOD, the IID setting refers to evaluation on test data drawn from the same distribution as the training data. The results show that our method achieves further improvements under the IID setting, indicating that increasing the diversity of training data can enhance model performance.

Efficiency and Scalability. We evaluate efficiency by analyzing runtime on small (ETTh1) and large (Traffic) datasets (Figure 4c, Figure 4d). Results show that model selection methods significantly reduce computational overhead compared to full fine-tuning. For example, on ETTh1, model selection methods typically require only 1,000 to 4,000 seconds, whereas fully fine-tuning each model takes approximately 4.97×10^4 seconds on the same GPU. This stark contrast highlights the critical importance of efficient model selection methods in practical applications. Moreover, the comparison across different dataset scales reveals that the runtime cost of SwiftTS remains relatively stable and is less sensitive to dataset size. In contrast, the time overhead of other model selection baselines increases dramatically as the dataset grows. Moreover, the cost of fine-tuning on the Traffic dataset reaches up to 3.46×10^6 seconds, making it prohibitively expensive for large-scale applications. Overall, SwiftTS achieves both superior performance and minimal time overhead.

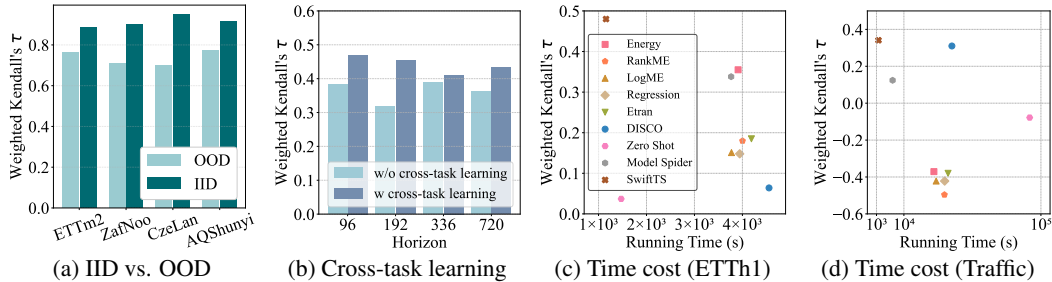


Figure 4: (a) Comparison of (a) average τ_ω for IID vs. OOD settings across four datasets and horizons, and (b) ablation study of cross-task learning, and method comparison w.r.t running time (second) and average τ_ω across horizons on (c) ETTh1 (small-scale) and (d) Traffic (large-scale).

6 CONCLUSION

This paper tackles the challenge of pre-trained model selection from a model hub for time series forecasting. We propose SwiftTS, a learning-guided framework with a lightweight dual-encoder architecture that independently embeds time series and candidate models, computing patchwise compatibility scores for efficient selection. To further enhance adaptability, SwiftTS incorporates a horizon-adaptive expert composition module for multi-task forecasting and leverages transferable cross-task learning to improve generalization across datasets and horizons. Extensive experiments show that SwiftTS achieves SOTA with high efficiency and scalability for real-world deployment.

ETHICS STATEMENT

This work is conducted entirely on publicly available benchmark datasets, as detailed in the paper, and does not involve the release of any personal or sensitive information. No human subjects are involved in this research, ensuring that our work complies with ethical standards in research integrity.

REPRODUCIBILITY STATEMENT

The performance of SwiftTS and the datasets used in our work are real, and all experimental results can be reproduced, as detailed in the paper. Details of model architecture, training procedures, and evaluation protocols are provided in the main text and appendix. To further facilitate reproducibility, we release the code and datasets at <https://anonymous.4open.science/r/SwiftTS-395C>.

REFERENCES

- Abdul Fatir Ansari, Lorenzo Stella, Ali Caner Türkmen, Xiyuan Zhang, Pedro Mercado, Huibin Shen, Oleksandr Shchur, Syama Sundar Rangapuram, Sebastian Pineda-Arango, Shubham Kapoor, Jasper Zschiegner, Danielle C. Maddix, Michael W. Mahoney, Kari Torkkola, Andrew Gordon Wilson, Michael Bohlke-Schneider, and Yuyang Wang. Chronos: Learning the language of time series. *CoRR*, abs/2403.07815, 2024.
- Peng Chen, Yingying Zhang, Yunyao Cheng, Yang Shu, Yihang Wang, Qingsong Wen, Bin Yang, and Chenjuan Guo. Pathformer: Multi-scale transformers with adaptive pathways for time series forecasting. In *The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024*. OpenReview.net, 2024. URL <https://openreview.net/forum?id=lJkOCMP2aW>.
- Abhimanyu Das, Weihao Kong, Rajat Sen, and Yichen Zhou. A decoder-only foundation model for time-series forecasting. In *ICML*, 2024.
- Tim Dettmers, Artidoro Pagnoni, Ari Holtzman, and Luke Zettlemoyer. Qlora: Efficient finetuning of quantized llms. In *NeurIPS*, 2023.
- Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *ICLR*, 2021.
- Vijay Ekambaram, Arindam Jati, Pankaj Dayama, Sumanta Mukherjee, Nam Nguyen, Wesley M. Gifford, Chandra Reddy, and Jayant Kalagnanam. Tiny time mixers (ttms): Fast pre-trained models for enhanced zero/few-shot forecasting of multivariate time series. In *NeurIPS*, 2024.
- Shanghua Gao, Teddy Koker, Owen Queen, Thomas Hartvigsen, Theodoros Tsiligkaridis, and Marinka Zitnik. Units: Building a unified time series model. *CoRR*, abs/2403.00131, 2024.
- Quentin Garrido, Randall Balestriero, Laurent Najman, and Yann LeCun. Rankme: Assessing the downstream performance of pretrained self-supervised representations by their rank. In *ICML*, volume 202 of *Proceedings of Machine Learning Research*, pp. 10929–10974. PMLR, 2023.

-
- Mohsen Gholami, Mohammad Akbari, Xinglu Wang, Behnam Kamranian, and Yong Zhang. Etran: Energy-based transferability estimation. In *ICCV*, pp. 18567–18576, 2023.
- Mononito Goswami, Konrad Szafer, Arjun Choudhry, Yifu Cai, Shuo Li, and Artur Dubrawski. MOMENT: A family of open time-series foundation models. In *ICML*, 2024.
- Aaron Hurst, Adam Lerer, Adam P. Goucher, Adam Perelman, Aditya Ramesh, Aidan Clark, AJ Ostrow, Akila Welihinda, Alan Hayes, Alec Radford, Aleksander Madry, Alex Baker-Whitcomb, Alex Beutel, Alex Borzunov, Alex Carney, Alex Chow, Alex Kirillov, Alex Nichol, Alex Paino, Alex Renzin, Alex Tachard Passos, Alexander Kirillov, Alexi Christakis, Alexis Conneau, Ali Kamali, Allan Jabri, Allison Moyer, Allison Tam, Amadou Crookes, Amin Tootoonchian, Ananya Kumar, Andrea Vallone, Andrej Karpathy, Andrew Braunstein, Andrew Cann, Andrew Codis-poti, Andrew Galu, Andrew Kondrich, Andrew Tulloch, Andrey Mishchenko, Angela Baek, Angela Jiang, Antoine Pelisse, Antonia Woodford, Anuj Gosalia, Arka Dhar, Ashley Pantuliano, Avi Nayak, Avital Oliver, Barret Zoph, Behrooz Ghorbani, Ben Leimberger, Ben Rossen, Ben Sokolowsky, Ben Wang, Benjamin Zweig, Beth Hoover, Blake Samic, Bob McGrew, Bobby Spero, Bogó Giertler, Bowen Cheng, Brad Lightcap, Brandon Walkin, Brendan Quinn, Brian Guarraci, Brian Hsu, Bright Kellogg, Brydon Eastman, Camillo Lugaresi, Carroll L. Wainwright, Cary Bassin, Cary Hudson, Casey Chu, Chad Nelson, Chak Li, Chan Jun Shern, Channing Conger, Charlotte Barette, Chelsea Voss, Chen Ding, Cheng Lu, Chong Zhang, Chris Beaumont, Chris Hallacy, Chris Koch, Christian Gibson, Christina Kim, Christine Choi, Christine McLeavey, Christopher Hesse, Claudia Fischer, Clemens Winter, Coley Czarnecki, Colin Jarvis, Colin Wei, Constantin Koumouzelis, and Dane Sherburn. Gpt-4o system card. *CoRR*, abs/2410.21276, 2024. doi: 10.48550/ARXIV.2410.21276. URL <https://doi.org/10.48550/arXiv.2410.21276>.
- Menglin Jia, Luming Tang, Bor-Chun Chen, Claire Cardie, Serge J. Belongie, Bharath Hariharan, and Ser-Nam Lim. Visual prompt tuning. In *ECCV*, volume 13693 of *Lecture Notes in Computer Science*, pp. 709–727. Springer, 2022.
- Ananya Kumar, Aditi Raghunathan, Robbie Matthew Jones, Tengyu Ma, and Percy Liang. Fine-tuning can distort pretrained features and underperform out-of-distribution. In *ICLR*, 2022.
- Guokun Lai, Wei-Cheng Chang, Yiming Yang, and Hanxiao Liu. Modeling long- and short-term temporal patterns with deep neural networks. In *SIGIR*, pp. 95–104, 2018.
- Quoc V. Le and Tomáš Mikolov. Distributed representations of sentences and documents. In *ICML*, volume 32, pp. 1188–1196, 2014.
- Xiaotong Li, Zixuan Hu, Yixiao Ge, Ying Shan, and Ling-Yu Duan. Exploring model transferability through the lens of potential energy. In *ICCV*, pp. 5406–5415, 2023.
- Yan Li, Xinjiang Lu, Yaqing Wang, and Dejing Dou. Generative time series forecasting with diffusion, denoise, and disentanglement. In *NeurIPS*, 2022.
- Zhe Li, Xiangfei Qiu, Peng Chen, Yihang Wang, Hanyin Cheng, Yang Shu, Jilin Hu, Chenjuan Guo, Aoying Zhou, Christian S. Jensen, and Bin Yang. Tsfm-bench: A comprehensive and unified benchmark of foundation models for time series forecasting, 2025.
- Yong Liu, Haoran Zhang, Chenyu Li, Xiangdong Huang, Jianmin Wang, and Mingsheng Long. Timer: Generative pre-trained transformers are large time series models. In *ICML*, 2024.
- Annamalai Narayanan, Mahinthan Chandramohan, Rajasekar Venkatesan, Lihui Chen, Yang Liu, and Shantanu Jaiswal. graph2vec: Learning distributed representations of graphs. *CoRR*, abs/1707.05005, 2017.
- Cuong V. Nguyen, Tal Hassner, Matthias W. Seeger, and Cédric Archambeau. LEEP: A new measure to evaluate transferability of learned representations. In *ICML*, volume 119 of *Proceedings of Machine Learning Research*, pp. 7294–7305, 2020.
- Yuqi Nie, Nam H. Nguyen, Phanwadee Sinthong, and Jayant Kalagnanam. A time series is worth 64 words: Long-term forecasting with transformers. In *The Eleventh International Conference on Learning Representations, ICLR 2023, Kigali, Rwanda, May 1-5, 2023*. OpenReview.net, 2023. URL <https://openreview.net/forum?id=Jbdc0vTOcol>.

-
- Michal Pándy, Andrea Agostinelli, Jasper R. R. Uijlings, Vittorio Ferrari, and Thomas Mensink. Transferability estimation using bhattacharyya class separability. In *CVPR*, pp. 9162–9172, 2022.
- Rafael Poyatos, Víctor Granda, Víctor Flo, Mark A. Adams, Balázs Adorján, David Aguadé, Marcos P. M. Aidar, Scott Allen, M. Susana Alvarado-Barrientos, Kristina J. Anderson-Teixeira, Luiza Maria Aparecido, M. Altaf Arain, Ismael Aranda, Heidi Asbjornsen, Robert Baxter, Eric Beamesderfer, Z. Carter Berry, Daniel Berveiller, Bethany Blakely, Johnny Boggs, Gil Bohrer, Paul V. Bolstad, Damien Bonal, Rosvel Bracho, Patricia Brito, Jason Brodeur, Fernando Casanoves, Jérôme Chave, Hui Chen, Cesar Cisneros, Kenneth Clark, Edoardo Cremonese, Hongzhong Dang, Jorge S. David, Teresa S. David, Nicolas Delpierre, Ankur R. Desai, Frederic C. Do, Michal Dohnal, Jean-Christophe Domec, Sebinasi Dziki, Colin Edgar, Rebekka Eichstaedt, Tarek S. El-Madany, Jan Elbers, Cleiton B. Eller, Eugénie S. Euskirchen, Brent Ewers, Patrick Fonti, Alicia Forner, David I. Forrester, Helber C. Freitas, Marta Galvagno, Omar Garcia-Tejera, Chandra Prasad Ghimire, Teresa E. Gimeno, John Grace, André Granier, Anne Griebel, Yan Guangyu, Mark B. Gush, Paul J. Hanson, Niles J. Hasselquist, Ingo Heinrich, Virginia Hernandez-Santana, Valentine Herrmann, Teemu Hölttä, Friso Holwerda, James Irvine, Supat Isarangkool Na Ayutthaya, Paul G. Jarvis, Hubert Jochheim, Carlos A. Joly, Julia Kaplick, Hyun Seok Kim, Leif Klemedtsson, Heather Kropp, Fredrik Lagergren, Patrick Lane, Petra Lang, Andrei Lapenas, Víctor Lechuga, Minsu Lee, Christoph Leuschner, Jean-Marc Limousin, Juan Carlos Linares, Maj-Lena Linderson, Anders Lindroth, Pilar Llorens, Álvaro López-Bernal, Michael M. Loranty, Dietmar Lüttschwager, Cate Macinnis-Ng, Isabelle Maréchaux, Timothy A. Martin, Ashley Matheny, Nate McDowell, Sean McMahon, Patrick Meir, Ilona Mészáros, Mirco Migliaavacca, Patrick Mitchell, Meelis Mölder, Leonardo Montagnani, Georgianne W. Moore, Ryogo Nakada, Furong Niu, Rachael H. Nolan, Richard Norby, Kimberly Novick, Walter Oberhuber, Nikolaus Obojes, A. Christopher Oishi, Rafael S. Oliveira, Ram Ören, Jean-Marc Ourcival, Teemu Paljakka, Oscar Perez-Priego, Pablo L. Peri, Richard L. Peters, Sebastian Pfautsch, William T. Pockman, Yakir Preisler, Katherine Rascher, George Robinson, Humberto Rocha, Alain Rocheteau, Alexander Röhl, Bruno H. P. Rosado, Lucy Rowland, Alexey V. Rubtsov, Santiago Sabaté, Yann Salmon, Roberto L. Salomón, Elisenda Sánchez-Costa, Karina V. R. Schäfer, Bernhard Schuldt, Alexandr Shashkin, Clément Stahl, Marko Stojanović, Juan Carlos Suárez, Ge Sun, Justyna Szatniewska, Fyodor Tatarinov, Miroslav Tesař, Frank M. Thomas, Pantana Tor-ngern, Josef Urban, Fernando Valladares, Christiaan van der Tol, Ilja van Meerveld, Andrej Varlagin, Holm Voigt, Jeffrey Warren, Christiane Werner, Willy Werner, Gerhard Wieser, Lisa Wingate, Stan Wullschleger, Koong Yi, Roman Zweifel, Kathy Steppe, Maurizio Mencuccini, and Jordi Martínez-Vilalta. Global transpiration data from sap flow measurements: the sapfluxnet database. *Earth System Science Data*, 13(6):2607–2649, June 2021. ISSN 1866-3516.
- Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J. Liu. Exploring the limits of transfer learning with a unified text-to-text transformer. *J. Mach. Learn. Res.*, 21:140:1–140:67, 2020.
- Benedek Rozemberczki, Oliver Kiss, and Rik Sarkar. Karate club: An API oriented open-source python framework for unsupervised learning on graphs. In *CIKM*, pp. 3125–3132. ACM, 2020.
- Wenqi Shao, Xun Zhao, Yixiao Ge, Zhaoyang Zhang, Lei Yang, Xiaogang Wang, Ying Shan, and Ping Luo. Not all models are equal: Predicting model transferability in a self-challenging fisher space. In *ECCV*, volume 13694 of *Lecture Notes in Computer Science*, pp. 286–302, 2022.
- Chao Song, Youfang Lin, Shengnan Guo, and Huaiyu Wan. Spatial-temporal synchronous graph convolutional networks: A new framework for spatial-temporal network data forecasting. In *AAAI*, pp. 914–921, 2020.
- Anh Tuan Tran, Cuong V. Nguyen, and Tal Hassner. Transferability and hardness of supervised classification tasks. In *ICCV*, pp. 1395–1405, 2019.
- Artur Trindade. Electricityloadaddiagrams20112014. <https://doi.org/10.24432/C58C86>, March 2015. Accessed on YYYY-MM-DD.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention is all you need. In Isabelle Guyon, Ulrike von Luxburg, Samy Bengio, Hanna M. Wallach, Rob Fergus, S. V. N. Vishwanathan, and Roman Garnett (eds.), *NeurIPS*, pp. 5998–6008, 2017.

-
- Yihang Wang, Yuying Qiu, Peng Chen, Kai Zhao, Yang Shu, Zhongwen Rao, Lujia Pan, Bin Yang, and Chenjuan Guo. ROSE: register assisted general time series forecasting with decomposed frequency learning. *CoRR*, abs/2405.17478, 2024.
- Zijian Wang, Yadan Luo, Liang Zheng, Zi Huang, and Mahsa Baktashmotlagh. How far pre-trained models are from neural collapse on the target dataset informs their transferability. In *ICCV*, pp. 5526–5535, 2023.
- Gerald Woo, Chenghao Liu, Akshat Kumar, Caiming Xiong, Silvio Savarese, and Doyen Sahoo. Unified training of universal time series forecasting transformers. In *ICML*, 2024.
- Haixu Wu, Jiehui Xu, Jianmin Wang, and Mingsheng Long. Autoformer: Decomposition transformers with auto-correlation for long-term series forecasting. In *NeurIPS*, pp. 22419–22430, 2021.
- Haixu Wu, Tengge Hu, Yong Liu, Hang Zhou, Jianmin Wang, and Mingsheng Long. Timesnet: Temporal 2d-variation modeling for general time series analysis. In *The Eleventh International Conference on Learning Representations, ICLR 2023, Kigali, Rwanda, May 1-5, 2023*. OpenReview.net, 2023. URL https://openreview.net/forum?id=ju_Uqw3840q.
- An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, Chujie Zheng, Dayiheng Liu, Fan Zhou, Fei Huang, Feng Hu, Hao Ge, Haoran Wei, Huan Lin, Jialong Tang, Jian Yang, Jianhong Tu, Jianwei Zhang, Jian Yang, Jiayi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keqin Bao, Kexin Yang, Le Yu, Lianghao Deng, Mei Li, Mingfeng Xue, Mingze Li, Pei Zhang, Peng Wang, Qin Zhu, Rui Men, Ruize Gao, Shixuan Liu, Shuang Luo, Tianhao Li, Tianyi Tang, Wenbiao Yin, Xingzhang Ren, Xinyu Wang, Xinyu Zhang, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yinger Zhang, Yu Wan, Yuqiong Liu, Zekun Wang, Zeyu Cui, Zhenru Zhang, Zhipeng Zhou, and Zihan Qiu. Qwen3 technical report. *CoRR*, abs/2505.09388, 2025. doi: 10.48550/ARXIV.2505.09388. URL <https://doi.org/10.48550/arXiv.2505.09388>.
- Kaichao You, Yong Liu, Jianmin Wang, and Mingsheng Long. Logme: Practical assessment of pre-trained models for transfer learning. In *ICML*, volume 139 of *Proceedings of Machine Learning Research*, pp. 12133–12143, 2021.
- Shuyi Zhang, Bin Guo, Anlan Dong, Jing He, Ziping Xu, and Song Xi Chen. Cautionary tales on air-quality improvement in beijing. *Proceedings of the Royal Society. A, Mathematical, physical, and engineering sciences*, 473(2205):20170457–20170457, 2017. ISSN 1364-5021.
- Tengxue Zhang, Yang Shu, Xinyang Chen, Yifei Long, Chenjuan Guo, and Bin Yang. Assessing pre-trained models for transfer learning through distribution of spectral components. In *AAAI*, pp. 22560–22568, 2025. URL <https://doi.org/10.1609/aaai.v39i21.34414>.
- Yi-Kai Zhang, Ting-Ji Huang, Yao-Xiang Ding, De-Chuan Zhan, and Han-Jia Ye. Model spider: Learning to rank pre-trained models efficiently. In *NeurIPS*, 2023.
- Haoyi Zhou, Shanghang Zhang, Jieqi Peng, Shuai Zhang, Jianxin Li, Hui Xiong, and Wancai Zhang. Informer: Beyond efficient transformer for long sequence time-series forecasting. In *AAAI*, pp. 11106–11115, 2021.

A APPENDIX

A.1 DATASETS

We evaluate SwiftTS on 14 multivariate time-series datasets spanning six distinct domains, including electricity, energy, traffic, environment, nature, and economics. (1) The ETT datasets (Zhou et al., 2021) contain 7 variables collected from two different power transformers between July 2016 and July 2018. The dataset consists of four subsets: ETTh1 and ETTh2, recorded hourly, and ETTm1 and ETTm2, recorded at 15-minute intervals. (2) Electricity (Trindade, 2015) records hourly electricity consumption from 321 customers over three years, from July 2016 to July 2019. (3) Solar (Lai et al., 2018) captures solar power generation from 137 photovoltaic plants in 2006, sampled every 10 minutes. (4) Wind (Li et al., 2022) consists of historical wind measurements (e.g., speed, direction). (5) PEMS08 (Song et al., 2020) contains three months of aggregated statistics on traffic flow, speed, and occupancy rate. (6) Traffic (Wu et al., 2021) contains hourly road occupancy rates measured by 862 sensors across freeways in the San Francisco Bay Area from 2015 to 2016. (7) Weather (Wu et al., 2021) includes 21 meteorological variables (e.g., temperature, humidity, and barometric pressure), recorded every 10 minutes across Germany in 2020. (8) AQShunyi (Zhang et al., 2017) provides hourly temperature measurements exhibiting strong seasonal patterns. (9) ZafNoo (Poyatos et al., 2021) is collected from the Sapflux data project and includes sap flow measurements and environmental variables. (10) CzeLan (Poyatos et al., 2021) is from the Sapflux data project, including sap flow measurements and environmental variables. (11) Exchange (Lai et al., 2018) comprises daily exchange rates for eight countries over a multi-year period. Table 4 summarizes the key statistics of these 14 multivariate time series datasets.

Table 4: Statistics of datasets.

Dataset	Domain	Frequency	Lengths	Dim	Split
ETTh1 (Zhou et al., 2021)	Electricity	1 hour	14,400	7	6:2:2
ETTh2 (Zhou et al., 2021)	Electricity	1 hour	14,400	7	6:2:2
ETTm1 (Zhou et al., 2021)	Electricity	15 mins	57,600	7	6:2:2
ETTm2 (Zhou et al., 2021)	Electricity	15 mins	57,600	7	6:2:2
Electricity (Trindade, 2015)	Electricity	1 hour	26,304	321	7:1:2
Solar (Lai et al., 2018)	Energy	10 mins	52,560	137	6:2:2
Wind (Li et al., 2022)	Energy	15 mins	48,673	7	7:1:2
PEMS08 (Song et al., 2020)	Traffic	5 mins	17,856	170	6:2:2
Traffic (Wu et al., 2021)	Traffic	1 hour	17,544	862	7:1:2
Weather (Wu et al., 2021)	Environment	10 mins	52,696	21	7:1:2
AQShunyi (Zhang et al., 2017)	Environment	1 hour	35,064	11	6:2:2
ZafNoo (Poyatos et al., 2021)	Nature	30 mins	19,225	11	7:1:2
CzeLan (Poyatos et al., 2021)	Nature	30 mins	19,934	11	7:1:2
Exchange (Lai et al., 2018)	Economic	1 day	7,588	8	7:1:2

A.2 BASELINES

In our study, we compare various model selection methods for pre-trained time series models, which can be broadly categorized into three paradigms: **(1) Feature-analytic methods:** These methods rely on intrinsic properties or statistical characteristics of features extracted by the pre-trained models to estimate their transferability. **RankME** (Garrido et al., 2023) evaluates the rank of feature matrices extracted from the model’s representations. **LogME** (You et al., 2021) computes the logarithm of the maximum label marginalized likelihood under a probabilistic model. **Regression** (Gholami et al., 2023) employs linear regression using Singular Value Decomposition (SVD) to approximate the mapping from model features to target labels. **Etran** (Gholami et al., 2023) combines both the energy score and the regression score into a unified metric. **DISCO** (Zhang et al., 2025) evaluates pre-trained models by analyzing the spectral distribution of their feature representations, enabling the assessment in both classification and regression tasks. **(2) Learning-based method: Model Spider** (Zhang et al., 2023) learns model representations and a similarity function through alignment with downstream task representations, facilitating model selection via the learned similarity. **(3)**

Brute-force method: Zero-shot measures a model’s ability to generalize to unseen tasks without any task-specific fine-tuning, offering valuable insights into its overall generalization capacity.

A.3 METRICS

Kendall’s τ measures the ordinal association between two rankings by evaluating the number of concordant and discordant pairs, which is defined as:

$$\tau = \frac{2}{K(K-1)} \sum_{1 \leq i < j \leq K} \text{sgn}(r^i - r^j) \text{sgn}(\hat{r}^i - \hat{r}^j) \quad (11)$$

where $\text{sgn}(x)$ is the sign function. However, in practical model selection scenarios, accurately identifying the top-performing models is often more critical than precisely ranking lower-performing ones. To reflect this priority, we employ the weighted version of Kendall’s τ , denoted as τ_ω . This variant adjusts the contribution of each pairwise comparison by assigning larger weights to higher-ranked models. A higher value of τ_ω indicates stronger consistency between estimated and actual rankings, reflecting the reliability of the evaluation metric in guiding model selection.

A.4 ALGORITHM

Algorithm 1: Pseudo-code of `SwiftTS`

```

1 Training:
2 Input: Meta-dataset  $\mathcal{D}_{\text{meta}} = \{D^i, Z, H^i, \mathbf{r}^i\}_{i=1}^N$ ; Total training epochs  $E$ ; Inner-loop
   learning rate  $\alpha$ ; Outer-loop learning rate  $\gamma$ ;
3 Randomly initialize the weights  $\theta$  of the whole model;
4 for  $e \leftarrow 1$  to  $E$  do
5    $\mathcal{T} \leftarrow$  Sample  $n$  tasks from  $\mathcal{D}_{\text{meta}}$ ;
6   for  $j \leftarrow 1$  to  $n$  do
7      $\text{support}_j, \text{query}_j \leftarrow$  Obtain support set and query set from  $\mathcal{T}_j$ ;
8      $\hat{\mathbf{r}} \leftarrow \text{ModelRanking}(Z, \text{support}_j, H)$ 
9     Evaluate  $\nabla_{\theta} \mathcal{L}_{\text{supp}}(\mathcal{T}_j; \theta)$  using 6;
10    Compute  $\theta'_j \leftarrow \theta - \alpha \nabla_{\theta} \mathcal{L}_{\text{supp}}(\mathcal{T}_j; \theta)$ ;
11    Update  $\theta \leftarrow \theta - \gamma \nabla_{\theta} \sum_{\mathcal{T}_j} \mathcal{L}_{\text{query}}(\mathcal{T}_j; \theta'_j)$ ;

12 Inference:
13 Input: An unseen downstream dataset  $X$ , model hub  $Z$ , forecasting horizon  $H$ ;
14  $\hat{\mathbf{r}} \leftarrow \text{ModelRanking}(Z, X, H)$ 
15 Return: The predicted ranking scores  $\hat{\mathbf{r}}$  for  $K$  candidate pre-trained models in the model
   hub.

16 Function  $\text{ModelRanking}(Z, X, H)$  :
17    $E_d, E_m \leftarrow E_D(X), E_M(Z)$ ;
18    $E_{\text{ca}} \leftarrow$  Compute patchwise CA using Eq. 4;
19    $w \leftarrow$  Compute the weights that router assigns to each expert using Eq. 7 with  $H$ ;
20   Return:  $\hat{\mathbf{r}} \leftarrow$  Rank pre-trained models using Eq. 8;

```

A.5 MORE EXPERIMENTAL RESULTS

Sensitivity Analysis. We examine the impact of $\lambda \in [0.0, 0.3, 0.5, 0.7, 0.9, 1.0]$ and report average results across all datasets. As shown in Figure 5a, the performance remains relatively stable due to the complementary roles of the two loss components: $\mathcal{L}_{\text{rank}}$ constrains the relative ranking among pre-trained models, while $\mathcal{L}_{\text{pred}}$ enforces an absolute constraint on the gap between predicted performance and actual fine-tuning results. The experimental results highlight the indispensable role of the prediction loss and suggest that the best performance is achieved when λ is set to 0.7.

Choices of Expert Numbers. We study the impact of expert number $G \in [2, 4, 6, 8, 10]$ (Figure 5b). A small G fails to adequately capture the differences across horizons, limiting its ability to perform

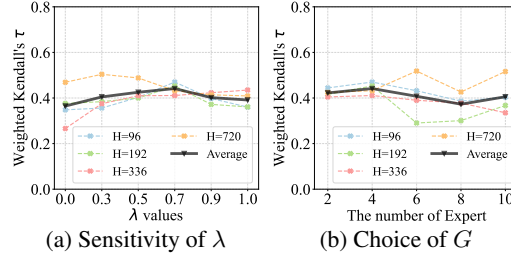


Figure 5: (a) Sensitivity analysis of the loss coefficient λ , (b) choice of the number of experts G .

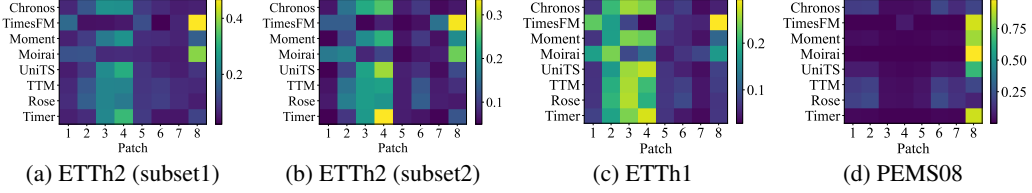


Figure 6: Visualization of patchwise cross-attention weights.

effective multi-task forecasting. Conversely, a large G increases computational cost and model complexity, which may lead to overfitting and degraded performance on downstream tasks. Overall, $G = 4$ provides the best trade-off between accuracy and efficiency, making it ideal for practical use.

Visualization. We visualize the patchwise cross-attention weights in Figure 6, which reveal several key observations: (1) Figure 6a and Figure 6b compare the data embeddings E_d obtained from different ETTh2 subsets. We observe that the weight distributions across different patches for the same model (same row) are quite similar. Furthermore, TTM and ROSE, the two best-performing models on ETTh2, also exhibit similar patchwise weight distributions. This suggests that models with comparable performance tend to share similar patchwise attention patterns within the same dataset. (2) When comparing E_d from two similar datasets, ETTh2 (Figure 6b) and ETTh1 (Figure 6c), we find that the patchwise weight distributions for the same model remain relatively consistent. This indicates that similar attention patterns are exhibited not only across different subsets of the same dataset, but also between similar datasets. (3) In contrast, when comparing two more distinct datasets, ETTh1 (Figure 6c) and PEMS08 (Figure 6d), significant differences emerge in the patchwise weight distributions of the same model. Moreover, among the top-performing models on PEMS08, such as Moirai, TimesFM, Moment, similar patchwise distributions are evident. The best model, Moirai, places the highest attention on the 8-th patch, whereas two weaker models, TTM and ROSE, show considerably lower weights for that patch. This indicates that models with different performance levels tend to focus on different patches.