

Supplementary Material of Frequency Matters: When Time Series Foundation Models Fail Under Spectral Shift

213 A Related Work

214 Following the success of foundation models in computer vision and natural language processing
215 (NLP), and with the rise of time series as a valid modality within the broader context of deep learning,
216 a growing number of foundation models have been proposed to tackle time series. While different
217 downstream applications are possible, the majority of these models have focused on forecasting.
218 Some approaches, such as GPT4TS [18] and Time-LLM [8], adapt existing large language models by
219 “reprogramming” them for the temporal domain, freezing most layers while finetuning lightweight,
220 time series-specific modules. In contrast, a wave of recent work has introduced models trained
221 entirely from scratch on large-scale, heterogeneous time series corpora.

222 Lag-Llama [12] presents a general-purpose foundation model for univariate probabilistic forecast-
223 ing, built on a decoder-only transformer. By explicitly modeling lagged dependencies, it adapts
224 autoregressive language modeling techniques to temporal patterns, achieving strong performance on
225 long-horizon forecasting tasks, though it is primarily designed for forecasting rather than broader
226 time series applications.

227 Chronos [1] adopts a probabilistic formulation by discretizing continuous time series into tokens
228 and applying autoregressive sequence modeling. Released as a family of five models ranging from
229 20M to 710M parameters, Chronos allows for flexible trade-offs between efficiency and accuracy.
230 While forecasting remains central, its probabilistic generation framework also enables extensions to
231 classification and anomaly detection via reinterpretation of generated trajectories.

232 Moirai [13] introduces a scalable temporal foundation model with hierarchical attention tailored to
233 long and irregular sequences. Through multi-resolution temporal representations, Moirai can capture
234 dependencies across a wide range of timescales, making it effective not only for forecasting but also
235 for classification and imputation tasks across domains such as healthcare and finance.

236 TimesFM [4] builds a general-purpose forecaster trained on a large collection of public and pro-
237 prietary datasets. With a transformer-based encoder and decoder backbone, it demonstrates strong
238 zero-shot performance across diverse application domains, including finance, energy, and traffic.
239 Beyond forecasting, TimesFM produces robust embeddings that transfer effectively to downstream
240 classification and anomaly detection tasks, particularly in low-label regimes.

241 MOMENT [14] proposes a universal time series foundation model that unifies forecasting with
242 representation learning across multiple modalities. Combining large-scale pretraining with adaptive
243 fine-tuning, MOMENT supports a broad suite of tasks, including forecasting, classification, anomaly
244 detection, and imputation. Its architecture is explicitly designed for multimodality, enabling the
245 integration of time series with contextual signals such as categorical features or text.

246 In terms of training methodology, these models share the same underlying philosophy of large-
247 scale pretraining on heterogeneous time series data, but diverge in their design choices. Lag-
248 Llama, Chronos, and TimesFM employ transformer-based encoders with objectives such as next-step
249 prediction and masked modeling; Chronos, in particular, leverages tokenization to enable language-
250 model style autoregression. On the other hand, Moirai emphasizes architectural scalability, using
251 hierarchical attention and multi-resolution sampling to pretrain on very long sequences. MOMENT
252 combines forecasting objectives with contrastive pretraining, producing general-purpose temporal
253 embeddings that are transferable across tasks.

254 Overall, a diverse set of methodologies and architectural choices has been proposed to address time
255 series downstream tasks ranging from forecasting to classification and anomaly detection. While these
256 models demonstrate strong performance on widely used benchmark datasets, their evaluation remains
257 confined mainly to controlled settings. In contrast, our work emphasizes real-world applications,

258 focusing on practical use cases that reflect the challenges and constraints encountered in industrial
 259 deployment scenarios, as detailed in the main paper.

260 B Spectral Analysis

261 To better understand why the time series foundation model MOMENT underperforms on our dataset,
 262 we analyze its dominant frequency components and compare them with those of the datasets used
 263 during pretraining. Figure 2 illustrates this comparison, focusing on the FordA and FaultDetectionA
 264 datasets from the UCR time series repository, which were part of MOMENT’s pretraining corpus.

265 As shown in the figure, our dataset exhibits fundamentally different dominant frequencies compared
 266 to those observed during pretraining. This discrepancy suggests that MOMENT was not exposed to
 267 such frequency patterns before, which likely contributed to its poor generalization and performance.

268 This analysis highlights the potential importance of frequency alignment in the generalization
 269 capability of time series foundation models. It also motivates the frequency-based hypothesis we
 270 explore further in Section 3.

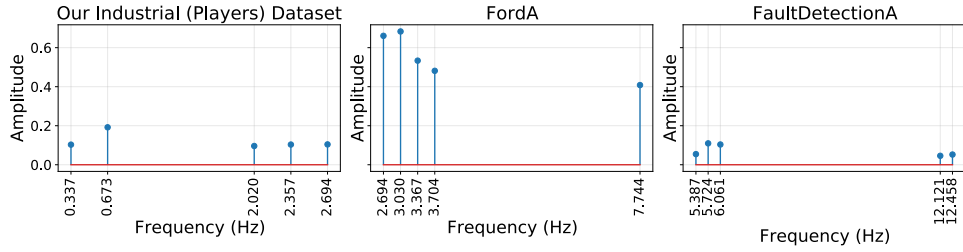


Figure 2: Analysis of the dominant frequencies in our dataset and the used dataset for pretraining MOMENT (FordA and FaultDetectionA).

271 C Data Generation

272 In the following section, and as a complement to Section 3, we provide additional details about the
 273 synthetic data generation process.

274 As previously explained, we constructed *seen* and *unseen* synthetic datasets from a set of considered
 275 datasets that were used to pretrain the considered foundation model. Specifically, for a real sequence
 276 $x(t)$, we compute its discrete Fourier transform (DFT) and retain the top- k dominant components,

$$\mathcal{F}\{x(t)\} \Rightarrow F(x) = \{(f_i, a_i)\}_{i=1}^k, \quad (1)$$

277 where f_i and a_i denote dominant frequencies and amplitudes. Based on the previous frequency extrac-
 278 tion, synthetic signals are generated by recombining these components with controlled perturbations
 279 (additive Gaussian noise and bounded amplitude scaling). This yields:

- 280 • **Seen samples:** dominant content constrained within dataset-level frequency bounds:

$$f_i \in [f_{\text{sim}}^{\text{low}}, f_{\text{sim}}^{\text{high}}],$$

281 with only mild perturbations;

- 282 • **Unseen samples:** at least one dominant component outside the empirical bounds,

$$f_j \notin [f_{\text{sim}}^{\text{low}} + \delta, f_{\text{sim}}^{\text{high}} + \delta],$$

283 or taken from complementary/shifted bands (low-only, high-only, or randomized). For
 284 instance, if the dominant range is $[0, 20]$ Hz, then the unseen dataset would be generated
 285 from $[20 + \delta, 40 + \delta]$ with δ being drawn at random subject to the constraint that the
 286 out-of-band upper limit does not exceed the maximum frequency present in the original
 287 series.

288 We note that the dataset-wide frequency bounds $[f_{\text{sim}}^{\text{low}}, f_{\text{sim}}^{\text{high}}]$ are estimated from the distribution of
 289 dominant frequencies across the dataset (e.g., via quantile summaries or min/max).

290 Based on the generated seen and unseen datasets, we generate a set of train/val/test splits for each
 291 dataset, obtained with a fixed 70/15/15 partition. Each split contains single-channel sequences of
 292 fixed length L (typically $L = 512$),

$$X^{(s)} \in \mathbb{R}^{N_s \times 1 \times L}, \quad s \in \{\text{train, val, test}\}. \quad (2)$$

293 **Regression targets.** Given the previously generated time series, we need to construct labels that
 294 could be used for the downstream task. In this perspective, each synthetic sample is paired with a
 295 continuous regression target derived directly from its generation metadata. Let $\mathcal{F}_{\text{used}} = \{f_j\}$ denote
 296 the set of frequencies used in the synthesis. The scalar target is defined as the sum of frequencies:

$$y = \sum_{f_j \in \mathcal{F}_{\text{used}}} f_j. \quad (3)$$

297 We note that while the previous labeling operation provides a frequency-aware label, it does indeed
 298 naturally increase with spectral displacement and distinguishes between seen and unseen samples.

299 Since such displacement can be used by the model to make the task easier, we use a normalization
 300 mechanism to optimize the labels and make them comparable in terms of scale:

$$\tilde{y}^{(s)} = \frac{y^{(s)} - \mu_y}{\sigma_y}, \quad y = \sigma_y \tilde{y}^{(s)} + \mu_y, \quad (4)$$

301 After this normalization operation, the labels for both seen and unseen datasets are within the same
 302 range, and therefore, we would expect the model to rather use temporal representations to extract the
 303 label.

304 In Figure 3 we present an example of the generation, where given an original sample, we analyze its
 305 original spectrum to generate the seen and unseen datasets.

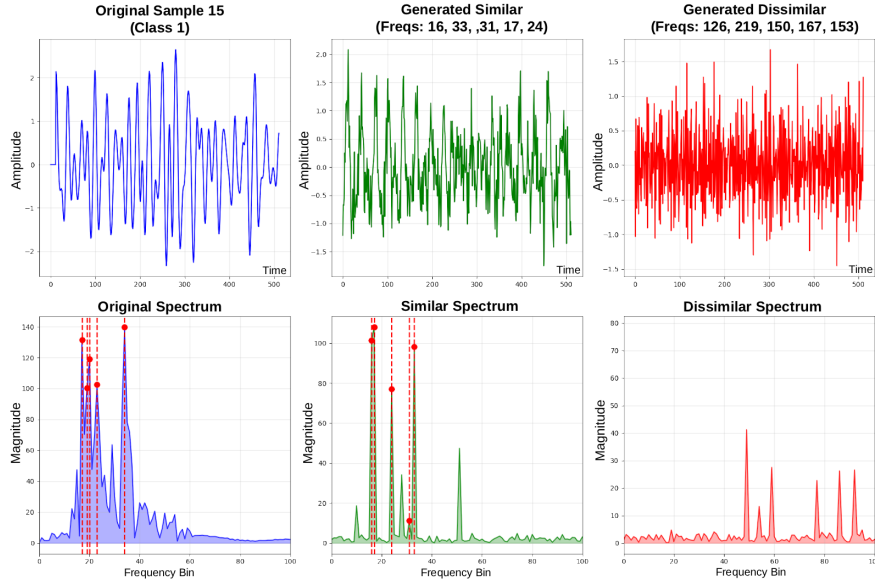


Figure 3: Synthetic series from seen vs. unseen frequency bands. Top: time domain; Bottom: spectra.

306 D Additional Results - Classification

307 To further illustrate our hypothesis beyond the regression task, we considered a classification task.
 308 Specifically, we bin adjacent frequency ranges (e.g., $[0, 10]$ vs. $[10, 20]$ Hz) and train a binary

309 classifier on top of the frozen encoder. The binning threshold is chosen as the median of the frequency
 310 ranges from the train split to ensure a balanced class distribution.

Table 3: Classification performance of a frozen TSFM encoder (MOMENT-small) with a trainable binary classifier on synthetic datasets.

Dataset	Test Accuracy		Test AUC	
	Seen ($\checkmark$)	Unseen ($\times$)	Seen ($\checkmark$)	Unseen ($\times$)
FordA	0.837 ± 0.004	0.829 ± 0.001	0.926 ± 0.002	0.901 ± 0.001
FordB	0.826 ± 0.000	0.838 ± 0.002	0.903 ± 0.002	0.926 ± 0.001
ElectricDevices	0.785 ± 0.003	0.650 ± 0.003	0.890 ± 0.001	0.716 ± 0.002
SmallKitchenAppliances	0.708 ± 0.010	0.756 ± 0.041	0.833 ± 0.018	0.859 ± 0.007
FaultDetectionA	0.766 ± 0.003	0.651 ± 0.003	0.858 ± 0.000	0.707 ± 0.003
FaultDetectionB	0.444 ± 0.048	0.472 ± 0.096	0.519 ± 0.306	0.476 ± 0.190
ECG5000	0.809 ± 0.020	0.591 ± 0.028	0.898 ± 0.007	0.692 ± 0.007
SwedishLeaf	0.689 ± 0.020	0.600 ± 0.040	0.796 ± 0.021	0.643 ± 0.028

311 Table 3 provides the mean and corresponding standard deviation of the Test accuracy and the Test
 312 AUC. As previously seen in the regression task, the performance is high for ranges within the
 313 presumed pretraining spectrum but drops for out-of-range comparisons (e.g., [40, 50] vs. [50, 60] Hz),
 314 reinforcing the frequency-mismatch explanation.

315 E Experimental Details

316 In all experiments, we use linear probing: the MOMENT backbone (i.e., the patch embedder and
 317 transformer encoder) is frozen, and only the regression or classification head is trained. The model
 318 is optimized using the Adam optimizer (learning rate 10^{-3}) for 50 epochs, with a mean squared
 319 error loss (for regression) or binary cross-entropy loss (for classification). We use the same hidden
 320 dimensionality as the one from the pretrained chosen backbone, corresponding to the predefined
 321 MOMENT configurations. Model selection is performed based on validation MSE, and the best
 322 checkpoint is then evaluated on the test set. We report all metrics in Table 1.

323 We use a combination of CPU, NVIDIA GPUs (L4 24GB, T4 16GB), and an Apple MPS device (M1
 324 MAX 32GB) to conduct the experiments in this paper. In the player engagement use case, we use
 325 CPU for the XGBoost model, a single GPU for TabNet, and PatchTST. We utilize distributed data
 326 parallelism with up to eight GPUs to accelerate the training process for the MOMENT model on the
 327 large mobile gaming dataset. In the frequency perspective experiments, we use a single GPU for
 328 regression experiments on the regression task and an Apple MPS device for the regression tasks.

329 F Reflection on the Role of Frequencies in TSFM Pretraining

330 In this work, we report an empirical observation about Time Series Foundation Models (TSFMs): a
 331 pretrained TSFM can underperform on domain-specific downstream tasks relative to models trained
 332 directly on in-domain time series. To probe this effect, we empirically evaluate how well a pretrained
 333 TSFM performs on synthetically generated time series whose dominant frequencies are varied. We
 334 observe a trend in frequency-alignment effect: the TSFM performs better on synthetic time series
 335 whose dominant frequencies overlap with those seen during pretraining than on synthetic time series
 336 samples from unseen or underrepresented frequency bands.

337 Our results suggest that, beyond standard self-supervised objectives for times pretraining (e.g.,
 338 masked modeling, forecasting, and contrastive learning), the spectral composition of the pretraining
 339 data may also play a critical role in the model’s generalization and robustness. One might consider
 340 adding an auxiliary task during pretraining to predict the dominant frequency bands or adopting
 341 frequency-based data augmentation or sampling to expand frequency coverage. We hypothesize that
 342 frequency-aware pretraining strategies, combined with standard time series pretraining methods, may
 343 be a fruitful future direction for improving the performance of TSFMs in zero-shot/few-shot settings.

References

- [1] Abdul Fatir Ansari, Lorenzo Stella, Caner Turkmen, Xiyuan Zhang, Pedro Mercado, Huibin Shen, Oleksandr Shchur, Syama Syndar 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. *Transactions on Machine Learning Research*, 2024.
- [2] Sercan Ö Arik and Tomas Pfister. Tabnet: Attentive interpretable tabular learning. In *Proceedings of the AAAI conference on artificial intelligence*, volume 35, pages 6679–6687, 2021.
- [3] Tianqi Chen and Carlos Guestrin. Xgboost: A scalable tree boosting system. In *Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining*, pages 785–794, 2016.
- [4] Abhimanyu Das, Weihao Kong, Rajat Sen, and Yichen Zhou. A decoder-only foundation model for time-series forecasting. In *Forty-first International Conference on Machine Learning*, 2024.
- [5] Hoang Anh Dau, Anthony Bagnall, Kaveh Kamgar, Chin-Chia Michael Yeh, Yan Zhu, Shaghayegh Gharghabi, Chotirat Ann Ratanamahatana, and Eamonn Keogh. The ucr time series archive. *IEEE/CAA Journal of Automatica Sinica*, 6(6):1293–1305, 2019.
- [6] Sofiane ENNADIR, Siavash Golkar, and Leopoldo Sarra. Joint embedding go temporal. In *NeurIPS Workshop on Time Series in the Age of Large Models*, 2024.
- [7] Kaiming He, Xinlei Chen, Saining Xie, Yanghao Li, Piotr Dollár, and Ross Girshick. Masked autoencoders are scalable vision learners. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 16000–16009, 2022.
- [8] Ming Jin, Shiyu Wang, Lintao Ma, Zhixuan Chu, James Y. Zhang, Xiaoming Shi, Pin-Yu Chen, Yuxuan Liang, Yuan-Fang Li, Shirui Pan, and Qingsong Wen. Time-LLM: Time series forecasting by reprogramming large language models. In *The Twelfth International Conference on Learning Representations*, 2024.
- [9] Christian Lessmeier, James Kuria Kimotho, Detmar Zimmer, and Walter Sextro. Condition monitoring of bearing damage in electromechanical drive systems by using motor current signals of electric motors: A benchmark data set for data-driven classification. In *PHM Society European Conference*, volume 3 1, 2016.
- [10] Jason Lines, Anthony Bagnall, Patrick Caiger-Smith, and Simon Anderson. Classification of household devices by electricity usage profiles. In Hujun Yin, Wenjia Wang, and Victor Rayward-Smith, editors, *Intelligent Data Engineering and Automated Learning - IDEAL 2011*, pages 403–412, Berlin, Heidelberg, 2011. Springer Berlin Heidelberg.
- [11] 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*, 2023.
- [12] Kashif Rasul, Arjun Ashok, Andrew Robert Williams, Arian Khorasani, George Adamopoulos, Rishika Bhagwatkar, Marin Biloš, Hena Ghonia, Nadhir Hassen, Anderson Schneider, Sahil Garg, Alexandre Drouin, Nicolas Chapados, Yuriy Nevmyvaka, and Irina Rish. Lag-llama: Towards foundation models for time series forecasting. In *R0-FoMo: Robustness of Few-shot and Zero-shot Learning in Large Foundation Models*, 2023.
- [13] Gerald Woo, Chenghao Liu, Akshat Kumar, Caiming Xiong, Silvio Savarese, and Doyen Sahoo. Unified training of universal time series forecasting transformers. In Ruslan Salakhutdinov, Zico Kolter, Katherine Heller, Adrian Weller, Nuria Oliver, Jonathan Scarlett, and Felix Berkenkamp, editors, *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pages 53140–53164. PMLR, 21–27 Jul 2024.
- [14] Gilbert Woo et al. Moment: A family of open time-series foundation models. *arXiv preprint arXiv:2310.10688*, 2023.

- 204 [15] George Zerveas, S Jayaraman, Anastasios Patel, Anastasios Bhamidipaty, and Carsten Eickhoff.
205 A transformer-based framework for multivariate time series representation learning. *Proceedings*
206 *of ACM SIGKDD*, 2021.
- 207 [16] Chaoning Zhang, Chenshuang Zhang, Junha Song, John Seon Keun Yi, Kang Zhang, and In So
208 Kweon. A survey on masked autoencoder for self-supervised learning in vision and beyond.
209 *arXiv:2208.00173*, 2022.
- 210 [17] Y Zhang, Y Yan, et al. Patchtst: A time series is worth 64 words. In *ICLR*, 2023.
- 211 [18] Tian Zhou, Ziqing Ma, Qingsong Wen, et al. One fits all: Power general time series analysis by
212 pretrained lm. *Advances in Neural Information Processing Systems*, 2023.