

LETS-C: Leveraging Text Embedding for Time Series Classification

Anonymous ACL submission

Abstract

Recent advancements in language modeling have shown promising results when applied to time series data. In particular, fine-tuning pre-trained large language models (LLMs) for time series classification tasks has achieved state-of-the-art (SOTA) performance on standard benchmarks. However, these LLM-based models have a significant drawback due to the large model size, with the number of trainable parameters in the millions. In this paper, we propose an alternative approach to leveraging the success of language modeling in the time series domain. Instead of fine-tuning LLMs, we utilize a text embedding model to embed time series and then pair the embeddings with a simple classification head composed of convolutional neural networks (CNN) and multilayer perceptron (MLP). We conducted extensive experiments on a well-established time series classification benchmark. We demonstrated LETS-C not only outperforms the current SOTA in classification accuracy but also offers a lightweight solution, using only 14.5% of the trainable parameters on average compared to the SOTA model. Our findings suggest that leveraging text embedding models to encode time series data, combined with a simple yet effective classification head, offers a promising direction for achieving high-performance time series classification while maintaining a lightweight model architecture. The implementation code is available at <https://anonymous.4open.science/r/LETS-C>.

1 Introduction

Time series classification (Bagnall et al., 2017; Abanda et al., 2019; Ismail Fawaz et al., 2019) has gained attention due to its applications in finance (Passalis et al., 2017), healthcare (Lipton et al., 2016), and activity recognition (Yang et al., 2015). The growing availability of time series data has driven demand for efficient and accurate classification methods. Advances in natural language processing (NLP) and large language mod-

els (LLMs) (Achiam et al., 2023) have demonstrated strong promises in modeling sequential data (Achiam et al., 2023). Inspired by this, researchers have explored applying LLMs to time series via prompting (Gruver et al., 2024; Liu et al., 2024; Merrill et al., 2024; Chu et al., 2024) or fine-tuning (Jin et al., 2023; Zhou et al., 2024), achieving state-of-the-art (SOTA) performance on well-established benchmarks for tasks including classification and forecasting.

However, LLMs for time series classification face a major drawback—their large size, often with billions of parameters, making them computationally expensive and impractical in resource-limited settings (Bommasani et al., 2021). Fine-tuning partially frozen pre-trained LLMs still requires millions of trainable parameters (Zhou et al., 2024). To mitigate this, we propose an alternative approach to leverage the success of language modeling in the time series domain. In particular, we propose a novel approach, LETS-C (Leveraging Text Embeddings for Time Series Classification), which utilizes off-the-shelf text embedding models instead of fine-tuning LLMs for time series classification. To the best of our knowledge, this work is the first to explore the potential of text embeddings in time series analysis, specifically classification, and demonstrate SOTA performance.

LETS-C combines text embeddings with a simple yet effective classification head composed of convolutional neural networks (CNNs) and a multilayer perceptron (MLP). By projecting time series data using text embedding models, we capture the intricate patterns and dependencies present in the temporal data. The embeddings and time series are then fed into the classification head, which learns to discriminate between different classes. Through extensive experiments on a well-established benchmark containing time series datasets from various domains, we demonstrated that LETS-C outperforms 27 baselines including the previous SOTA

method. Moreover, LETS-C is significantly more efficient, using much less trainable parameters than the previous SOTA. Our key contributions are:

- **Text Embeddings for Time Series:** We introduce LETS-C, the first work to leverage text embeddings for time series analysis, specifically for classification tasks.
- **State-of-the-Art Performance:** LETS-C achieves SOTA performance in classification accuracy on a well-established benchmark containing time series datasets across different domains, surpassing 27 baseline models.
- **Computationally Lightweight:** LETS-C is significantly more efficient, achieving higher accuracy while using much fewer trainable parameters (14.5%) than the existing SOTA method.

In addition, we conducted comprehensive analyses to showcase the effectiveness of LETS-C:

- **Intrinsic Discriminative Power of Time Series Embeddings:** We demonstrate the advantage of text embeddings for time series, showing that embeddings of time series from the same class are more similar than those from different classes, explaining the boost in classification accuracy.
- **Generalization Across Various Text Embedding Models:** LETS-C with different text embedding models consistently outperforms previous SOTA with much fewer trainable parameters, further validating the effectiveness of our approach.
- **Optimizing Model Size with Minimal Accuracy Loss:** LETS-C retains a high percentage of accuracy even as the classification model size shrinks considerably, enhancing computational efficiency with minimal impact on performance.

2 Related Work

We review the related works in three key areas: time series classification, the application of language models to time series data, and text embeddings. See Appendix A for an extended review.

Time Series Classification Time series classification has been an active research area for decades. Early methods focused on distance-based approaches (Abanda et al., 2019), such as Dynamic Time Warping (Berndt and Clifford, 1994) and distance kernels with Support Vector Machines (Kam-pouraki et al., 2008). Others used feature extraction with classifiers like eXtreme Gradient Boosting (XGBoost) (Chen and Guestrin, 2016) and

LightGBM (Ke et al., 2017). Deep learning-based approaches, including CNNs (Wu et al., 2022a; Zhao et al., 2017), MLPs (Zhang et al., 2022a), and Recurrent Neural Networks (RNNs) like Long Short-Term Memory (LSTM) (Lai et al., 2018), later gained popularity for learning complex patterns and handling long sequences. More recently, Transformer-based models (Vaswani et al., 2017) have been adapted from NLP to time series (Nie et al., 2023; Zhou et al., 2022; Zhang et al., 2022b; Eldele et al., 2021; Wu et al., 2021; Zhou et al., 2021; Koval et al., 2024), leveraging self-attention for long-range dependencies. Additionally, unsupervised representation learning methods pre-train models with masked time series modeling to minimize reconstruction error, followed by fine-tuning for downstream tasks like classification (Franceschi et al., 2019; Tonekaboni et al., 2021; Yue et al., 2022; Eldele et al., 2021; Zerveas et al., 2021; Goswami et al., 2024). However, these complex models often require larger sizes and higher computational costs, particularly for training.

Language Models for Time Series The success of LLMs in NLP has inspired their application in time series analysis. Surveys (Zhang et al., 2024; Jiang et al., 2024) provide insights into key methodologies, challenges, and future directions. Recent studies explore integrating time series and language, including structured time series-text modeling (Khadanga et al., 2019; Deznabi et al., 2021), natural language descriptions of time series (Murakami et al., 2017; Jhamtani and Berg-Kirkpatrick, 2021; Fons et al., 2024), and various applications (Li et al., 2024a; Drinkall et al., 2024; Kavarada et al., 2024). Pre-trained LLMs have been used for time series forecasting via prompting (Gruver et al., 2024; Liu et al., 2024; Cao et al., 2023), while (Yu et al., 2023) explored their potential for generating explainable financial forecasts. Time-LLM (Jin et al., 2023) maps time series to the language embedding space, enabling LLM-based forecasting (Yang et al., 2021). More importantly, recent work, OneFitsAll (Zhou et al., 2024) achieved SOTA performance on various time series tasks by fine-tuning LLMs like GPT (Radford et al., 2019). In contrast, we propose leveraging text embeddings instead of directly using LLMs for time series analysis. Our approach establishes new SOTA performance on time series classification tasks with significantly fewer trainable parameters, making it a more efficient alternative.

Text Embeddings Text embeddings are crucial in NLP, mapping words or sentences into dense vector spaces to capture semantic and syntactic information. Text embedding techniques range from word-level embeddings like Word2Vec (Mikolov et al., 2013) and GloVe (Pennington et al., 2014) to contextualized embeddings from pre-trained models like BERT (Devlin et al., 2019) and RoBERTa (Liu et al., 2019b). In time series, unsupervised methods have been proposed for learning embeddings (Franceschi et al., 2019; Eldele et al., 2021; Zerveas et al., 2021; Yue et al., 2022; Sun et al., 2023; Goswami et al., 2024). However, large-scale datasets are scarcer in time series than in NLP, making it more challenging to learn embeddings from scratch. To our knowledge, we are the **first** to leverage well-trained text embeddings from NLP for time series classification.

3 Methodology

Given a time series classification dataset $\mathcal{D} = \{(\mathbf{x}_i, y_i)_{i=1}^N\}$, where $\mathbf{x}_i$ is a multivariate time series sample, and $y_i \in \{1, 2, \dots, C\}$ is the corresponding class label, the goal is to learn a classifier that accurately predicts the class label $\hat{y}_i$ for each time series. As illustrated in Figure 1, we propose LETS-C framework that harnesses text embeddings for time series classification tasks. Specifically, we 1) initially preprocess the time series data to normalize it, then 2) subsequently generate text embeddings from the normalized time series, 3) fuse embeddings with the time series data, and finally 4) feed the fused representation to a classification head that consists of CNNs and MLP. The choice of a simple classification head is intentional, we aim to test the hypothesis that the text embeddings of the time series provide sufficiently powerful representations for effective classification.

Preprocessing To ensure consistent scales across all model inputs, each feature dimension of time series $\mathbf{x}_i$ is min-max normalized to the range $[0, 1]$ based on the minimum and maximum feature values of each dimension across the training data.

Text Embedding of Time Series It is crucial to carefully format the preprocessed time series into strings before using text embeddings, as the tokenization of numerical strings can significantly affect the embeddings. (Liu and Low, 2023) has shown that tokenization impacts a model’s arithmetic abilities, with commonly used subword tokenization methods like Byte Pair Encoding (BPE)

arbitrarily subdividing numbers, causing similar numbers to appear very differently. To mitigate this, we adopted a digit-space tokenization strategy, as suggested by (Gruver et al., 2024), where each digit is spaced, commas are added to separate time steps, and decimal points are omitted for fixed precision. For instance, a series to be formatted with a precision of two decimal places, such as 0.645, 6.45, 64.5, 645.0, would be converted to "6 4 , 6 4 5 , 6 4 5 0 , 6 4 5 0 0" prior to tokenization. This method ensures separate tokenization of each digit, preserving numerical integrity and enhancing pattern recognition in language models.

Next, we utilized the text-embedding-3-large model (OpenAI, 2024) to embed the formatted time series into the embedding space. It is important to note that we are using only the text embedding model, unlike the LLMs used in previous works (Jin et al., 2023; Zhou et al., 2024; Gruver et al., 2024). This model was selected for several reasons: it is highly ranked on the Massive Text Embedding Benchmark (MTEB) leaderboard (Muennighoff et al., 2022; MTEB, 2024), known for its effectiveness in a variety of downstream tasks such as text search and sentence similarity; it supports a high maximum token length of 8191, accommodating our time series datasets; and it offers a high-dimensional vector space of 3072 dimensions. This large dimensionality captures a broad spectrum of temporal features, additionally the model allows for truncation to reduce dimensions as needed for specific applications without substantial loss of semantic information (Kusupati et al., 2022). This capability to truncate dimensions is particularly advantageous for optimizing efficiency while maintaining robust performance, aligning well with our goal of a lightweight framework.

In particular, we generate a text embedding for each dimension of $\mathbf{x}_i$, transforming $\mathbf{x}_i \in \mathbb{R}^{d \times l_x}$ into embeddings $\mathbf{e}_i \in \mathbb{R}^{d \times l_e}$. Here, d is the number of dimensions in the multivariate time series, l_x is the series length, and l_e is the embedding length. Specifically, each dimension of $\mathbf{x}_i$ (of length l_x) is first converted into a string, resulting in d separate strings. Each string is then independently embedded using an embedding model, as studies (Nie et al., 2023; Zhou et al., 2024; Goswami et al., 2024) suggest that channel-wise modeling is an effective strategy for multivariate time series. Our approach differs slightly depending on the type of embedding model used: 1) *Proprietary model*: For

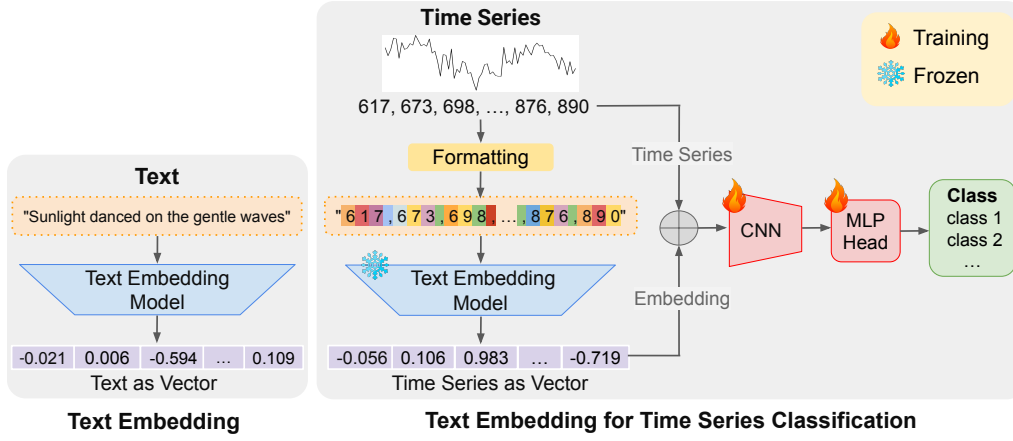


Figure 1: **Left:** Conventional text embedding. **Right:** Our proposed LETS-C framework normalizes time series and formats them to tokenize each digit separately. It then embeds the time series, fuses the embeddings with the original series via element-wise addition, and uses a simple CNN-MLP classification head to perform classification. Only the lightweight CNN and MLP head are trained. For illustration, we show a single-dimensional time series $\mathbf{x}_i \in \mathbb{R}^{l_x}$, transformed into an embedding vector $\mathbf{e}_i \in \mathbb{R}^{l_e}$.

OpenAI’s text-embedding-3-large, the API directly returns an embedding vector of size l_e for each string; 2) *Open-weight models*: We extract the last token’s embedding from the model’s final hidden states to represent the entire sequence, effectively transforming the dimension from l_x to l_e . This approach captures the final contextual representation by condensing the sequence into a single embedding derived from the last token. Embeddings for each dimension are then combined into a $d \times l_e$ matrix, effectively transforming the input from $d \times l_x$ to $d \times l_e$. This transformation preserves independent contextual representations for each dimension while ensuring a consistent embedding structure. Note that the embedding computation in LETS-C is a one-time pass, then the computed embeddings of the time series can be stored and reused for further training, in contrast to the persistent computational cost caused by fine-tuning parts of LLMs. To validate the suitability of off-the-shelf text embeddings for time series classification, we compared cosine similarities between text embeddings of time series from the same and different classes. Intra-class similarity was consistently higher, demonstrating their effectiveness (see Section 5.2).

Fusing Embedding and Time Series Next, we fuse the embedding with the preprocessed time series using element-wise addition, applying zero padding for dimensional consistency. The time series embedding vector serves as a feature representation of temporal patterns present in the entire time series, and adding the raw time series integrates both information sources. This approach is well-established, particularly in ResNet (He et al.,

2016), where raw data is combined with feature representations via shortcut connections. Additionally, fusing embeddings from different modalities is widely supported (Manzoor et al., 2023; Guo et al., 2019; Poria et al., 2018; Baltrušaitis et al., 2018; Jabri et al., 2016), with element-wise addition being a common and effective method for combining multimodal embeddings. This fusion allows both time series and text embeddings to contribute meaningfully to the final representation, improving the model’s overall performance. As shown in Section 5.3, using either the text embedding or the raw time series alone, as well as alternative fusion methods to addition, is less effective. While element-wise addition is an empirical design choice in our architecture, our experiments demonstrate that it achieves SOTA performance while maintaining minimal model complexity compared to alternatives such as concatenation, or fusion networks, or cross-attention.

Lightweight Classification Head Lastly, we pair the fused time series representation with a simple classification head composed of 1D CNNs and an MLP for time series classification. The output from CNNs are flattened and fed through the final MLP head with a softmax activation, which outputs a vector of the probabilities of each time series class. As detailed in the Appendix I, hyperparameter search determines the number of convolutional blocks in CNN, the number of linear layers in MLP, and the use of batch normalization, dropout, activation, and pooling. With a simple classification head, our model is lightweight and requires much less trainable parameters compared to the existing SOTA built on transformers, as detailed in Section 5.1.

4 Experimental Protocol and Details

Datasets and Evaluation Metrics We evaluated LETS-C against a well-established benchmark for multivariate time series classification, commonly adopted for comparison in various studies (Zhou et al., 2024; Li et al., 2024b; Wu et al., 2022a; Zerveas et al., 2021). This benchmark, selected from the UEA Archive (Bagnall et al., 2018), is purposefully curated to cover challenging and diverse domains, including the following datasets: Ethanol-Concentration (Large et al., 2018), FaceDetection (Rik Henson, 2023), Handwriting (Shokoohi-Yekta et al., 2017), Heartbeat (Liu et al., 2016), JapaneseVowels (Kudo et al., 1999), PEMS-SF (Cuturi, 2011), SelfRegulationSCP1 (Birbaumer et al., 1999), SelfRegulationSCP2 (Birbaumer et al., 1999), SpokenArabicDigits (Bedda and Hammami, 2010), and UWaveGestureLibrary (Liu et al., 2009). This commonly adopted benchmark offers a comprehensive testing environment, with multivariate dimensions ranging from 3 to 963, time series lengths up to 1751, and up to 26 classes. See Appendix B for details on each dataset.

To assess the classifiers, we used metrics including classification accuracy and *AvgWins*. *AvgWins* is defined as the average number of times that a method outperforms other methods across benchmarked datasets, with ties also being counted towards this average. Additionally, we analyzed models’ computational efficiency in terms of trainable model parameters, training, and inference time.

Baselines We included 27 baseline models to ensure a comprehensive comparison. We utilize the same supervised learning baselines outlined by (Zhou et al., 2024; Wu et al., 2022a), namely: **Classical methods:** 1) Dynamic Time Warping (DTW) (Berndt and Clifford, 1994), 2) eXtreme Gradient Boosting (XGBoost) (Chen and Guestrin, 2016), and 3) RandOm Convolutional KERNel Transform (ROCKET) (Dempster et al., 2020); **MLP-based methods:** 4) LightTS (Zhang et al., 2022a), and 5) DLinear (Zeng et al., 2023); **RNN-based models:** 6) Long Short-Term Memory (LSTM) (Hochreiter and Schmidhuber, 1997), 7) Long- and Short-term Time-series Network (LSTNet) (Lai et al., 2018), and 8) Linear State Space Layer (LSSL) (Gu et al., 2021); **CNN-based models:** 9) Temporal Convolutional Network (TCN) (Franceschi et al., 2019), and 10) TimesNet (Wu et al., 2022a); **Transformer-based models:** 11) Transformer (Vaswani et al., 2017), 12) Reformer (Kitaev et al., 2020), 13) In-

former (Zhou et al., 2021), 14) Pyraformer (Liu et al., 2021), 15) Autoformer (Wu et al., 2021), 16) Non-stationary Transformer (Liu et al., 2022), 17) FEDformer (Zhou et al., 2022), 18) ETSformer (Woo et al., 2022), 19) Flowformer (Wu et al., 2022b), 20) Patch Time Series Transformer (PatchTST) (Nie et al., 2023); and **LLM-based model:** 21) OneFitsAll (Zhou et al., 2024). For details, see Appendix C. Note that some of these methods were adapted from forecasting to classification tasks, without altering the core design of the models (Appendix D). Additionally, we included the unsupervised representation learning baselines outlined by (Goswami et al., 2024), namely: **CNN-based methods:** 22) T-Loss (Franceschi et al., 2019), 23) Temporal Neighborhood Coding (TNC) (Tonekaboni et al., 2021), 24) TS2Vec (Yue et al., 2022); **Transformer-based models:** 25) Time Series representation learning framework via Temporal and Contextual Contrasting (TS-TCC) (Eldede et al., 2021), 26) Time Series Transformer (TST) (Zerveas et al., 2021), and 27) MOMENT (Goswami et al., 2024). For details, refer to Appendices E (unsupervised baselines) and F (adapting unsupervised models for classification). Reproduction details for baselines and LETS-C implementation details are in Appendices G and H, resp.

5 Results and Analysis

5.1 Performance and Efficiency

Comparison to State-of-the-art Table 1 and Figure 3 present a comparative analysis of LETS-C against 27 baseline models on the benchmark introduced above. We observe that LETS-C consistently demonstrates robust performance across all datasets, achieving the highest average accuracy of 76.16% and *AvgWins* of 40%, compared to 27 benchmark models. This includes the most recent SOTA model OneFitsAll (accuracy: 73.97%, *AvgWins*: 20%) and an older SOTA TimesNet (accuracy: 73.57%, *AvgWins*: 0%). Notably, LETS-C surpasses OneFitsAll on six out of ten datasets by a significant margin. LETS-C is particularly effective on challenging datasets like PEMS-SF and EthanolConcentration (EC). PEMS-SF is characterized by exceptionally high dimensionality with 963 features, and EC contains an extremely long time series at length of 1751. These results showcase the competitive edge of LETS-C against the previous SOTA methods, thus establishing a new benchmark for time series classification.

Table 1: Comparison of classification accuracy (%) and AvgWins (%). **Red:** Best, **Blue:** Second best. **Abbreviations:** EC: Ethanol Concentration, FD: Face Detection, HW: Handwriting, HB: Heartbeat, JV: Japanese Vowels, SCP1: Self-Regulation SCP1, SCP2: Self-Regulation SCP2, SAD: Spoken Arabic Digits, UW: UWave Gesture Library.

Model/Dataset		EC	FD	HW	HB	JV	PEMS-SF	SCP1	SCP2	SAD	UW	Average ↑	AvgWins % ↑
<i>Supervised Learning Methods</i>													
Classical methods	DTW (Berndt and Clifford, 1994)	32.3	52.9	28.6	71.7	94.9	71.1	77.7	53.9	96.3	90.3	66.97	0%
	XGBoost (Chen and Guestrin, 2016)	43.7	63.3	15.8	73.2	86.5	98.3	84.6	48.9	69.6	75.9	65.98	10%
	ROCKET (Dempster et al., 2020)	45.2	64.7	58.8	75.6	96.2	75.1	90.8	53.3	71.2	94.4	72.53	20%
MLP	LightTS (Zhang et al., 2022a)	29.7	67.5	26.1	75.1	96.2	88.4	89.8	51.1	100	80.3	70.42	10%
	DLinear (Zeng et al., 2023)	32.6	68	27	75.1	96.2	75.1	87.3	50.5	81.4	82.1	67.53	0%
RNN	LSTM (Hochreiter and Schmidhuber, 1997)	32.3	57.7	15.2	72.2	79.7	39.9	68.9	46.6	31.9	41.2	48.56	0%
	LSTNet (Lai et al., 2018)	39.9	65.7	25.8	77.1	98.1	86.7	84	52.8	100	87.8	71.79	10%
	LSSL (Gu et al., 2021)	31.1	66.7	24.6	72.7	98.4	86.1	90.8	52.2	100	85.9	70.85	10%
CNN	TCN (Franceschi et al., 2019)	28.9	52.8	53.3	75.6	98.9	68.8	84.6	55.6	95.6	88.4	70.25	0%
	TimesNet (Wu et al., 2022a)	35.7	68.6	32.1	78	98.4	89.6	91.8	57.2	99	85.3	73.57	0%
Transformer	Transformer (Vaswani et al., 2017)	32.7	67.3	32	76.1	98.7	82.1	92.2	53.9	98.4	85.6	71.9	0%
	Reformer (Kitaev et al., 2020)	31.9	68.6	27.4	77.1	97.8	82.7	90.4	56.7	97	85.6	71.52	0%
	Informer (Zhou et al., 2021)	31.6	67	32.8	80.5	98.9	81.5	90.1	53.3	100	85.6	72.13	20%
	Pyrformer (Liu et al., 2021)	30.8	65.7	29.4	75.6	98.4	83.2	88.1	53.3	99.6	83.4	70.75	0%
	Autoformer (Wu et al., 2021)	31.6	68.4	36.7	74.6	96.2	82.7	84	50.6	100	85.9	71.07	10%
	Non-stationary Transformer (Liu et al., 2022)	32.7	68	31.6	73.7	99.2	87.3	89.4	57.2	100	87.5	72.66	20%
	FEDformer (Zhou et al., 2022)	31.2	66	28	73.7	98.4	80.9	88.7	54.4	100	85.3	70.66	10%
	ETSformer (Woo et al., 2022)	28.1	66.3	32.5	71.2	95.9	86	89.6	55	100	85	70.96	10%
	Flowformer (Wu et al., 2022b)	33.8	67.6	33.8	77.6	98.9	83.8	92.5	56.1	98.8	86.6	72.95	0%
	PatchTST (Nie et al., 2023)	26.2	68.5	25.9	66.8	96.0	87.9	85.7	53.3	97.2	85.0	69.25	0%
LLM	OneFitsAll (Zhou et al., 2024)	34.2	69.2	32.7	77.2	98.6	87.9	93.2	59.4	99.2	88.1	73.97	20%
<i>Unsupervised Representation Learning Methods</i>													
CNN	T-Loss (Franceschi et al., 2019)	20.5	51.3	45.1	74.1	98.9	67.6	84.3	53.9	90.5	87.5	67.37	0%
	TNC (Tonekaboni et al., 2021)	29.7	53.6	24.9	74.6	97.8	69.9	79.9	55.0	93.4	75.9	65.47	0%
	TS2Vec (Yue et al., 2022)	30.8	50.1	51.5	68.3	98.4	68.2	81.2	57.8	98.8	90.6	69.57	0%
Transformer	TS-TCC (Eldele et al., 2021)	28.5	54.4	49.8	75.1	93.0	73.4	82.3	53.3	97.0	75.3	68.21	0%
	TST (Zerveas et al., 2021)	26.2	53.4	22.5	74.6	97.8	74.0	75.4	55.0	92.3	57.5	62.87	0%
	MOMENT (Goswami et al., 2024)	35.7	63.3	30.8	72.2	71.6	89.6	84.0	47.8	98.1	90.9	68.4	0%
LETS-C (Ours)		52.9	68.9	23.8	78	99.2	93.1	93.2	62.8	99.2	90.6	76.17	40%

Table 2: Comparison of trainable parameters (millions) for LETS-C vs. OneFitsAll. The Ratio (%) = $100 \times \frac{\text{\# of trainable parameters in LETS-C}}{\text{\# of trainable parameters in OneFitsAll}}$, quantifying the efficiency of LETS-C relative to OneFitsAll.

Model/Dataset		EC	FD	HW	HB	JV	PEMS-SF	SCP1	SCP2	SAD	UW	Average ↓
Trainable parameters (M)	LETS-C (Ours)	0.28	0.003	0.15	0.04	0.14	0.56	0.30	0.33	0.14	0.26	0.22
	OneFitsAll	1.42	2.37	1.73	2.03	1.32	10.23	0.98	1.04	1.82	1.0	2.39
	Ratio (%)	19.89	0.16	8.89	2.28	11.19	5.51	30.83	32.06	7.77	26.21	14.48

Computational Efficiency Next, we aim to assess how well LETS-C balances performance with computational efficiency, which is crucial for usage in resource-constrained environments. Table 2 provides a detailed analysis of the trainable parameters associated with LETS-C compared to the previous SOTA model, OneFitsAll. Our method achieved higher performance with only 14.48% of the trainable model parameters on average, compared to OneFitsAll. Despite the advantage of OneFitsAll over other leading models like TimesNet and FEDformer on its reduced parameter count, OneFitsAll still requires much more trainable parameters than our approach. Detailed training and inference time comparisons are reported in Appendix K. We show that LETS-C offers a lightweight approach to time series classification while achieving the SOTA accuracy. It’s important to note that in LETS-C, the

text embedding computation is a one-time operation, unlike models like OneFitsAll, which continue to incur computational costs while fine-tuning partially frozen LLMs.

5.2 Effectiveness of LETS-C

Intrinsic Discriminative Power of Text Embeddings on Time Series To analyze text embeddings extracted from time series and assess their effectiveness, we compared embeddings of time series from the same and different classes. This study revealed the built-in power of text embeddings to readily distinguish time series from different classes. Specifically, we computed the average cosine similarity for time series pairs within the same class and across different classes. We average similarities from multiple channels to handle multivariate time series effectively. This approach

helps us quantify how closely time series in the embedding space are related, both within each class (intra-class) and across distinct classes (inter-class). Figure 2 (left) visualizes these embedding similarities through heatmaps. The heatmaps are scaled using min-max normalization, where warmer colors in the heatmap represent higher similarities and cooler shades indicate lower similarities. Diagonal entries show intra-class similarities, while off-diagonal entries reveal inter-class similarities. We observe that intra-class similarity consistently exceeds inter-class similarity, demonstrating that text embeddings effectively retain and convey significant information from the underlying time series.

Generalization Across Various Text Embedding Models

To assess the generalization of our approach across various text embedding models beyond text-embedding-3-large, we evaluate LETS-C using three additional embeddings: e5-mistral-7b-instruct (Wang et al., 2024), gte-large-en-v1.5 (Li et al., 2023), and nomic-embed-text-v1 (Nussbaum et al., 2024). These models were selected for their high rankings in the MTEB overall leaderboard and their variations in embedding dimensions, maximum token lengths, and model sizes. Details on embedding models are in Appendix O. Table 3 presents detailed accuracy metrics and trainable parameters for these various embedding models in the LETS-C framework. We observe that LETS-C consistently outperforms current baselines in terms of average classification accuracy and AvgWins, while utilizing only a fraction of the trainable model parameters across all explored text embedding models. Specifically, text-embedding-3-large achieves an average accuracy of 76.17%, while e5-mistral-7b-instruct, gte-large-en-v1.5, & nomic-embed-text-v1 achieve accuracies of 74.89%, 76.02%, and 75.12% respectively. These results not only surpass the previous SOTA accuracy of 73.97% but also require significantly fewer trainable model parameters—just 14.48%, 15.62%, 9.24%, and 5.31% compared to OneFitsAll. Among these text embedding variations, higher text embedding dimensionality leads to more trainable parameters on average. e5-mistral-7b-instruct requires the most trainable parameters (0.25M) due to its 4096 dimensions, followed by text-embedding-3-large (3072D, 0.22M), gte-large-en-v1.5 (1024D, 0.19M), and

nomic-embed-text-v1 (768D, 0.09M). Consequently, our approach generalizes across diverse text embedding models, demonstrates superior performance, with the benefit of being lightweight.

Optimizing Model Size with Minimal Accuracy Loss

Next, we examine the trade-offs between model accuracy and model size in our approach. To vary model size, we adjust the number of linear and convolution layers in the classification head, ranging from 1 to 5 layers each, creating model variants of different sizes. Across all datasets, we find that significant reductions in model parameters result in only a minimal loss of accuracy. Figure 2 (right) illustrates this trade-off, showing accuracy and parameter retention relative to the optimal performance of LETS-C (see Tables 1 and 2 for reference optimal values). The results highlight the efficiency of LETS-C, demonstrating its ability to maintain high accuracy with significantly fewer parameters across all datasets. While the trade-off between model size and accuracy is data-dependent, in general, substantial parameter reductions lead to only slight accuracy decreases. For detailed results, see Appendix P and Table 10. This analysis confirms that while reducing parameters can lower accuracy, the impact remains manageable, making LETS-C suitable for applications requiring efficient inference.

5.3 Additional Analysis

Ablation Study To empirically evaluate the benefits of fusing text embeddings with time series over variants using only one modality, we conduct an ablation study. As shown in Appendix Q and Table 11, we observe that the combination of both embeddings and time series achieves the highest average accuracy, at 76.17%, and the best number of AvgWins, compared to the ablated versions. These demonstrate the significant performance gains from fusing embeddings with time series data, which are essential for optimal model accuracy.

Alternative Methods for Fusing Time Series with Embeddings

We explore two additional methods for fusing time series and embeddings beyond simple addition. The first method involves a *Fusion network* that first processes embeddings and time series data through convolutional and dense layers in two separate branches, then merges the features from both branches into a final dense network. The second method employs *Concatenation*, where the time series and embeddings are concatenated and

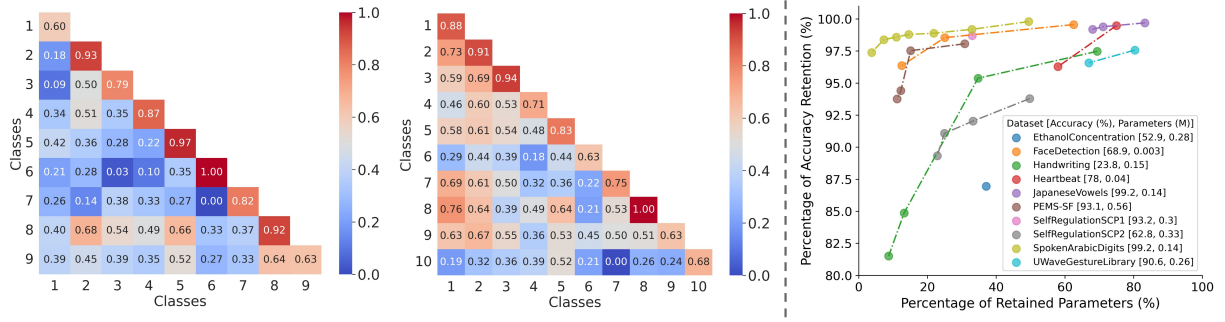


Figure 2: **Left:** Heatmaps illustrating within- and between-class cosine similarities of text embeddings derived from the training time series in JapaneseVowels (**far left**) with 9 classes and SpokenArabicDigits (**middle**) with 10 classes. **Right:** Trade-off between the percentage of accuracy retention and model parameter retention relative to LETS-C’s optimal values. The optimal LETS-C accuracy (%) and parameters (M) for each dataset are in the legend.

Table 3: Comparison of LETS-C with various embedding models against OneFitsAll. Performances surpassing OneFitsAll are in **bold**, with the best in **Red**. AvgWins scores above 50% indicate consistent superiority, calculated as 1 for outperforming OneFitsAll and 0 otherwise. Ratio (%) = $100 \times \frac{\text{\# of trainable parameters in LETS-C}}{\text{\# of trainable parameters in OneFitsAll}}$ measures computational efficiency relative to OneFitsAll.

Accuracy ↑													
Method/Dataset	EC	FD	HW	HB	JV	PEMS-SF	SCP1	SCP2	SAD	UW	Average ↑	AvgWins % ↑	
OneFitsAll	34.2	69.2	32.7	77.2	98.6	87.9	93.2	59.4	99.2	88.1	73.97	—	
LETS-C	text-embedding-3-large	52.9	68.9	23.8	78	99.2	93.1	93.2	62.8	99.2	90.6	76.17	80%
	e5-mistral-7b-instruct	55.5	68.7	23.3	77.6	99.2	84.4	93.9	59.4	98.5	88.4	74.89	60%
	gte-large-en-v1.5	57.8	68.8	24.7	77.6	98.4	91.3	94.2	60	99	88.4	76.02	60%
	nomi-c-embed-text-v1	52.9	68	24.8	76.6	99.2	88.4	93.9	59.4	98.6	89.4	75.12	60%
Trainable Parameters (M) ↓													
Method/Dataset	EC	FD	HW	HB	JV	PEMS-SF	SCP1	SCP2	SAD	UW	Average ↓		
OneFitsAll	1.42	2.37	1.73	2.03	1.32	10.23	0.98	1.04	1.82	1.0	2.39		
LETS-C	text-embedding-3-large	0.28	0.003	0.15	0.04	0.14	0.56	0.30	0.33	0.14	0.26	0.22	
	Ratio %	19.89	0.16	8.89	2.28	11.19	5.51	30.83	32.06	7.77	26.21	14.48	
	e5-mistral-7b-instruct	0.13	0.40	0.39	0.17	0.16	0.30	0.24	0.24	0.33	0.16	0.25	
	Ratio %	9.48	16.93	22.78	8.54	12.55	2.97	25.06	23.59	18.39	15.93	15.62	
	gte-large-en-v1.5	0.18	0.31	0.31	0.12	0.04	0.56	0.03	0.07	0.22	0.09	0.19	
	Ratio %	12.91	13.44	18.27	6.11	3.36	5.51	3.77	7.65	12.34	9.00	9.24	
	nomi-c-embed-text-v1	0.03	0.03	0.36	0.07	0.05	0.19	0.02	0.10	0.06	0.02	0.09	
	Ratio %	2.21	1.31	20.99	3.58	3.99	1.85	3.02	10.03	3.34	2.78	5.31	

processed through a lightweight classification head. Despite cross-attention being another alternative for fusing different modalities, we didn’t include it in this study due to the computational complexity it adds to the model. Appendix R and Table 12 present details on classification accuracy and trainable model parameters for these variations. We observe that the addition approach in the LETS-C architecture achieves the highest average classification accuracy (76.11%) compared to the fusion network (73.40%) and concatenation (74.22%). It also records the best AvgWins at 70%. Further, the number of parameters increases with both the concatenation and fusion network approaches, resulting in additional complexity compared to the addition method. As our goal was to develop a lightweight model, we opted for the addition approach due to its simpler parameter structure.

6 Conclusion

We introduced LETS-C, a novel approach that leverages text embeddings for time series classification. To our knowledge, this is the first exploration of using off-the-shelf text embeddings in time series analysis, particularly for classification. By projecting time series through text embedding models and employing a simple yet effective classification head, LETS-C achieves SOTA performance on a well-established benchmark covering various domains, surpassing 27 baselines. Additionally, LETS-C is significantly more lightweight than previous SOTA methods, achieving higher accuracy with fewer trainable parameters, as well as less training and inference time. Our comprehensive analysis highlights LETS-C’s robustness across different text embedding models, the advantage of using text embeddings for time series classification, and the trade-off between accuracy and model size.

7 Limitations and Broader Perspective

Limitations and future work: Our approach has a few key limitations. One is the maximum sequence length constraint of text embedding models, which restricts the handling of longer time series and may necessitate truncation or downsampling, potentially affecting the capture of long-term dependencies. Additionally, while we adhered to established benchmarks for consistency, we did not evaluate our method on a newly constructed or strictly held-out dataset to check for potential data leakage into the OpenAI embeddings. Testing on a novel dataset could further validate our results and ensure that the approach’s effectiveness is not influenced by pre-existing knowledge. Another limitation is the monetary cost associated with using OpenAI’s text embeddings, which we discuss in more detail in Appendix L.

Future research could explore alternative time series tokenization strategies, as well as extensions to other time series tasks. We believe our findings will inspire further exploration of text embeddings in the time series domain, paving the way for more powerful and efficient methods for various time series tasks.

Broader perspective: The embeddings generated may lack interpretability, making it difficult to understand model decisions in high-stakes applications. Additionally, the resources accompanying this study will be responsibly released for research purposes only.

Datasets and code: The benchmarks used in this study are publicly available from the UEA Archive (Bagnall et al., 2018) and were curated by previous research. Specifically, they include the following datasets: EthanolConcentration (Large et al., 2018), FaceDetection (Rik Henson, 2023), Handwriting (Shokoohi-Yekta et al., 2017), Heartbeat (Liu et al., 2016), JapaneseVowels (Kudo et al., 1999), PEMS-SF (Cuturi, 2011), SelfRegulation-SCP1 (Birbaumer et al., 1999), SelfRegulation-SCP2 (Birbaumer et al., 1999), SpokenArabicDigits (Bedda and Hammami, 2010), and UWaveGestureLibrary (Liu et al., 2009). We abide by their terms of use. The implementation code is available at <https://anonymous.4open.science/r/LETS-C>. To support reproducibility and facilitate future research, we will make the resources associated with this study available upon acceptance.

References

- Amaia Abanda, Usue Mori, and Jose A Lozano. 2019. A review on distance based time series classification. *Data Mining and Knowledge Discovery*, 33(2):378–412.
- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Anthony Bagnall, Hoang Anh Dau, Jason Lines, Michael Flynn, James Large, Aaron Bostrom, Paul Southam, and Eamonn Keogh. 2018. The uea multi-variate time series classification archive, 2018. *arXiv preprint arXiv:1811.00075*.
- Anthony Bagnall, Jason Lines, Aaron Bostrom, James Large, and Eamonn Keogh. 2017. The great time series classification bake off: a review and experimental evaluation of recent algorithmic advances. *Data mining and knowledge discovery*, 31:606–660.
- Claus Bahlmann, Bernard Haasdonk, and Hans Burkhardt. 2002. Online handwriting recognition with support vector machines-a kernel approach. In *Proceedings eighth international workshop on frontiers in handwriting recognition*, pages 49–54. IEEE.
- Shaojie Bai, J Zico Kolter, and Vladlen Koltun. 2018. An empirical evaluation of generic convolutional and recurrent networks for sequence modeling. *arXiv preprint arXiv:1803.01271*.
- Tadas Baltrušaitis, Chaitanya Ahuja, and Louis-Philippe Morency. 2018. Multimodal machine learning: A survey and taxonomy. *IEEE transactions on pattern analysis and machine intelligence*, 41(2):423–443.
- Mouldi Bedda and Nacereddine Hammami. 2010. Spoken Arabic Digit. UCI Machine Learning Repository. DOI: <https://doi.org/10.24432/C52C9Q>.
- Donald J Berndt and James Clifford. 1994. Using dynamic time warping to find patterns in time series. In *Proceedings of the 3rd international conference on knowledge discovery and data mining*, pages 359–370.
- Niels Birbaumer, Nimr Ghanayim, Thilo Hinterberger, Iver Iversen, Boris Kotchoubey, Andrea Kübler, Juri Perelmouter, Edward Taub, and Herta Flor. 1999. A spelling device for the paralysed. *Nature*, 398(6725):297–298.
- Rishi Bommasani, Drew A Hudson, Ehsan Adeli, Russ Altman, Simran Arora, Sydney von Arx, Michael S Bernstein, Jeannette Bohg, Antoine Bosselut, Emma Brunskill, et al. 2021. On the opportunities and risks of foundation models. *arXiv preprint arXiv:2108.07258*.
- Defu Cao, Furong Jia, Sercan O Arik, Tomas Pfister, Yixiang Zheng, Wen Ye, and Yan Liu. 2023.

732	Tempo: Prompt-based generative pre-trained trans-	Elizabeth Fons, Rachneet Kaur, Soham Palande, Zhen	788
733	former for time series forecasting. <i>arXiv preprint</i>	Zeng, Tucker Balch, Manuela Veloso, and Svitlana	789
734	<i>arXiv:2310.04948</i> .	Vyetrenko. 2024. Evaluating large language models	790
735	Tianqi Chen and Carlos Guestrin. 2016. Xgboost: A	on time series feature understanding: A comprehen-	791
736	scalable tree boosting system. In <i>Proceedings of</i>	sive taxonomy and benchmark . In <i>Proceedings of the</i>	792
737	<i>the 22nd acm sigkdd international conference on</i>	<i>2024 Conference on Empirical Methods in Natural</i>	793
738	<i>knowledge discovery and data mining</i> , pages 785–	<i>Language Processing (EMNLP)</i> , pages 21598–21634,	794
739	794.	Miami, Florida, USA. Association for Computational	795
740	Zheng Chu, Jingchang Chen, Qianglong Chen, Weijiang	Linguistics.	796
741	Yu, Haotian Wang, Ming Liu, and Bing Qin. 2024.	Jean-Yves Franceschi, Aymeric Dieuleveut, and Martin	797
742	TimeBench: A comprehensive evaluation of tempo-	Jaggi. 2019. Unsupervised scalable representation	798
743	ral reasoning abilities in large language models . In	learning for multivariate time series. <i>Advances in</i>	799
744	<i>Proceedings of the 62nd Annual Meeting of the Asso-</i>	<i>neural information processing systems</i> , 32.	800
745	<i>ciation for Computational Linguistics (ACL) (Volume</i>	Jonas Geiping and Tom Goldstein. 2023. Cramming:	801
746	<i>1: Long Papers)</i> , pages 1204–1228, Bangkok, Thai-	Training a language model on a single gpu in one day.	802
747	land. Association for Computational Linguistics.	In <i>International Conference on Machine Learning</i> ,	803
748	Marco Cuturi. 2011. Fast global alignment kernels. In	pages 11117–11143. PMLR.	804
749	<i>Proceedings of the 28th international conference on</i>	Aidan N Gomez, Mengye Ren, Raquel Urtasun, and	805
750	<i>machine learning (ICML-11)</i> , pages 929–936.	Roger B Grosse. 2017. The reversible residual net-	806
751	Tri Dao, Dan Fu, Stefano Ermon, Atri Rudra, and	work: Backpropagation without storing activations.	807
752	Christopher Ré. 2022. Flashattention: Fast and	<i>Advances in neural information processing systems</i> ,	808
753	memory-efficient exact attention with io-awareness.	30.	809
754	<i>Advances in Neural Information Processing Systems</i> ,	Mononito Goswami, Konrad Szafer, Arjun Choudhry,	810
755	35:16344–16359.	Yifu Cai, Shuo Li, and Artur Dubrawski. 2024. Mo-	811
756	Angus Dempster, François Petitjean, and Geoffrey I	ment: A family of open time-series foundation mod-	812
757	Webb. 2020. Rocket: exceptionally fast and accurate	els . In <i>Forty-first International Conference on Ma-</i>	813
758	time series classification using random convolutional	<i>chine Learning</i> .	814
759	kernels. <i>Data Mining and Knowledge Discovery</i> ,	Nate Gruver, Marc Finzi, Shikai Qiu, and Andrew G	815
760	34(5):1454–1495.	Wilson. 2024. Large language models are zero-shot	816
761	Jacob Devlin, Ming-Wei Chang, Kenton Lee, and	time series forecasters. <i>Advances in Neural Informa-</i>	817
762	Kristina Toutanova. 2019. BERT: Pre-training of	<i>tion Processing Systems</i> , 36.	818
763	deep bidirectional transformers for language under-	Albert Gu, Tri Dao, Stefano Ermon, Atri Rudra, and	819
764	standing . In <i>Proceedings of the 2019 Conference of</i>	Christopher Ré. 2020. Hippo: Recurrent mem-	820
765	<i>the North American Chapter of the Association for</i>	ory with optimal polynomial projections. <i>Advances</i>	821
766	<i>Computational Linguistics (NAACL): Human Lan-</i>	<i>in neural information processing systems</i> , 33:1474–	822
767	<i>guage Technologies, Volume 1 (Long and Short Pa-</i>	1487.	823
768	<i>pers)</i> , pages 4171–4186, Minneapolis, Minnesota.	Albert Gu, Karan Goel, and Christopher Ré. 2021. Effi-	824
769	Association for Computational Linguistics.	ciently modeling long sequences with structured state	825
770	Iman Deznabi, Mohit Iyyer, and Madalina Fiterau. 2021.	spaces. <i>arXiv preprint arXiv:2111.00396</i> .	826
771	Predicting in-hospital mortality by combining clin-	Wenzhong Guo, Jianwen Wang, and Shiping Wang.	827
772	ical notes with time-series data. In <i>Findings of</i>	2019. Deep multimodal representation learning: A	828
773	<i>the association for computational linguistics: ACL-</i>	survey. <i>Ieee Access</i> , 7:63373–63394.	829
774	<i>IJCNLP 2021</i> , pages 4026–4031.	Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian	830
775	Felix Drinkall, Eghbal Rahimikia, Janet Pierrehumbert,	Sun. 2016. Deep residual learning for image recog-	831
776	and Stefan Zohren. 2024. Time machine GPT . In	nition. In <i>Proceedings of the IEEE conference on</i>	832
777	<i>Findings of the Association for Computational Lin-</i>	<i>computer vision and pattern recognition</i> , pages 770–	833
778	<i>guistics: NAACL 2024</i> , pages 3281–3292, Mexico	778.	834
779	City, Mexico. Association for Computational Lin-	Sepp Hochreiter and Jürgen Schmidhuber. 1997. Long	835
780	guistics.	short-term memory. <i>Neural computation</i> , 9(8):1735–	836
781	Emadeldeen Eldele, Mohamed Ragab, Zhenghua Chen,	1780.	837
782	Min Wu, Chee Keong Kwoh, Xiaoli Li, and Cuntai	Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan	838
783	Guan. 2021. Time-series representation learning via	Allen-Zhu, Yanzhi Li, Shean Wang, Lu Wang,	839
784	temporal and contextual contrasting. In <i>Proceedings</i>	and Weizhu Chen. 2021. Lora: Low-rank adap-	840
785	<i>of the Thirtieth International Joint Conference on Ar-</i>	tation of large language models. <i>arXiv preprint</i>	841
786	<i>tificial Intelligence</i> . International Joint Conferences	<i>arXiv:2106.09685</i> .	842
787	on Artificial Intelligence Organization.		

843	Hassan Ismail Fawaz, Germain Forestier, Jonathan We-	<i>in Natural Language Processing and the 9th Inter-</i>	898
844	ber, Lhassane Idoumghar, and Pierre-Alain Muller.	<i>national Joint Conference on Natural Language Pro-</i>	899
845	2019. Deep learning for time series classification:	<i>cessing (EMNLP-IJCNLP)</i> , pages 6432–6437, Hong	900
846	a review. <i>Data mining and knowledge discovery</i> ,	Kong, China. Association for Computational Linguis-	901
847	33(4):917–963.	tics.	902
848	Allan Jabri, Armand Joulin, and Laurens Van	Taesung Kim, Jinhee Kim, Yunwon Tae, Cheonbok	903
849	Der Maaten. 2016. Revisiting visual question answer-	Park, Jang-Ho Choi, and Jaegul Choo. 2021. Re-	904
850	ing baselines. In <i>European conference on computer</i>	versible instance normalization for accurate time-	905
851	<i>vision</i> , pages 727–739. Springer.	series forecasting against distribution shift. In <i>Inter-</i>	906
852	Young-Seon Jeong, Myong K Jeong, and Olufemi A	<i>national Conference on Learning Representations</i> .	907
853	Omitaomu. 2011. Weighted dynamic time warping	Nikita Kitaev, Łukasz Kaiser, and Anselm Levskaya.	908
854	for time series classification. <i>Pattern recognition</i> ,	2020. Reformer: The efficient transformer. <i>arXiv</i>	909
855	44(9):2231–2240.	<i>preprint arXiv:2001.04451</i> .	910
856	Harsh Jhamtani and Taylor Berg-Kirkpatrick. 2021.	Ross Koval, Nicholas Andrews, and Xifeng Yan. 2024.	911
857	Truth-conditional captions for time series data . In	Financial forecasting from textual and tabular time	912
858	<i>Proceedings of the 2021 Conference on Empirical</i>	series . In <i>Findings of the Association for Computa-</i>	913
859	<i>Methods in Natural Language Processing (EMNLP)</i> ,	<i>tional Linguistics: EMNLP 2024</i> , pages 8289–8300,	914
860	pages 719–733, Online and Punta Cana, Dominican	Miami, Florida, USA. Association for Computational	915
861	Republic. Association for Computational Linguistics.	Linguistics.	916
862	Albert Q Jiang, Alexandre Sablayrolles, Arthur Men-	Mineichi Kudo, Jun Toyama, and Masaru Shimbo. 1999.	917
863	sch, Chris Bamford, Devendra Singh Chaplot, Diego	Japanese Vowels. UCI Machine Learning Repository.	918
864	de las Casas, Florian Bressand, Gianna Lengyel, Guil-	DOI: https://doi.org/10.24432/C5NS47 .	919
865	laume Lample, Lucile Saulnier, et al. 2023. Mistral	Aditya Kusupati, Gantavya Bhatt, Aniket Rege,	920
866	7b. <i>arXiv preprint arXiv:2310.06825</i> .	Matthew Wallingford, Aditya Sinha, Vivek Ramanu-	921
867	Yushan Jiang, Zijie Pan, Xikun Zhang, Sahil Garg, An-	jan, William Howard-Snyder, Kaifeng Chen, Sham	922
868	derson Schneider, Yuriy Nevmyvaka, and Dongjin	Kakade, Prateek Jain, et al. 2022. Matryoshka repre-	923
869	Song. 2024. Empowering time series analysis	sensation learning. <i>Advances in Neural Information</i>	924
870	with large language models: A survey . <i>ArXiv</i> ,	<i>Processing Systems</i> , 35:30233–30249.	925
871	abs/2402.03182.	Guokun Lai, Wei-Cheng Chang, Yiming Yang, and	926
872	Ming Jin, Shiyu Wang, Lintao Ma, Zhixuan Chu,	Hanxiao Liu. 2018. Modeling long-and short-term	927
873	James Y Zhang, Xiaoming Shi, Pin-Yu Chen, Yuxuan	temporal patterns with deep neural networks. In	928
874	Liang, Yuan-Fang Li, Shirui Pan, et al. 2023. Time-	<i>The 41st international ACM SIGIR conference on re-</i>	929
875	llm: Time series forecasting by reprogramming large	<i>search & development in information retrieval</i> , pages	930
876	language models. <i>arXiv preprint arXiv:2310.01728</i> .	95–104.	931
877	Argyro Kampouraki, George Manis, and Christophoros	James Large, E Kate Kemsley, Nikolaus Wellner, Ian	932
878	Nikou. 2008. Heartbeat time series classification	Goodall, and Anthony Bagnall. 2018. Detecting	933
879	with support vector machines. <i>IEEE transactions on</i>	forged alcohol non-invasively through vibrational	934
880	<i>information technology in biomedicine</i> , 13(4):512–	spectroscopy and machine learning. In <i>Pacific-Asia</i>	935
881	518.	<i>Conference on Knowledge Discovery and Data Min-</i>	936
882	Masayuki Kawarada, Tatsuya Ishigaki, Goran Topić,	ing, pages 298–309. Springer.	937
883	and Hiroya Takamura. 2024. Demonstration selec-	Mosh Levy, Alon Jacoby, and Yoav Goldberg. 2024.	938
884	tion strategies for numerical time series data-to-text .	Same task, more tokens: the impact of input length on	939
885	In <i>Findings of the Association for Computational</i>	the reasoning performance of large language models.	940
886	<i>Linguistics: EMNLP 2024</i> , pages 7378–7392, Mi-	<i>arXiv preprint arXiv:2402.14848</i> .	941
887	ami, Florida, USA. Association for Computational	Nian Li, Chen Gao, Mingyu Li, Yong Li, and Qingmin	942
888	Linguistics.	Liao. 2024a. EconAgent: Large language model-	943
889	Guolin Ke, Qi Meng, Thomas Finley, Taifeng Wang,	empowered agents for simulating macroeconomic	944
890	Wei Chen, Weidong Ma, Qiwei Ye, and Tie-Yan Liu.	activities . In <i>Proceedings of the 62nd Annual Meet-</i>	945
891	2017. Lightgbm: A highly efficient gradient boost-	<i>ing of the Association for Computational Linguistics</i>	946
892	ing decision tree. <i>Advances in neural information</i>	<i>(ACL) (Volume 1: Long Papers)</i> , pages 15523–15536,	947
893	<i>processing systems</i> , 30.	Bangkok, Thailand. Association for Computational	948
894	Swaraj Khadanga, Karan Aggarwal, Shafiq Joty, and	Linguistics.	949
895	Jaideep Srivastava. 2019. Using clinical notes with	Zehan Li, Xin Zhang, Yanzhao Zhang, Dingkun Long,	950
896	time series data for ICU management . In <i>Proceed-</i>	Pengjun Xie, and Meishan Zhang. 2023. Towards	951
897	<i>ings of the 2019 Conference on Empirical Methods</i>	general text embeddings with multi-stage contrastive	952
		learning. <i>arXiv preprint arXiv:2308.03281</i> .	953

954	Zekun Li, Shiyang Li, and Xifeng Yan. 2024b. Time	Muhammad Arslan Manzoor, Sarah Albarri, Ziting	1008
955	series as images: Vision transformer for irregularly	Xian, Zaiqiao Meng, Preslav Nakov, and Shangsong	1009
956	sampled time series. <i>Advances in Neural Information</i>	Liang. 2023. Multimodality representation learning:	1010
957	<i>Processing Systems</i> , 36.	A survey on evolution, pretraining and its applica-	1011
		tions. <i>ACM Transactions on Multimedia Computing,</i>	1012
958	Zachary C Lipton, David Kale, and Randall Wetzel.	<i>Communications and Applications</i> , 20(3):1–34.	1013
959	2016. Directly modeling missing data in sequences		
960	with rnns: Improved classification of clinical time se-	Mike A Merrill, Mingtian Tan, Vinayak Gupta, Thomas	1014
961	ries. In <i>Machine learning for healthcare conference</i> ,	Hartvigsen, and Tim Althoff. 2024. Language mod-	1015
962	pages 253–270. PMLR.	els still struggle to zero-shot reason about time series .	1016
		In <i>Findings of the Association for Computational</i>	1017
963	Chengyu Liu, David Springer, Qiao Li, Benjamin	<i>Linguistics: EMNLP 2024</i> , pages 3512–3533, Mi-	1018
964	Moody, Ricardo Abad Juan, Francisco J Chorro,	ami, Florida, USA. Association for Computational	1019
965	Francisco Castells, José Millet Roig, Ikaro Silva,	Linguistics.	1020
966	Alistair EW Johnson, et al. 2016. An open access		
967	database for the evaluation of heart sound algorithms.	Tomas Mikolov, Ilya Sutskever, Kai Chen, Greg S Cor-	1021
968	<i>Physiological measurement</i> , 37(12):2181.	rado, and Jeff Dean. 2013. Distributed representa-	1022
		tions of words and phrases and their compositionality.	1023
969	Haoxin Liu, Zhiyuan Zhao, Jindong Wang, Harshavard-	<i>Advances in neural information processing systems</i> ,	1024
970	han Kamarthi, and B. Aditya Prakash. 2024. LST-	26.	1025
971	Prompt: Large language models as zero-shot time		
972	series forecasters by long-short-term prompting . In	MTEB. 2024. Massive Text Embedding Benchmark	1026
973	<i>Findings of the Association for Computational Lin-</i>	(MTEB) Leaderboard. https://huggingface.co/	1027
974	<i>guistics: ACL 2024</i> , pages 7832–7840, Bangkok,	spaces/mteb/leaderboard . Accessed: 2024-05-	1028
975	Thailand. Association for Computational Linguistics.	13.	1029
976	Jiayang Liu, Lin Zhong, Jehan Wickramasuriya, and	Niklas Muennighoff, Nouamane Tazi, Loïc Magne, and	1030
977	Venu Vasudevan. 2009. uwave: Accelerometer-based	Nils Reimers. 2022. Mteb: Massive text embedding	1031
978	personalized gesture recognition and its applications.	benchmark . <i>arXiv preprint arXiv:2210.07316</i> .	1032
979	<i>Pervasive and Mobile Computing</i> , 5(6):657–675.		
		Soichiro Murakami, Akihiko Watanabe, Akira	1033
980	Liyuan Liu, Haoming Jiang, Pengcheng He, Weizhu	Miyazawa, Keiichi Goshima, Toshihiko Yanase, Hi-	1034
981	Chen, Xiaodong Liu, Jianfeng Gao, and Jiawei Han.	roya Takamura, and Yusuke Miyao. 2017. Learning	1035
982	2019a. On the variance of the adaptive learning rate	to generate market comments from stock prices.	1036
983	and beyond. <i>arXiv preprint arXiv:1908.03265</i> .	In <i>Proceedings of the 55th Annual Meeting of the</i>	1037
		<i>Association for Computational Linguistics (ACL)</i>	1038
984	Shizhan Liu, Hang Yu, Cong Liao, Jianguo Li, Weiya-	<i>(Volume 1: Long Papers)</i> , pages 1374–1384.	1039
985	Lin, Alex X Liu, and Schahram Dustdar. 2021.		
986	Pyraformer: Low-complexity pyramidal attention for	Yuqi Nie, Nam H Nguyen, Phanwadee Sinthong, and	1040
987	long-range time series modeling and forecasting. In	Jayant Kalagnanam. 2023. A time series is worth 64	1041
988	<i>International conference on learning representations</i> .	words: Long-term forecasting with transformers. In	1042
		<i>The Eleventh International Conference on Learning</i>	1043
989	Tiedong Liu and Bryan Kian Hsiang Low. 2023. Goat:	<i>Representations</i> .	1044
990	Fine-tuned llama outperforms gpt-4 on arithmetic		
991	tasks. <i>arXiv preprint arXiv:2305.14201</i> .	Zach Nussbaum, John X. Morris, Brandon Duderstadt,	1045
		and Andriy Mulyar. 2024. Nomic embed: Training a	1046
992	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-	reproducible long context text embedder . <i>Preprint</i> ,	1047
993	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,	arXiv:2402.01613 .	1048
994	Luke Zettlemoyer, and Veselin Stoyanov. 2019b.		
995	Roberta: A robustly optimized bert pretraining ap-	OpenAI. 2024. OpenAI Embedding Models.	1049
996	proach. <i>arXiv preprint arXiv:1907.11692</i> .	https://platform.openai.com/docs/guides/	1050
		embeddings , https://openai.com/index/	1051
997	Yong Liu, Haixu Wu, Jianmin Wang, and Mingsheng	new-embedding-models-and-api-updates/ .	1052
998	Long. 2022. Non-stationary transformers: Exploring	Accessed: 2024-05-13.	1053
999	the stationarity in time series forecasting. <i>Advances</i>		
1000	<i>in Neural Information Processing Systems</i> , 35:9881–	Victor Pan. 2017. Fast approximate computations with	1054
1001	9893.	cauchy matrices and polynomials. <i>Mathematics of</i>	1055
		<i>Computation</i> , 86(308):2799–2826.	1056
1002	Xueguang Ma, Liang Wang, Nan Yang, Furu Wei, and	Victor Y Pan. 2012. <i>Structured matrices and polynomi-</i>	1057
1003	Jimmy Lin. 2024. Fine-tuning llama for multi-stage	<i>als: unified superfast algorithms</i> . Springer Science	1058
1004	text retrieval. In <i>Proceedings of the 47th Inter-</i>	& Business Media.	1059
1005	<i>national ACM SIGIR Conference on Research and</i>		
1006	<i>Development in Information Retrieval</i> , pages 2421–	Nikolaos Passalis, Avraam Tsantekidis, Anastasios	1060
1007	2425.	Tefas, Juho Kannianen, Moncef Gabbouj, and	1061

1062	Alexandros Iosifidis. 2017. Time-series classification using neural bag-of-features. In <i>2017 25th European Signal Processing Conference (EUSIPCO)</i> , pages 301–305. IEEE.	1116
1063		1117
1064		1118
1065		1119
1066	Jeffrey Pennington, Richard Socher, and Christopher D Manning. 2014. Glove: Global vectors for word representation. In <i>Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)</i> , pages 1532–1543.	1120
1067		1121
1068		1122
1069		1123
1070		1124
1071	Soujanya Poria, Navonil Majumder, Devamanyu Hazarika, Erik Cambria, Alexander Gelbukh, and Amir Hussain. 2018. Multimodal sentiment analysis: Addressing key issues and setting up the baselines. <i>IEEE Intelligent Systems</i> , 33(6):17–25.	1125
1072		
1073		
1074		
1075		
1076	Jacob Portes, Alexander Trott, Sam Havens, Daniel King, Abhinav Venigalla, Moin Nadeem, Nikhil Sardana, Daya Khudia, and Jonathan Frankle. 2024. Mosaibert: a bidirectional encoder optimized for fast pretraining. <i>Advances in Neural Information Processing Systems</i> , 36.	1126
1077		1127
1078		1128
1079		1129
1080		1130
1081		
1082	Ofir Press, Noah A Smith, and Mike Lewis. 2020. Shortformer: Better language modeling using shorter inputs. <i>arXiv preprint arXiv:2012.15832</i> .	1131
1083		1132
1084		1133
1085	Alec Radford, Jeff Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. 2019. Language models are unsupervised multitask learners.	1134
1086		1135
1087		
1088	UEA Rik Henson. 2023. DecMeg2014 - Decoding the Human Brain. https://www.kaggle.com/c/decoding-the-human-brain/data . Accessed: 2024-04-15.	1136
1089		1137
1090		1138
1091		1139
1092	Tim Salimans and Durk P Kingma. 2016. Weight normalization: A simple reparameterization to accelerate training of deep neural networks. <i>Advances in neural information processing systems</i> , 29.	1140
1093		1141
1094		1142
1095		1143
1096	Noam Shazeer. 2020. Glu variants improve transformer. <i>arXiv preprint arXiv:2002.05202</i> .	1144
1097		1145
1098	Hiroshi Shimodaira, Ken-ichi Noma, Mitsuru Nakai, and Shigeki Sagayama. 2001. Dynamic time-alignment kernel in support vector machine. <i>Advances in neural information processing systems</i> , 14.	1146
1099		
1100		
1101		
1102	Mohammad Shoeybi, Mostofa Patwary, Raul Puri, Patrick LeGresley, Jared Casper, and Bryan Catanzaro. 2019. Megatron-lm: Training multi-billion parameter language models using model parallelism. <i>arXiv preprint arXiv:1909.08053</i> .	1147
1103		1148
1104		1149
1105		1150
1106		1151
1107	Mohammad Shokoohi-Yekta, Bing Hu, Hongxia Jin, Jun Wang, and Eamonn Keogh. 2017. Generalizing dtw to the multi-dimensional case requires an adaptive approach. <i>Data mining and knowledge discovery</i> , 31:1–31.	1152
1108		1153
1109		1154
1110		1155
1111		
1112	Jianlin Su, Murtadha Ahmed, Yu Lu, Shengfeng Pan, Wen Bo, and Yunfeng Liu. 2024. Roformer: Enhanced transformer with rotary position embedding. <i>Neurocomputing</i> , 568:127063.	1156
1113		1157
1114		1158
1115		1159
		1160
	Chenxi Sun, Yaliang Li, Hongyan Li, and Shenda Hong. 2023. Test: Text prototype aligned embedding to activate llm’s ability for time series. <i>arXiv preprint arXiv:2308.08241</i> .	1161
		1162
		1163
		1164
	Christian Szegedy, Wei Liu, Yangqing Jia, Pierre Sermanet, Scott Reed, Dragomir Anguelov, Dumitru Erhan, Vincent Vanhoucke, and Andrew Rabinovich. 2015. Going deeper with convolutions. In <i>Proceedings of the IEEE conference on computer vision and pattern recognition</i> , pages 1–9.	1165
		1166
		1167
		1168
		1169
	Sana Tonekaboni, Danny Eytan, and Anna Goldenberg. 2021. Unsupervised representation learning for time series with temporal neighborhood coding. In <i>International Conference on Learning Representations (ICLR)</i> .	1170
		1171
		1172
		1173
		1174
		1175
		1176
		1177
		1178
		1179
		1180
		1181
		1182
		1183
		1184
		1185
		1186
		1187
		1188
		1189
		1190
		1191
		1192
		1193
		1194
		1195
		1196
		1197
		1198
		1199
		1200
		1201
		1202
		1203
		1204
		1205
		1206
		1207
		1208
		1209
		1210
		1211
		1212
		1213
		1214
		1215
		1216
		1217
		1218
		1219
		1220
		1221
		1222
		1223
		1224
		1225
		1226
		1227
		1228
		1229
		1230
		1231
		1232
		1233
		1234
		1235
		1236
		1237
		1238
		1239
		1240
		1241
		1242
		1243
		1244
		1245
		1246
		1247
		1248
		1249
		1250
		1251
		1252
		1253
		1254
		1255
		1256
		1257
		1258
		1259
		1260
		1261
		1262
		1263
		1264
		1265
		1266
		1267
		1268
		1269
		1270
		1271
		1272
		1273
		1274
		1275
		1276
		1277
		1278
		1279
		1280
		1281
		1282
		1283
		1284
		1285
		1286
		1287
		1288
		1289
		1290
		1291
		1292
		1293
		1294
		1295
		1296
		1297
		1298
		1299
		1300
		1301
		1302
		1303
		1304
		1305
		1306
		1307
		1308
		1309
		1310
		1311
		1312
		1313
		1314
		1315
		1316
		1317
		1318
		1319
		1320
		1321
		1322
		1323
		1324
		1325
		1326
		1327
		1328
		1329
		1330
		1331
		1332
		1333
		1334
		1335
		1336
		1337
		1338
		1339
		1340
		1341
		1342
		1343
		1344
		1345
		1346
		1347
		1348
		1349
		1350
		1351
		1352
		1353
		1354
		1355
		1356
		1357
		1358
		1359
		1360
		1361
		1362
		1363
		1364
		1365
		1366
		1367
		1368
		1369
		1370
		1371
		1372
		1373
		1374
		1375
		1376
		1377
		1378
		1379
		1380
		1381
		1382
		1383
		1384
		1385
		1386
		1387
		1388
		1389
		1390
		1391
		1392
		1393
		1394
		1395
		1396
		1397
		1398
		1399
		1400
		1401
		1402
		1403
		1404
		1405
		1406
		1407
		1408
		1409
		1410
		1411
		1412
		1413
		1414
		1415
		1416
		1417
		1418
		1419
		1420
		1421
		1422
		1423
		1424
		1425
		1426
		1427
		1428
		1429
		1430
		1431
		1432
		1433
		1434
		1435
		1436
		1437
		1438
		1439
		1440
		1441
		1442
		1443
		1444
		1445
		1446
		1447
		1448
		1449
		1450
		1451
		1452
		1453
		1454
		1455
		1456
		1457
		1458
		1459
		1460
		1461
		1462
		1463
		1464
		1465
		1466
		1467
		1468
		1469
		1470
		1471
		1472
		1473
		1474
		1475
		1476
		1477
		1478
		1479
		1480
		1481
		1482
		1483
		1484
		1485
		1486
		1487
		1488
		1489
		1490
		1491
		1492
		1493
		1494
		1495
		1496
		1497
		1498
		1499
		1500
		1501
		1502
		1503
		1504
		1505
		1506
		1507
		1508

1170	Chao-Han Huck Yang, Yun-Yun Tsai, and Pin-Yu Chen.	long sequence time-series forecasting. In <i>Proceed-</i>	1225
1171	2021. Voice2series: Reprogramming acoustic mod-	<i>ings of the AAAI conference on artificial intelligence</i> ,	1226
1172	els for time series classification. In <i>International</i>	volume 35, pages 11106–11115.	1227
1173	<i>conference on machine learning</i> , pages 11808–11819.		
1174	PMLR.		
1175	Jianbo Yang, Minh Nhut Nguyen, Phyo Phyo San, Xi-	Tian Zhou, Ziqing Ma, Qingsong Wen, Xue Wang,	1228
1176	aoli Li, and Shonali Krishnaswamy. 2015. Deep	Liang Sun, and Rong Jin. 2022. Fedformer: Fre-	1229
1177	convolutional neural networks on multichannel time	quency enhanced decomposed transformer for long-	1230
1178	series for human activity recognition. In <i>Ijcai</i> , vol-	term series forecasting. In <i>International conference</i>	1231
1179	ume 15, pages 3995–4001. Buenos Aires, Argentina.	<i>on machine learning</i> , pages 27268–27286. PMLR.	1232
1180	Xinli Yu, Zheng Chen, and Yanbin Lu. 2023. Harness-	Tian Zhou, Peisong Niu, Liang Sun, Rong Jin, et al.	1233
1181	ing LLMs for temporal data - a study on explainable	2024. One fits all: Power general time series analysis	1234
1182	financial time series forecasting . In <i>Proceedings of</i>	by pretrained lm. <i>Advances in neural information</i>	1235
1183	<i>the 2023 Conference on Empirical Methods in Natu-</i>	<i>processing systems</i> , 36.	1236
1184	<i>ral Language Processing (EMNLP): Industry Track</i> ,		
1185	pages 739–753, Singapore. Association for Compu-	Appendix	1237
1186	tational Linguistics.	Table of Contents	1238
1187	Zhihan Yue, Yujing Wang, Juanyong Duan, Tianmeng	A Related Works	15 1239
1188	Yang, Congrui Huang, Yunhai Tong, and Bixiong Xu.	B Datasets	15 1240
1189	2022. Ts2vec: Towards universal representation of	C Comparison Baselines: Supervised	1241
1190	time series. In <i>Proceedings of the AAAI Conference</i>	Learning Methods	17 1242
1191	<i>on Artificial Intelligence</i> , volume 36, pages 8980–	D Adapting Supervised Forecasting Base-	1243
1192	8987.	line Models for Classification	25 1244
1193	Ailing Zeng, Muxi Chen, Lei Zhang, and Qiang Xu.	E Comparison Baselines: Unsupervised	1245
1194	2023. Are transformers effective for time series fore-	Representation Learning Methods	25 1246
1195	casting? In <i>Proceedings of the AAAI conference</i>	F Adapting Unsupervised Representation	1247
1196	<i>on artificial intelligence</i> , volume 37, pages 11121–	Learning Models for Classification	27 1248
1197	11128.	G Reproduction Details for Baselines	27 1249
1198	George Zerveas, Srideepika Jayaraman, Dhaval Patel,	H Implementation Details for LETS-C	27 1250
1199	Anuradha Bhamidipaty, and Carsten Eickhoff. 2021.	I Hyperparameter Settings for LETS-C	27 1251
1200	A transformer-based framework for multivariate time	J Model Performance	28 1252
1201	series representation learning. In <i>Proceedings of</i>	K Computational Cost Analysis	28 1253
1202	<i>the 27th ACM SIGKDD conference on knowledge</i>	L Monetary Costs Analysis	28 1254
1203	<i>discovery & data mining</i> , pages 2114–2124.	M Assessing Text Embeddings with Cosine	1255
1204	Tianping Zhang, Yizhuo Zhang, Wei Cao, Jiang Bian,	Similarity	29 1256
1205	Xiaohan Yi, Shun Zheng, and Jian Li. 2022a. Less	N Numerical Precision for Tokenization	30 1257
1206	is more: Fast multivariate time series forecasting	O Generalization Across Various Text Em-	1258
1207	with light sampling-oriented mlp structures. <i>arXiv</i>	bedding Models	30 1259
1208	<i>preprint arXiv:2207.01186</i> .	P Trade-offs: Model Accuracy vs. Parame-	1260
1209	Xiang Zhang, Ziyuan Zhao, Theodoros Tsiligkaridis,	ter Complexity	32 1261
1210	and Marinka Zitnik. 2022b. Self-supervised		
1211	contrastive pre-training for time series via time-		
1212	frequency consistency. <i>Advances in Neural Infor-</i>		
1213	<i>mation Processing Systems</i> , 35:3988–4003.		
1214	Xiyuan Zhang, Ranak Roy Chowdhury, Rajesh K.		
1215	Gupta, and Jingbo Shang. 2024. Large language mod-		
1216	els for time series: A survey . <i>ArXiv</i> , abs/2402.01801.		
1217	Bendong Zhao, Huanzhang Lu, Shangfeng Chen, Jun-		
1218	liang Liu, and Dongya Wu. 2017. Convolutional neu-		
1219	ral networks for time series classification. <i>Journal</i>		
1220	<i>of Systems Engineering and Electronics</i> , 28(1):162–		
1221	169.		
1222	Haoyi Zhou, Shanghang Zhang, Jieqi Peng, Shuai		
1223	Zhang, Jianxin Li, Hui Xiong, and Wancai Zhang.		
1224	2021. Informer: Beyond efficient transformer for		

Q Ablation Study

34

R Alternative Methods for Fusing Time Series with Embeddings

34

A Related Works

Time series classification has been an active research area for decades. Early approaches investigated distance-based approaches (Abanda et al., 2019) for time series classification. Some built nearest neighbor classifiers based on explicit time series distance measures such as Dynamic Time Warping (DTW) (Wang et al., 2013; Jeong et al., 2011; Berndt and Clifford, 1994). Others have used distance kernels instead and learned Support Vector Machines (SVMs) (Kampouraki et al., 2008; Bahlmann et al., 2002; Shimodaira et al., 2001), or extracted features and learned linear (Dempster et al., 2020) or tree-based classifiers such as eXtreme Gradient Boosting (XGBoost) (Chen and Guestrin, 2016).

Later, deep learning-based approaches are widely adopted because of their ability to learn complex patterns. Convolutional Neural Networks (CNNs) have proven to be successful in learning local patterns in time series data (Wu et al., 2022a; Franceschi et al., 2019; Zhao et al., 2017). Similarly, Multilayer Perceptron (MLP) can provide simple but effective time series classifiers (Zhang et al., 2022a). Recurrent Neural Networks (RNNs), such as Long Short-Term Memory (LSTM) effectively handle long sequence modeling (Gu et al., 2021; Lai et al., 2018; Hochreiter and Schmidhuber, 1997).

More recently, transformer-based models (Vaswani et al., 2017) have revolutionized the NLP domain, and these models have been adapted to the time series domain (Zhou et al., 2022; Wu et al., 2021; Zhou et al., 2021). The self-attention mechanism in transformers is known for modeling long-range dependencies in sequence data. However, the increasing complexity of these models often comes with larger model sizes and higher computational costs, especially for training.

B Datasets

We benchmark our model using the following 10 multivariate datasets from the UEA Time Series Classification Archive Bagnall et al. (2018). See Table 4 for their data characteristics.

B.1 EthanolConcentration

EthanolConcentration (Large et al., 2018) comprises raw spectra from water-and-ethanol solutions contained within 44 unique, real whisky bottles, featuring ethanol concentrations of 35%, 38%, 40%, and 45%. Scotch Whisky regulations require a minimum alcohol content of 40%, a standard that producers adhere to in order to comply with labeling specifications. The dataset presents a classification task to identify the ethanol concentration from spectral readings of any given bottle. Each record includes three spectral readings from the same bottle and batch, obtained by positioning the bottle between a light source and a spectroscope. These spectral readings, which cover wavelengths from 226nm to 1101.5nm at a 0.5nm resolution, were recorded over a one-second integration time using a StellarNet BLACKComet-SR spectrometer. The methodology deliberately avoids optimizing for clarity or consistency in the spectral path, aiming to simulate the varied conditions typical of rapid screening tests that may be performed on batches of spirits for quality assurance.

B.2 FaceDetection

The FaceDetection dataset originates from a 2014 Kaggle competition (Rik Henson, 2023). The challenge involves identifying whether a subject is viewing a picture of a face or a scrambled image using magnetoencephalography (MEG) data, independent of the individual subject. This dataset specifically includes only the training portion from the competition, organized by patient. It comprises data from 10 training subjects (subject01 to subject10) and 6 testing subjects (subject11 to subject16). Each subject has approximately 580 to 590 trials, resulting in a total of 5,890 training trials and 3,524 test trials. Each trial features 1.5 seconds of MEG data, initiated 0.5 seconds before the stimulus is presented, and is associated with a class label—Face (class 1) or Scramble (class 0). The data were down-sampled to 250Hz and subjected to a high-pass filter at 1Hz, producing 62 observations per channel.

B.3 Handwriting

The Handwriting dataset (Shokoohi-Yekta et al., 2017) consists of motion data captured from a smartwatch while subjects wrote the 26 letters of the alphabet. Developed at the University of California, Riverside (UCR), this dataset includes 150

Table 4: Dataset Characteristics. **Abbreviations:** EC: EthanolConcentration, FD: FaceDetection, HW: Handwriting, HB: Heartbeat, JV: JapaneseVowels, PEMS: PEMS-SF, SCP1: SelfRegulationSCP1, SCP2: SelfRegulationSCP2, SAD: SpokenArabicDigits, UW: UWaveGestureLibrary

Characteristic	EC	FD	HW	HB	JV	PEMS-SF	SCP1	SCP2	SAD	UW
Train Size	261	5890	150	204	270	267	268	200	6599	120
Test Size	263	3524	850	205	370	173	293	180	2199	320
Number of Dimensions	3	144	3	61	12	963	6	7	13	3
Series Length	1751	62	152	405	29	144	896	1152	93	315
Number of Classes	4	2	26	2	9	7	2	2	10	8
Type	Spectro	EEG	Motion/Human Activity Recognition	Audio	Audio	Occupancy rate	EEG	EEG	Speech	EEG

training cases and 850 test cases. It features six dimensions, comprising three accelerometer readings and three gyroscope readings.

B.4 Heartbeat

The Heartbeat dataset originates from the PhysioNet/CinC Challenge 2016 [Liu et al. \(2016\)](#) and consists of cardiac sound recordings from a diverse pool of participants, both healthy individuals and patients with cardiac conditions. Recordings were made in various settings, clinical and non-clinical, and captured from multiple body locations including the aortic, pulmonic, tricuspid, and mitral areas, among up to nine potential sites. The dataset categorizes these sounds into two primary classes: normal and abnormal. Normal heart sounds were obtained from healthy subjects, while abnormal sounds were recorded from patients diagnosed with cardiac ailments, predominantly heart valve defects such as mitral valve prolapse, mitral regurgitation, aortic stenosis, and post-valvular surgery conditions, as well as coronary artery disease.

The audio recordings, inclusive of contributions from both children and adults, were uniformly truncated to five seconds. Spectrograms of each truncated audio were generated using a window size of 0.061 seconds with a 70% overlap. This multi-variate dataset is structured with each dimension representing a frequency band derived from the spectrogram. There are 113 instances in the normal class and 296 in the abnormal class.

B.5 JapaneseVowels

The Japanese Vowels dataset [Kudo et al. \(1999\)](#), sourced from the UCI Machine Learning Repository, comprises recordings from nine male speakers who pronounced the Japanese vowels ‘a’ and ‘e’. Each utterance was analyzed using a 12-degree linear prediction to extract a 12-dimensional time-

series representation, with lengths varying originally from 7 to 29. For consistency, all instances in the dataset have been padded to the maximum length of 29. The objective of the classification task is to identify the speaker; hence each 12-by-29 instance matrix is associated with a single class label, ranging from 1 to 9. This dataset serves as a benchmark for assessing the efficacy of time-series classification models in distinguishing speakers based on LPC cepstrum coefficients obtained from their speech patterns.

The dataset includes a total of 640 time-series instances. A training set consists of 30 utterances per speaker, totaling 270 instances. The test set, however, comprises 370 instances and varies in distribution—ranging from 24 to 88 instances per speaker—owing to external factors such as timing and availability during the experimental setup.

B.6 PEMS-SF

The PEMS-SF dataset [Cuturi \(2011\)](#) contains 15 months of daily data sourced from the California Department of Transportation. This dataset details the occupancy rates, ranging from 0 to 1, across various car lanes on the freeways of the San Francisco Bay Area. The data spans from January 1, 2008, to March 30, 2009, with measurements taken every 10 minutes. Each day is treated as an individual time series with a dimension of 963, corresponding to the number of sensors that consistently functioned throughout the observation period. The length of each time series is 144 data points (6 per hour x 24 hours). The dataset excludes public holidays and two anomalous days (March 8, 2009, and March 9, 2008) when sensors recorded no data between 2:00 and 3:00 AM, resulting in a total of 440 valid time series. The classification task involves identifying the day of the week for each series, labeling them with integers from 1 (Monday) to 7 (Sun-

day). Each attribute within a record reflects the occupancy rate recorded by a sensor at a specific timestamp throughout the day.

B.7 SelfRegulationSCP1

The SelfRegulationSCP1 dataset, sourced from [Birbaumer et al. \(1999\)](#), involves recordings from a healthy subject who was instructed to control a cursor on a screen using cortical potentials. This process was facilitated by tracking the subject’s slow cortical potentials (Cz-Mastoids), where cortical positivity resulted in downward cursor movements and cortical negativity caused it to move upward. Each trial, lasting six seconds, was designed to capture these dynamics, with visual feedback provided between the second 2 and 5.5 of the trial. During each trial, a goal was visually indicated at either the top or bottom of the screen starting from 0.5 seconds to the end of the trial, guiding the subject to generate negative or positive potentials correspondingly. The usable data for each trial, however, spans only 3.5 seconds—from the second 2 to 5.5—corresponding to 896 samples per channel given the sampling rate of 256 Hz.

Data capture involved a PsyLab EEG8 amplifier and a PCIM-DAS1602/16 A/D converter, recording over channels positioned according to the 10/20 system. The dataset includes a training set of 268 trials—168 from the first day and 100 from the second, mixed randomly—and 293 test instances, with class labels indicating positivity or negativity.

B.8 SelfRegulationSCP2

The SelfRegulationSCP2 dataset [Birbaumer et al. \(1999\)](#) includes data from an artificially respirated ALS patient who was tasked with controlling a cursor on a computer screen using cortical potentials. Auditory and visual cues were used to guide the patient, with slow cortical potentials measured at the Cz-Mastoids. A positive potential moved the cursor downward, whereas a negative potential moved it upward. Each trial lasted 8 seconds, with the cursor movement direction (up for negativity, down for positivity) indicated both visually and auditorily from the 0.5 to 7.5 second marks. Auditory instructions were given precisely at the 0.5-second mark, and visual feedback was available from seconds 2 to 6.5. Only the data from this 4.5-second feedback period, translating to 1152 samples per channel at a 256 Hz sampling rate, are used for training and testing.

EEG data were collected from several sites according to the 10/20 system and included channels for detecting vertical eye movements (vEOG). The EEG signals were not corrected for EOG artifacts, providing a raw view of the cortical activity. The dataset comprises 200 trials for training, evenly split between two classes, and an additional 180 trials for testing, recorded on the same day but after the training session data. Each trial spans 7 dimensions and a series length of 1152.

B.9 Spoken Arabic Digits

The Spoken Arabic Digits dataset [Bedda and Hammami \(2010\)](#) consists of 8,800 time series data entries derived from the vocal utterances of 88 native Arabic speakers (44 males and 44 females, aged between 18 and 40). Each dataset entry represents one of ten Arabic digits, spoken ten times by each speaker. The dataset captures 13 Mel Frequency Cepstral Coefficients (MFCCs) for each sound snippet, which are extracted under the following audio processing conditions:

- **Sampling rate:** 11025 Hz
- **Bit depth:** 16 bits
- **Window function:** Hamming
- **Pre-emphasis filter:** $1 - 0.97Z^{-1}$

Each line in the database corresponds to one frame of analysis, listing the 13 MFCCs separated by spaces. These coefficients effectively capture the spectral properties essential for recognizing spoken digits. This structured approach facilitates robust time-series analysis for speech recognition tasks involving Arabic numerals.

B.10 UWaveGestureLibrary

The UWaveGestureLibrary [Liu et al. \(2009\)](#) comprises a set of eight simple gestures, each generated from accelerometer data. The dataset records the X, Y, and Z coordinates corresponding to each gesture’s motion. Every time series within this dataset consists of 315 data points.

C Comparison Baselines: Supervised Learning Methods

To provide a thorough comparison, we evaluate our approach against 21 supervised baseline models for time series classification. These baselines can be categorized into Classical methods, models based

on MLPs, RNNs, CNNs, transformers, and LLMs. The details of which are provided below.

C.1 Classical methods

- 1) **Dynamic Time Warping (DTW)** [Berndt and Clifford \(1994\)](#) is a method for measuring similarity between two time series, $X = (x_1, x_2, \dots, x_M)$ and $Y = (y_1, y_2, \dots, y_N)$, which may differ in length and are sampled at equidistant points in time. DTW identifies the best alignment between these series by minimizing the effects of distortion and shifting in time, allowing for the comparison of similar shapes across different phases.

The core of DTW is the construction of a local cost matrix, $C \in \mathbb{R}^{M \times N}$, with entries $c_{i,j} = \|x_i - y_j\|$ for $i \in [1, M]$ and $j \in [1, N]$, which represents the pairwise distances between points in the two series. The objective is to find a warping path $p = (p_1, p_2, \dots, p_L)$ where each $p_l = (p_i, p_j)$ lies within $[1, M] \times [1, N]$. This path aligns the series by following the route that minimizes cumulative distance, adhering to several constraints: it must start and end at $p_1 = (1, 1)$ and $p_L = (M, N)$ (boundary condition), maintain the temporal ordering of points $m_1 \leq m_2 \leq \dots \leq m_L$ and $n_1 \leq n_2 \leq \dots \leq n_L$ (monotonicity condition), and prevent large temporal jumps (step size condition). The optimal warping path is identified through a recursive process aimed at minimizing the total cost associated with p , calculated as $c_p(X, Y) = \sum_{l=1}^L c(x_{m_l}, y_{n_l})$. This path, P^* , where $c_{P^*} = \min_{p \in P} c_p(X, Y)$, defines the DTW distance, quantifying the similarity between the series. However, the computational cost of this process is $O(MN)$, where M and N are the lengths of the two series, rendering it computationally demanding for large datasets.

For classifying time series, the DTW distance is integrated with the k-nearest neighbors (k-NN) algorithm. This approach computes the DTW distance of a target series to all other series in a training dataset and assigns a classification based on the most common class among the k-nearest neighbors. Thus, DTW effectively accommodates series with time shifts, providing a robust, distance-based method for classifying time series.

- 2) **eXtreme Gradient Boosting (XGBoost)** [\(Chen and Guestrin, 2016\)](#) is a state-of-the-art machine learning algorithm that primarily uses decision trees as base learners to construct a robust ensemble model. XGBoost sequentially builds a series of weak learners—typically decision trees—and enhances each successive tree by correcting errors made by its predecessors, a technique known as boosting. This process involves an additive model where each new decision tree is improved by leveraging the cumulative knowledge of the trees that came before it, optimizing for maximum information gain at each split using a greedy algorithm. To prevent overfitting, XGBoost incorporates regularization directly into its loss function and employs shrinkage to moderate the learning rate. Additionally, the optimization method, which is gradient-based, minimizes a cost function by iteratively adjusting the model’s parameters in response to the gradients of the errors. XGBoost also refines the decision tree construction process by using a Similarity Score and Gain to determine the most effective node splits, further improving the model’s accuracy and efficiency.
- 3) **RandOm Convolutional KErnel Transform (ROCKET)** [\(Dempster et al., 2020\)](#) is a method for time series classification that employs random convolutional kernels to transform series data, which is then used to train a linear classifier. Unlike traditional convolutional neural networks (CNNs) that rely on learned kernels, ROCKET utilizes a broad array of random kernels. Each kernel is uniquely characterized by random properties such as length, weights, biases, dilation, and padding. This configuration forms a single-layer convolutional neural network, where the randomized kernel weights contribute to generating input for a softmax layer, thus optimizing the feature extraction process. ROCKET also efficiently scales for large datasets due to its linear complexity relative to the length of the time series and the number of training samples. The key advantages of ROCKET include:
 - Number of Kernels: ROCKET utilizes a substantial number of kernels in a single layer. The non-learned nature of these kernels reduces computational costs sig-

nificantly, enabling the use of numerous kernels without substantial overhead.

- Variety of Kernels: Unlike typical CNNs, where kernels might share characteristics, each ROCKET kernel is distinct in its attributes, enhancing the diversity and the ability to detect various patterns.
- Kernel Dilation: Dilation in ROCKET is randomly assigned to each kernel, differing from the exponential increase with depth seen in traditional CNNs. This randomness is vital for identifying patterns across different scales and frequencies.
- Feature Extraction: Beyond employing global max pooling techniques through the maximum value of feature maps, ROCKET utilizes a novel metric—the proportion of positive values. This metric allows classifiers to assess the prevalence of patterns more accurately within the time series.

C.2 MLP-based methods

- 4) **LightTS** (Zhang et al., 2022a) is an MLP-based time series forecasting model that employs simple MLP structures to manage both short-term and long-term temporal dependencies. This model includes two downsampling strategies: interval sampling, which targets long-term dependencies, and continuous sampling, which focuses on short-term local patterns. These strategies are based on the principle that downsampling typically preserves the majority of a time series’ crucial information, thus maintaining model efficiency. LightTS utilizes an MLP-based framework on top of these downsampling techniques, enabling effective information exchange among different down-sampled subsequences and time steps. This configuration allows LightTS to adaptively select relevant information for forecasting and to efficiently handle very long input sequences by processing only a fraction of the data after downsampling.
- 5) **DLinear** (Zeng et al., 2023) is a recent non-transformer model developed in response to the difficulty transformer-based models face in capturing ordering information within time series. DLinear integrates a decomposition scheme similar to those used in Autoformer and FEDformer but relies on linear layers for

processing. Initially, it decomposes a raw data input into a trend component using a moving average kernel and a remainder (seasonal) component. Subsequently, two single-layer linear layers are applied independently to each component. The outputs from these layers are then summed to produce the final prediction. This approach allows DLinear to enhance performance over a standard linear model by explicitly handling trends within the data.

C.3 RNN-based models

- 6) **Long Short-Term Memory (LSTM)** Hochreiter and Schmidhuber (1997) addresses the vanishing gradient problem that plagues vanilla RNNs in processing longer sequences. This issue arises as the network propagates forward, and the small weight values in the hidden layers are multiplied repeatedly, causing the gradients to diminish rapidly. As a result, the weights in the initial layers become increasingly difficult to train, which impacts the training of subsequent weights, making RNNs challenging to train overall. LSTM mitigates this problem by incorporating a memory cell equipped with various gates that regulate the flow of information into and out of the cell, enabling it to handle long-short term dependencies effectively. An LSTM unit utilizes a cell state and three gates—input, forget, and output—to manage information. Each gate includes a sigmoid layer σ that outputs values between 0 and 1, representing the proportion of information allowed through the gate, and a point-wise multiplication operation. Specifically, the forget gate $f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f)$ determines which information to discard from the previous cell state c_{t-1} by analyzing the current input x_t and the previous hidden state h_{t-1} . The input gate $i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i)$ decides which new information to update, and the update to the cell state $\tilde{c}_t = \tanh(W_c \cdot [h_{t-1}, x_t] + b_c)$ is computed. The cell state is then updated to $c_t = f_t \cdot c_{t-1} + i_t \cdot \tilde{c}_t$. Finally, the output gate $o_t = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o)$ determines what portion of the cell state to output, with the output hidden state given by $h_t = o_t \cdot \tanh(c_t)$.
- 7) **Long- and Short-term Time-series Network (LSTNet)** (Lai et al., 2018) incorporates both CNN and RNN components to perform com-

prehensive time series analysis. The CNN extracts short-term local dependency patterns from multi-dimensional input variables, while the RNN is tasked with capturing complex long-term dependencies in time series trends. To address the issue of scale insensitivity commonly found in neural network models, LSTNet integrates a traditional autoregressive model. Additionally, LSTNet features a Recurrent-skip structure designed to effectively capture very long-term dependence patterns and to facilitate easier optimization, leveraging the periodic properties of the input time series signals. Further enhancing its robustness, LSTNet also employs a traditional autoregressive linear model in parallel with its nonlinear neural network components. This dual approach makes the network particularly adept at managing time series that exhibit significant scale variations.

- 8) **Linear State Space Layer (LSSL)** (Gu et al., 2021), which is part of the structured state space sequence model, introduces a parameterization for state space models (SSM) designed to enhance computational efficiency. LSSL modifies the structured state matrices by decomposing them into a low-rank and a normal term (Gu et al., 2020; Voelker et al., 2019). This modification facilitates the computation of the truncated generating function in frequency space, rather than expanding the standard SSM in coefficient space. Furthermore, LSSL employs the Woodbury identity to adjust the low-rank term, and the normal term is stably diagonalized. This approach simplifies the computations by involving processes associated with a Cauchy kernel (Pan, 2012, 2017), known for its stability in theoretical contexts. These modifications allow LSSL to efficiently manage both computational and memory resources.

C.4 CNN-based models

- 9) **Temporal Convolutional Network (TCN)** (Bai et al., 2018; Franceschi et al., 2019) is a CNN-based architecture designed to capture extended historical data with long memory capabilities and requires minimal tuning in practice. TCNs utilize dilated causal convolutions to ensure that predictions do not prematurely incorporate future data. The dilations signif-

icantly expand the network’s receptive field, enabling it to cover a broader range of historical context. Furthermore, TCNs integrate residual connections to facilitate the effective training of deeper models.

A TCN model comprises a series of n TCN residual blocks, where n is a hyperparameter. Each block contains two dilated causal convolutional layers, which are fully convolutional to ensure that the output size is consistent with the input size. These causal convolutions guarantee that the output at any given time t depends only on inputs from time t and earlier. The convolutional layers are applied with a stride of 1, and padding adjustments maintain the convolutional nature of the network. Each convolutional layer applies a dilation factor d , typically set as $d = 2^i$ for the i -th block, to exponentially increase the receptive field as the network deepens. Mathematically, a convolution with dilation factor d on an element x of a 1D input g with a filter f of length k is computed as $(g *_d f)(x) = \sum_{j=0}^{k-1} f(j) \cdot g(x - d \cdot j)$.

Following the convolutional layers, the sequence of operations includes weight normalization (Salimans and Kingma, 2016), ReLU activation, and a dropout layer. Weight normalization improves gradient conditioning and accelerates convergence by reparameterizing each weight vector $\mathbf{w}$ as $\mathbf{w} = \frac{g}{\|\mathbf{v}\|} \mathbf{v}$, where $\mathbf{v}$ has a fixed norm and g is a scalar. This process decouples the magnitude of the weight vector from its direction, enhancing network optimizability. Each block concludes with an element-wise addition of the block’s input and the estimated residual mapping, followed by a ReLU activation. Dimension alignment for this addition is achieved with a 1×1 convolution.

- 10) **TimesNet** (Wu et al., 2022a) is a method that overcomes the limitations of traditional 1D time series representation by transforming these series into 2D tensors organized across multiple periods. It captures short-term intraperiod variations and long-term interperiod trends by embedding them into the columns and rows of the 2D tensors, respectively. This transformation allows for more efficient modeling of temporal variations using 2D convolu-

tional kernels, thereby extending the analysis into a more comprehensive 2D space. TimesNet ensures simultaneous representation of both intraperiod and interperiod variations, with modules specifically tailored to emphasize the unique temporal patterns of each period.

The central component of TimesNet, the TimesBlock, is a versatile and adaptive structure designed to detect multiperiodicity and extract complex temporal patterns from these 2D tensors. Its parameter-efficient Inception-block (Szegedy et al., 2015) based architecture boosts the model’s analytical capabilities, facilitating a detailed examination of distinct temporal variations associated with different periods. This approach allows TimesNet to transcend the constraints of 1D representations, enabling a unified and thorough analysis of temporal variations.

C.5 Transformer-based models

- 11) **Transformer** (Vaswani et al., 2017) comprises a dual-component architecture with both an encoder and a decoder, each containing stacked self-attention and point-wise, fully connected layers. Each component consists of six identical layers. In the encoder, each layer includes a multi-head self-attention mechanism and a position-wise fully connected feed-forward network, complemented by a residual connection and layer normalization. The decoder replicates this configuration but adds a third sub-layer for multi-head attention on the encoder’s output. It also modifies its self-attention mechanism to prevent forward-looking attention, thus preserving the model’s auto-regressive properties for sequential generation.

The Transformer utilizes multi-head attention to enable nuanced interactions between its encoder and decoder and to facilitate detailed processing across different positions in the sequence. This attention mechanism projects queries, keys, and values through multiple linear transformations, enabling diverse representation and integration of information across subspaces. It is applied in three distinct forms: encoder-decoder attention allows decoder queries to attend to all positions in the encoder output; self-attention within the

encoder lets each position process information from all preceding positions; and self-attention within the decoder restricts attention to prevent future positions from influencing the sequence, maintaining sequential integrity. This structured approach helps the Transformer effectively manage and process long sequence dependencies, making it adaptable for a broad range of sequence-based applications.

- 12) **Reformer** (Kitaev et al., 2020) enhances the efficiency of the transformer model with two significant modifications, making it more suitable for processing long sequences. Firstly, it replaces the traditional dot-product attention with a locality-sensitive hashing mechanism. This change reduces the computational complexity from $O(L^2)$ to $O(L \log L)$, where L is the sequence length. Secondly, Reformer employs reversible residual layers, which regenerate the activations of any layer from the subsequent layer’s activations using only the model parameters. This approach eliminates the need for storing multiple copies of activations for each layer, substantially reducing memory usage. The reversible layers, introduced in (Gomez et al., 2017), require only a single set of activations to be stored for the entire model, significantly reducing the memory cost usually multiplied by the number of layers (N). Moreover, the Reformer processes activations in chunks within feed-forward layers, further decreasing memory demands. These adjustments, along with the use of locality-sensitive hashing for attention computation, not only minimize memory and computational overhead but also maintain the model’s performance on par with the traditional transformer model for long sequences.
- 13) **Informer** (Zhou et al., 2021) is a transformer-based model designed specifically to tackle challenges in long sequence time-series forecasting, such as quadratic time complexity, high memory usage, and constraints of the traditional encoder-decoder architecture. The Informer introduces several modifications to enhance efficiency and effectiveness in processing long sequences:
 - ProbSparse Self-Attention Mechanism: This mechanism replaces the conven-

tional self-attention in Transformers. It is engineered to achieve $O(L \log L)$ in both time complexity and memory usage, efficiently managing dependency alignments without compromising performance.

- Self-Attention Distilling: This process improves attention management by concentrating on dominant attention scores and halving the input size for each cascading layer. This method effectively manages extremely long input sequences and significantly lowers the space complexity to $O((2 - \epsilon)L \log L)$.
- Generative Style Decoder: Diverging from the typical step-by-step decoding process, the generative style decoder in Informer predicts long time-series sequences in a single forward operation. This approach accelerates inference for long-sequence predictions and reduces the propagation of cumulative errors during the inference phase.

The Informer leverages these enhancements—ProbSparse self-attention for efficient processing, self-attention distilling to focus on important attention scores, and a generative style decoder for rapid sequence generation—to improve its performance in forecasting long sequences and capturing long-range dependencies between extensive time-series inputs and outputs.

- 14) **Pyraformer** (Liu et al., 2021) is a transformer-based model that employs a pyramidal attention module (PAM) for efficient management of time series data. This module uses inter-scale and intra-scale connections to summarize features at different resolutions and capture temporal dependencies across various ranges. The design integrates a tree structure for inter-scale connections and neighboring intra-scale connections to achieve a multi-resolution representation of time series. This architecture ensures that Pyraformer scales linearly with the input series length, optimizing computational efficiency while maintaining a constant signal path length relative to the sequence length L , and keeping both time and space complexity linear with L .

Pyraformer operates by first embedding ob-

served data, covariates, and positions in a manner similar to the Informer (Zhou et al., 2021). It then constructs a multi-resolution C -ary tree through a coarser-scale construction module (CSCM), where each coarser scale node aggregates information from C finer scale nodes. This structure allows Pyraformer to model temporal dependencies efficiently across different scales through sparse intra-scale connections, thus reducing computational overhead. Depending on the specific needs of different downstream tasks, Pyraformer adapts its output structure to effectively meet the requirements of diverse time series analyses.

- 15) **Autoformer** (Wu et al., 2021) is a transformer-based model that incorporates time series decomposition, drawing inspiration from classical time series analysis methods. Unlike traditional transformers that rely on self-attention mechanisms to capture long-range dependencies, Autoformer introduces an auto-correlation mechanism as an alternative. This change addresses the issues that traditional models face with intricate temporal patterns and long-term forecasting, where identifying reliable dependencies can be challenging. To enhance efficiency in handling long series, traditional transformers sometimes use sparse versions of self-attention, which can limit the effective use of information.

Autoformer integrates decomposition blocks directly into its structure, moving away from the typical preprocessing approach of series decomposition. These blocks are specifically designed to progressively isolate long-term trends from the data during the forecasting process, allowing the model to refine and decompose the data iteratively. This setup improves the handling of complex time series. The auto-correlation mechanism, inspired by stochastic process theory and based on the periodicity of the series, focuses on identifying dependencies and aggregating representation at the sub-series level, effectively capturing and utilizing similarities derived from underlying periodic patterns.

The architecture of Autoformer adheres to a residual and encoder-decoder framework. The encoder eliminates long-term trend-cyclical components through the series decomposition

blocks and focuses on modeling seasonal patterns. In contrast, the decoder accumulates the trend component extracted from the hidden variables, using past seasonal information to enhance forecasting accuracy. Additionally, Autoformer incorporates a moving average within its decomposition blocks to smooth out periodic fluctuations and highlight long-term trends, thereby facilitating a more targeted analysis of stable trend components within the time series.

- 16) **Non-stationary Transformer** (Liu et al., 2022) addresses the challenges posed by non-stationary real-world data, where the joint distribution changes over time, often leading to the degradation of transformer performance. Non-stationary Transformer consists of two interdependent modules: series stationarization and de-stationary attention. Series stationarization normalizes the input data to unify its statistical properties, enhancing predictability, and adjusts the output to restore the original statistics. This module utilizes a straightforward normalization approach without additional parameters. De-stationary attention, on the other hand, aims to counteract the potential over-normalization by reintroducing the intrinsic non-stationary characteristics of the data into the model’s temporal dependencies. The de-stationary attention module approximates how attention mechanisms would function on unnormalized data and integrates these insights back into the model to maintain crucial temporal dynamics. This setup allows stationformer to balance the predictability benefits gained from normalized data with the rich, detailed patterns inherent in the raw non-stationary series.

Structurally, Non-stationary Transformer adapts the traditional encoder-decoder setup. The encoder extracts information from past observations, while the decoder aggregates this information to refine predictions. The framework modifies the standard transformer by applying series stationarization to both the input and output of the model, and replaces the conventional self-attention mechanism with de-stationary attention. This adaptation aims to enhance the model’s ability to predict non-stationary series by effectively managing the challenges associated with data variability

over time.

- 17) **Frequency Enhanced Decomposed Transformer (FEDformer)** (Zhou et al., 2022) is a transformer-based method that incorporates a time series decomposition scheme to address the limitations of traditional transformers, particularly their high computational demands and challenges in capturing global time series trends. By combining transformers with the seasonal-trend decomposition method, FEDformer aims to separate the broad trends from more detailed fluctuations in time series data. This allows the transformer component to focus on more granular details while the decomposition handles the overall profile of the series.

The architecture of FEDformer includes specialized blocks such as Fourier-enhanced and Wavelet-enhanced blocks within the transformer framework. These blocks serve as substitutes for the conventional self-attention and cross-attention mechanisms, and they enable the model to analyze important structures through frequency domain mapping. FEDformer employs a selective approach to incorporating Fourier components, which helps keep the computational complexity and memory usage linear in relation to the length of the time series. Specifically, FEDformer is structured as a deep decomposition architecture that integrates Frequency Enhanced Block (FEB), Frequency Enhanced Attention (FEA) connecting the encoder and decoder, and the Mixture Of Experts Decomposition block (MOEDecomp). This setup leverages both seasonal-trend decomposition and distribution analysis to facilitate the processing of time series data.

- 18) **ETSformer** (Woo et al., 2022) is a transformer-based architecture tailored for time series forecasting, integrating exponential smoothing techniques. ETSformer uses two novel mechanisms, Exponential Smoothing Attention (ESA) and Frequency Attention (FA), which are designed to replace the traditional self-attention mechanism in standard transformers. ESA utilizes attention scores based on relative time lags, enabling efficient handling of growth components with a computational complexity of $O(L \log L)$ for a

2129	length- L lookback window. Similarly, FA em-	attentions found in typical attention mecha-	2179
2130	employs Fourier transformations to identify dom-	anisms. Flowformer embeds flow conservation	2180
2131	inant seasonal patterns, selecting bases with	within its attention mechanism to streamline	2181
2132	the highest amplitudes in the frequency do-	the process of information aggregation and	2182
2133	main to achieve the same level of complexity.	refinement. By integrating these elements,	2183
2134		Flowformer seeks to provide an alternative	2184
2135	ETSformer is structured with modular decom-	approach that could potentially handle large	2185
2136	position blocks that allow it to dissect time-	datasets and complex time series data more	2186
2137	series data into distinct, interpretable compo-	efficiently than traditional transformers, with-	2187
2138	nents such as level, growth, and seasonality.	out the computational complexity typically	2188
2139	This design facilitates layer-wise decompo-	associated with these models.	2189
2140	sition of the time series into these compo-		
2141	nents, enhancing the model’s ability to capture	20) Patch Time Series Transformer (PatchTST)	2190
2142	and represent complex temporal dynamics.	(Nie et al., 2023) is a transformer-based model	2191
2143	The architecture systematically extracts latent	for multivariate time series forecasting. It	2192
2144	growth and seasonal patterns through a deep,	is built on two key principles: (i) segmenta-	2193
2145	multi-layered approach, where each layer pro-	tion of time series into subseries-level patches,	2194
2146	gressively refines the extraction of these tem-	which serve as input tokens to the transformer,	2195
2147	poral features. The final forecast generated	and (ii) channel-independence, where each	2196
2148	by ETSformer integrates these decomposed	channel contains a single univariate time se-	2197
2149	elements—level, trend, and seasonality—into	ries that shares the same embedding and trans-	2198
2150	a cohesive output that is both practical and	former weights across all series. The patching	2199
2151	interpretable for human analysts. By empha-	mechanism provides three main advantages:	2200
2152	sizing recent observations, the model aligns	it retains local semantic information in embed-	2201
2153	with the principles of exponential smoothing,	dings, significantly reduces the computational	2202
2154	ensuring that more recent trends carry greater	and memory complexity of attention maps	2203
2155	weight in the forecast.	for a given look-back window, and enables	2204
2156		the model to attend to longer historical de-	2205
2157	19) Flowformer (Wu et al., 2022b) modifies the	pendencies. The model operates as follows:	2206
2158	traditional transformer architecture by inte-	multivariate time series data is divided into	2207
2159	grating flow network theory to tackle the scal-	independent channels, each processed sep-	2208
2160	ability issues typical of standard transform-	arately while sharing the same transformer	2209
2161	ers. The conventional attention mechanism in	backbone. Each univariate time series under-	2210
2162	transformers is known for its quadratic com-	goes instance normalization before being seg-	2211
2163	plexity, which limits their ability to process a	mented into patches, which then serve as in-	2212
2164	large number of tokens and scale to larger	put tokens for the transformer. PatchTST em-	2213
2165	models effectively. To address this, Flow-	employs masked self-supervised learning, where	2214
2166	former introduces the Flow-Attention mecha-	patches are randomly selected and set to zero,	2215
2167	nism, grounded in the principles of flow con-	requiring the model to reconstruct the miss-	2216
2168	servation. This mechanism reimagines atten-	ing values. The training objective minimizes	2217
2169	tion as information flowing from sources (val-	Mean Squared Error (MSE) loss between the	2218
2170	ues) to sinks (results) via learned flow capaci-	predicted and ground truth values, demonstrat-	2219
2171	ties (attentions), aiming to achieve linear com-	ing the effectiveness of transformers when	2220
2172	plexity.	time series data is segmented into patches at	2221
2173		the subseries level. For downstream time se-	2222
2174	The Flow-Attention mechanism manages the	ries tasks, the same transformer encoder is	2223
2175	flow of information by regulating incoming	extended and trained with a linear layer for	2224
2176	flow at sinks to initiate source competition	task-specific predictions.	2225
2177	and outgoing flow at sources for sink allo-		
2178	cation. This management of flow helps in	C.6 LLM-based models	2226
	aggregating relevant information without re-	21) OneFitsAll (Zhou et al., 2024) employs pre-	2227
	lying on specific inductive biases and aims	trained language and computer vision models,	2228
	to prevent the common issue of degenerated		

developed from billions of tokens, for time series analysis. This approach, termed the Frozen Pretrained Transformer (FPT), retains the original self-attention and feedforward (FFN) layers of the pre-trained models, which hold the majority of the learned knowledge. The model is fine-tuned for various time series classification tasks, with adjustments made only to the positional embeddings and layer normalization layers to adapt to specific downstream tasks. Additionally, to more effectively manage local semantic information, OneFitsAll incorporates a patching technique as described by (Nie et al., 2023). This method aggregates adjacent time steps into a single patch-based token, thereby increasing the historical time horizon that can be processed by the model without increasing the token length, reducing information redundancy within transformer models.

D Adapting Supervised Forecasting Baseline Models for Classification

It is important to note that some of our supervised learning-based baselines were adapted from forecasting to classification tasks without altering the core design of the models. We employ a lightweight multi-layer perceptron (MLP)-based projection layer on top of the model’s encoded features (originally designed for forecasting) to map them to classification labels. The parameters of this layer are learned during training, enabling the model to adapt effectively for classification tasks. This approach aligns with methods used in prior works such as OneFitsAll (Zhou et al., 2024) and TimesNet (Wu et al., 2022a) when adapting forecasting models for multivariate classification.

To further clarify how each baseline model is used for classification, all baseline codes are available in our implementation repository at <https://anonymous.4open.science/r/LETS-C>.

E Comparison Baselines: Unsupervised Representation Learning Methods

E.1 CNN-based models

- 1) **T-Loss** (Franceschi et al., 2019) is an unsupervised method for learning general-purpose representations of multivariate time series while handling variations in length. It trains a scalable encoder, structured as a deep convolutional neural network with dilated con-

volution, to produce fixed-length vector representations regardless of input length. The loss function is designed as a triplet loss with time-based negative sampling, leveraging the encoder’s ability to process time series of unequal lengths. The objective is to ensure that similar time series obtain similar representations without requiring supervision to learn such similarity.

2) Temporal Neighborhood Coding (TNC)

(Tonekaboni et al., 2021) is a self-supervised approach that leverages the local smoothness of signals to define temporal neighborhoods and learn generalizable representations for time series. It builds on the smoothness of the generative process of signals to capture meaningful representations for time series windows by ensuring that, in the representation space, signals close in time are distinguishable from those farther apart. In other words, temporal proximity remains identifiable in the encoding space. For each sample window, a neighborhood distribution is first defined. The encoder learns the distribution of windows in the representation space, and samples from these distributions are fed into a discriminator, which predicts the probability of the windows belonging to the same neighborhood. TNC is a general framework, making it agnostic to both the nature of the time series and the architecture of the encoder. The encoder can be any parametric model suited to the signal properties. For the discriminator, a simple multi-headed binary classifier is used.

3) TS2Vec (Yue et al., 2022) is a universal contrastive learning framework for time series representation learning across all semantic levels.

It performs hierarchical contrastive learning over augmented context views, enabling robust contextual representations for each timestamp while capturing multi-resolution temporal dependencies. To learn representations at different granularity, TS2Vec applies a simple aggregation mechanism. The representation of an arbitrary sub-sequence is obtained via max pooling over the corresponding timestamps, allowing the model to generate fine-grained representations at any resolution. The framework hierarchically discriminates positive and negative samples across both instance-

wise and temporal dimensions, reinforcing the ability to capture contextual relationships in time series data.

The contrastive objective is based on augmented context views, ensuring that representations of the same sub-series in different augmented contexts remain consistent. To achieve this, two overlapping subseries are randomly sampled from an input time series, and consistency of contextual representations is encouraged on the common segment. The encoder is optimized jointly with temporal and instance-wise contrastive loss, with the total loss summed across multiple scales in a hierarchical manner. The encoder consists of three key components: an input projection layer, a timestamp masking module, and a dilated CNN module. The input projection layer maps observations at each timestamp into a high-dimensional latent space via a fully connected layer. The timestamp masking module then masks latent vectors at randomly selected timestamps to generate an augmented context view. Importantly, masking is applied at the latent vector level rather than raw input values, as time series values may be unbounded, making it impractical to use a special token for raw data.

E.2 Transformer-based models

- 4) **Time-Series Representation Learning via Temporal and Contextual Contrasting (TS-TCC)** (Eldele et al., 2021) is an unsupervised framework for learning time-series representations from unlabeled data using contrastive learning. It transforms raw time-series data into two different but correlated views through weak and strong augmentations. TS-TCC consists of two key modules: a temporal contrasting module and a contextual contrasting module. The temporal contrasting module enhances robustness by enforcing a cross-view prediction task, where past latent features from one augmentation predict the future of another. This design encourages the model to learn transformation-invariant representations while handling perturbations from different timesteps and augmentations. The contextual contrasting module further refines representation learning by maximizing similarity among different contexts of the same sample while minimizing similarity across

samples. By building on the representations learned from temporal contrasting, this module enhances the model’s ability to capture discriminative features. TS-TCC employs simple yet effective augmentations applicable to any time-series data, ensuring adaptability across diverse datasets.

- 5) **Time Series Transformer (TST)** (Zerveas et al., 2021) is a transformer-based framework for unsupervised representation learning of multivariate time series. It employs a masked Mean Squared Error (MSE) loss to train a transformer model, extracting dense vector representations through an autoregressive input denoising objective. Specifically, parts of the input are masked (set to zero), and the model is trained to predict the missing values. Since transformers are inherently insensitive to input order, TST incorporates positional encoding to capture the sequential nature of time series data. Only the predictions on the masked values contribute to the MSE loss. The pre-trained model can be applied to various downstream tasks, including classification. For classification, the final representation vectors from all time steps are concatenated into a single vector, which serves as input to a linear output layer with parameters corresponding to the number of classes.
- 6) **MOMENT** (Goswami et al., 2024) is an open-source foundation model for general-purpose time series analysis. It follows a masked time series modeling approach, where the goal is to reconstruct masked input time series using a lightweight reconstruction head. First, a univariate time series is segmented into disjoint fixed-length sub-sequences (patches). Before patching, reversible instance normalization (Kim et al., 2021) is applied to the time series. Each patch is then mapped to an embedding, using a trainable linear projection if all time steps are observed, and a designated learnable mask embedding if some patches are masked. During pretraining, patches are randomly masked by replacing their patch embeddings with a special learnable mask embedding. These patch embeddings are then fed into a transformer encoder to learn patch representations. The transformed patch embeddings are used to reconstruct both masked

and unmasked time series patches through a lightweight reconstruction head. Pretraining ensures that embeddings and high-level representations of the time series are effectively learned, with the objective of minimizing the masked reconstruction error (Mean Squared Error) between the ground truth and predicted patches.

To process multivariate time series, MOMENT operates independently on each channel, aligning with recent studies (Nie et al., 2023; Zhou et al., 2024) that support channel-wise modeling as an effective strategy. Unlike traditional architectures that use a decoder of similar size to the encoder, MOMENT employs a lightweight reconstruction head, enabling task-specific fine-tuning with minimal trainable parameters while preserving the encoder’s high-level learned features. For classification, MOMENT operates in a zero-shot setting by replacing its reconstruction head with a linear classification head, mapping patch representations to logits corresponding to class labels.

F Adapting Unsupervised Representation Learning Models for Classification

There are two approaches to adapting unsupervised representation learning baseline models for classification:

- **Two-stage approach:** First, sequence-level representations for each time series are obtained without access to labels, using the unsupervised representation learning model. In the second stage, a machine learning classifier (e.g., a Support Vector Machine with an RBF kernel) is trained on these representations with labeled data. The classifier applies a softmax function to generate a probability distribution over classes, with cross-entropy loss computed against the categorical ground truth labels. This approach is used for all unsupervised baselines except MOMENT, specifically for T-Loss, TNC, TS2Vec, TS-TCC, and TST.
- **Direct adaptation approach:** Instead of a two-stage process, the reconstruction head in the unsupervised representation learning model is replaced with a linear classification

head that maps patch representations to logits corresponding to class labels. MOMENT adopts this approach by replacing its reconstruction head with a linear head that outputs logits equal to the number of classes.

G Reproduction Details for Baselines

To facilitate the reproduction of all baselines and clarify how each baseline model is used, all baseline codes are available in our implementation repository at <https://anonymous.4open.science/r/LETS-C>. Additionally, Table 5 lists the original code sources for each baseline model. All models, except MOMENT, were trained on each dataset individually—either with labels for supervised learning methods or without labels for unsupervised representation learning methods.

H Implementation Details for LETS-C

The experiments were conducted on a Linux machine equipped with an NVIDIA T4 Tensor Core GPU with 16GB of memory. We utilized the PyTorch v2.4.0 deep learning platform running on Python 3.11 for all models. We set a fixed random seed for all experiments to ensure reproducibility. All configurations employed the RAdam optimizer Liu et al. (2019a), with its default hyperparameters settings $(\beta_1, \beta_2) = (0.9, 0.999)$. Further, we explored LETS-C’s performance across three different text embedding models: e5-mistral-7b-instruct (Wang et al., 2024), gte-large-en-v1.5 (Li et al., 2023), and nomic-embed-text-v1 (Nussbaum et al., 2024), in addition to text-embedding-3-large in Section 5.3. Exploratory hyperparameter optimization was conducted, revealing that 1-3 1D convolutional layers and 1-2 linear layers are optimal for the performance of LETS-C across all datasets. For a detailed description of the hyperparameters, see Appendix Section I. For tokenization, we found that maintaining a precision of one decimal place optimizes performance (see Appendix Section N for details), thus we adopted that precision throughout the experiments.

I Hyperparameter Settings for LETS-C

I.1 Hyperparameter Search Space

Our hyperparameter search for the lightweight classification head included determining the optimal number of convolutional layers, whether to use dropout or pooling within the convolutional blocks,

Table 5: Code source for each baseline.

Model	Source
<i>Supervised Learning Methods</i>	
DTW	https://github.com/markdregan/K-Nearest-Neighbors-with-Dynamic-Time-Warping/tree/master
XGBoost	https://github.com/dmlc/xgboost
ROCKET	https://github.com/angus924/rocket/tree/master
LightTS	https://github.com/thuml/Time-Series-Library/blob/main/models/LightTS.py
DLinear	https://github.com/thuml/Time-Series-Library/blob/main/models/DLinear.py
LSTM	https://anonymous.4open.science/r/LETS-C/baselines/models/supervised/lstm.py
LSTNet	https://github.com/laiguokun/LSTNet
LSSL	https://github.com/state-spaces/s4
TCN	https://github.com/White-Link/UnsupervisedScalableRepresentationLearningTimeSeries
TimesNet	https://github.com/thuml/Time-Series-Library/blob/main/models/TimesNet.py
Transformer	https://github.com/thuml/Time-Series-Library/blob/main/models/Transformer.py
Reformer	https://github.com/thuml/Time-Series-Library/blob/main/models/Reformer.py
Informer	https://github.com/thuml/Time-Series-Library/blob/main/models/Informer.py
Pyraformer	https://github.com/thuml/Time-Series-Library/blob/main/models/Pyraformer.py
Autoformer	https://github.com/thuml/Time-Series-Library/blob/main/models/Autoformer.py
Non-stationary Transformer	https://github.com/thuml/Time-Series-Library/blob/main/models/Nonstationary_Transformer.py
FEDformer	https://github.com/thuml/Time-Series-Library/blob/main/models/FEDformer.py
ETSformer	https://github.com/thuml/Time-Series-Library/blob/main/models/ETSformer.py
Flowformer	https://github.com/thuml/Flowformer/tree/main/Flowformer_TimeSeries
PatchTST	https://github.com/yuqinie98/PatchTST
One Fits All	https://github.com/DAMO-DI-ML/NeurIPS2023-One-Fits-All
<i>Unsupervised Representation Learning Methods</i>	
T-Loss	https://github.com/White-Link/UnsupervisedScalableRepresentationLearningTimeSeries
TNC	https://github.com/sanatonek/TNC_representation_learning
TS2Vec	https://github.com/zhihanyue/ts2vec
TS-TCC	https://github.com/emadeldeen24/TS-TCC
TST	https://github.com/gzerveas/mvts_transformer
MOMENT	https://github.com/moment-timeseries-foundation-model/moment
Ours	https://anonymous.4open.science/r/LETS-C

the number of dense layers, the type of activation function for each layer, and whether to employ batch normalization. We explored 1 to 5 convolutional layers and 1 to 4 linear layers. We also tested various learning rates, ranging from 0.001 to 0.05, different activation functions such as ReLU, GELU, and tanh, and truncation lengths for the text-embedding-large from 64 to 1024.

I.2 Optimal Hyperparameters

Table 6 presents the optimal configurations for our experiments across various datasets. All configurations employed tokenization with a precision of one decimal place and utilized the text-embedding-3-large embeddings. It is crucial to note that the count of linear layers includes the output layer; therefore, a configuration with one linear layer means that this single layer functions as the output layer within the model architecture. Furthermore, all convolutional layers utilized a kernel size of 3. Additionally, each configuration used the RAdam optimizer Liu et al. (2019a) with its default hyperparameter settings $(\beta_1, \beta_2) = (0.9, 0.999)$.

J Model Performance

Figure 3 displays a comparison of models based on the average classification accuracy across all datasets, as summarized in Table 4.

K Computational Cost Analysis

Table 7 discusses the training and inference times for our model as compared to the SOTA OneFitsAll (Zhou et al., 2024) across the 10 benchmark datasets. We observe that our approach significantly reduces the total training time to just 14.66% and the inference time to 31.67% of those required by OneFitsAll.

L Monetary Costs Analysis

Our approach incurs monetary costs when utilizing OpenAI’s text-embedding-3-large model (OpenAI, 2024). To compute the cost of generating text embeddings for our datasets, we followed the pricing structure of \$0.130 per million tokens. Each time series sample, including all time steps, was embedded separately for each dimension. To estimate the number of tokens per sample, we con-

Table 6: Optimal hyperparameter configuration for LETS-C.

LETS-C						
Dataset/Configuration	Model Hyperparameter				Training Process	
	Embedding dimension	Conv layers	Linear layers	Activation	Learning rate	Batch size
EthanolConcentration	64	2	1	tanh	0.007	64
FaceDetection	64	1	1	tanh	0.007	128
Handwriting	16	1	2	tanh	0.007	64
Heartbeat	512	1	1	tanh	0.007	64
Japanese Vowels	512	1	2	GELU	0.007	64
PEMS-SF	1024	3	1	tanh	0.007	64
Self-Regulation SCP1	64	1	1	tanh	0.007	64
Self-Regulation SCP2	1024	2	1	GELU	0.007	64
Spoken Arabic Digits	512	1	1	tanh	0.001	64
UWave Gesture Library	1024	1	1	tanh	0.001	64

Table 7: Comparison of training and inference times (in seconds) per batch for OneFitsAll (Zhou et al., 2024) versus our model across the 10 benchmark datasets. The Ratio (%) is calculated using the formula $100 \times \frac{\text{time of LETS-C}}{\text{time of OneFitsAll}}$, quantifying the computational efficiency of our model relative to OneFitsAll.

Model/Dataset		EC	FD	HW	HB	JV	PEMS-SF	SCP1	SCP2	SAD	UW	Average ↓
Time												
Training time (s)	LETS-C (Ours)	0.01	0.01	0.009	0.01	0.01	0.11	0.004	0.01	0.01	0.01	0.02
	OneFitsAll	0.35	0.23	0.14	0.07	0.02	0.28	0.11	0.13	0.97	0.07	0.24
	Ratio (%)	3.17	8.02	6.71	16.18	42.85	40.49	4.30	8.40	1.21	15.30	14.66
Inference time (s)	LETS-C (Ours)	0.01	0.05	0.01	0.01	0.01	0.10	0.004	0.01	0.02	0.01	0.02
	OneFitsAll	0.16	0.10	0.06	0.04	0.02	0.15	0.06	0.08	0.15	0.04	0.09
	Ratio (%)	10.58	50.50	14.75	31.15	68.41	66.10	7.02	19.59	18.00	30.55	31.67

sidered the numerical format of the time series values, where each time step consists of a three-digit number. These digits are separated by spaces, and consecutive time steps are delimited by commas, resulting in an estimated six tokens per time step. The total token count for each dataset was then computed as:

$$\text{Total Tokens} = \text{Total Samples} \times \text{Num. Dimensions} \times \text{Tokens per Series},$$

where the estimated tokens per series were derived from the series length multiplied by six. The dataset characteristics, including train and test sizes, number of dimensions, and series lengths, are detailed in Table 4. Using this approach, the total number of tokens across all datasets was approximately 1.05 billion, leading to an estimated embedding cost of \$137.07. It is important to note that API costs may fluctuate over time, leading to potentially different or outdated values. Furthermore, cost remains a common limitation for all LLM-based approaches that rely on closed-source

paid models.

M Assessing Text Embeddings with Cosine Similarity

Figure 4 visualizes the within-class and between-class cosine similarities of text embeddings derived from the testing time series. Each matrix entry is scaled using min-max normalization to range from 0 to 1, where warmer colors in the heatmap represent higher similarities and darker shades indicate lower similarities. Diagonal entries show within-class similarities, highlighting intra-class cohesion, while off-diagonal entries reveal between-class relationships

Similar to the training set, we observe that within-class similarity consistently exceeds across-class similarity, thereby validating the hypothesis that text embeddings effectively retain and convey significant information from the underlying time series data.

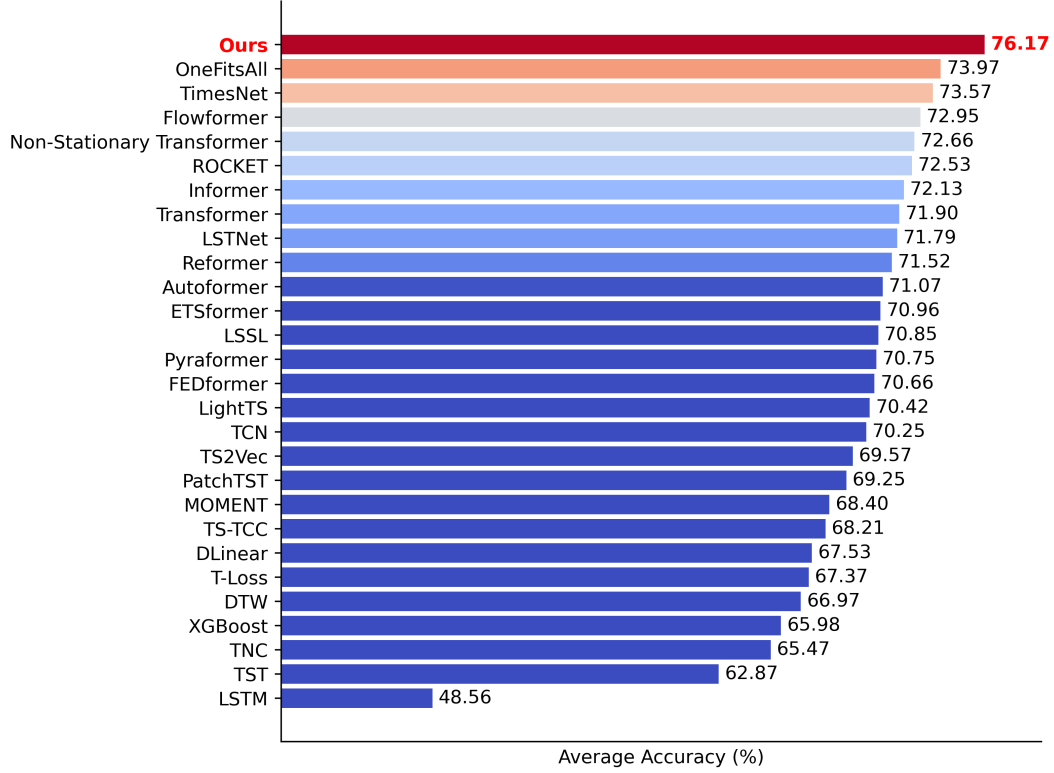


Figure 3: Model comparison based on classification accuracy, averaged across all datasets listed in Table 4. For detailed results, refer to Table 1 in the main paper.

N Numerical Precision for Tokenization

To explore the impact of numerical precision on the computation of embeddings, we analyzed classification accuracy across precisions 1 to 6 using four datasets: Handwriting, Heartbeat, JapaneseVowels, and UWaveGestureLibrary. We selected these datasets because they were the smaller ones among the 10 available, making the computation of embeddings more computationally affordable. Figure 5 illustrates the average classification accuracy (%) across these numerical precisions. Detailed results can be found in Table 8.

Studies in the NLP domain have shown that longer inputs do not perform well with language models (Press et al., 2020; Levy et al., 2024). Our empirical analysis supports this claim, revealing a decrease in classification accuracy with increased numerical precision when computing text embeddings. The average accuracy starts at 72.8% with precision 1 and declines to 69% at precision 6. Additionally, the percentage of AvgWins is highest at precision 1. As numerical precision increases, so does the length of the time series and the input to the text embedding model. Note that the maximum token length for text-embedding-3-large

embeddings is 8191, and thus keeping precision of 1 ensures that context length doesn’t exceed the maximum permissible token length. This issue is especially problematic for datasets with longer time series, such as the EthanolConcentration dataset, which includes 1751 time steps per sample. This finding led us to opt for precision 1 in our study.

Note that these precision results also depend on the type of tokenization selected, and defining an appropriate tokenization for time series is one of the potential future directions for this work.

O Generalization Across Various Text Embedding Models

To evaluate the generalization capabilities of our approach across different text embedding models beyond text-embedding-3-large (OpenAI, 2024), the performance of LETS-C was tested using three alternative embedding models: e5-mistral-7b-instruct (Wang et al., 2024), gte-large-en-v1.5 (Li et al., 2023), and nomic-embed-text-v1 (Nussbaum et al., 2024). A summary of the embedding models utilized in this study is provided in Table 9.

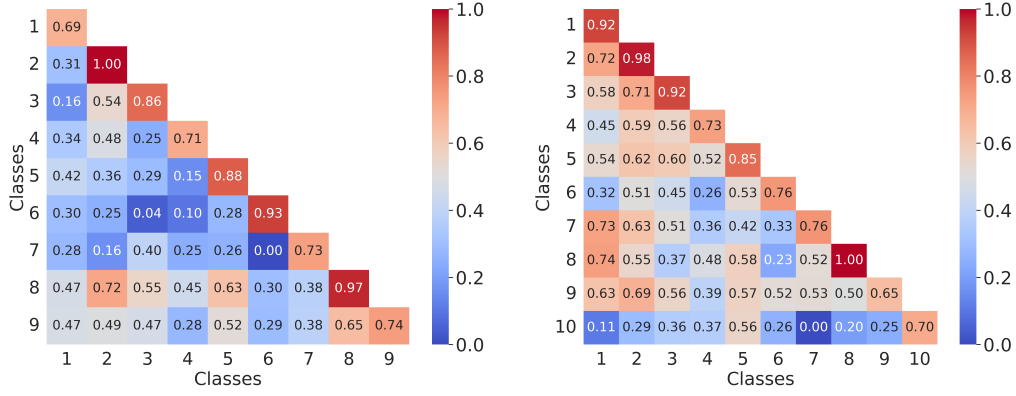


Figure 4: Heatmaps illustrating within-class and between-class cosine similarities of text embeddings derived from the testing time series data in Japanese Vowels (**left**) with 9 classes and Spoken Arabic Digits (**right**) with 10 classes. On both axes, x and y represent different classes. Diagonal entries indicate within-class similarities, and off-diagonal entries represent between-class similarities. Warmer colors signify higher cosine similarities, while cooler colors suggest lower similarities.

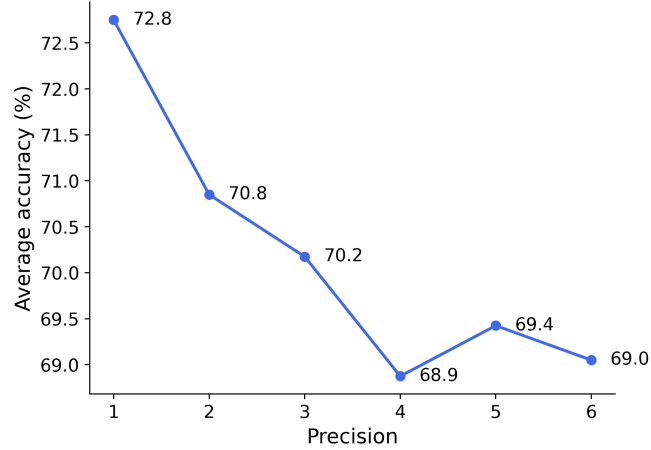


Figure 5: Average classification accuracy (%) across numerical precisions 1 to 6. Results are averaged from four datasets: Handwriting, Heartbeat, JapaneseVowels, and UWaveGestureLibrary. See Table 8 in the Appendix for detailed results.

O.1 e5-mistral-7b-instruct

The e5-mistral-7b-instruct model (Wang et al., 2024) is based on the pretrained Mistral-7b checkpoint (Jiang et al., 2023) and was fine-tuned using the RankLLaMA training methodology (Ma et al., 2024). It employed LoRA (Hu et al., 2021) with a rank of 16 on a mixture of multilingual datasets, which included both synthetic data and a collection of 13 public datasets, yielding approximately 1.8 million examples after sampling. Proprietary LLMs, such as GPT-4, were prompted to generate diverse synthetic data with instructions in multiple languages. Leveraging the strong language understanding capabilities of the Mistral model, e5-mistral-7b-instruct achieved state-

of-the-art results across nearly all task categories on the competitive MTEB benchmark. Techniques such as gradient checkpointing, mixed precision training, and DeepSpeed ZeRO-3 were applied to further reduce GPU memory requirements.

O.2 gte-large-en-v1.5

The gte-large-en-v1.5 model (Li et al., 2023) is a general text embedding (GTE) model that utilizes contrastive learning on an open-source large-scale dataset comprising unsupervised text pairs extracted from various sources. To enhance the quality of the learned text representations, high-quality text pairs with human labels from multiple sources were employed for contrastive fine-tuning.

Table 8: Numerical Precision for Tokenization: This table reports classification accuracy (%). **Red**: Best performance.

Dataset / Precision	1	2	3	4	5	6
Handwriting	23.2	15.9	11.6	11.1	12.0	11.2
Heartbeat	78.0	78.5	79.0	78.5	79.5	78.0
JapaneseVowels	99.2	98.4	97.6	96.8	95.9	95.1
UWaveGestureLibrary	90.6	90.6	92.5	89.1	90.3	91.9
Average $\uparrow$	72.75	70.85	70.175	68.875	69.425	69.05
AvgWins % $\uparrow$	50%	0%	25%	0%	25%	0%

Table 9: Summary of selected embedding models used in the study.

MTEB Rank	Model	Embedding Dimensions	Max Token Length
15	text-embedding-3-large (OpenAI, 2024)	3072	8191
6	e5-mistral-7b-instruct (Wang et al., 2024)	4096	32768
9	gte-large-en-v1.5 (Li et al., 2023)	1024	8192
35	nomic-embed-text-v1 (Nussbaum et al., 2024)	768	8192

The model utilizes a multi-stage contrastive learning approach and benefits from a diverse training data mixture, enabling it to achieve strong generalization performance for single-vector embeddings.

O.3 nomic-embed-text-v1

The nomic-embed-text-v1 model (Nussbaum et al., 2024) is a fully reproducible, open-source English text embedding model with an 8192 context length that outperforms both OpenAI Ada-002 and OpenAI text-embedding-3-small on short and long-context tasks. To accommodate long sequence lengths, the model adapts the BERT architecture (Devlin et al., 2019) with several optimizations: replacing absolute positional embeddings with rotary positional embeddings (Su et al., 2024), using SwiGLU activation instead of GeLU (Shazeer, 2020), implementing Flash Attention (Dao et al., 2022), setting Dropout to 0 (Geiping and Goldstein, 2023), and ensuring the vocabulary size is a multiple of 64 (Portes et al., 2024; Shoybi et al., 2019). These modifications result in a 137M parameter encoder.

P Trade-offs: Model Accuracy vs. Parameter Complexity

Table 10 illustrates the trade-off between model accuracy and the complexity of training parameters in our model, which utilizes both embeddings and time series data as inputs to a lightweight frame-

work. To vary model size, we adjust the number of linear and convolution layers in the classification head, ranging from 1 to 5 layers each, creating model variants of different sizes.

Across all datasets, we find that significant reductions in model parameters result in only a minimal loss of accuracy. This trade-off between model accuracy and parameter complexity is data-dependent. However, we generally observe a trend where a reduction in parameters leads to only a slight decrease in accuracy. Next, let’s take a closer look at three datasets: Heartbeat, PEMS-SF, and Spoken Arabic Digits, to understand how these trade-offs manifest in different contexts.

- **Heartbeat:** In the Heartbeat dataset, we retain 99.48% of the optimal model’s accuracy using only 75% of its trainable parameters. Specifically, the optimal model achieves an accuracy of 78% with 46,426 trainable parameters, while the second-best model achieves 77.6% accuracy with 34,820 parameters. Moreover, we retain 96.28% accuracy with a further reduction to 57.95% of the parameters.
- **PEMS-SF:** In the PEMS-SF dataset, the optimal model starts with an accuracy of 93.1% and 564,231 trainable parameters. Reducing the parameters to 173,866 (30.81% of the original), the model maintains 98.06% of its optimal accuracy at 91.3%. Further param-

Table 10: Trade-off between model accuracy and the complexity of training parameters. The accuracy and parameters of the best model are highlighted in **bold**. The accuracy difference is calculated as the raw difference between the accuracies of the reduced model and the best model. The % Delta in accuracy and parameters is defined separately for each as $100 \times \frac{\text{Accuracy of the reduced model}}{\text{Accuracy of the optimal model}}$ and $100 \times \frac{\text{Parameters of the reduced model}}{\text{Parameters of the optimal model}}$, quantifying the accuracy and computational efficiency of the reduced model relative to our best model.

Dataset	Accuracy (%) $\uparrow$	Trainable Parameters $\downarrow$	Difference % Delta in Accuracy $\uparrow$	% Delta in Parameters $\downarrow$
EthanolConcentration	52.9	283950	-	-
	46	105344	-6.9 86.95	37.09
FaceDetection	68.9	3842	-	-
	68.6	2402	-0.3 99.56	62.51
	67.9	962	-1.0 98.54	25.03
	66.4	482	-2.5 96.37	12.54
Handwriting	23.8	154526	-	-
	23.2	107226	-0.6 97.47	69.39
	22.7	53626	-1.1 95.37	34.70
	20.2	20394	-3.6 84.87	13.19
	19.4	13426	-4.4 81.51	8.68
Heartbeat	78	46426	-	-
	77.6	34820	-0.4 99.48	75.00
	75.1	26908	-2.9 96.28	57.95
Japanese Vowels	99.2	148233	-	-
	98.9	123401	-0.3 99.69	83.24
	98.6	105353	-0.6 99.39	71.07
	98.4	100857	-0.8 99.19	68.03
PEMS-SF	93.1	564231	-	-
	91.3	173866	-1.8 98.06	30.81
	90.8	85210	-2.3 97.52	15.10
	87.9	69077	-5.2 94.41	12.24
	87.3	62901	-5.8 93.77	11.14
Self-Regulation SCP1	93.2	302626	-	-
	92.5	99657	-0.7 99.24	32.93
Self-Regulation SCP2	62.8	334402	-	-
	58.9	166306	-3.9 93.78	49.73
	57.8	111106	-5.0 92.03	33.22
	57.2	83330	-5.6 91.08	24.91
	56.1	76386	-6.7 89.33	22.84
Spoken Arabic Digits	99.2	141790	-	-
	99	70066	-0.2 99.79	49.41
	98.4	46658	-0.8 99.19	32.90
	98.1	30964	-1.1 98.89	21.83
	98	20646	-1.2 98.79	14.56
	97.8	15487	-1.4 98.58	10.92
	97.6	10328	-1.6 98.38	7.28
	96.6	5308	-2.6 97.37	3.74
UWave Gesture Library	90.6	263338	-	-
	88.4	211556	-2.2 97.57	80.33
	87.5	176298	-3.1 96.57	66.94

eter reductions to 85,210 (15.10%), 69,077 (12.24%), and 62,901 (11.14%) result in accuracies of 90.8%, 87.9%, and 87.3%, respectively. These reductions illustrate that even significant reductions in parameters only lead to a slight decrease in performance.

- **SpokenArabicDigits:** For the Spoken Arabic Digits dataset, reducing the number of trainable parameters generally correlates with a minor decline in accuracy, though the trade-off is modest. The optimal model, achieving an accuracy of 99.2% with 141,790 trainable parameters, shows that even with substantial reductions to 70,066 parameters (49.41% of the original), the accuracy remains high at 99%, retaining 99.79% of the original model’s accuracy. Further reductions to 46,658 (32.90%), 30,964 (21.83%), 20,646 (14.56%), 15,487 (10.92%), 10,328 (7.28%), and 5,308 (3.74%) yield accuracies of 98.4%, 98.1%, 98%, 97.8%, 97.6%, and 96.6% respectively.

These examples highlight that efficiency in terms of trainable parameters does not linearly correspond to a loss in model accuracy across various datasets. While fewer parameters generally lead to a lower accuracy, the decrement is often proportional and manageable, making these models highly suitable for deployment in resource-constrained environments or for applications requiring rapid processing with minimal computational overhead.

Q Ablation Study

Table 11 presents the results of this study, comparing the performance of our default approach, which adds embeddings to time series, to variants that exclude either embeddings or the time series.

R Alternative Methods for Fusing Time Series with Embeddings

We explored two additional methods for fusing time series and embeddings beyond simple addition. The first method involves a *Fusion network* that first processes embeddings and time series data through convolutional and dense layers in two separate branches, then merges the features from both branches into a final dense network. The second method employs *Concatenation*, where the time series and embeddings are concatenated and processed through a lightweight classification head.

Despite cross-attention being another alternative for fusing different modalities, we didn’t include it in this study due to the computational complexity it adds to the model. Table 12 presents the classification accuracy and trainable model parameters for these variations.

We observe that the addition approach in the LETS-C architecture achieves the highest average classification accuracy (76.11%) compared to the fusion network (73.40%) and concatenation (74.22%). Notably, when compared to all baselines, concatenation emerges as the second-best method after addition, which achieves the highest accuracy. The fusion approach ranks fifth overall, with only OneFitsAll and TimesNet outperforming it, following addition and concatenation, in terms of average accuracy.

Table 11: Comparison of classification accuracy (%) for different configurations: ours (both embeddings and time series), embeddings only, and time series only. **Red**: Best.

Method/Dataset	EC	FD	HW	HB	JV	PEMS-SF	SCP1	SCP2	SAD	UW	Average $\uparrow$	AvgWins % $\uparrow$
LETS-C (Ours)	52.9	68.9	23.8	78	99.2	93.1	93.2	62.8	99.2	90.6	76.17	60%
embedding only	38	59.4	20.1	80	99.2	92.5	88.7	63.9	99.2	88.4	72.94	40%
time series only	42.6	69.6	25.1	76.1	98.1	89	93.2	58.3	98.9	83.8	73.47	30%

Table 12: Comparison of classification accuracy (%) and trainable model parameters (millions) for alternative methods of fusing time series with embeddings. Higher AvgWins and accuracy averages indicate superior performance, while lower model parameter averages suggest greater computational efficiency. **Red**: Best performance.

Accuracy $\uparrow$												
Method/Dataset	EC	FD	HW	HB	JV	PEMS-SF	SCP1	SCP2	SAD	UW	Average $\uparrow$	AvgWins % $\uparrow$
LETS-C (Addition)	52.9	68.9	23.8	78	99.2	93.1	93.2	62.8	99.2	90.6	76.17	70%
Fusion network	44.1	66.5	23.6	76.6	98.1	86.1	92.8	56.1	99.2	90.9	73.40	20%
Concatenation	43	65.1	22.5	79	98.9	93.1	93.9	58.9	99	88.8	74.22	30%

Trainable Parameters (M) $\downarrow$												
Method/Dataset	EC	FD	HW	HB	JV	PEMS-SF	SCP1	SCP2	SAD	UW	Average $\downarrow$	
LETS-C (Addition)	0.28	0.003	0.15	0.04	0.14	0.56	0.30	0.33	0.14	0.26	0.22	
Fusion Network	0.19	0.35	0.33	0.28	0.16	5.54	0.37	0.27	0.23	0.04	0.78	
Concatenation	0.42	0.009	0.26	0.17	0.11	0.28	0.23	0.39	0.09	0.46	0.24	