

Time-MQA: Time Series Multi-Task Question Answering with Context Enhancement

Anonymous ACL submission

Abstract

Time series data are foundational in finance, healthcare, and energy domains. However, most existing methods and datasets remain focused on a narrow spectrum of tasks, such as forecasting or anomaly detection. To bridge this gap, we introduce Time Series Multi-Task Question Answering (Time-MQA), a unified framework that enables natural language queries across multiple time series tasks - numerical analytical tasks and open-ended question answering with reasoning. Central to Time-MQA is the TSQA dataset, a large-scale dataset containing $\sim 200k$ question-answer pairs derived from diverse time series spanning environment, traffic, etc. This comprehensive resource covers various time series lengths and promotes robust model development. We further demonstrate how continually pre-training large language models (Mistral 7B, Llama-3 8B, and Qwen-2.5 7B) on the TSQA dataset enhanced time series reasoning capabilities, moving beyond mere numeric tasks and enabling more advanced and intuitive interactions with temporal data. The complete TSQA dataset, models, executable codes, user study questionnaires for evaluation, and results have all been open-sourced here¹.

1 Introduction

Time series analysis has long been fundamental to real-world applications in finance, healthcare, energy, and other domains (Nie et al., 2024; Xu et al., 2023; Yang et al., 2021b; Chen et al., 2024b). Before the emergence of large language models (LLMs), most research in this area focused on numerical analytical tasks such as forecasting and anomaly detection, with methods typically constrained to a single objective. Recently, the community has begun to explore ways to enhance these conventional time series approaches by incorporating LLMs (Fons et al., 2024; Zhang et al., 2024b);

¹<https://anonymous.4open.science/r/Time-MQA-6085/>

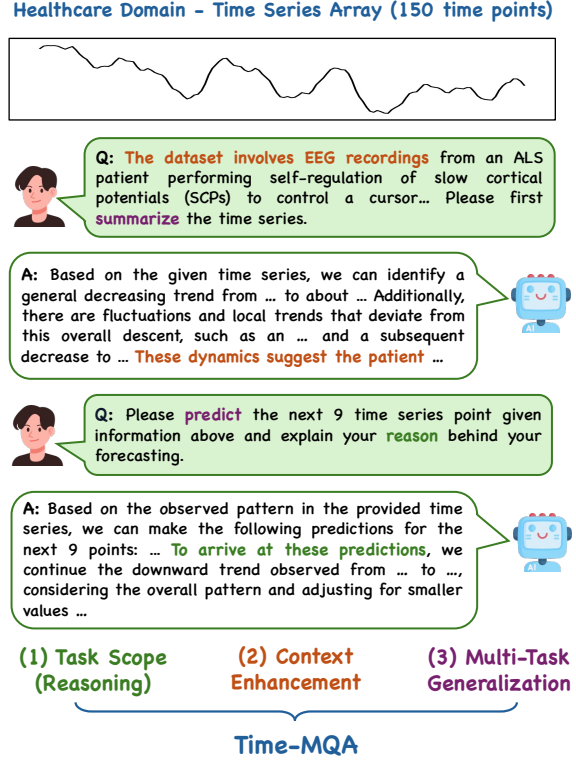


Figure 1: Example of Time-MQA with context enhancement. The LLM answers the user’s questions based on the provided context and the input time series. It can conduct reasoning and handle multiple time series tasks.

however, while some efforts do leverage contextual information, they typically focus on a single task, such as forecasting, leaving a gap in broader, multi-task reasoning and inference capabilities (Merrill et al., 2024; Ansari et al., 2024; Frisoni et al., 2024; Jin et al., 2024a). To bridge this gap, we propose a unified *Time Series Multi-task Question Answering (Time-MQA)* framework that integrates diverse tasks with natural language queries.

Time-MQA is crucial for advanced reasoning and inference, as it enables models to interpret temporal data through natural language queries and uncover deeper insights beyond mere classi-

cal numeric tasks (Chow et al., 2024; Xu et al., 2024). However, there is a notable lack of relevant datasets — specifically, paired language with time series — which severely limits the development of models capable of dialogue and reasoning in this domain (Jin et al., 2024b; Kong et al., 2025). This significant gap underscores the urgent need for comprehensive resources that integrate multiple time series tasks under a question-answering framework, allowing users to query these tasks in natural language.

In this paper, we introduce *Time-MQA*, a new framework for multi-task time series question answering with context enhancement. As illustrated in Figure 1, *Time-MQA* expands beyond numerical analytical tasks by consolidating open-ended questions and classical time series tasks, such as forecasting, into a single framework. Users can pose queries in natural language, enabling a more intuitive and flexible interface for a wide range of time series analyses.

To support *Time-MQA*, we proposed *TSQA*, a large-scale dataset with approximately 200k question-answer pairs spanning multiple domains, including healthcare, environment, AIOps, machine sensors, finance, energy, traffic, IoT, nature, transport, human activities, and the web. This dataset covers various time series lengths and tasks, ensuring broad coverage and robustness. Notably, *TSQA* features open-ended reasoning questions with more elaborate text-based explanations. To the best of our knowledge, this is the first large-scale QA dataset in the time series domain that spans multiple domains and tasks, effectively bridging the gap between classical time series analysis and modern LLM-driven approaches.

We further demonstrate the utility of the *TSQA* dataset by employing continual pre-training techniques on representative LLMs such as Mistral 7B (Jiang et al., 2023), Llama-3 8B (Dubey et al., 2024), and Qwen-2.5 7B (Yang et al., 2024). Our experiments reveal that these models trained on the *TSQA* dataset can effectively acquire time series knowledge and reasoning abilities, enabling more advanced capabilities beyond basic numeric handling of time series data.

In summary, our contributions include:

- We propose *Time-MQA*, a multi-task time series question answering framework that leverages contextual enhancement to extend beyond traditional numerical analysis tasks.

- We construct *TSQA*, a dataset comprising $\sim 200k$ question-answer pairs across over twelve domains (e.g., healthcare, finance, and energy) and five tasks (e.g., forecasting, anomaly detection, and open-ended reasoning question answering).
- We demonstrate that fine-tuning LLMs on the *TSQA* dataset equips them with time series-specific knowledge, enabling natural language queries for comprehensive time series analysis.

2 Related Works

2.1 Classical Time Series Tasks

Time series analysis has been extensively studied in various real-world applications, such as finance, healthcare, climate, electricity, AIOps, and industrial system maintenance (Nie et al., 2024; Guo et al., 2024; Ma et al., 2024; Yang et al., 2023a, 2021a). Classical time series tasks focus on extracting insights from these time series and addressing challenges associated with temporal patterns (Fuller, 2009; Hamilton, 2020).

Forecasting is one the most fundamental tasks. It predicts future time points based on historical values and features (Lim and Zohren, 2021). Depending on the temporal horizon, forecasting can be categorized into short-term forecasting, which captures immediate fluctuations, and long-term forecasting, which models more complex temporal dependencies (Wang et al., 2024b). Common methods include statistical methods (e.g., ARIMA and exponential smoothing) and deep learning-based methods (e.g., RNN, LSTM, Transformer-based architectures) (Wen et al., 2022; Miller et al., 2024).

Anomaly detection seeks to identify abnormal patterns or deviations from expected behavior in time-series data (Zamanzadeh Darban et al., 2024). Classical approaches rely on statistical models like z-score analysis and dynamic thresholding, while contemporary methods incorporate deep learning-based frameworks, such as autoencoders, transformers, and graph neural networks, to capture complex dependencies and temporal correlations (Han et al., 2022; Yang et al., 2023b).

Imputation addresses the issue of missing or corrupted data in sequence. It is essential for ensuring data integrity in downstream analysis (Du et al., 2024). Traditional imputation techniques, such as interpolation, have been widely employed, whereas deep learning-based methods, such as variational autoencoders (VAE), generative adversarial

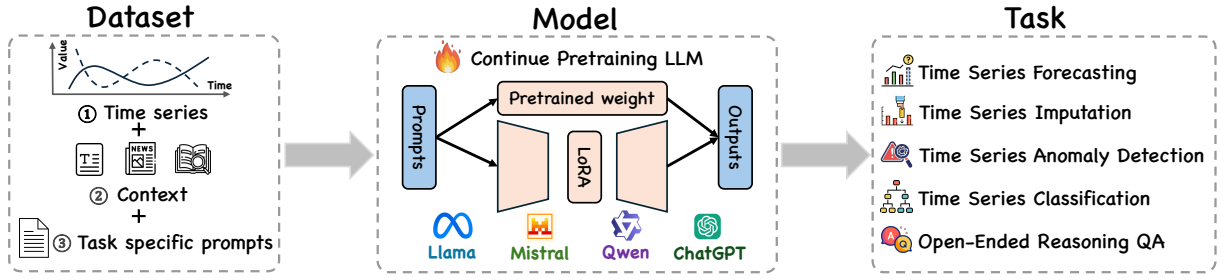


Figure 2: The overview of the proposed Time-MQA framework.

networks (GANs), and diffusion models, have recently demonstrated promising results in learning complex missing patterns and improving imputation performance (Wang et al., 2025).

Beyond these tasks, time series classification/regression, generation, augmentation, and decomposition are also frequently used in real-world scenarios (Mohammadi Foumani et al., 2024; Wen et al., 2020; Zhang et al., 2024a).

2.2 Text-Enhanced Time Series Tasks

Recent advancements in time series analysis have demonstrated the potential of incorporating textual information based on LLMs to enhance time series tasks (Jin et al., 2024b; Wang et al., 2024a). Unlike classical approaches that rely solely on numerical data, text-enhanced time series analysis leverages domain-specific textual descriptions, contextual metadata, or associated reports to improve the cognitive understanding and modeling of time-dependent patterns (Liu et al., 2024c). This hybrid approach mitigates the limitations of unimodal time series models by integrating additional semantic and contextual cues that are useful for decision-making (Kong et al., 2025; Singh et al., 2024).

In detail, text-enhanced time series forecasting and anomaly detection tasks benefit from textual information by incorporating expert reports to refine predictions and provide anomaly causal explanations (Hollmann et al., 2025; Chen et al., 2023). Similarly, classification, imputation, and generation tasks can be enhanced by leveraging textual descriptions as auxiliary supervision or describing missing values (Bernardini et al., 2023; Moor et al., 2023). It will help models distinguish subtle variations across different categories and generate more informed reconstructions.

Recent research has explored various techniques for integrating textual and time-series data, including LLM-based alignment, cross-modal attention

mechanisms, and contrastive learning strategies that jointly encode text and time-series representations (Jin et al., 2023b; Zhang et al., 2023; Liu et al., 2024a). Some approaches, such as Time-LLM (Jin et al., 2023a), directly adapt LLMs to process text and time-series data, whereas others, like Time-MMD (Liu et al., 2024b), employ weighted fusion methods to combine textual embeddings with deep time-series backbones. By enhancing classical time series tasks with textual information, text-enhanced time series models offer greater robustness and richer interpretability across diverse applications (Jin et al., 2024b; Kong et al., 2025).

2.3 Language Question Answering

Question Answering (QA) in Natural Language Processing (NLP) involves systems that interpret human language queries to retrieve or generate accurate answers (Biancofiore et al., 2024; Chen et al., 2024c). It evolves from rule-based systems to neural architectures driven by LLMs like GPT-4 (Achiam et al., 2023) and Llama (Touvron et al., 2023). These models leverage massive text corpora and large-scale datasets for end-to-end pre-training, fine-tuned via supervised learning or reinforcement learning with human feedback (RLHF) to align responses with factual and contextual relevance (Liu et al., 2023). Innovations such as retrieval-augmented generation (RAG) combine parametric knowledge with external data sources, while benchmarks like SQuAD, HotpotQA, MuSiQue, FinTextQA, SyllabusQA, and ToolQA drive progress (Trivedi et al., 2022; Ho et al., 2020; Yang et al., 2018; Zhuang et al., 2024; Chen et al., 2024a; Fernandez et al., 2024). However, challenges remain, such as handling ambiguous queries, ensuring the accuracy of generated answers, and maintaining efficiency in processing large volumes of data. Ongoing research focuses on enhancing the reasoning capabilities of QA sys-

tems, improving their ability to handle complex and nuanced questions, and expanding their applicability across diverse domains (Singh et al., 2025).

3 Methodology

3.1 The Time-MQA Framework

The Time-MQA framework broadens traditional analysis by integrating diverse objectives - such as forecasting, imputation, classification, anomaly detection, and notably, open-ended queries - into a unified question-answering paradigm (as shown in Figure 2. Formally, let $\mathbf{X} = \{x_1, x_2, \dots, x_T\}$ denote a time series input, where $x_t \in \mathbb{R}^d$ represents a d -dimensional observation at timestep t . Let $\mathbf{C}$ represent additional contextual information (e.g., textual metadata, domain-specific knowledge, or other modalities). For a question Q expressed in natural language, the goal is to generate an answer A conditioned on both X and contextual information C . Time-MQA aims to learn a function $f : (\mathbf{X}, \mathbf{C}, Q) \rightarrow A$, where A is the correct answer to the query Q . Depending on the nature of Q , A can take diverse forms, such as predicted values, classification labels, a set of anomalous timestamps, and textual explanations. Specifically, the model component of Time-MQA is based on continued pre-trained LLMs (i.e., Mistral 7B (Jiang et al., 2023), Llama-3 8B (Dubey et al., 2024), and Qwen-2.5 7B (Yang et al., 2024)), using the prepared TSQA dataset. To optimize parameter usage, Time-MQA employs Parameter-Efficient Fine-Tuning (PEFT) with a LoRA adapter. The supported various tasks of the Time-MQA with context enhancement are demonstrated with details in Figure 3.

Key Distinctions. There are three main differences between Time-MQA and traditional time series analysis tasks:

- (1) **Task Scope:** Traditional tasks focus on singular objectives (e.g., forecasting future values or classifying the time series). In contrast, Time-MQA unifies these under a question-driven paradigm, enabling both conventional tasks (e.g., “Forecast the next 5 values”) and complex queries (e.g., “Why did the temperature drop abruptly at hour 10?”) that require joint reasoning across detection, explanation, and contextual knowledge.
- (2) **Context Enhancement:** Traditional methods rely solely on the time series $\mathbf{X}$. Time-

MQA integrates auxiliary context $\mathbf{C}$ to resolve ambiguities and improve robustness. For instance, even when analyzing the same time series, differences in a dataset’s background information can yield contrasting predictions (Williams et al., 2024).

- (3) **Multi-Task Generalization:** Unlike single-purpose models, Time-MQA dynamically adapts to diverse question types through a unified architecture, eliminating the need for task-specific pipelines. This flexibility allows it to handle various time series tasks within a single framework, promoting knowledge sharing via shared representations that enhance both performance and interpretability.

3.2 The TSQA Dataset

This section will introduce our self-constructed 200k-level text-enhanced time series dataset, i.e., TSQA. From perspectives such as the data collection process and types across multiple domains and tasks. We compare it with other datasets and demonstrate the advantages of TSQA.

3.2.1 Dataset Composition and Categorization

The raw data we use comes from various classic publicly available datasets in the time series domain, covering multiple tasks and application areas. Innovatively, we have incorporated textual descriptions, including background information, feature descriptions, etc. From a task perspective, our dataset can be categorized into the following types. Some examples of the proposed TSQA dataset can be found in the Appendix.

Forecasting. In the forecasting task, we utilize UTSD datasets (Liu et al.), publicly available time series forecasting datasets, such as ETTh1, ETTh2, ETTm1, ECL, Weather, etc. (Zhou et al., 2021), and the latest text-enhanced time series datasets as the raw time-series data. Additionally, for the first time, we incorporate financial datasets with earnings call transcripts (The Motley Fool, 2024). To ensure the generalization ability of algorithms and models, we set the input length and prediction length to random values between 64–256 and 8–32, respectively. Furthermore, we enhance all data with background information, feature descriptions, and task descriptions as textual information based on the original data sources. In summary, the forecasting dataset includes 42,557 data instances spanning application domains such as energy, environment,

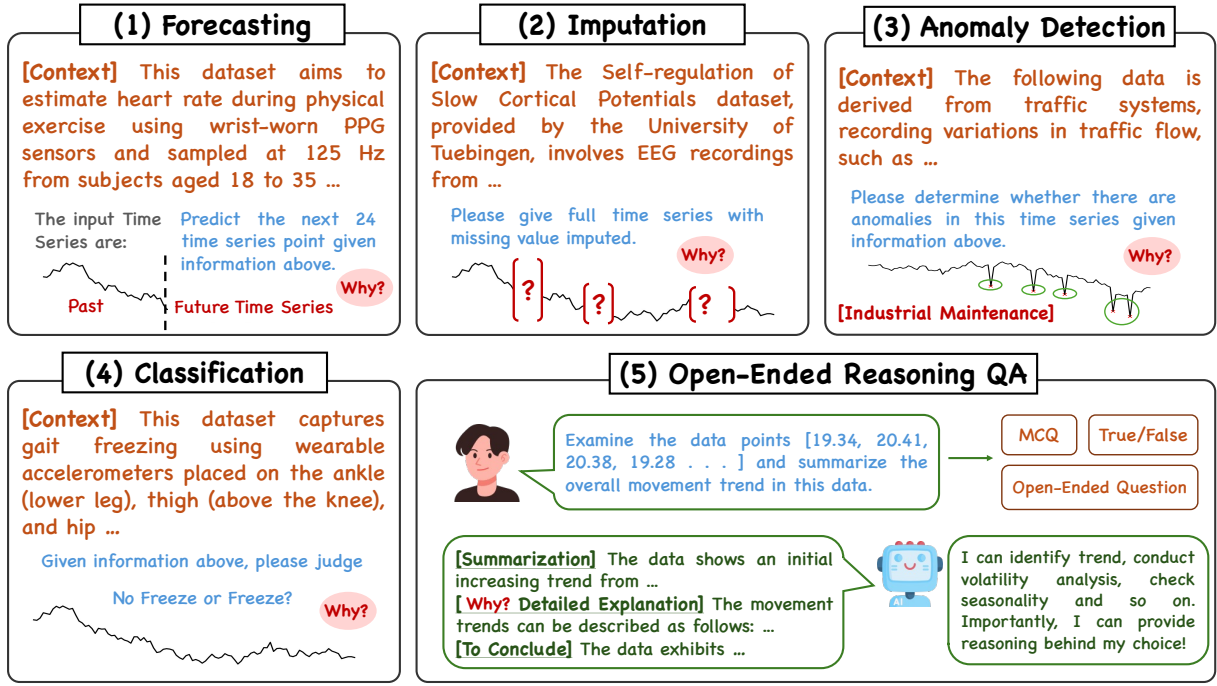


Figure 3: The demonstration of the Time-MQA with context enhancement.

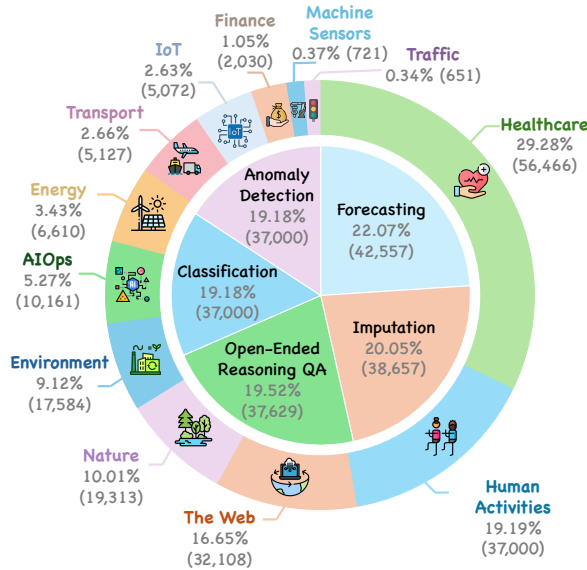


Figure 4: Distribution of data statistics in TSQA dataset.

health, IoT, nature, transport, Web, AIOps, etc.

Imputation. The imputation task shares the same original time series data sources as the forecasting task. Additionally, we randomly set the input length between 64 and 256, randomly removed 4 to 12 values, and replaced them with "X". The imputation dataset consists of a total of 38,657 instances from multiple application domains.

Anomaly Detection. For the anomaly detection task, we utilize commonly used public datasets in the field, such as the NASA Series dataset, UCR, ECG, KPI, MGAB, NAB, SensorScope, and Yahoo, among others (Zamanzadeh Darban et al., 2024). To enhance the generalization capability of algorithms and models, we randomly set the input length between 8 and 256 while ensuring a balanced distribution between anomalous and normal data. The anomaly detection dataset comprises 37,000 instances from various application domains, including AIOps, Health, Finance, Machinery, Industrial Sensors, Environment, and Traffic.

Classification. For the classification task, the data sources are relatively straightforward, primarily derived from the human activity recognition application domain (Kwapisz et al., 2011; Bachlin et al., 2009). The dataset includes both binary and multi-class classification tasks. To ensure class balance, we collected a total of 37,000 data instances. Additionally, the input sequence length is randomly set between 8 and 32 to maintain diversity in data representation.

Open-Ended Reasoning QA. We used different parts of the UTSD dataset (Liu et al.) - distinct from those employed in our forecasting tasks - to avoid data leakage. To generate open-ended reasoning QA, we utilized GPT-4o, instructing it

Table 1: Comparison of text-enhanced time series datasets. ☉ and ✓ indicate having only real data and having both synthetic and real data.

Dataset	Numerical	Cognitive	Source	Domain	Size
TS-Insights	✓	✗	☉	4	~10k
ChatTS	✗	✓	✓	4	~2.2k
CiK	✓	✗	☉	7	~2.9k
TimeMMD	✓	✗	☉	9	~16k
TSQA	✓	✓	✓	12	~200k

to create questions covering various topics such as trends, seasonality, cyclical patterns, summarization, volatility, anomalies, structural breaks, and other statistical properties (without limiting it strictly to these areas). We also incorporated multiple question types, including multiple-choice, true/false, and open-ended formats. The prompts used to generate these QA pairs are provided in the Appendix. In total, we generated 37,629 data instances, from which we manually reviewed and selected 1,400 QA pairs for use in the continual pretraining step.

3.2.2 Data Statistics

The TSQA dataset comprises 192,843 (~200k) question-answer pairs spanning twelve domains - healthcare, finance, energy, traffic, environment, IoT, nature, transport, human activities, machine sensors, AIOps, and the web - and five task types: forecasting, imputation, anomaly detection, classification, and open-ended reasoning (see Figure 4 for distributions). Within the open-ended reasoning QA, the dataset includes 6,919 true/false questions, 11,281 multiple-choice questions, and 12,510 open-ended questions, offering a broad and diverse range of question formats.

3.2.3 Comparison with Existing Datasets

We summarize several existing datasets compared to our proposed TSQA dataset in Table 1. Unlike prior datasets (TS-Insights (Zhang et al., 2023), ChatTS (Xie et al., 2024), CiK (Williams et al., 2024), and TimeMMD (Liu et al., 2024b)), which focus on either classical numerical analytical time series tasks (e.g., forecasting, anomaly detection, imputation, classification) or text-enhanced cognitive tasks (e.g., reasoning, QA), our TSQA dataset supports both, making it the most comprehensive benchmark. It also covers the widest range of application domains (12 vs. a maximum of 9) and is significantly larger (~200k instances vs. ~10k–16k). Furthermore, TSQA includes both real and syn-

Hyperparameter	Assignment
Base model	Mistral 7B
Computing infrastructure	1*A100-80GB GPU
Max steps	4000
Warm-up steps	1000
Batch size per device	4
Gradient accumulation steps	8
Learning rate	5e-5
Embedding learning rate	1e-5
Optimizer	AdamW (8-bit)
Learning rate scheduler	Cosine
Weight decay	0.1
LoRA rank (r)	16
LoRA alpha	16
LoRA dropout	0.0
LoRA target modules	q_proj, k_proj, v_proj, o_proj, gate_proj, up_proj, down_proj
Training time	~1 Day

Table 2: Hyper-parameters and training time of fine-tuning the Mistral 7B based on the TSQA dataset.

thetic data, ensuring greater diversity and robustness for text-enhanced time series analysis.

4 Experiment and Result

4.1 Experimental Settings

In the experiments, we considered the ratio of application domains and data sources by randomly selecting 1,400 QA pairs for each task type - forecasting, imputation, anomaly detection, and classification - and by manually reviewing and selecting 1,400 QA pairs for open-ended reasoning, leading to a total of 7,000 QA pairs. We then followed the setting from (Cheng et al., 2024). Specifically, to ensure the model adequately learns the time-series domain, we then combined our dataset with a general QA corpus sourced from OpenOrca (Lian et al., 2023) at a 70% to 30% ratio, resulting in 10k QA pairs overall. Finally, we formatted all QA pairs so that the question and answer were clearly labeled, and then we tokenized the text. An example of a formatted QA pair from the pre-tokenized text is provided in the Appendix. All training runs were conducted on a single A100 80GB GPU. For a detailed overview of the hyperparameters used and training time, please refer to Table 2, using Mistral 7B for example.

4.2 Results

Main Results. Table 3 presents the performance of three fine-tuned models, alongside the GPT-4o (Achiam et al., 2023) model and Doubao (Volcengine, 2023) results, across multiple tasks. For

Table 3: Comparison of our three fine-tuned models, GPT-4o, and Doubao across diverse tasks. Forecasting and imputation tasks were evaluated using average MSE, while anomaly detection, classification, and open-ended reasoning tasks (including multiple-choice questions (MCQs) and true-false judgment) were measured by accuracy. * Doubao uses simple mean forecasting, which outputs the same value for all forecasting.

Backbone	Classical Numerical Task				Open-Ended Reasoning QA	
	Forecasting ↓	Imputation ↓	Anomaly Detection ↑	Classification ↑	Judgment ↑	MCQ ↑
Doubao	—*	0.018	0.52	0.44	0.78	0.56
GPT-4o	1.79	0.018	0.64	0.32	0.72	0.58
Llama-3 8B	2.01	0.020	0.54	0.24	0.74	0.48
Qwen-2.5 7B	1.82	0.016	0.68	0.52	0.82	0.54
Mistral 7B	1.35	0.014	0.58	0.44	0.80	0.64

evaluation, we randomly selected 50 QA pairs for each task type (or question format), taking into account their respective application domains.

Overall, our fine-tuned models demonstrated improved performance across multiple tasks. For open-ended reasoning QA, the fine-tuned Qwen model achieved an accuracy of 82% on judgment questions, while the Mistral model reached 64% on multiple-choice questions. In forecasting tasks, the MSE values were relatively high, likely due to the long time series in our dataset, which can be challenging for models to handle. Nevertheless, our fine-tuned Mistral model outperformed GPT-4o.

For imputation tasks, the MSE remained low, possibly because having more extensive time series data provides additional information that supports better imputation. Our fine-tuned Mistral model generally showed stronger performance. These results suggest that by continuously pre-training on QA time series pairs, models can learn and leverage time series patterns more effectively.

Forecasting with Long Time Series. Figure 5 illustrates an example of forecasting results produced by our fine-tuned Mistral 7B model and GPT-4o on a long time-series input. Notably, the Mistral 7B model’s forecasts outperform those generated by GPT-4o. Moreover, the fine-tuned model can provide the rationale behind its predictions. When we posed follow-up questions in natural language, the model offered explanations of its reasoning. This illustrates that the model trained using the TSQA dataset can effectively learn complex time-series patterns and seasonal trends.

TSQA-Tuned vs. Zero-Shot Model. To further demonstrate the effectiveness of our fine-tuned model on the proposed TSQA dataset, we compared the performance of the tuned Mistral 7B with that of the original pretrained model on open-

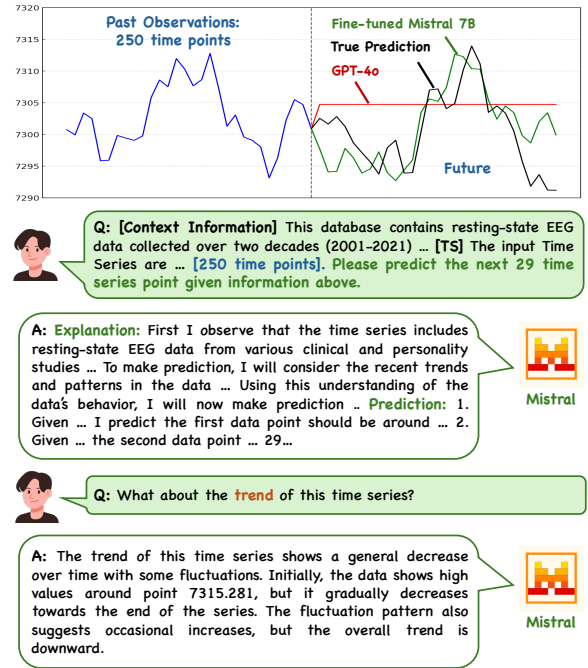


Figure 5: Example of long time series forecasting.

Table 4: Performance of TSQA-tuned and Zero-shot Mistral 7B on two tasks (evaluated at accuracy).

Model	Judgment ↑	MCQ ↑
Zero-shot	0.78	0.60
TSQA-tuned	0.80	0.64

ended reasoning QA tasks (including judgment and multiple-choice questions). The results, shown in Table 4, indicate that the TSQA-tuned model performs slightly better than the zero-shot model, suggesting the value of our TSQA dataset for improving question-answering performance.

Open-Ended Reasoning QA. Figure 6 presents an example of an open-ended reasoning question answering and compares the responses generated by our fine-tuned models with those from GPT-

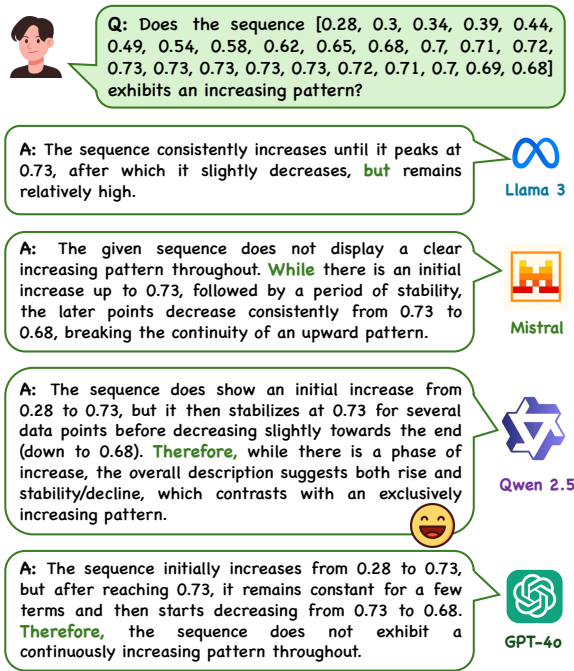


Figure 6: Example of open-ended reasoning QA.

40. As shown, our fine-tuned models effectively justify the reasons behind their answers. Notably, the Qwen-2.5 7B model provides a more thorough and comprehensive reasoning process compared to the other models.

4.3 User Study on Time Series QA

Given the subjectivity of time series question answering and reasoning tasks, we conducted a user study. We recruited 78 participants from diverse domains, such as data science researchers, financial practitioners, and students of liberal arts, who participated in this study anonymously. The four optional models are A) Llama-3 8B, B) Mistral 7B, C) Qwen-2.5 7B, and D) GPT-4o sequentially. Participants assessed responses based on accuracy, completeness, clarity, and user preference. The complete result and questionnaire of the user study can be found in the Appendix, and structured summaries of the findings are as follows:

Accuracy and Completeness. Mistral achieved the highest accuracy in tasks requiring numerical precision. For cyclical pattern identification (Q4), it received 80.8% user preference, and 69.2% selected its response as most comprehensive for volatility analysis (Q6). However, in trend analysis (Q12), 51.3% preferred Qwen over Mistral (37.2%). In summary, Mistral’s and Qwen’s answers to the numerical questions are generally considered to be

more accurate as well as more comprehensive.

Explanation Clarity. Explanation clarity was measured by asking participants to rate the understandability of the models’ reasoning. In Question 7 (“Considering the data points [60.0, 28.0, ...], do you see any seasonal patterns?”), more participants found Mistral’s explanation clear, praising its logical structure and ease of understanding. Although Qwen’s responses were also well received, they trailed slightly behind in these clarity ratings.

User Preference. In terms of overall user preference, when asked which answer they preferred (as in Question 11), mistral was selected by 70.5% of participants, while Qwen received 32.1% of the votes. This reflects a general tendency among users to favor responses that combine detailed data analysis with clear, logical explanations. In contrast to the above two models, Llama and GPT-4o are slightly less well accepted.

Overall, from the user study, Qwen and Mistral emerged as the top models for accuracy-driven tasks, while Mistral excelled in generating thorough explanations and more preferences. This highlights a trade-off between precision and interpretability in open-ended time series reasoning.

5 Conclusion

In this paper, we introduce *Time-MQA*, a multi-task framework that unifies time series analysis through natural language question answering, overcoming the limitations of single-task approaches. By integrating forecasting, imputation, anomaly detection, classification, and open-ended reasoning, *Time-MQA* enables flexible interactions with time-series data. The introduced *TSQA* dataset comprised of ~200k question-answer pairs across twelve diverse domains and varied time series lengths. It is a vital resource for advancing time series question answering and reasoning in LLMs. Experimental results show that continually pre-training models such as Mistral 7B, Llama 8B, and Qwen-2.5 7B on *TSQA* enhance their ability to interpret temporal patterns—going beyond numeric tasks to generate contextually rich insights. This work bridges the gap between traditional time series analysis and modern LLM capabilities, making temporal information more accessible and context-aware. Future directions include expanding the *TSQA* dataset to dynamic real-world data streams and exploring strategies to further refine time series reasoning.

Limitations

While Time-MQA advances multi-task time series question answering, it has limitations. Our proposed TSQA dataset, though diverse, may not cover all real-world scenarios, particularly highly irregular or domain-specific time series. But data with those properties may be common in real scenarios (Mulayim et al., 2024). Our approach relies on pre-training LLM architectures, which may not be optimized for long numerical sequences, which may cause hallucination problems (Yao et al., 2023). Evaluating open-ended reasoning remains challenging, as current automated metrics may not fully capture response quality (Kong et al., 2025). Additionally, computational costs for fine-tuning large models are high, which may limit some real-time applications (Jin et al., 2024b).

Ethics Statement

Except for the pre-training corpora in the finance domain, all datasets and language models used in this work are publicly available. The TSQA dataset was constructed from well-established, publicly accessible time series benchmarks and synthetic data generation followed ethical guidelines to minimize biases and ensure data quality. Our approach does not involve any personally identifiable information (PII) or sensitive data, and we adhere to responsible AI practices by following the ACL Code of Ethics.

References

Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.

Abdul Fatir Ansari, Lorenzo Stella, Caner Turkmen, Xiyuan Zhang, Pedro Mercado, Huibin Shen, Oleksandr Shchur, Syama Sundar Rangapuram, Sebastian Pineda Arango, Shubham Kapoor, et al. 2024. Chronos: Learning the language of time series. *arXiv preprint arXiv:2403.07815*.

Marc Bachlin, Daniel Roggen, Gerhard Troster, Meir Plotnik, Noit Inbar, Inbal Meidan, Talia Herman, Marina Brozgol, Eliya Shaviv, Nir Giladi, et al. 2009. Potentials of enhanced context awareness in wearable assistants for parkinson’s disease patients with the freezing of gait syndrome. In *2009 International Symposium on Wearable Computers*, pages 123–130. IEEE.

Michele Bernardini, Anastasiia Doynychko, Luca Romeo, Emanuele Frontoni, and Massih-Reza Amini.

2023. A novel missing data imputation approach based on clinical conditional generative adversarial networks applied to ehr datasets. *Computers in Biology and Medicine*, 163:107188.

Giovanni Maria Biancofiore, Yashar Deldjoo, Tommaso Di Noia, Eugenio Di Sciascio, and Fedelucio Narducci. 2024. Interactive question answering systems: Literature review. *ACM Computing Surveys*, 56(9):1–38.

Jian Chen, Peilin Zhou, Yining Hua, Yingxin Loh, Kehui Chen, Ziyuan Li, Bing Zhu, and Junwei Liang. 2024a. Fintextqa: A dataset for long-form financial question answering. *Association for Computational Linguistics(ACL)*.

Wei Chen, Yuxuan Liang, Yuanshao Zhu, Yanchuan Chang, Kang Luo, Haomin Wen, Lei Li, Yanwei Yu, Qingsong Wen, Chao Chen, et al. 2024b. Deep learning for trajectory data management and mining: A survey and beyond. *arXiv preprint arXiv:2403.14151*.

Xiaoyang Chen, Ben He, Hongyu Lin, Xianpei Han, Tianshu Wang, Boxi Cao, Le Sun, and Yingfei Sun. 2024c. Spiral of silences: How is large language model killing information retrieval?—a case study on open domain question answering. *Association for Computational Linguistics(ACL)*.

Zhuo Chen, Wen Zhang, Yufeng Huang, Mingyang Chen, Yuxia Geng, Hongtao Yu, Zhen Bi, Yichi Zhang, Zhen Yao, Wenting Song, et al. 2023. Teleknowledge pre-training for fault analysis. In *2023 IEEE 39th International Conference on Data Engineering (ICDE)*, pages 3453–3466. IEEE.

Daixuan Cheng, Yuxian Gu, Shaohan Huang, Junyu Bi, Minlie Huang, and Furu Wei. 2024. Instruction pre-training: Language models are supervised multitask learners. *arXiv preprint arXiv:2406.14491*.

Winnie Chow, Lauren Gardiner, Haraldur T Hallgrímsson, Maxwell A Xu, and Shirley You Ren. 2024. Towards time series reasoning with llms. *arXiv preprint arXiv:2409.11376*.

Wenjie Du, Jun Wang, Linglong Qian, Yiyuan Yang, Zina Ibrahim, Fanxing Liu, Zepu Wang, Haoxin Liu, Zhiyuan Zhao, Yingjie Zhou, et al. 2024. Tsi-bench: Benchmarking time series imputation. *arXiv preprint arXiv:2406.12747*.

Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*.

Nigel Fernandez, Alexander Scarlatos, and Andrew Lan. 2024. Syllabusqa: A course logistics question answering dataset. *arXiv preprint arXiv:2403.14666*.

665	Elizabeth Fons, Rachneet Kaur, Soham Palande, Zhen	Ming Jin, Yifan Zhang, Wei Chen, Kexin Zhang, Yux-	718
666	Zeng, Tucker Balch, Manuela Veloso, and Svitlana	uan Liang, Bin Yang, Jindong Wang, Shirui Pan, and	719
667	Vyetrenko. 2024. Evaluating large language mod-	Qingsong Wen. 2024b. Position: What can large	720
668	els on time series feature understanding: A compre-	language models tell us about time series analysis.	721
669	hensive taxonomy and benchmark. <i>arXiv preprint</i>	In <i>Forty-first International Conference on Machine</i>	722
670	<i>arXiv:2404.16563</i> .	<i>Learning</i> .	723
671	Giacomo Frisoni, Alessio Cocchieri, Alex Presepi, Gi-	Yaxuan Kong, Yiyuan Yang, Shiyu Wang, Chenghao	724
672	anluca Moro, and Zaiqiao Meng. 2024. To generate	Liu, Yuxuan Liang, Ming Jin, Stefan Zohren, Dan	725
673	or to retrieve? on the effectiveness of artificial con-	Pei, Yan Liu, and Qingsong Wen. 2025. Position:	726
674	texts for medical open-domain question answering.	Empowering time series reasoning with multimodal	727
675	<i>Association for Computational Linguistics(ACL)</i> .	llms. <i>arXiv preprint arXiv:2502.01477</i> .	728
676	Wayne A Fuller. 2009. <i>Introduction to statistical time</i>	Jennifer R Kwapisz, Gary M Weiss, and Samuel A	729
677	<i>series</i> . John Wiley & Sons.	Moore. 2011. Activity recognition using cell phone	730
678	Yizhen Guo, Tian Zhou, Wanyi Jiang, Bo Wu, Liang	accelerometers. <i>ACM SigKDD Explorations Newslet-</i>	731
679	Sun, and Rong Jin. 2024. Maximizing the impact	<i>ter</i> , 12(2):74–82.	732
680	of deep learning on subseasonal-to-seasonal climate	Wing Lian, Bleys Goodson, Eugene Pentland, Austin	733
681	forecasting: The essential role of optimization. <i>arXiv</i>	Cook, Chanvichet Vong, and "Teknium". 2023.	734
682	<i>preprint arXiv:2411.16728</i> .	Openorca: An open dataset of gpt augmented flan	735
683	James D Hamilton. 2020. <i>Time series analysis</i> . Prince-	reasoning traces. https://https://huggingface.co/Open-Orca/OpenOrca .	736
684	ton university press.		737
685	Songqiao Han, Xiyang Hu, Hailiang Huang, Minqi	Bryan Lim and Stefan Zohren. 2021. Time-series	738
686	Jiang, and Yue Zhao. 2022. Adbench: Anomaly de-	forecasting with deep learning: a survey. <i>Philo-</i>	739
687	tection benchmark. <i>Advances in Neural Information</i>	<i>sophical Transactions of the Royal Society A</i> ,	740
688	<i>Processing Systems</i> , 35:32142–32159.	379(2194):20200209.	741
689	Xanh Ho, Anh-Khoa Duong Nguyen, Saku Sugawara,	Haoxin Liu, Chenghao Liu, and B Aditya Prakash.	742
690	and Akiko Aizawa. 2020. Constructing a multi-hop	2024a. A picture is worth a thousand numbers: En-	743
691	qa dataset for comprehensive evaluation of reasoning	abling llms reason about time series via visualization.	744
692	steps. <i>arXiv preprint arXiv:2011.01060</i> .	<i>NAACL</i> .	745
693	Noah Hollmann, Samuel Müller, Lennart Purucker,	Haoxin Liu, Shangqing Xu, Zhiyuan Zhao, Ling kai	746
694	Arjun Krishnakumar, Max Körfer, Shi Bin Hoo,	Kong, Harshavardhan Kamarthi, Aditya B Sasanur,	747
695	Robin Tibor Schirrmester, and Frank Hutter. 2025.	Megha Sharma, Jiaming Cui, Qingsong Wen, Chao	748
696	Accurate predictions on small data with a tabular	Zhang, et al. 2024b. Time-mmd: Multi-domain mul-	749
697	foundation model. <i>Nature</i> , 637(8045):319–326.	timodal dataset for time series analysis. In <i>The Thirty-</i>	750
698	Albert Q Jiang, Alexandre Sablayrolles, Arthur Men-	<i>eight Conference on Neural Information Processing</i>	751
699	sch, Chris Bamford, Devendra Singh Chaplot, Diego	<i>Systems Datasets and Benchmarks Track</i> .	752
700	de las Casas, Florian Bressand, Gianna Lengyel, Guil-	Haoxin Liu, Zhiyuan Zhao, Jindong Wang, Harshavard-	753
701	laume Lample, Lucile Saulnier, et al. 2023. Mistral	han Kamarthi, and B. Aditya Prakash. 2024c. LST-	754
702	7b. <i>arXiv preprint arXiv:2310.06825</i> .	Prompt: Large language models as zero-shot time	755
703	Chuanyang Jin, Yutong Wu, Jing Cao, Jiannan Xiang,	series forecasters by long-short-term prompting. In	756
704	Yen-Ling Kuo, Zhiting Hu, Tomer Ullman, Antonio	<i>Findings of the Association for Computational Lin-</i>	757
705	Torralba, Joshua B Tenenbaum, and Tianmin Shu.	<i>guistics: ACL 2024</i> , pages 7832–7840. Association	758
706	2024a. Mmtom-qa: Multimodal theory of mind ques-	for Computational Linguistics.	759
707	tion answering. <i>arXiv preprint arXiv:2401.08743</i> .	Yiheng Liu, Tianle Han, Siyuan Ma, Jiayue Zhang,	760
708	Ming Jin, Shiyu Wang, Lintao Ma, Zhixuan Chu,	Yuanyuan Yang, Jiaming Tian, Hao He, Antong Li,	761
709	James Y Zhang, Xiaoming Shi, Pin-Yu Chen, Yuxuan	Mengshen He, Zhengliang Liu, et al. 2023. Summary	762
710	Liang, Yuan-Fang Li, Shirui Pan, et al. 2023a. Time-	of chatgpt-related research and perspective towards	763
711	llm: Time series forecasting by reprogramming large	the future of large language models. <i>Meta-Radiology</i> ,	764
712	language models. <i>arXiv preprint arXiv:2310.01728</i> .	page 100017.	765
713	Ming Jin, Qingsong Wen, Yuxuan Liang, Chaoli Zhang,	Yong Liu, Haoran Zhang, Chenyu Li, Xiangdong	766
714	Siqiao Xue, Xue Wang, James Zhang, Yi Wang,	Huang, Jianmin Wang, and Mingsheng Long. Timer:	767
715	Haifeng Chen, Xiaoli Li, et al. 2023b. Large models	Generative pre-trained transformers are large time se-	768
716	for time series and spatio-temporal data: A survey	ries models. In <i>Forty-first International Conference</i>	769
717	and outlook. <i>arXiv preprint arXiv:2310.10196</i> .	<i>on Machine Learning</i> .	770
		Ziqing Ma, Wenwei Wang, Tian Zhou, Chao Chen,	771
		Bingqing Peng, Liang Sun, and Rong Jin. 2024. Fu-	772
		sionsf: Fuse heterogeneous modalities in a vector	773

Yiyuan Yang, Yi Li, Taojia Zhang, Yan Zhou, and Haifeng Zhang. 2021a. Early safety warnings for long-distance pipelines: A distributed optical fiber sensor machine learning approach. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pages 14991–14999.

Yiyuan Yang, Chaoli Zhang, Tian Zhou, Qingsong Wen, and Liang Sun. 2023b. Dcdetector: Dual attention contrastive representation learning for time series anomaly detection. In *Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, pages 3033–3045.

Yiyuan Yang, Haifeng Zhang, and Yi Li. 2021b. Long-distance pipeline safety early warning: a distributed optical fiber sensing semi-supervised learning method. *IEEE sensors journal*, 21(17):19453–19461.

Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Bengio, William W Cohen, Ruslan Salakhutdinov, and Christopher D Manning. 2018. Hotpotqa: A dataset for diverse, explainable multi-hop question answering. *arXiv preprint arXiv:1809.09600*.

Jia-Yu Yao, Kun-Peng Ning, Zhen-Hui Liu, Mu-Nan Ning, Yu-Yang Liu, and Li Yuan. 2023. Llm lies: Hallucinations are not bugs, but features as adversarial examples. *arXiv preprint arXiv:2310.01469*.

Zahra Zamanzadeh Darban, Geoffrey I Webb, Shirui Pan, Charu Aggarwal, and Mahsa Salehi. 2024. Deep learning for time series anomaly detection: A survey. *ACM Computing Surveys*, 57(1):1–42.

Kexin Zhang, Qingsong Wen, Chaoli Zhang, Rongyao Cai, Ming Jin, Yong Liu, James Y Zhang, Yuxuan Liang, Guansong Pang, Dongjin Song, et al. 2024a. Self-supervised learning for time series analysis: Taxonomy, progress, and prospects. *IEEE Transactions on Pattern Analysis and Machine Intelligence*.

Weiqi Zhang, Jiexia Ye, Ziyue Li, Jia Li, and Fugee Tsung. 2024b. Dualtime: A dual-adaptor multimodal language model for time series representation. *arXiv preprint arXiv:2406.06620*.

Yunkai Zhang, Yawen Zhang, Ming Zheng, Kezhen Chen, Chongyang Gao, Ruian Ge, Siyuan Teng, Amine Jelloul, Jinneng Rao, Xiaoyuan Guo, Chiang-Wei Fang, Zeyu Zheng, and Jie Yang. 2023. *Insight miner: A large-scale multimodal model for insight mining from time series*. In *NeurIPS 2023 AI for Science Workshop*.

Haoyi Zhou, Shanghang Zhang, Jieqi Peng, Shuai Zhang, Jianxin Li, Hui Xiong, and Wancai Zhang. 2021. Informer: Beyond efficient transformer for long sequence time-series forecasting. In *Proceedings of the AAAI conference on artificial intelligence*, volume 35, pages 11106–11115.

Yuchen Zhuang, Yue Yu, Kuan Wang, Haotian Sun, and Chao Zhang. 2024. Toolqa: A dataset for llm question answering with external tools. *Advances in Neural Information Processing Systems*, 36.

A Examples of TSQA Dataset

Below are several examples from the TSQA dataset.

A.1 Forecasting.

Question: The Self-regulation of Slow Cortical Potentials dataset, provided by the University of Tuebingen, involves EEG recordings from a healthy subject controlling a cursor on a screen through cortical potentials. The sampling frequency of this dataset is 256 Hz. The input Time Series are *[Time Series Data Points]*. Please predict the next nine time series points given the information above.

Answer: Based on the given information, the predictions are [-11.69, -10.72, -9.84, -9.88, -9.03, -8.78, -8.41, -8.28, -7.59].

Question: From March 31, 2014 to June 04, 2014, the weighted ILI rate was 1.74542, and the ili total count reached 14637 across 1866 providers serving a total of 797863 patients, especially 0-4 age group was 3582, 5-24 age group was 5659, 25-49 age group was 3153, 50-64 age group was 1436, 65+ age group was 807, and 25-64 age group was unavailable. The internet search results for the same period are summarized as follows: ... The prior historical average weighted ILI rate during this period was 1.286616667 with a standard deviation of 0.073690536. Please analyze the data and provide a target result on June 04, 2014 with a summary of the findings.

Answer: Based on the provided textual information, I predict that the influenza infection situation in the United States for the upcoming period is likely to be moderate to severe, with a potential for increased risk of complications and co-infections, particularly among vulnerable populations such as those with cystic fibrosis and critically ill patients. Therefore, the predicted weighted ILI rate for June 04, 2014 is estimated to be 1.83453.

A.2 Imputation.

Question: This dataset, sourced from the Kaggle web traffic forecasting competition, contains daily time series data representing the number of hits for 145,063 Wikipedia pages from July 1, 2015, to September 10, 2017. The input Time Series are *[Time Series Data Points with Missing Values]*. Please give full-time series with missing value imputed. The missing values are represented by 'X'.

Answer: Based on the given information, the full-time series with missing value imputed are *[Time Series Data Points]*.

A.3 Classification.

Question: Following data provides accelerometer data for activity recognition research. The dataset has a sampling rate of 20Hz and records accelerometer data for six activity states: walking, jogging, sitting, standing, upstairs, and downstairs. Each sample includes acceleration values for the X, Y, and Z axes, ranging from -20 to 20, where 10 represents 1g and 0 indicates no acceleration. The recorded acceleration includes gravitational acceleration, so when the phone is stationary on a flat surface, the vertical axis registers approximately ± 10 . We provide 10 timestamps of accelerometer data, with each timestamp containing X, Y, and Z values, for a total of 30 values. The recorded Time Serie is *[Time Series Data Points]*. Please judge whether this data corresponds to 'Walking' or 'Jogging' or 'Upstairs' or 'Downstairs' or 'Sitting' or 'Standing' given the information above.

Answer: Based on the given information, the activity is Walking.

A.4 Anomaly Detection.

Question: The following data is derived from AWS CloudWatch services and is primarily used for monitoring and managing metrics related to cloud infrastructure and application performance. These include CPU usage, memory consumption, network throughput, and disk I/O metrics. This data can be utilized to identify performance issues in cloud services, such as traffic spikes, resource exhaustion, or network anomalies, and to perform root cause analysis for rapid problem resolution. The input Time Serie is *[Time Series Data Points]*. Please determine whether there are anomalies in this time series given the information above.

Answer: Based on the given information, this time series includes Anomaly Points.

A.5 Open-Ended Reasoning QA.

Question: Summarize the overall behavior observed in the data points *[Time Series Data Points]*.

Answer: The data initially shows a slow increase followed by a decreasing trend. The middle part of the data points represents a peak with few higher fluctuations and later shows a slightly decline indicating a volatile pattern.

Question: Determine if there's a structural break in the data points *[Time Series Data Points]*.
A) Yes, between 0.71 and 0.72 B) Yes, between 0.65 and 0.63 C) No structural break

Answer: C) No structural break. The data follows a gradual upward trend and subsequent flattening with slight variability, but no abrupt change in the level or direction indicating a structural break.

Question: The data points *[Time Series Data Points]* exhibit a clear seasonal pattern. True or False?

Answer: False. While there are fluctuations and some repeated values, no clear cyclical pattern emerges that would indicate seasonality. The variations appear to be more random than driven by any seasonal influences.

B Open-Ended Reasoning QA Generation

Figure 7 illustrated the prompt for generating open-ended reasoning QA pairs.

Prompt: Given the time series data: {target_series_str}. Please select 24 time series points and generate 10 unique instruction-tuning tasks for an LLM we are tuning for time series question-and-answer tasks.

For each question, you can select different 24 time series points from given time series. The questions can vary in format, including multiple-choice, true/false, or open-ended question types. You may also ask questions related to trends, seasonality, cyclical patterns, summarization, volatility, anomalies, structural breaks or statistical property, but don't be restricted to these topics. Please also ensure diversity in your question-answer pairs.

Your response should be formatted as a JSON object with four fields: "task_type": This field should indicate the category or topic type, which provides a clear classification for the question content. "question_format": This field should specify the type of question, such as multiple-choice, true/false, or another question format. "question": This field should contain the question related to the given data. If it's a multiple-choice question, include the available options as well. "answer with reasoning": This field should firstly provide the correct answer to the input question, based on the given data. Then, it should explain the process used to arrive at the answer. If a mathematical operation is involved, mention the formula and fill it in with the relevant numbers.

Expected Output: "task_type": summarization. "question_format": open ended question. "question": Examine the data points [-40.66, -41.47, -41.06, -40.09, -40.59, -41.38, -40.16, -40.31, -38.84, -37.19, -35.19, -33.72, -33.41, -31.09, -29.22, -30.34, -28.91, -30.59, -31.16, -32.31, -32.88, -32.66, -32.72, -30.84] and summarize the overall movement trend in this data. "answer with reasoning": Initially stable with minor fluctuations, followed by a significant decrease. The data starts off stable with minor ups and downs and later shows a significant decrease towards the end. Thank you!

Figure 7: Prompt for open-ended reasoning QA generation.

C Training Data Format

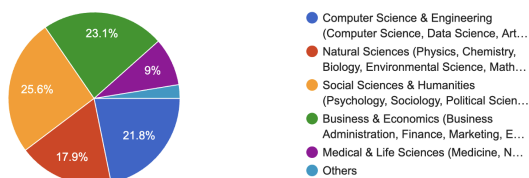
We format our question-and-answer pairs using a specifically designed template to clearly separate questions from answers. The template is structured as follows: <QUE> {Question} <ANS> {Answer} </END>. Additionally, for continual pre-training of the Llama model, we prepend <|begin_of_text|> at the start of each sample and append <|end_of_text|> at the end. For the Mistral model, we use <s> at the beginning of each sample and </s> at the end. In the case of the

1072 Qwen model, only <|endoftext|> is added at the
1073 end of each sample.

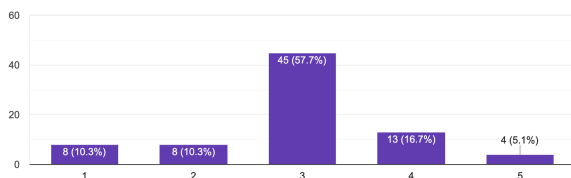
1074 **D User Study**

1075 A total of 78 surveys were collected, and the results
1076 for each question are as follows:

1077 **Question 1.** What is your field of study?



1078 **Question 2.** How familiar are you with time se-
1079 ries analysis? (1 = Not at all, 5 = Very familiar)



1080 **Question 3.** How accurate is the answer with
1081 respect to the question?

1082 Analyse and summarise the trend and pat-
1083 tern in the data set [68.0, 83.0, 95.0, 103.5, 108.5,
1084 112.0, 113.0, 114.0, 114.5, 115.5, 116.0, 116.0,
1085 115.0, 114.0, 113.5, 111.5, 108.0, 102.5, 100.5,
1086 93.5, 88.0, 82.0, 78.0, 73.5].

Question 3 (Answers):

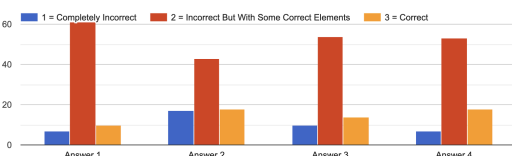
Question 3: Analyse and summarise the trend and pattern in the data set (68.0, 83.0, 95.0, 103.5, 108.5, 112.0, 113.0, 114.0, 114.5, 115.5, 116.0, 115.0, 114.0, 113.5, 111.5, 108.0, 102.5, 100.5, 93.5, 88.0, 82.0, 78.0, 73.5).

Answer 1: The data set exhibits a general downward trend after an initial period of increase. Initially, the values rise from 68.0 to 116.0, indicating an upward trend, and then begin to decline, dropping from 116.0 to 73.5. This suggests a downward trend. The pattern is characterized by an initial peak followed by a decline. The reason for these changes may be related to seasonal or cyclical factors, but without further context, it is impossible to determine the specific causes.

Answer 2: The data shows a clear upward trend initially, reaching a peak at 116.0, followed by a gradual decline towards the end of the series. Overall, it demonstrates a strong upward trend followed by a decline. The pattern is characterized by an initial peak followed by a decline. The reason for these changes may be related to seasonal or cyclical factors, but without further context, it is impossible to determine the specific causes.

Answer 3: The data shows an increasing trend initially, peaking at 116.0 before starting a gradual decline towards the end of the series. Overall, it demonstrates a strong upward trend followed by a decline. The pattern is characterized by an initial peak followed by a decline. The reason for these changes may be related to seasonal or cyclical factors, but without further context, it is impossible to determine the specific causes.

Answer 4: Overall Trend: The data initially increases from 68.0 to a peak of 116.0 around the middle of the sequence, after reaching the peak, the values decline, dropping back down to 73.5 at the end. The values rise from 68.0 to 116.0, suggesting a steady increase. This upward trend is disrupted between 112.0 and 115.5, then drops slightly as it reaches the peak. The highest values (112.0-116.0) are stable for a short period before declining. After reaching 116.0, the numbers start decreasing, indicating a downward trend. The decline is initially gradual but becomes steeper after 114.0, with a significant drop to 73.5. This dataset exhibits a rise-peak-fall pattern. The values increase steadily, peak, stabilize briefly, and then decline in a symmetrical manner. This could indicate a seasonal trend, a business cycle, or a time-based natural fluctuation.



Question 4. Which do you think is the most accurate answer with respect to the open-ended reasoning question?

Identify any cyclical patterns in the data points [0.28, 0.27, 0.26, 0.25, 0.24, 0.23, 0.24, 0.25, 0.26, 0.28, 0.3, 0.33, 0.36, 0.39, 0.43, 0.46, 0.5, 0.53, 0.57, 0.6, 0.63, 0.65, 0.67, 0.68].

Question 4 (Answers):

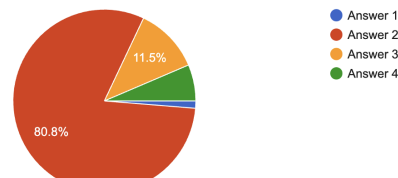
Question 4: Identify any cyclical patterns in the data points [0.28, 0.27, 0.26, 0.25, 0.24, 0.23, 0.24, 0.25, 0.26, 0.28, 0.3, 0.33, 0.36, 0.39, 0.43, 0.46, 0.5, 0.53, 0.57, 0.6, 0.63, 0.65, 0.67, 0.68].

Answer 1: There is no clear cyclical pattern in the data points provided. The data shows a steady increase which could indicate an upward trend but not a clear cycle.

Answer 2: The data does not clearly demonstrate any distinct cyclical patterns. Instead, it shows a general upward trend with minor fluctuations.

Answer 3: There is a cyclical pattern in the data where values initially decrease from 0.39 to 0.23, then start increasing again to reach 0.68.

Answer 4: To identify cyclical patterns in your data, we analyze the trends and fluctuations. Initial Decline (0.28 - 0.23): The first six values show a decreasing trend. Initial Upward Fluctuation (0.24 - 0.39): From the 7th value onwards, the data exhibits a mostly increasing trend. A minor oscillation around 0.25-0.26 suggests a small cycle within the initial decline phase. A strong upward trend follows, resembling a possible seasonal or trend-driven cycle. Apply Fourier analysis or autocorrelation to detect precise cycles. If manually analyzing, observe for repeated peaks and troughs over fixed intervals.



Question 5. How complete is the answer? Did it address all parts of the question?

Examine the data points [19.34, 20.41, 20.38, 19.28, 19.75, 22.84, 25.09, 24.97, 24.75, 26.25, 27.0, 25.03, 22.41, 21.78, 22.44, 21.19, 19.44, 19.62, 20.62, 20.16, 18.03, 17.19, 18.31, 19.91] and summarize the overall movement trend in this data.

Question 5 (Answers):

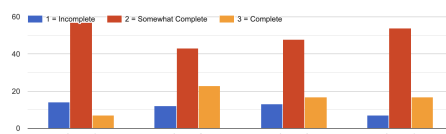
Question 5: Examine the data points [19.34, 20.41, 20.38, 19.28, 19.75, 22.84, 25.09, 24.97, 24.75, 26.25, 27.0, 25.03, 22.41, 21.78, 22.44, 21.19, 19.44, 19.62, 20.62, 20.16, 18.03, 17.19, 18.31, 19.91].

Answer 1: The data shows an initial increasing trend from 19.34 to a peak of 27.0, followed by a general downward trend towards the end of the series. Overall, it demonstrates a strong upward trend followed by a decline. The pattern is characterized by an initial peak followed by a decline. The reason for these changes may be related to seasonal or cyclical factors, but without further context, it is impossible to determine the specific causes.

Answer 2: The data shows an initial increasing trend from 19.34 to a peak of 27.0, followed by a general downward trend towards the end of the series. Overall, it demonstrates a strong upward trend followed by a decline. The pattern is characterized by an initial peak followed by a decline. The reason for these changes may be related to seasonal or cyclical factors, but without further context, it is impossible to determine the specific causes.

Answer 3: The data points show a general upward trend from 19.34 to a peak of 27.0, followed by a general downward trend towards the end of the series. Overall, it demonstrates a strong upward trend followed by a decline. The pattern is characterized by an initial peak followed by a decline. The reason for these changes may be related to seasonal or cyclical factors, but without further context, it is impossible to determine the specific causes.

Answer 4: To analyze the overall movement trend in the given data points, we first observe the initial values (19.34 to 20.41) and the final values (18.03 to 19.91). The data exhibits a slight downward trend. The values start around 19.34 and fluctuate slightly before rising significantly to a peak of 27.0. After reaching the peak, the values decline, showing a clear downward trend.



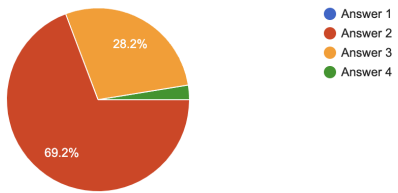
1105
1106
1107
1108
1109
1110
1111
1112
1113
1114
1115
1116

Question 6. Which answer is the most comprehensive in relation to the question?

Considering the data points [47.0, 51.0, 51.0, 20.0, 48.0, 44.0, 41.0, 28.0, 48.0, 38.0, 67.0, 66.0, 70.0, 54.0, 63.0, 56.0, 76.0, 58.0, 60.0, 58.0, 69.0, 56.0, 39.0, 50.0], which statement best describes the volatility of this time series? A) The volatility is low with consistent data points. B) There is moderate volatility with regular patterns. C) High volatility with sudden peaks and drops. D) The time series exhibits a perfectly linear trend.

Question 6 (Answers):

Question 6: Considering the data points [47.0, 51.0, 51.0, 20.0, 48.0, 44.0, 41.0, 28.0, 48.0, 38.0, 67.0, 66.0, 70.0, 54.0, 63.0, 56.0, 76.0, 58.0, 60.0, 58.0, 69.0, 56.0, 39.0, 50.0], which statement best describes the volatility of this time series? A) The volatility is low with consistent data points. B) There is moderate volatility with regular patterns. C) High volatility with sudden peaks and drops. D) The time series exhibits a perfectly linear trend.			
Answer 1 The correct answer is C). High volatility with sudden peaks and drops. The data shows significant fluctuations, with values jumping from 20.0 to 76.0, then rising to 70.0 before dropping to 39.0. These sharp changes signify high volatility.	Answer 2 C) High volatility with sudden peaks and drops. Justification: The data points fluctuate significantly from as low as 20.0 to as high as 76.0, with several data points exhibiting high volatility in the series without a consistent linear trend. Comments: A) The data points exhibit a low level of volatility and are not significantly fluctuating, with values jumping from 20.0 to 76.0, then rising to 70.0 before dropping to 39.0. These sharp changes signify high volatility. B) While there is some volatility with a few rising and falling patterns, the deviation from the mean is substantial to identify the volatility as moderate. D) The data does not display a perfectly linear trend but rather exhibits considerable fluctuations, making this response incorrect.	Answer 3 The correct answer is C) High volatility with sudden peaks and drops. The data contains several high peaks (67, 66, 70) and low values (20, 39), indicating significant volatility in the data. Additionally, there are rapid changes between these extremes. Therefore, the Python code snippet related to this question is: (Please provide the data and summary statistics) Given these observations, the best answer is C) High volatility with sudden peaks and drops. This is because the series has considerable fluctuations, with notable jumps and drops rather than smooth or consistent changes.	Answer 4 To assess the volatility of the given time series, we examine the range, fluctuations, and presence of sudden peaks or drops. The data points range from 20.0 to 76.0, which is a significant spread. There are some outlier values, such as 20.0 and 76.0, which are not followed by high volatility or consistent patterns. Given these observations, the best answer is C) High volatility with sudden peaks and drops. This is because the series has considerable fluctuations, with notable jumps and drops rather than smooth or consistent changes.



1117
1118
1119
1120
1121
1122
1123
1124
1125

Question 7. How clear or understandable is the reasoning provided in the answer?

Considering the data points [60.0, 28.0, 46.0, 51.0, 37.0, 46.0, 41.0, 35.0, 32.0, 66.0, 10.0, 21.0, 67.0, 72.0, 83.0, 61.0, 11.0, 31.0, 62.0, 50.0, 56.0, 76.0, 53.0, 86.0], do you see any seasonal patterns? A) Yes, there is a seasonal pattern. B) No, there is no identifiable seasonal pattern.

Question 8. How clearly does the response explain its reasoning? (1 = Not Clear At All, 5 = Really Clear)

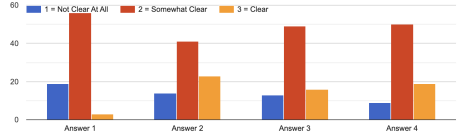
What is the volatility level in the dataset [0.83, 0.8, 0.76, 0.75, 0.75, 0.76, 0.78, 0.8, 0.8, 0.8, 0.81, 0.8, 0.8, 0.8, 0.8, 0.79, 0.8, 0.82, 0.83, 0.82, 0.82, 0.82, 0.82, 0.81]? A) High B) Moderate C) Low D) No volatility"

Answer: C) Low Explanation: The dataset shows moderate stability with minor fluctuations,

1137

Question 7 (Answers):

Question 7: Considering the data points [60.0, 28.0, 46.0, 51.0, 37.0, 46.0, 41.0, 35.0, 32.0, 66.0, 10.0, 21.0, 67.0, 72.0, 83.0, 61.0, 11.0, 31.0, 62.0, 50.0, 56.0, 76.0, 53.0, 86.0], do you see any seasonal patterns? A) Yes, there is a seasonal pattern. B) No, there is no seasonal pattern. C) Yes, there is a seasonal pattern. D) No, there is no seasonal pattern.			
Answer 1 A) Yes, there is a seasonal pattern. The data exhibits a repeating pattern every 10 units, suggesting a seasonal component.	Answer 2 B) No, there is no identifiable seasonal pattern. The data points seem to fluctuate irregularly without showing a consistent periodicity or cycle that would suggest seasonality.	Answer 3 B) No, there is no identifiable seasonal pattern. The data points seem to fluctuate irregularly without showing a consistent periodicity or cycle that would suggest seasonality.	Answer 4 To determine if there is a seasonal pattern in the given data points, we need to analyze the data for any repeating cycles or trends that would indicate seasonality. A seasonal pattern typically repeats at regular intervals.
Evaluation: The data points show varying levels of volatility, with some points exhibiting significant fluctuations (e.g., 10.0, 21.0, 67.0, 72.0, 83.0, 61.0, 11.0, 31.0, 62.0, 50.0, 56.0, 76.0, 53.0, 86.0). These fluctuations suggest a non-linear trend, which is not consistent with a seasonal pattern. Therefore, the data points do not exhibit a seasonal pattern.			
Based on this, the most reasonable answer is B) No, there is no identifiable seasonal pattern. Therefore, the data points do not exhibit a seasonal pattern.			



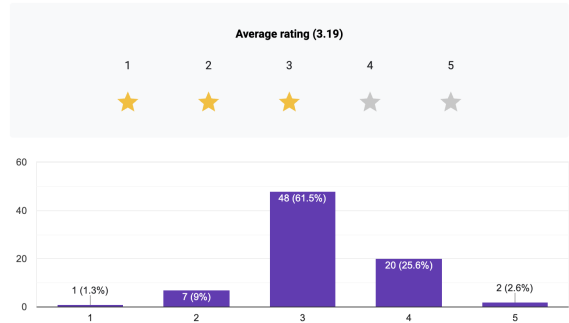
indicating a low volatility level. The changes between consecutive points are generally small and consistent rather than erratic or extreme. Thinking process: The dataset has some movement, but it is not highly variable or erratic. There are no significant spikes or drops across the data points. The changes are steady and minor, suggesting a low level of volatility. Comparing to other options:

A) High: The dataset does not show high levels of volatility with major fluctuations.

B) Moderate: While there is some movement, it is not as severe as to classify it as moderate volatility.

D) No volatility: There is some movement, but it is not so minimal that it would be considered no volatility.

Therefore, the appropriate response and explanation for this question is C) Low volatility.



Question 9. Does the explanation logically follow from the data or the question's context?

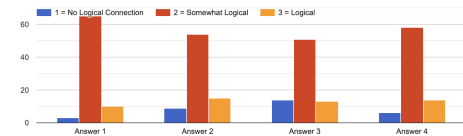
1160
1161
1162

1163
1164
1165
1166
1167

Summarise the overall movement trend in the data points [2.16, 1.86, 1.84, 2.49, 1.98, 1.96, 1.73, 1.79, 2.13, 1.43, 1.49, 1.31, 0.55, -0.16, -0.12, 0.63, 0.66, 0.51, 0.28, -1.14, -2.81, -3.21, -2.26, -1.65].

Question 9 (Answers):

Answer 1 The overall movement trend in the data points is a downward trend with some minor recovery towards the end. The data shows a downward trend overall, starting from a positive value and gradually decreasing to negative values by the end. It's indicative of a declining pattern across the data points. Explanation: To summarize the movement trend, I focused on the general direction of the data points. While the values are mixed, the overall trend is downward, starting from a positive value and ending with negative values. The initial values are relatively positive, but they decrease steadily, with a notable drop around 0.55, followed by a slight recovery, and then a sharp decline towards the end. Step-by-step justification: Considering the starting and ending values: The starting value is 2.16, and the ending value is -1.65. From the general direction, the values decrease from positive to negative, showing an overall downward movement. Highlight key shifts: There is a consistent decline without significant recovery or change in direction, supporting the trend. The overall movement trend is a decline or downward trend across the data points.	Answer 2 The data shows a downward trend overall, starting from a positive value and gradually decreasing to negative values by the end. It's indicative of a declining pattern across the data points. Explanation: To summarize the movement trend, I focused on the general direction of the data points. While the values are mixed, the overall trend is downward, starting from a positive value and ending with negative values. The initial values are relatively positive, but they decrease steadily, with a notable drop around 0.55, followed by a slight recovery, and then a sharp decline towards the end. Step-by-step justification: Considering the starting and ending values: The starting value is 2.16, and the ending value is -1.65. From the general direction, the values decrease from positive to negative, showing an overall downward movement. Highlight key shifts: There is a consistent decline without significant recovery or change in direction, supporting the trend. The overall movement trend is a decline or downward trend across the data points.	Answer 3 Based on the provided data points, the overall trend appears to be downward with some minor fluctuations in the middle section. The values start around 2.16 and generally decrease towards -3.21, indicating a general decline. Explanation: The first few values (around 2.16, 1.86, 1.84, 2.49) are relatively high, but they quickly drop to lower values (1.98, 1.96, 1.73, 1.79). The values then fluctuate slightly, but the overall trend is downward, with a notable drop around 0.55, followed by a slight recovery, and then a sharp decline towards the end. Step-by-step justification: Considering the starting and ending values: The starting value is 2.16, and the ending value is -1.65. From the general direction, the values decrease from positive to negative, showing an overall downward movement. Highlight key shifts: There is a consistent decline without significant recovery or change in direction, supporting the trend. The overall movement trend is a decline or downward trend across the data points.	Answer 4 The overall movement trend in the data points shows an initial fluctuation around 0.5, followed by a general decline, with a sharp drop in negative values towards the end. The data points show a sharp drop in negative values towards the end, indicating a sharp decline. Explanation: The first few values (around 2.16, 1.86, 1.84, 2.49) are relatively high, but they quickly drop to lower values (1.98, 1.96, 1.73, 1.79). The values then fluctuate slightly, but the overall trend is downward, with a notable drop around 0.55, followed by a slight recovery, and then a sharp decline towards the end. Step-by-step justification: Considering the starting and ending values: The starting value is 2.16, and the ending value is -1.65. From the general direction, the values decrease from positive to negative, showing an overall downward movement. Highlight key shifts: There is a consistent decline without significant recovery or change in direction, supporting the trend. The overall movement trend is a decline or downward trend across the data points.
---	--	--	---



1168
1169
1170
1171
1172
1173
1174
1175
1176
1177
1178
1179
1180
1181

Question 10. How would you rate the quality of the explanation and reasoning with respect to the question? (1 = Really Poor, 5 = Excellent Reasoning)

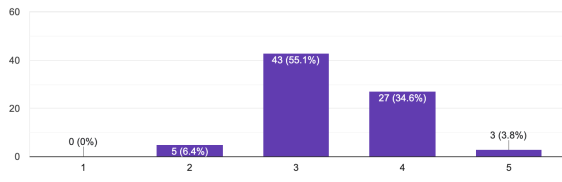
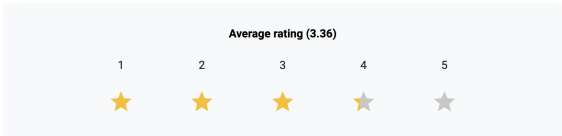
In the time series data [0.29, 0.29, 0.29, 0.29, 0.29, 0.29, 0.29, 0.28, 0.28, 0.28, 0.28, 0.28, 0.28, 0.28, 0.29, 0.3, 0.32, 0.34, 0.37, 0.4, 0.44, 0.47, 0.51], is there evidence of any structural breaks?

Answer: There is a subtle structural break between the stability around 0.28 and the upward shift starting around 0.3.

Explanation: A structural break in time series data is a sudden change in the data pattern or mean level. In the given data, after initially fluctuating around 0.29, there is a persistent drop to 0.28, and then a change occurs towards an upward trend from 0.3 onwards. Although subtle, this transition from 0.28 to 0.3 indicates a structural break where the mean level shifts. It signals the end of one data pattern and the beginning of another distinct phase.

Reasoning: The data stability around 0.29 suggests a constant level, followed by a sudden shift to 0.28, and then an upward trend from 0.3 onwards. This progressive change in behavior indicates there is indeed evidence of a structural break in the sequence.

1198



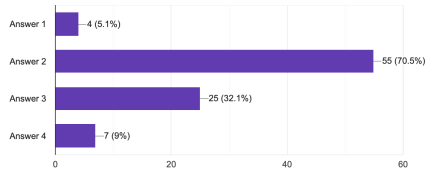
Question 11. Overall, which answer do you prefer for this question? (You can choose one or more)

1199
1200
1201
1202
1203
1204
1205
1206

Given the data points [6.64, 1.71, 27.49, 1.7, 0.2, 6.05, 33.49, 21.13, 0.0, 0.99, 5.34, 10.74, 2.6, 11.04, 14.36, 9.08, 3.69, 0.26, 8.4, 6.37, 2.35, 6.02, 114.07, 49.47], describe the level of volatility found in this data.

Question 11 (Answers):

Answer 1 The data exhibits high volatility, as indicated by the wide range of values from 0.0 to 114.07. The volatility is particularly evident in the presence of such an extreme value. Explanation: In this series, we see values ranging from as low as 0.0 to as high as 114.07, indicating large fluctuations or changes in value from one point to the next. These values suggest a high level of volatility. The presence of an outlier (114.07) significantly impacts the overall volatility, making it difficult to discern a consistent pattern from the other data points. To quantify the volatility, we can calculate the standard deviation. However, due to the presence of the outlier, the standard deviation would be heavily influenced by that value. Thus, we should focus on the range and the presence of the outlier to describe the level of volatility. Highlight key shifts: The data shows a sharp increase in volatility, particularly around the outlier value of 114.07, which is significantly higher than the rest of the data points.	Answer 2 The data shows a high level of volatility, with significant fluctuations between low values and a large outlier at 114.07. The volatility is particularly evident in the presence of such an extreme value. Explanation: In this series, we see values ranging from as low as 0.0 to as high as 114.07, indicating large fluctuations or changes in value from one point to the next. These values suggest a high level of volatility. The presence of an outlier (114.07) significantly impacts the overall volatility, making it difficult to discern a consistent pattern from the other data points. To quantify the volatility, we can calculate the standard deviation. However, due to the presence of the outlier, the standard deviation would be heavily influenced by that value. Thus, we should focus on the range and the presence of the outlier to describe the level of volatility. Highlight key shifts: The data shows a sharp increase in volatility, particularly around the outlier value of 114.07, which is significantly higher than the rest of the data points.	Answer 3 The data exhibits high volatility, as indicated by the wide range of values from 0.0 to 114.07. The volatility is particularly evident in the presence of such an extreme value. Explanation: In this series, we see values ranging from as low as 0.0 to as high as 114.07, indicating large fluctuations or changes in value from one point to the next. These values suggest a high level of volatility. The presence of an outlier (114.07) significantly impacts the overall volatility, making it difficult to discern a consistent pattern from the other data points. To quantify the volatility, we can calculate the standard deviation. However, due to the presence of the outlier, the standard deviation would be heavily influenced by that value. Thus, we should focus on the range and the presence of the outlier to describe the level of volatility. Highlight key shifts: The data shows a sharp increase in volatility, particularly around the outlier value of 114.07, which is significantly higher than the rest of the data points.	Answer 4 The data exhibits high volatility, as indicated by the wide range of values from 0.0 to 114.07. The volatility is particularly evident in the presence of such an extreme value. Explanation: In this series, we see values ranging from as low as 0.0 to as high as 114.07, indicating large fluctuations or changes in value from one point to the next. These values suggest a high level of volatility. The presence of an outlier (114.07) significantly impacts the overall volatility, making it difficult to discern a consistent pattern from the other data points. To quantify the volatility, we can calculate the standard deviation. However, due to the presence of the outlier, the standard deviation would be heavily influenced by that value. Thus, we should focus on the range and the presence of the outlier to describe the level of volatility. Highlight key shifts: The data shows a sharp increase in volatility, particularly around the outlier value of 114.07, which is significantly higher than the rest of the data points.
---	--	---	---



Question 12. Which answer do you consider is the best?

1207
1208
1209
1210
1211
1212
1213
1214
1215
1216
1217
1218
1219
1220
1221
1222
1223

The sequence [0.28, 0.3, 0.34, 0.39, 0.44, 0.49, 0.54, 0.58, 0.62, 0.65, 0.68, 0.7, 0.71, 0.72, 0.73, 0.73, 0.73, 0.73, 0.73, 0.72, 0.71, 0.7, 0.69, 0.68] exhibits an increasing pattern?

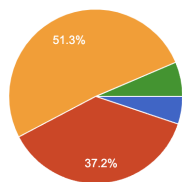
Answer 1: The sequence consistently increases until it peaks at 0.73, after which it slightly decreases, but remains relatively high.

Answer 2: The given sequence does not display a clear increasing pattern throughout. While there is an initial increase up to 0.73, followed by a period of stability, the later points decrease consistently from 0.73 to 0.68, breaking the

continuity of an upward pattern.

Answer 3: The sequence does show an initial increase from 0.28 to 0.73, but it then stabilizes at 0.73 for several data points before decreasing slightly towards the end (down to 0.68). Therefore, while there is a phase of increase, the overall description suggests both rise and stability/decline, which contrasts with an exclusively increasing pattern. Hence, the correct response is False regarding continuous increase throughout.

Answer 4: The sequence initially increases from 0.28 to 0.73, but after reaching 0.73, it remains constant for a few terms and then starts decreasing from 0.73 to 0.68. Therefore, the sequence does not exhibit a continuously increasing pattern throughout.



- Answer 1: The sequence consistently increases until it peaks at 0.73, after which it slightly decreases, but remain...
- Answer 2: The given sequence does not display a clear increasing pattern throughout. While there is an initial inc...
- Answer 3: The sequence does show an initial increase from 0.28 to 0.73, but it...
- Answer 4: The sequence initially increases from 0.28 to 0.73, but after r...