

FLAIRR-TS – Forecasting LLM-Agents with Iterative Refinement and Retrieval for Time Series

Gunjan Jalori

Google

gunjanjalori@google.com

Preetika Verma *

Carnegie Mellon University, USA

preetikv@andrew.cmu.edu

Sercan Ö Arık

Google

soarik@google.com

Abstract

Time series Forecasting with large language models (LLMs) requires bridging numerical patterns and natural language. Effective forecasting on LLM often relies on extensive pre-processing and fine-tuning. Recent studies show that a frozen LLM can rival specialized forecasters when supplied with a carefully engineered natural-language prompt, but crafting such a prompt for each task is itself onerous and ad-hoc. We introduce FLAIRR-TS, a test-time prompt optimization framework that utilizes an agentic system: a Forecaster-agent generates forecasts using an initial prompt, which is then refined by a refiner agent, informed by past outputs and retrieved analogs. This adaptive prompting generalizes across domains using creative prompt templates and generates high-quality forecasts without intermediate code generation. Experiments on benchmark datasets show improved accuracy over static prompting and retrieval-augmented baselines, approaching the performance of specialized prompts. FLAIRR-TS provides a practical alternative to tuning, achieving strong performance via its agentic approach to adaptive prompt refinement and retrieval.

1 Introduction

Recent studies demonstrate that LLMs can leverage their vast pre-trained knowledge to achieve competitive zero-shot and few-shot time-series forecasting (TSF) performance, often rivaling specialized models through direct prompting alone (Xue and Salim, 2024). The efficacy of LLMs in TSF is often stymied by the **prompt engineering bottleneck**. The performance of a frozen, pre-trained LLM is critically dependent on the precise natural language prompt it receives. Crafting optimal prompts is currently a laborious, ad-hoc process requiring significant domain expertise and iterative

manual tuning for each new dataset or scenario, limiting scalability and robust generalization (Niu et al., 2024). This challenge has spurred research into more sophisticated prompting strategies (Liu et al., 2024; Tang et al., 2024) and test-time methods without altering weights (Jin et al., 2024).

Given that LLMs can iteratively refine their outputs through feedback (as in Madaan et al. (2023) and Chen and others (2025)), we explore their capability to autonomously refining their prompts at test time to enhance TSF capabilities. We introduce **FLAIRR-TS - Forecasting LLM-Agents with Iterative Refinement and Retrieval**, a framework designed to enhance TSF capabilities of LLMs without any training. This approach aims to mitigate the manual prompt engineering burden while simultaneously improving prediction accuracy by grounding forecasts in relevant historical context. FLAIRR-TS integrates a **Forecaster agent** (F) for initial predictions, a **Refiner Agent** for Iterative Refinement Tuning (IRT), and a **Retrieval agent** (R) that sources semantically similar historical time series segments, akin to Retrieval Augmented Generation (RAG) adapted for TSF (Han et al., 2023). This entire cycle of prompt adaptation and forecast refinement occurs without any weight updates, offering a compelling alternative to costly tuning. Beyond the capabilities of FLAIRR-TS for general applicability, we also investigate the upper bounds for performance with judiciously-designed prompts. Inspired by (Sahoo et al., 2025), we introduce **Architected Strategy Prompts** (ASPs), a set of specialized prompts, which include directives for specific analytical procedures or induce particular cognitive approaches. While FLAIRR-TS excels at automated, test-time prompt refinement without prior domain-specific tuning, ASPs allow exploring the performance when, manual strategy-driven design is employed for prompt improvement.

Our main contributions are summarized as:

- We propose **FLAIRR-TS**, a novel prompting

*Work done at Google

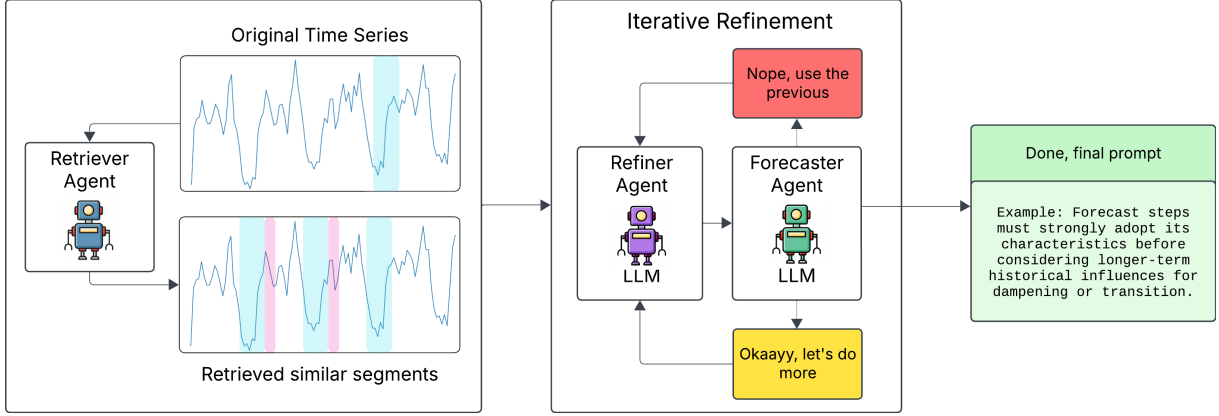


Figure 1: Flowchart of the the proposed method framework, consisting Retrieval, Forecaster and Refiner agents.

and test-time optimization framework for TSF with iterative refinement and retrieval.

- We utilize retrieval augmentation for TSF with LLMs with the introduced **ASPs**, developed via systematic prompt design, to reveal the significant impact of specialized instructions and to serve as high-performance benchmarks.
- We demonstrate that FLAIRR-TS consistently improves forecasting accuracy across diverse datasets without model fine-tuning, outperforming static domain agnostic prompting and a non-iterative retrieval-augmented baseline

2 Methodology

2.1 Overall Agentic Architecture

We propose FLAIRR-TS, a framework combining test-time optimization for iterative refinement via prompting by an agentic system, and retrieval-augmented context to enhance TSF with pre-trained LLMs. As illustrated in Figure 1 and formally detailed in Algorithm 1, it operates as a multi-agent system. The **Forecaster Agent** generates predictions using a prompt that is dynamically improved by the **Refiner Agent** during an **Iterative Tuning** phase. This process is enriched by the **Retrieval Agent** that provides the relevant historical context and augments it to the input provided to the forecaster. The core iterative cycle (Alg. 1, lines 7-20) involves forecasting, evaluating the forecast against recent ground truth (e.g., via a metric like MAE), and refining the prompt. The Refiner agent can signal early termination if the forecast improvement falls below a defined threshold, τ_{stop} (e.g., a 5% reduction in MAE). Otherwise, if maximum iterations (N_{iter}) are reached, the system defaults

to the prompt that yielded the best observed MAE. This adaptive optimization occurs at test-time without any model training.

2.2 Core Agent Descriptions

Retrieval Agent. Inspired by RAFT (Han et al., 2023), this agent (Alg. 1, line 8) enhances the Forecaster Agent’s inputs by retrieving M historical time series segments (S_{retr}) from a historical database. This database is constructed by applying a sliding window of length L across the entire training split of the dataset. Similarity between the current context window (X_{Ctx}) and the historical segments is measured using **Pearson’s correlation**, with the top- M most similar segments being retrieved. These segments, along with their actual outcomes, provide illustrative examples of past pattern evolutions, directly augmenting the context (C_{aug}) given to the Forecaster-agent (see Appendix D for data formatting details).

Refiner-agent (R). Functioning as a meta-optimizer (Alg. 1, line 12), the Refiner Agent is stateful and analyzes the entire history of the current refinement session. As shown in the algorithm, this history, $\{(P^{(i)}, mae^{(i)})\}_{i=0}^{k-1}$, includes all previously attempted prompts and their resulting errors. By observing this full trajectory, the agent can make non-myopic, informed decisions about the next prompt modification (P_{next}). It provides a done_signal if the forecast quality meets the pre-defined termination criterion (i.e., MAE improvement is less than τ_{stop}). Its detailed reasoning, guided by a specific prompt structure (see Appendix B), might yield feedback such as, Pay closer attention to sudden changes in the last 10% of the input sequence.

Algorithm 1 FLAIRR-TS Algorithm

Require: Training data X , Historical series $X_{1:t-1}$, Horizon H , Initial prompt P_0 , Context length L , #Segments M , Max iterations N_{iter} , Recent ground truth $X_{t:t+H}$, Stopping threshold τ_{stop}

Ensure: Selected prompt P_{out}

```
1:  $P_{\text{curr}} \leftarrow P_0$ ;  $P_{\text{best}} \leftarrow P_0$ ;  $\text{mae}_{\text{min}} \leftarrow \infty$ ;  $\hat{X}_{\text{best}} \leftarrow \text{nil}$ ;  $\text{early\_stop} \leftarrow \text{false}$ 
2:  $X_{\text{HistDB}} \leftarrow X_{1:t-L-1}$ ;  $X_{\text{Ctx}} \leftarrow X_{t-L:t}$   $\triangleright$  Setup context and historical DB
3: for  $k \leftarrow 1$  to  $N_{\text{iter}}$  do
4:    $S_{\text{retr}} \leftarrow \text{RETRIEVESEGMENTS}(X_{\text{HistDB}}, X_{\text{Ctx}}, M)$ 
5:    $C_{\text{aug}} \leftarrow \text{AUGMENTCONTEXT}(X_{\text{Ctx}}, S_{\text{retr}})$ 
6:    $\hat{X}_{\text{cand}} \leftarrow \text{FORECASTERLLM}(P_{\text{curr}}, C_{\text{aug}}, H)$ 
7:    $\text{mae}_{\text{curr}} \leftarrow \text{CALCULATEMAE}(\hat{X}_{\text{cand}}, X_{t:t+H})$ 
8:   if  $\text{mae}_{\text{curr}} < \text{mae}_{\text{min}}$  then
9:      $\text{mae}_{\text{min}} \leftarrow \text{mae}_{\text{curr}}$ ;  $P_{\text{best}} \leftarrow P_{\text{curr}}$ ;  $\hat{X}_{\text{best}} \leftarrow \hat{X}_{\text{cand}}$ 
10:  end if
11:   $(P_{\text{next}}, \text{done\_signal}) \leftarrow \text{REFINERLLM}(\{(P^{(i)}, \text{mae}^{(i)})\}_{i=0}^{k-1}, P_{\text{curr}}, \tau_{\text{stop}})$   $\triangleright$  Refiner sees history
    of prompts and errors
12:  if  $\text{done\_signal}$  then
13:     $P_{\text{out}} \leftarrow P_{\text{curr}}$ ;  $\text{early\_stop} \leftarrow \text{true}$ ; break
14:  end if
15:   $P_{\text{curr}} \leftarrow P_{\text{next}}$ 
16: end for
17: if not  $\text{early\_stop}$  then  $\triangleright$  Fallback to best MAE if max iterations reached
18:    $P_{\text{out}} \leftarrow P_{\text{best}}$ 
19: end if
20: return  $P_{\text{out}}$ 
```

Forecaster-agent (F). This agent (Algorithm 1, line 10) is responsible for generating the time series forecast ($\hat{X}_{\text{cand}}$). It uses the current prompt (P_{curr})—either the initial prompt P_0 or the one refined by the Refiner Agent—along with the augmented context (C_{aug}) provided by the Retriever Agent. FLAIRR-TS allows for the utilization of a potentially more compact LLM as this agent, with its behavior shaped by dynamically optimized prompts. The structure of the prompts is detailed in Appendix C.

2.3 Architected Strategy Prompts (ASP)

We present some judiciously-designed prompts inspired by the results of FLAIRR. Instead of merely asking an LLM to predict future values, we aim to induce more complex, imaginative reasoning. If manual prompt engineering is permitted, we aim to improve accuracy further by carefully editing the best prompts from FLAIRR. ASPs are developed by building upon the results achieved by FLAIRR. Some examples include:

Analytical

Deep STL analysis: (inspired by (Zhou et al., 2024)) perform an STL decomposition, forecast each component, then recombine them via STL addition.

Thinking–Inductive

Monte-Hall Prompting: frame forecasting as a decision game so the model evaluates several scenarios before committing.

Imaginative

(a) **Many-Worlds Reasoning:** simulate multiple plausible futures and aggregate them.

(b) **D&D Dungeon-Master:** forecast a character’s hit-point trajectory over upcoming turns.

More details about ASP prompts are in Appendix E.

3 Experiments

Our experiments utilize the Informer (Zhou et al., 2021) benchmark datasets¹: ETT (ETTh1, ETTh2,

¹Full experimental parameters and any dataset-specific preprocessing are in the Appendix.

Dataset	Horizon	Supervised				PTMs		Prompt			
		Informer	DLinear	FEDformer	PatchTST	TTM	Time-LLM	LSTP	FLAIRR (Ours)	ASP(G2.5P) (Ours)	ASP(G2.0F) (Ours)
ETTh1	96	0.76	0.39	0.58	0.41	0.36	0.46	0.15	<u>0.101</u>	0.078	0.118
	192	0.78	0.41	0.64	0.49	0.39	0.54	<u>0.22</u>	0.246	0.208	<u>0.223</u>
traffic	96	0.69	0.28	0.56	0.25	0.46	0.25	0.32	<u>0.145</u>	0.143	0.184
	192	0.58	0.28	0.58	<u>0.26</u>	0.49	0.25	0.31	0.326	0.324	0.296

Table 1: Performance comparison (MAE) of supervised models and zero-shot methods on benchmark datasets. FLAIRR (Ours), ASP(G2.5P) (Ours), and ASP(G2.0F) (Ours) are our proposed/evaluated methods.

Dataset	Horizon	Supervised				Prompt			
		Informer	AutoFormer	FedFormer	PatchTST	LSTP	FLAIRR (Ours)	ASP(G2.5P) (Ours)	ASP(G2.0F) (Ours)
ILI	4	1.54	1.24	2.54	0.43	0.38	0.271	<u>0.264</u>	0.189
	12	2.33	1.82	2.67	0.43	0.39	0.249	0.183	<u>0.197</u>
	20	2.12	1.90	1.75	1.26	0.73	<u>0.589</u>	0.564	0.867
	24	3.99	1.79	1.50	1.72	1.55	<u>0.724</u>	0.722	1.004
	24	1.45	1.38	1.95	1.55	0.17	<u>0.110</u>	0.084	0.125
Weather	48	1.57	1.43	1.67	1.56	0.24	<u>0.160</u>	0.142	0.238
	96	1.48	1.67	1.96	1.12	0.39	0.290	<u>0.257</u>	0.243
	120	1.90	1.74	2.02	1.31	0.51	0.383	0.309	<u>0.369</u>

Table 2: Performance comparison (MAE) on datasets whose test periods post-date the Gemini 2.5 Pro knowledge cut-off. FLAIRR and both ASP variants are ours; *Informer-PatchTST* are supervised baselines; *LSTP* is a prior prompt-based method.

ETTh1, ETTh2), Electricity, and Traffic. We also benchmark on several newer datasets, including Weather and ILINet, and we test on 2025 data to ensure the test period is after the knowledge cutoff date of Gemini. More details are in Appendix F. The characteristics of all datasets (domains, frequencies, evaluated horizons H) and our approach to data integrity are provided in Section F.

LLM Backbone: We implement FLAIRR using Gemini 2.5 Pro and ASP using both Gemini 2.5 Pro and Gemini 2 Flash. For our ablation study, we replicate these experiments with DeepSeek-V3.

Data & Execution: To ensure robust results, we normalize inputs using standard scaling, control for numerical precision in prompts, and report the median performance over five independent runs for each experiment.

Results: Results are presented in Table 1 for long-horizon datasets and Table 2 for short-horizon datasets. We use the Mean Absolute Error (MAE) metric. We compare our work with the most recent prompt-based method, LSTPrompt (Liu et al., 2024) (using a frozen Gemini as its backbone), and two of the best-performing PTMs—TTM (Ekambaram et al., 2024) and Time-LLM (Jin et al., 2024). We also compare against non-LLM supervised methods like DLinear (Zeng et al., 2022).

Analysis: Across 20 distinct scenarios, LAIRR and ASP, outperform all competing models in 14 cases, including every smaller horizon task, and consistently surpass the LSTP baseline on all datasets.

3.1 Ablations

We disentangle the impact of *Retrieval* and *Iterative Refinement (IR)* by successively activating them on top of a *Simple Prompt*. Figure 2 shows the MAE on ETTh2 for Gemini 2.5 Pro, Gemini

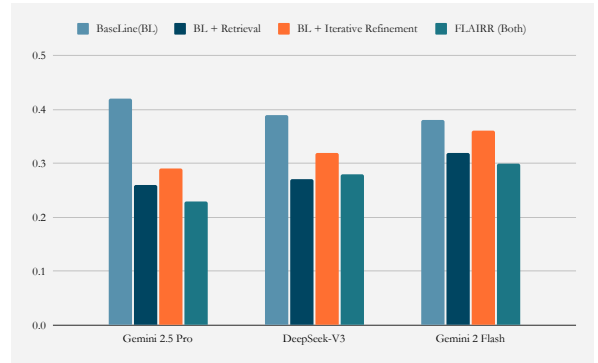


Figure 2: Ablation results, average MAE. Lower MAE is better.

2 Flash, and the open-source DeepSeek-V3.

Observations: Retrieval lowers error by grounding forecasts in analogous history, while IR refines outputs through on-the-fly prompt correction. Their combination (**FLAIRR-TS**) delivers the lowest MAE across all three backbones. Crucially, the same trend holds for DeepSeek-V3, demonstrating that our gains are architecture-agnostic.

4 Conclusion

FLAIRR-TS advances the approach to prompting for time-series forecasting. Its core contribution is not to surpass every hand-tuned prompt, but to significantly reduce the burden of manual tuning. The framework provides a systematic and automated process for refining prompts, ensuring consistently high performance from even simple starting instructions. Through its agentic, feedback-driven interactions, FLAIRR-TS offers a scalable pathway to unlocking the full potential of LLMs for forecasting across any dataset or horizon.

5 Limitations and Future Work

- **Evaluation coverage.** It'd be important to extend empirical validation to cover robustness to irregular sampling, regime shifts, or domain drifts.
- **Analogue-retrieval assumption.** FLAIRR-TS assumes the presence of semantically similar historical segments. When they do not exist (e.g. for novel events or cold-start prediction scenarios), the refinement loop would come with more risks for compounding errors.
- **Numerical fidelity of LLMs.** LLMs exhibit limited precision on long or out-of-range sequences, and might hallucinate trends under noise or scale shifts, constraining their reliability. Our method is expected to take advantage of future improvements in LLMs for long input-output modeling and numerical understanding.
- **Inference cost.** Iterative prompting adds multiple LLM calls per forecast. Latency and energy consumption might be prohibitive for real-time, high-frequency settings, motivating for approaches like LLM distillation.

References

- Zhaofeng Chen and others. 2025. SETS: Self-verification and self-correction for improved test-time scaling. Anticipated for International Conference on Machine Learning (ICML). Placeholder entry. Details may need updating upon actual publication. Search for preprint by Zhaofeng Chen on SETS.
- Vijay Ekambaram, Arindam Jati, Pankaj Dayama, Sumanta Mukherjee, Nam H. Nguyen, Wesley M. Gifford, Chandra Reddy, and Jayant Kalagnanam. 2024. Tiny time mixers (tms): Fast pre-trained models for enhanced zero/few-shot forecasting of multivariate time series. *Preprint*, arXiv:2401.03955.
- Nate Gruver, Marc Finzi, Shikai Qiu, and Andrew Gordon Wilson. 2023. Large language models are zero-shot time series forecasters. In *Advances in Neural Information Processing Systems 36 (NeurIPS 2023)*, pages 24013–24034. ArXiv:2310.07820.
- Seungone Han, Peiyuan Liao, Poming P. Chiu, Jennifer Hobbs, Sungtae An, Min hwan Oh, Vikas K. Garg, Caiming Xiong, and Yoonkey Kim. 2023. Retrieval augmented time series forecasting. In *Advances in Neural Information Processing Systems 36 (NeurIPS 2023)*, pages 73654–73670. ArXiv:2310.16227.
- Ming Jin, Shiyu Wang, Lintao Ma, Zhixuan Chu, James Y. Zhang, Xiaoming Shi, Pin-Yu Chen, Yuxuan Liang, Yuan-Fang Li, Shirui Pan, and Qingsong Wen. 2024. Time-LLM: Time series forecasting by reprogramming large language models. In *The Twelfth International Conference on Learning Representations (ICLR)*. ArXiv:2310.01728.
- Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. 2020. Retrieval-augmented generation for knowledge-intensive NLP tasks. In *Advances in Neural Information Processing Systems 33 (NeurIPS 2020)*, pages 9459–9474.
- Haoxin Liu, Zhiyuan Zhao, Jindong Wang, Harshvardhan Kamarthi, and B. Aditya Prakash. 2024. Lst-prompt: Large language models as zero-shot time series forecasters by long-short-term prompting. *Preprint*, arXiv:2402.16132.
- Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler Hallinan, Luyu Gao, Sarah Wiegrefe, Uri Alon, Nouha Dziri, Shrimai Prabhumoye, Yiming Yang, Shashank Gupta, Bodhisattwa Prasad Majumder, Katherine Hermann, Sean Welleck, Amir Yazdanbakhsh, and Peter Clark. 2023. Self-Refine: Iterative refinement with Self-Feedback. arXiv preprint arXiv:2303.17651. *Preprint*, arXiv:2303.17651.
- Yuqi Nie, Nam H. Nguyen, Phanwadee Sinthong, and Jayant Kalagnanam. 2023. A time series is worth 64 words: Long-term forecasting with transformers. *Preprint*, arXiv:2211.14730.
- Peisong Niu, Tian Zhou, Xue Wang, Liang Sun, and Rong Jin. 2024. Understanding the role of textual prompts in llm for time series forecasting: an adapter view. *Preprint*, arXiv:2311.14782.
- Pranab Sahoo, Ayush Kumar Singh, Sriparna Saha, Vinija Jain, Samrat Mondal, and Aman Chadha. 2025. A systematic survey of prompt engineering in large language models: Techniques and applications. *Preprint*, arXiv:2402.07927.
- Jingyi Tang, Zongyao Zhang, Daksh Minhas, Chengzhang Li, Haomin Chen, Minghuan Tan, Chetan Shah, and Joyce C. Ho. 2024. Prompting medical large vision-language models to diagnose pathologies by visual question answering. arXiv preprint arXiv:2407.21368. *Preprint*, arXiv:2407.21368. Verified from.[53] Placeholder ‘tang-et al-2024-enrichingprompts’ resolved.
- Yu-Hsiang Lin Wan, Akshita Agrawal, Chiyu Max Jiang, Eunsol Choi, and Graham Neubig. 2024. Self-supervised prompting for cross-lingual in-context learning in low-resource languages. arXiv preprint arXiv:2406.18880. *Preprint*, arXiv:2406.18880. Verified from.[55] Placeholder ‘wan-et al-2024-incontextemplars’ resolved.
- Hao Xue and Flora D. Salim. 2024. Promptcast: A new prompt-based learning paradigm for time series forecasting. *IEEE Transactions on Knowledge and Data Engineering*, 36(11):6851–6864.

- Ailing Zeng, Muxi Chen, Lei Zhang, and Qiang Xu. 2022. [Are transformers effective for time series forecasting?](#) *Preprint*, arXiv:2205.13504.
- 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 *The Thirty-Fifth AAAI Conference on Artificial Intelligence, AAAI 2021, Virtual Conference*, volume 35, pages 11106–11115. AAAI Press.
- Tian Zhou, Ziqing Ma, Qingsong Wen, Xue Wang, Liang Sun, and Rong Jin. 2022. [FEDformer: Frequency enhanced decomposed transformer for long-term series forecasting](#). In *Proceedings of the 39th International Conference on Machine Learning (ICML 2022)*, volume 162 of *Proceedings of Machine Learning Research*, pages 27268–27286. PMLR.
- Wendi Zhou, Xiao Li, Lin Geng Foo, Yitan Wang, Harold Soh, Caiming Xiong, and Yoonkey Kim. 2024. [TEMPO: Temporal representation prompting for large language models in time-series forecasting](#). arXiv preprint arXiv:2405.18384. Anticipated for NeurIPS 2024. *Preprint*, arXiv:2405.18384.

A Related Work

Time Series Forecasting with LLMs: Traditional time series forecasting has relied on models explicitly trained for the task, from statistical methods to deep architectures like RNN variants and temporal CNNs, up through recent Transformer-based models (e.g. FEDformer (Zhou et al., 2022) and PatchTST ((Nie et al., 2023))) tailored for long-range sequences. These approaches require substantial training on each target dataset. In contrast, emerging research explores using pre-trained LLMs as general-purpose forecasters via prompting at inference time only, without gradient-based fine-tuning. Xue and Salim (2024) pioneered this direction with PromptCast, formulating forecasting as a prompt-completion task: historical values are encoded into a textual prompt (possibly with instructions) and the LLM’s next-token predictions are decoded as forecasts. Gruver et al. (2023) similarly represent numerical time series as token sequences and treat extrapolation as language modeling, finding that GPT-3 and LLaMA-2 can zero-shot extrapolate time series with accuracy comparable to or exceeding specialized trained models. TNotably, these LLM-based approaches leverage the models’ strong sequence modeling and few-shot generalization for competitive benchmark results, without requiring abilities to achieve competitive results on standard benchmarks without any task-specific training data. Nevertheless, naive prompt formulations might overlook important temporal dynamics and patterns. Recent works therefore propose more advanced test-time prompting strategies. Liu et al. (2024) introduce LSTPrompt, which splits the prediction into short- and long-term sub-tasks and guides the LLM through a chain-of-thought reasoning process; this method outperforms earlier prompt baselines and even approaches the accuracy of dedicated TS models. Tang et al. (2024) report that enriching prompts with external knowledge (e.g. known seasonal periods or contextual clues) and using natural language rephrasings of the input can significantly improve an LLM’s forecasting accuracy. Another technique, Time-LLM (Jin et al., 2024), reprograms a frozen LLM by mapping time-series data into textual “patches” and prepending learned prompt tokens, allowing the model to output forecasts that outperform state-of-the-art specialized forecasters without any fine-tuning of the LLM’s weights. On the other hand, Zeng et al. (2022) offer a cautionary perspective:

through extensive ablations, they found that removing the LLM or replacing it with a simple attention-based network in these pipelines often does not hurt performance (and sometimes improves it), calling into question how much current LLM-for-TS methods truly benefit from the pre-trained language model. To push LLM-based forecasting further, researchers are drawing on insights from prompt optimization and test-time reasoning. For example, Wan et al. (2024) show that intelligently selecting and reusing in-context exemplars can yield larger gains than optimizing instructions alone, suggesting that careful few-shot prompt design is crucial. Chen and others (2025) propose a self-verification and self-correction framework (SETS) that lets the model iteratively refine its outputs at inference, achieving better accuracy scaling on complex reasoning tasks. Incorporating such techniques into zero-shot forecasting prompts is an exciting direction. In summary, the literature demonstrates a nascent but growing paradigm of using pre-trained LLMs directly for time series forecasting, with multiple studies showing that, given the right prompts, foundation models can attain forecast accuracy rivaling traditional specialized models. While these methods demonstrate progress in leveraging LLMs for forecasting, the dynamic and optimal design of prompts—especially those needing to integrate complex reasoning, external knowledge, and iterative feedback—remains a key challenge. Our work, FLAIRR-TS, aims to address this by structuring the forecasting process around specialized agents for dynamic prompt adaptation and refinement.

Agentic Frameworks with Iterative Refinement

The concept of employing multiple interacting agents or distinct processing roles for complex problem-solving has gained traction in AI. Such agentic systems can distribute tasks, specialize functionalities, and enable more sophisticated reasoning or generation processes. Iterative refinement, where an output is progressively improved through feedback loops, is a common characteristic of these systems and is also seen in self-correction mechanisms within single LLMs (e.g., Self-Refine by Madaan et al. (2023)). For instance, systems might involve a generator agent and a critic agent, or distinct agents for planning, execution, and verification. FLAIRR-TS draws inspiration from these paradigms by structuring its operation around specialized agents: a Forecaster-agent for initial prediction, a retriever agent for sourcing relevant con-

text, and a refiner agent for iterative prompt refinement. This agentic decomposition facilitates more targeted and adaptable modifications to distinct aspects of the forecasting prompt through these specialized roles. Crucially, unlike traditional multi-agent systems where agents might be independently trained or involve complex coordination protocols, FLAIRR-TS implements these roles using LLMs at test time to dynamically adapt the prompting strategy itself. The "refinement" occurs in the textual instructions and contextual information fed to the LLM, rather than through updates to model weights, distinguishing it from model distillation or training paradigms. This focus on inference-time prompt adaptation through an agentic perspective is a key aspect of our approach. This structured approach also aims to ensure that the LLM's reasoning and generative capabilities are a core component of the forecasting process, addressing concerns about their actual contribution in some prior LLM-for-TS pipelines.

Retrieval Augmented Generation: Retrieval Augmented Generation (RAG) (Lewis et al., 2020) has become a standard technique for enhancing LLMs in knowledge-intensive NLP tasks. RAG systems retrieve relevant documents or passages from an external corpus and provide them as additional context to the LLM, improving factual grounding and reducing hallucination. Recently, Han et al. (2023) adapted this concept to time series forecasting with their Retrieval Augmented Time Series Forecasting (RAFT) approach. RAFT retrieves historical time series segments similar to the current input window and uses them to augment the context provided to a forecasting model (in their case, an LLM). Our work directly builds upon and integrates the RAFT principle within the Retrieval agent component of FLAIRR-TS. We hypothesize that the effectiveness of RAFT can be further enhanced by optimizing the prompt that instructs the LLM on how to utilize the retrieved historical context, which is precisely what the agentic interaction within FLAIRR-TS aims to achieve.

B Refiner Agent

You are an expert Time-Series-Forecasting **Prompt Engineer** acting as a **Refiner Agent**. Your goal is to analyze a set of forecasting attempts made by a **Forecaster Agent** and provide specific, actionable **Learnings** on how to improve the *initial forecasting prompt* used by the Forecaster. The Forecaster Agent uses a

base prompt and adds new forecasting instructions to it based on your learnings.

Key Information for Your Analysis for this Iteration {it + 1}:

1. Current Forecasting Instructions Under Review: {current_instructions_under_review}
2. Overall Mean Absolute Error (MAE) for this batch of samples: {mae_to_report_to_teacher}

You will also be given a batch of individual samples, where each sample includes:

1. The full Prompt the Forecaster Agent used (includes the instructions above).
2. The Forecaster Agent's Predictions for the OT variable.
3. The Ground-Truth OT values.

Your Analysis Task:

1. **Identify error patterns.** Compare Predictions with Ground Truths. Look for systematic errors (e.g., over/under-prediction, lagging, volatility mis-handling).
2. **Correlate errors with prompts and instructions.** Check whether the current instructions are ambiguous, misleading, too complex, or otherwise harmful.
3. **Formulate "Learnings".** Give concrete, generalizable improvements (e.g., adjust look-back horizon, drop STL decomposition, add weekday feature).
4. **Determine "Done" status.**
 - If the percentage reduction in MAE is less than the pre-defined stopping threshold (τ_{stop}), output Done: True.
 - Otherwise output Done: False.

Output Format—exactly this template

Learnings: <your concise, actionable suggestions here>

Done: <True or False>

Confidence in output: <High | Medium | Low> — one-line rationale.

C Forecaster Agent

Prompt-Synthesis Instructions

Example: Forecasting-Instruction Refinement

You are an intelligent agent that synthesizes forecasting prompts based on expert feedback. You will receive *Learnings* from a **Refiner Agent** that suggest improvements to an initial time-series forecasting prompt. Your task is to turn these learnings into concise and effective *prompt-forecasting instructions*. These instructions will be appended to a base forecasting prompt.

The forecasting instructions should:

- Be a short set of guiding principles (maximum 3 actionable items).
- Directly address the issues and suggestions in the Learnings.
- Be clearly phrased for another LLM to follow.
- ****Do not include placeholders such as {previous_data} or {prediction_data}.**
- ****Do not change the output format or the forecasting task itself.**
- If no actionable learnings exist, output a safe generic set—or state:

No specific new instructions generated due to lack of actionable learnings.

Example (Refiner Agent said “focus on recent volatility”):

Learnings: The model often misses sudden spikes; the prompt should ask the forecaster to pay more attention to recent volatility and its effect on the next step.

Your Output (forecasting instructions): “Critically assess the volatility in the most recent data points. Your forecast for the next step should reflect whether this volatility is expected to continue, increase, or decrease. Explain this assumption in your reasoning.”

Learnings you received:

{current_learnings}

Based on these learnings, generate only the refined prompt-forecasting instructions below (no extra commentary).

Refined Prompt Forecasting Instructions:

<model prediction here>

D Prompt template

Objective

Provide a well-reasoned forecast for the {target_variable} value in the next row of the dataset, given the historical data.

Dataset Instructions

- **Dataset:** data_name, data_description
- **Variable to Predict:** {target_variable}.
- **Task:** Predict the {target_variable} values for the next {prediction_length} steps using the historical data.
- **Constraints:**
 - Adhere strictly to the specified output format.

If instructions:

Forecasting Instructions: {instructions}

If raft_context: The {raft_context} contains the M retrieved historical segments. Each segment and its

corresponding ground-truth outcome are formatted as comma-separated text strings.

Input Data

- Historical Data: {previous_sequence_length_data}

Output Format — exactly this

Predicted Values: [predicted_value_1, ...]

Reasoning: [Your detailed reasoning]

Certainty Estimate: [Percentage certainty]

Certainty Reasoning: [reasoning]

E Prompt Library

The following library of prompts was designed to test different cognitive pathways of the LLM. The strategies are grouped into categories: Analytical prompts that enforce a structured decomposition of the problem; Thinking-Inductive prompts that encourage probabilistic or scenario-based reasoning; and Imaginative prompts that use metaphor and creative framing to elicit novel patterns.

teacher-student-loop

ACT I – REFINER AGENT Propose a first-pass forecast for the next {sequence_length} steps.

ACT II – FORECASTER AGENT Evaluate the Refiner’s forecast against the most recent known data and suggest corrections.

ACT III – REFINER AGENT Incorporate feedback and provide the refined forecast.

self-verification-sets

Step 1 – Generate candidate forecast A for {sequence_length} steps. Step 2 – Generate independent candidate forecast B. Step 3 – For each horizon h, if the two differ beyond an acceptable tolerance, reconcile them (e.g., by averaging). Provide only the reconciled forecast.

meta-prompt-conf-bands

Forecast {sequence_length} steps and include 68% and 95% confidence bands. Briefly explain the uncertainty assumptions before the numbers.

imaginary-python-repl

You are **ForecastPy**, a mental Python REPL. Think then “run code in your head” that derives the forecast for the next {sequence_length} steps.

synesthetic-soundtrack

Interpret the past sequence as MIDI velocity (0–127) and compose the next {sequence_length} beats that extend the melody. Provide both the MIDI integers and the values rescaled to original units.

color-gradient-canvas

Map each value to an RGB triplet on a blue-to-red gradient. Produce a grid of HEX colours that encodes the next {sequence_length} points.

dungeon-master

You are a D&D Dungeon Master. The party's HP over the last turns is shown. Forecast HP for the next {sequence_length} turns, assuming no boss fights and only mild potion use.

micro-essay-poisson

Write a ≤ 60 -word micro-abstract describing the generative mechanism, then list {sequence_length} λ parameters for a Poisson baseline.

reverse-sudoku

Think of the next {sequence_length} points as filling a 9×11 Sudoku-like grid whose row sums match the recent history. Provide the grid and a flattened list.

many-worlds-ensemble

Create forecasts for four parallel universes (A–D) shifted by -2σ , -1σ , $+1\sigma$, $+2\sigma$, each {sequence_length} steps long, then provide a consensus median forecast.

haiku-seeded

Compose a three-line haiku that metaphorically describes the upcoming pattern, then list the {sequence_length} numeric forecasts, one per line.

F Datasets

Experiments were performed on a diverse set of widely-used time-series-forecasting (TSF) benchmark datasets spanning multiple domains, sampling frequencies, and statistical characteristics (e.g., seasonality, trend, noise levels). All datasets are normalized with StandardScaling from sklearn package. The datasets are:

- **ETT** (ETTh1, ETTh2, ETTm1, ETTm2) – Electricity Transformer Temperature data recorded at hourly (h) or 15-minute (m) intervals; widely used for long-sequence forecasting with OT as target variable (ETTh: 17,420 total data points, ETTm: 69,680 total data points)
- **Electricity** – Hourly household electricity data of customers with electricity consumption as target variable (26,304 total data points)
- **Traffic** – Hourly occupancy rates from California road-traffic sensors (2021–2025 March) with traffic volume as target variable (17,544 data points)
- **ILINet** – Weekly Influenza-Like-Illness counts from the CDC (2002–2025 April) with total ILI patients as target variable (1,441 total data points)²
- **Weather** – Hourly weather data from Chicago with temperature as target variable (35,052 total data points)³

F.1 Data Integrity

A significant consideration when utilizing Large Language Models (LLMs) for time series forecasting is the potential for the model's pre-training data to inadvertently include samples from the test set, which could lead to an overestimation of predictive performance. To rigorously uphold data integrity in this study, we employed ILINet and weather datasets as benchmarks, with a specific focus on temporal data separation. Our experimental design ensures that all data samples within the test set originate from dates strictly subsequent to the known training data cut-off date of the LLM employed for inference. This chronological separation mitigates the risk of test data contamination, providing a robust and fair evaluation of the LLM's ability to generalize and forecast genuinely unseen future values.

F.2 Evaluation Metrics

Forecasting performance was assessed with two standard error metrics:

$$\text{MAE} = \frac{1}{H} \sum_{i=1}^H \left| \hat{X}_{t+i} - X_{t+i} \right|, \quad (1)$$

²<https://gis.cdc.gov/grasp/fluview/fluportaldashboard.html>

³<https://www.kaggle.com/datasets/curiel/chicago-weather-database>

Where H is the prediction horizon, $\hat{X}_{t+i}$ is the predicted value, and X_{t+i} is the ground-truth value. Lower values indicate better performance for both metrics. These metrics were computed directly from the experimental results.

F.3 Hyperparameter Settings

Here we detail the key hyperparameters used for the FLAIRR-TS framework in our experiments.

ts_max_iter Value: 5. The maximum number of refinement iterations allowed in the iterative refinement loop.

ts_stopping_criteria Value: 5%. The refinement process is stopped early if the percentage reduction in Mean Absolute Error (MAE) between iterations falls below this threshold.

ts_sample_size Value: 3. The number of validation samples used within each refinement iteration to evaluate the quality of a candidate prompt.

raft_m_retrieval Value: 2. The number of similar historical segments retrieved by the Retrieval Agent to be used as context.

G Future directions

There are several avenues for future work. One direction is to incorporate quantitative validation in the loop: currently, the Refiner-agent’s feedback quality is not directly measured. If we had a small hold-out set or could use the model’s own likelihood of the data, we might select or weight feedback. This leans towards techniques in automatic prompt optimization where a reward is defined. Additionally, while FLAIRR-TS currently uses natural language for feedback from the Refiner-agent, one could imagine hybrid approaches where the Refiner-agent suggests pseudo-code or formulaic adjustments (if the LLM agents are equipped with a calculator tool). That could improve handling of scale and magnitude issues. On the retrieval side, exploring more advanced analog search (perhaps using learned embeddings or matching not just on raw values but pattern descriptors) might yield even more relevant cases to show the Refiner-agent, especially for complex multivariate data.

From an application perspective, deploying FLAIRR-TS in an interactive forecasting system would be very interesting. Because FLAIRR-TS’s intermediate steps (the prompts, the retrieved

analog, the feedback) are human-readable, a human analyst could intervene in the loop – agreeing or disagreeing with the Refiner-agent’s critique, or adding their own feedback. This could turn forecasting into a collaborative dialog between human, Forecaster-agent, and Refiner-agent. In settings like supply chain or epidemiology forecasting, such a system could help build trust as well, since each refinement step can be scrutinized.