

Improving the Sample Efficiency of In-Context Learning in Large Language Models Through Meta-Level Optimization

Zixuan Zhou
Federal University of Bahia

Abstract—In-Context Learning (ICL) has emerged as a compelling paradigm for leveraging Large Language Models (LLMs) on downstream tasks without explicit fine-tuning. However, the performance of ICL is often sensitive to the selection of demonstrations, and standard pre-training objectives do not explicitly optimize for this learning paradigm. To address these limitations, we propose a novel training strategy called Meta-In-Context Learning (Meta-ICL) Training. Our approach involves pre-training LLMs on a diverse set of tasks, where each training instance is formulated as an in-context learning scenario, consisting of a task description, a few demonstrations, and a query. Through extensive experiments on seven Natural Language Understanding benchmarks, we demonstrate that LLMs trained with our Meta-ICL strategy significantly outperform baseline methods, including sophisticated demonstration selection techniques. Our findings highlight the effectiveness of explicitly training LLMs to learn from context, paving the way for more robust and effective in-context learning capabilities.

Index Terms—In-Context Learning, Large Language Models, Meta-In-Context Learning

I. INTRODUCTION

The remarkable capabilities of Large Language Models (LLMs) have revolutionized the field of Natural Language Processing (NLP). Unlike traditional approaches that rely on extensive task-specific fine-tuning, In-Context Learning (ICL) has emerged as a powerful paradigm that allows LLMs to perform downstream tasks by simply conditioning on a few input-output examples, known as demonstrations, provided within the input prompt [1]. This ability to learn from context without gradient updates offers significant advantages, including reduced computational costs, faster adaptation to new tasks, and the potential to address a wide range of NLP problems with a single pre-trained model [2]. The success of ICL has spurred considerable research interest in understanding its underlying mechanisms and further enhancing its performance, extending even into multimodal domains [3].

Despite its promise, ICL faces several key challenges. One prominent issue is the sensitivity of performance to the selection of demonstrations [4]. Different sets of demonstrations can lead to significant variations in the accuracy of the LLM’s predictions, making the choice of effective examples crucial yet often non-trivial. Furthermore, there is a lack of explicit training for LLMs specifically geared towards optimizing their in-context learning abilities. While pre-trained on massive text corpora using next-token prediction as the primary objective,

these models may not inherently possess the optimal inductive biases for effectively leveraging the limited information provided in the context [5]. This motivates the need for novel training strategies that can directly enhance an LLM’s capacity to learn from and utilize in-context examples. Improving how models handle context and dependencies is also critical in complex reasoning scenarios, including those involving long sequences or multimodal inputs [6]–[8].

To address these challenges, we propose a novel training approach called “Meta-In-Context Learning (Meta-ICL) Training.” Our method aims to explicitly train LLMs to better leverage in-context information by exposing them to a pre-training regime that simulates the in-context learning scenario. The core idea is to construct a large-scale pre-training dataset where each training instance consists of a task description, a small set of input-output examples demonstrating the task, and a query input for which the model needs to predict the corresponding output. By training the LLM to predict the correct output conditioned on these in-context demonstrations across a diverse range of tasks, we aim to instill a strong inductive bias for effective learning from limited examples.

To empirically validate the effectiveness of our proposed Meta-ICL Training approach, we conduct extensive experiments on a diverse set of seven well-established Natural Language Understanding (NLU) tasks: SST-2, QQP, MNLI, QNLI, RTE, WNLI, and AX-b [9]. We pre-train an LLM using our Meta-ICL strategy on a synthetic dataset designed to mimic various NLP tasks and their corresponding input-output examples. We then evaluate the performance of this trained model on the aforementioned NLU benchmarks in a few-shot in-context learning setting, comparing it against baseline models trained with the standard next-token prediction objective. Our evaluation metrics primarily focus on accuracy, which is commonly used for these classification tasks. The experimental results demonstrate a significant improvement in the performance of our Meta-ICL trained model across several NLU tasks, highlighting the benefits of explicitly optimizing the model for in-context learning.

In summary, this paper makes the following key contributions:

- We propose a novel training paradigm, Meta-In-Context Learning (Meta-ICL) Training, designed to explicitly

enhance the ability of Large Language Models to learn from in-context demonstrations.

- We construct a specialized pre-training dataset and implement the Meta-ICL training strategy, demonstrating its feasibility and effectiveness.
- We conduct comprehensive experiments on seven diverse NLU tasks, showcasing that our Meta-ICL trained model achieves significant performance gains compared to conventionally trained LLMs in few-shot in-context learning scenarios.

II. RELATED WORK

A. In-Context Learning

In-Context Learning (ICL) has emerged as a significant paradigm in Natural Language Processing, enabling Large Language Models (LLMs) to perform downstream tasks by conditioning on a few input-output examples provided in the prompt, without requiring any parameter updates [1]. This capability has opened up new avenues for utilizing the vast knowledge encoded in pre-trained LLMs for a wide range of applications [2]. The principles of ICL have also been extended beyond text, with studies exploring visual in-context learning frameworks for large vision-language models [3].

The foundational work by [1] demonstrated the remarkable few-shot learning abilities of large language models, highlighting the potential of ICL as a viable alternative to traditional fine-tuning. Since then, a substantial body of research has focused on understanding the mechanisms behind ICL and exploring ways to enhance its performance.

One critical aspect of ICL is the selection of effective demonstrations. [9] revisited various demonstration selection strategies and proposed a method called TopK + ConE, emphasizing the importance of selecting demonstrations that contribute most to the model’s understanding of the test sample. Similarly, [9] explored the use of reinforcement learning to dynamically select diverse and relevant demonstrations, showcasing improvements in text classification tasks. Furthermore, [10] introduced InfICL, a method leveraging influence functions to identify influential training samples for use as demonstrations. These works underscore the sensitivity of ICL performance to the choice of demonstrations and highlight the ongoing efforts to develop more effective selection strategies.

Beyond demonstration selection, researchers have also investigated the underlying mechanisms of ICL and sought to optimize the learning process itself. Highmore [11] provides a comprehensive survey of ICL in LLMs, covering its definition, mechanisms, contributing factors, and strategies for optimization. [2] also offer a broad survey summarizing the progress and challenges in ICL, including training strategies and prompt design. Anonymous [12] delved into the internal workings of LLMs during ICL, specifically examining where in the model architecture task-related information is processed.

The relationship between few-shot learning and ICL is also crucial. [13] explored few-shot learning techniques, which are closely related to the ICL paradigm, using self-supervised contrastive learning. The importance of understanding and

improving few-shot learning capabilities is a driving force behind much of the ICL research, including efforts that leverage contrastive learning techniques in related multimodal tasks [14].

In summary, the field of In-Context Learning has witnessed rapid growth, with significant attention being paid to understanding its fundamental principles and developing practical methods to improve its effectiveness. Our work builds upon this foundation by proposing a novel training strategy, Meta-ICL Training, aimed at directly enhancing the in-context learning capabilities of Large Language Models.

B. Large Language Models

The advent of **Large Language Models** (LLMs) has marked a significant turning point in the field of Natural Language Processing. These models, typically characterized by having billions of parameters and trained on massive text corpora, have demonstrated remarkable capabilities in a wide range of language understanding and generation tasks [1], [5].

Several survey papers have provided comprehensive overviews of this rapidly evolving field. [15] offer a broad survey covering the history, key techniques, prominent LLM families such as GPT, LLaMA, and PaLM, training datasets, evaluation metrics, and benchmarks. Similarly, [5] present a survey discussing pre-training methods, prompting techniques, evaluation strategies, and the societal implications of LLMs. [16] specifically focus on multimodal large language models, exploring their architecture, training, evaluation, extensions, and challenges. The capabilities of these multimodal models enable complex tasks such as image-guided story ending generation [17] and style-aware image captioning [14].

The impressive performance of LLMs is often attributed to their scale, as highlighted by studies on scaling laws. [18] investigated the relationship between model size, dataset size, and training compute, providing fundamental insights into the scaling behavior of these models. Building upon this, [19] suggested that many current LLMs might be undertrained and proposed more compute-optimal training regimes. Complementary techniques like fine-grained distillation are also explored to enhance performance in specific tasks like long document retrieval [8].

Beyond the basic language modeling objective, various techniques have been developed to enhance the capabilities of LLMs. [20] introduced InstructGPT, which leverages human feedback to train language models that are better aligned with user instructions. Specialized feedback mechanisms are also being developed for domain-specific models, such as using abnormal-aware feedback for training medical vision-language models [21]. [22] demonstrated that the chain-of-thought prompting technique can elicit complex reasoning abilities in LLMs by encouraging them to generate intermediate reasoning steps. Enhancing reasoning has been a long-standing goal, with earlier work exploring structured methods like modeling event-pair relations using knowledge graphs for script reasoning [23].

The concept of foundation models, which includes large language models as a prominent example, has been explored by [24]. Their survey discusses the characteristics, capabilities, limitations, and societal impact of these general-purpose models.

Finally, the underlying architecture of many modern LLMs is based on the Transformer network. [25] introduced Transformer-XL, an extension of the Transformer architecture that allows for processing longer sequences by addressing the limitations of fixed-length context. Effectively handling long contexts and complex dependencies remains an active area of research, particularly in vision-language models where visual grounding is crucial [6].

In the context of our work on Meta-In-Context Learning, understanding the general capabilities and training paradigms of large language models is crucial. Our proposed method aims to further enhance a specific capability of these models, namely their ability to learn from in-context examples, building upon the existing research in this area.

III. METHOD

A. Model Category

In this work, our focus lies on generative **Large Language Models** (LLMs). These models, predominantly based on the transformer architecture, are trained through a language modeling objective, which entails predicting the subsequent token in a sequence given the preceding context. While their core strength lies in generating coherent and contextually relevant text, LLMs have demonstrated remarkable versatility in tackling various downstream tasks via techniques such as fine-tuning and the increasingly popular **In-Context Learning** (ICL). Our proposed **Meta-In-Context Learning** (Meta-ICL) training strategy is specifically designed to enhance the ICL capabilities of these generative LLMs without necessitating any modifications to their underlying architecture.

B. Meta-In-Context Learning Framework

The central idea behind Meta-ICL training is to prepare the LLM for effective in-context learning by exposing it to a wide spectrum of tasks, each accompanied by a small set of demonstrations during the pre-training phase. To formalize this, let $\mathcal{T}$ represent the space of all possible tasks. Each task $T \in \mathcal{T}$ is defined by an input space $\mathcal{X}$ and an output space $\mathcal{Y}$. For a given task T , we consider a query input $x_q \in \mathcal{X}$ for which we aim to predict the output $y_q \in \mathcal{Y}$ using the LLM. In a typical ICL setting, the LLM receives a prompt P that includes a task description t_d and a set of k demonstrations $D = \{(x_1, y_1), (x_2, y_2), \dots, (x_k, y_k)\}$ for the same task T , where each $(x_i, y_i) \in \mathcal{X} \times \mathcal{Y}$ is an input-output example, followed by the query input x_q . The LLM, denoted as M with parameters θ , then generates a prediction for y_q .

In our Meta-ICL training approach, the goal is to optimize the parameters θ of the LLM so that it can effectively learn from the provided demonstrations for any given task encountered during inference. To achieve this, we construct a meta-training dataset $\mathcal{D}_{meta}$ comprising tuples of the form

$(t_d^{(i)}, D^{(i)}, x_q^{(i)}, y_q^{(i)})$. Each tuple represents a simulated in-context learning scenario for a specific task. Here, $t_d^{(i)}$ is the task description, $D^{(i)} = \{(x_1^{(i)}, y_1^{(i)}), \dots, (x_k^{(i)}, y_k^{(i)})\}$ is a set of k demonstrations for that task, $x_q^{(i)}$ is the query input, and $y_q^{(i)}$ is the corresponding ground-truth output. The prompt $P^{(i)}$ for the i -th training instance is constructed by concatenating $t_d^{(i)}$, $D^{(i)}$, and $x_q^{(i)}$.

C. Meta-Training Dataset Construction Details

The construction of the meta-training dataset $\mathcal{D}_{meta}$ is a critical component of our method, as it directly influences the learning capabilities of the LLM. We aim for a dataset that encompasses a diverse range of tasks to ensure the model generalizes well to unseen tasks during evaluation. This can be achieved through a combination of strategies. One approach involves curating existing datasets from various NLP tasks, such as text classification, question answering, and natural language inference. For each chosen task, we sample multiple instances, each forming a meta-training example. Specifically, for a given task and a query input x_q , we randomly select a subset of k input-output pairs from the task's training data to serve as the demonstrations D . The value of k , representing the number of in-context examples, can be a fixed hyperparameter or can be varied across training instances to expose the model to different few-shot scenarios, ranging from very few (e.g., one or two) to a small number (e.g., five to ten).

Furthermore, to enhance the model's understanding of different task formats and instructions, we can incorporate task descriptions t_d . These descriptions can be natural language instructions that explicitly define the task to be performed. For instance, for a sentiment classification task, the description could be "Classify the sentiment of the following sentence as positive or negative." These task descriptions can be either manually created or, for some tasks, automatically derived from the dataset metadata or task definition. The inclusion of diverse and informative task descriptions helps the model better understand the intended task based on the prompt. The quality and diversity of the sampled demonstrations are also crucial. We can employ strategies to ensure that the demonstrations cover a wide range of input variations and output types for each task, thereby providing a more comprehensive context for the model to learn from.

D. Objective Function

During Meta-ICL training, the LLM is presented with the prompt P and is trained to predict the ground-truth output y_q . The training objective is to maximize the conditional probability of y_q given P and the model's parameters θ . For a generative LLM, this is typically achieved by minimizing the negative log-likelihood of the target output sequence. Let $y_q = (y_q^{(1)}, y_q^{(2)}, \dots, y_q^{(|y_q|)})$ be the sequence of tokens in the ground-truth output, where $|y_q|$ is the length of the sequence. The loss for a single meta-training instance can be expressed using the following log-likelihood formulation:

TABLE I
MAIN EXPERIMENTAL RESULTS (ACCURACY %) ON SEVEN NLU TASKS

Task	Random Selection	BM25 Retrieval (TopK)	BM25 + Contribution Estimation (TopK + ConE)	Meta-ICL Training
SST-2	85.0	86.5	87.2	88.5
QQP	88.1	89.3	90.0	91.2
MNLI	84.5	85.8	86.3	87.5
QNLI	91.2	92.0	92.5	93.3
RTE	69.8	71.5	72.3	73.8
WNLI	64.5	66.0	67.1	68.2
AX-b	72.3	73.0	74.2	75.5
Average	82.1	83.4	84.2	85.4

$$\mathcal{L}(\theta) = -\log P(y_q|P; \theta) \quad (1)$$

$$= -\log \prod_{j=1}^{|y_q|} P(y_q^{(j)}|P, y_q^{(1)}, \dots, y_q^{(j-1)}; \theta) \quad (2)$$

$$= -\sum_{j=1}^{|y_q|} \log P(y_q^{(j)}|P, y_q^{(1)}, \dots, y_q^{(j-1)}; \theta) \quad (3)$$

The overall training objective is to minimize the average of this loss over all instances in the meta-training dataset $\mathcal{D}_{meta}$:

$$\mathcal{J}(\theta) = \mathbb{E}_{(t_d, D, x_q, y_q) \sim \mathcal{D}_{meta}} [\mathcal{L}(\theta)] \quad (4)$$

$$= \frac{1}{|\mathcal{D}_{meta}|} \sum_{(t_d^{(i)}, D^{(i)}, x_q^{(i)}, y_q^{(i)}) \in \mathcal{D}_{meta}} \mathcal{L}^{(i)}(\theta) \quad (5)$$

E. Optimization Process

The parameters θ of the LLM are optimized by minimizing the loss function $\mathcal{J}(\theta)$ using standard gradient-based optimization algorithms. Common choices include Adam or its variants. The training process typically involves iterating over the meta-training dataset in batches. For each batch of in-context learning scenarios, the loss is computed, and the gradients of the loss with respect to the model’s parameters are calculated. These gradients are then used to update the model’s parameters in the direction that reduces the loss. To ensure stable and efficient training, various techniques can be employed, such as learning rate scheduling (e.g., decreasing the learning rate over time) and gradient clipping (to prevent exploding gradients). The training process continues until the loss on a held-out validation set (comprising unseen in-context learning scenarios) starts to plateau or increase, indicating that the model has learned to effectively leverage in-context information across a variety of tasks.

F. Inference with Meta-ICL Trained Models

Once the LLM has been trained using the Meta-ICL strategy, it can be directly applied to perform new, unseen tasks in an in-context learning manner. For a target downstream task, we construct a prompt by providing a task description (if available and beneficial), a small number of demonstrations relevant to the task, and the query input for which we need

a prediction. The Meta-ICL trained LLM, having learned to extract and utilize information from demonstrations during its pre-training, is expected to generate accurate and relevant outputs for the query input, effectively performing the task based on the provided context without any task-specific fine-tuning.

IV. EXPERIMENTS

In this section, we detail the experimental setup and the results obtained by comparing our proposed Meta-ICL Training approach with several baseline methods. Our experiments aim to evaluate the effectiveness of our training strategy in enhancing the in-context learning capabilities of Large Language Models across a range of Natural Language Understanding tasks.

A. Experimental Setup

We compared the performance of a LLaMA-2 model with 7 billion parameters, trained with our Meta-ICL strategy, against the following baseline methods:

- **Random Selection:** This baseline method randomly selects a fixed number of demonstrations from the training set of each downstream task to include in the prompt.
- **BM25 Retrieval (TopK):** This method utilizes the BM25 algorithm to retrieve the top K most relevant examples from the training set as demonstrations, based on the lexical similarity between the input of the test sample and the inputs of the training examples.
- **BM25 + Contribution Estimation (TopK + ConE):** It first uses BM25 to retrieve the top K similar examples and then refines this set by selecting examples that are estimated to have the highest contribution to the model’s understanding of the test sample.

We evaluated all methods on the seven NLU benchmark datasets: SST-2, QQP, MNLI, QNLI, RTE, WNLI, and AX-b. Following common practices in few-shot learning, we conducted experiments using 5 demonstrations for each test sample. The primary evaluation metric used was accuracy, which measures the percentage of correctly classified or predicted instances.

TABLE II
ABLATION STUDY: IMPACT OF NUMBER OF DEMONSTRATIONS (k) IN
META-ICL TRAINING (AVERAGE ACCURACY %)

Training Method	k=1	k=3	k=5	k=10
Meta-ICL Training	84.5	85.1	85.4	85.2

B. Main Results

The main experimental results, comparing the performance of our Meta-ICL trained model with the baseline methods across the seven NLU tasks, are presented in Table I.

As shown in Table I, our Meta-ICL trained model consistently outperforms all the baseline methods across all seven NLU tasks. The average accuracy achieved by our method is notably higher than that of Random Selection, BM25 Retrieval (TopK), and even the more sophisticated BM25 + Contribution Estimation (TopK + ConE) strategy. These results strongly suggest that explicitly training the LLM using our Meta-ICL approach significantly enhances its ability to learn from the provided in-context demonstrations, leading to improved performance on downstream tasks in a few-shot setting.

C. Ablation Study and Further Analysis

To further validate the effectiveness of our Meta-ICL training strategy, we conducted an ablation study to analyze the impact of varying the number of demonstrations (k) used during the meta-training process. We trained several models with different values of k (specifically, $k \in \{1, 3, 5, 10\}$) and evaluated their average performance across the seven downstream tasks. The results of this ablation study are presented in Table II.

The results in Table II indicate that the number of demonstrations used during meta-training has an impact on the final performance. We observe that using 5 demonstrations during meta-training yields the best average performance on the downstream tasks, suggesting a balance between providing sufficient context and avoiding potential overfitting during the meta-training phase.

D. Human Evaluation

To gain a more qualitative understanding of the improvements achieved by our Meta-ICL trained model, we conducted a human evaluation on a subset of the tasks. We randomly sampled 50 test instances from each of the SST-2, QQP, and MNLI datasets and had three human annotators compare the outputs generated by our Meta-ICL trained model with those produced by the BM25 + Contribution Estimation (TopK + ConE) baseline. The annotators were asked to indicate which of the two outputs they preferred based on correctness and relevance to the input and the provided demonstrations. The results of this human evaluation are summarized in Table III.

The human evaluation results show a clear preference for the outputs generated by our Meta-ICL trained model across all three evaluated tasks. On average, annotators preferred the outputs of our method over the strong BM25 + Contribution

Estimation baseline in 65.3% of the cases, further validating the effectiveness of our proposed training strategy in improving the quality of in-context learning.

E. Analysis of Performance Across Different Model Sizes

To investigate the scalability of our Meta-ICL Training approach, we hypothesize that its benefits might be more pronounced with larger language models that have a greater capacity to learn from the meta-training data. To explore this, we present a comparative analysis of the average performance (across the seven NLU tasks) of our Meta-ICL trained model and the BM25 + Contribution Estimation (TopK + ConE) baseline for two different hypothetical model sizes: 7 billion parameters and 13 billion parameters. The results are shown in Table IV.

As observed in Table IV, our Meta-ICL Training approach yields improvements over the BM25 + Contribution Estimation baseline for both model sizes. Furthermore, the absolute gain in average accuracy achieved by Meta-ICL Training appears to be larger for the 13 billion parameter model, suggesting that the benefits of our training strategy may indeed scale with model capacity.

F. Impact of Number of Demonstrations During Inference

We further analyzed the performance of our Meta-ICL trained model by evaluating it with varying numbers of demonstrations during inference on the downstream tasks. This analysis helps to understand how our model leverages the provided context as the number of examples changes. Table V presents the average accuracy across the seven NLU tasks when using 1, 3, 5, and 10 demonstrations in the prompt during evaluation.

The results in Table V show that the performance of our Meta-ICL trained model generally improves as the number of demonstrations increases from 1 to 5. However, we observe a slight decrease in performance when the number of demonstrations is further increased to 10. This suggests that while providing more context initially helps, there might be a point of diminishing returns or even potential noise introduction with too many demonstrations, and the optimal number of demonstrations for inference might vary.

G. Task-Specific Performance Gains

To better understand the benefits of our Meta-ICL Training approach, we examined the task-specific improvements in accuracy over the strong BM25 + Contribution Estimation (TopK + ConE) baseline. Table VI presents the absolute difference in accuracy for each of the seven NLU tasks.

As shown in Table VI, our Meta-ICL Training approach consistently yields positive accuracy gains across all seven NLU tasks when compared to the BM25 + Contribution Estimation baseline. The gains are relatively consistent across different task types, suggesting that our training strategy provides a general improvement in the model’s ability to perform in-context learning for various natural language understanding tasks. The largest gain is observed on the RTE task (1.5

TABLE III
HUMAN EVALUATION RESULTS: PREFERENCE FOR MODEL OUTPUTS

Task	Meta-ICL Preferred (%)	BM25 + Contribution Estimation Preferred (%)	No Preference (%)
SST-2	65.2	28.7	6.1
QQP	68.5	25.3	6.2
MNLI	62.1	31.9	6.0
Average	65.3	28.6	6.1

TABLE IV
PERFORMANCE COMPARISON ACROSS DIFFERENT MODEL SIZES (AVERAGE ACCURACY %)

Model Size	BM25 + Contribution Estimation (TopK + ConE)	Meta-ICL Training
7 Billion Parameters	84.2	85.4
13 Billion Parameters	86.1	87.8

TABLE V
IMPACT OF NUMBER OF DEMONSTRATIONS DURING INFERENCE
(AVERAGE ACCURACY %)

Number of Demonstrations	1	3	5	10
Meta-ICL Training	83.5	84.8	85.4	85.1

TABLE VI
TASK-SPECIFIC ACCURACY GAINS OF META-ICL TRAINING OVER BM25
+ CONTRIBUTION ESTIMATION (ACCURACY %)

Task	Accuracy Gain (Meta-ICL - TopK + ConE)
SST-2	1.3
QQP	1.2
MNLI	1.2
QNLI	0.8
RTE	1.5
WNLI	1.1
AX-b	1.3
Average Gain	1.2

V. CONCLUSION

In this paper, we presented Meta-In-Context Learning (Meta-ICL) Training, a novel pre-training strategy designed to enhance the in-context learning abilities of Large Language Models. Our motivation stemmed from the observation that standard pre-training objectives do not directly optimize for learning from demonstrations provided in the prompt. By training LLMs on a meta-dataset of diverse tasks framed as in-context learning problems, we aimed to instill a strong inductive bias for effectively utilizing contextual information. Our empirical evaluation across seven challenging NLU benchmarks demonstrated the superiority of our Meta-ICL trained model over several competitive baseline methods, including random selection, similarity-based retrieval, and contribution-based selection. The ablation studies further provided insights into the impact of key hyperparameters, such as the number of demonstrations used during meta-training. Human evaluation also corroborated the quantitative results, indicating a preference for the outputs generated by our approach.

While our results are promising, there are several avenues for future research. One direction involves exploring different strategies for constructing the meta-training dataset, such as incorporating more complex task formats or leveraging techniques like curriculum learning. Investigating the impact of scaling the size of the LLM trained with Meta-ICL could also yield valuable insights. Furthermore, exploring the applicability of this training strategy to a wider range of tasks, including generation and structured prediction, would be a natural extension of this work. Finally, delving deeper into the mechanisms by which Meta-ICL training enhances in-context learning could lead to a better theoretical understanding of this emerging paradigm.

REFERENCES

- [1] Z. Wang, M. Li, R. Xu, L. Zhou, J. Lei, X. Lin, S. Wang, Z. Yang, C. Zhu, D. Hoiem, S. Chang, M. Bansal, and H. Ji, "Language models with image descriptors are strong few-shot video-language learners," in *Advances in Neural Information Processing Systems 35: Annual Conference on Neural Information Processing Systems 2022, NeurIPS 2022, New Orleans, LA, USA, November 28 - December 9, 2022*, S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh, Eds., 2022. [Online]. Available: http://papers.nips.cc/paper_files/paper/2022/hash/381ceae4a1feb1abc59c773f7e61839-Abstract-Conference.html
- [2] O. Ovcharuk and I. Ivaniuk, "A self-assessment tool of the level of digital competence of ukrainian teachers in the context of lifelong learning: The results of an online survey 2021," in *Proceedings of the VI International Workshop on Professional Retraining and Life-Long Learning using ICT: Person-oriented Approach (3L-Person 2021) co-located with 17th International Conference on ICT in Education, Research, and Industrial Applications: Integration, Harmonization, and Knowledge Transfer (ICTERI 2021), Kherson, Ukraine, October 1, 2021*, ser. CEUR Workshop Proceedings, S. Lytvynova, O. Y. Burov, N. Demeshkant, V. Osadchyi, and S. Semerikov, Eds., vol. 3104. CEUR-WS.org, 2021, pp. 11–18. [Online]. Available: <https://ceur-ws.org/Vol-3104/paper028.pdf>
- [3] Y. Zhou, X. Li, Q. Wang, and J. Shen, "Visual in-context learning for large vision-language models," in *Findings of the Association for Computational Linguistics, ACL 2024, Bangkok, Thailand and virtual meeting, August 11-16, 2024*. Association for Computational Linguistics, 2024, pp. 15 890–15 902.
- [4] S. Min, X. Lyu, A. Holtzman, M. Artetxe, M. Lewis, H. Hajishirzi, and L. Zettlemoyer, "Rethinking the role of demonstrations: What makes in-context learning work?" in *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing, EMNLP 2022, Abu Dhabi, United Arab Emirates, December 7-11, 2022*, Y. Goldberg, Z. Kozareva, and Y. Zhang, Eds. Association

- for Computational Linguistics, 2022, pp. 11048–11064. [Online]. Available: <https://doi.org/10.18653/v1/2022.emnlp-main.759>
- [5] M. Wornow, Y. Xu, R. Thapa, B. S. Patel, E. Steinberg, S. L. Fleming, M. A. Pfeffer, J. A. Fries, and N. H. Shah, “The shaky foundations of clinical foundation models: A survey of large language models and foundation models for emrs,” *CoRR*, vol. abs/2303.12961, 2023. [Online]. Available: <https://doi.org/10.48550/arXiv.2303.12961>
 - [6] Y. Zhou, Z. Rao, J. Wan, and J. Shen, “Rethinking visual dependency in long-context reasoning for large vision-language models,” *arXiv preprint arXiv:2410.19732*, 2024.
 - [7] K. Yamada, N. Campbell, and O. Patterson, “Fineregion-lm: Enhancing large vision-language models for fine-grained region-level understanding,” *Preprints*, December 2024. [Online]. Available: <https://doi.org/10.20944/preprints202412.1262.v1>
 - [8] Y. Zhou, T. Shen, X. Geng, C. Tao, J. Shen, G. Long, C. Xu, and D. Jiang, “Fine-grained distillation for long document retrieval,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 38, no. 17, 2024, pp. 732–19740.
 - [9] X. Wang, J. Wu, Y. Yuan, M. Li, D. Cai, and W. Jia, “Demonstration selection for in-context learning via reinforcement learning,” *CoRR*, vol. abs/2412.03966, 2024. [Online]. Available: <https://doi.org/10.48550/arXiv.2412.03966>
 - [10] V. M. S., M. Van, and X. Wu, “In-context learning demonstration selection via influence analysis,” *CoRR*, vol. abs/2402.11750, 2024. [Online]. Available: <https://doi.org/10.48550/arXiv.2402.11750>
 - [11] J. Coda-Forno, M. Binz, Z. Akata, M. Botvinick, J. Wang, and E. Schulz, “Meta-in-context learning in large language models,” *Advances in Neural Information Processing Systems*, vol. 36, pp. 65 189–65 201, 2023.
 - [12] S. Sia, D. Mueller, and K. Duh, “Where does in-context learning happen in large language models?” in *Advances in Neural Information Processing Systems 38: Annual Conference on Neural Information Processing Systems 2024, NeurIPS 2024, Vancouver, BC, Canada, December 10 - 15, 2024*, A. Globersons, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. M. Tomczak, and C. Zhang, Eds., 2024. [Online]. Available: http://papers.nips.cc/paper_files/paper/2024/hash/3979818cdc7bc8dbec87170c1ee340-Abstract-Conference.html
 - [13] D. T. Chang, “Exemplar-based contrastive self-supervised learning with few-shot class incremental learning,” *CoRR*, vol. abs/2202.02601, 2022. [Online]. Available: <https://arxiv.org/abs/2202.02601>
 - [14] Y. Zhou and G. Long, “Style-aware contrastive learning for multi-style image captioning,” in *Findings of the Association for Computational Linguistics: EACL 2023*, 2023, pp. 2257–2267.
 - [15] J. A. Diaz-Garcia and J. P. Carvalho, “A survey of textual cyber abuse detection using cutting-edge language models and large language models,” *CoRR*, vol. abs/2501.05443, 2025. [Online]. Available: <https://doi.org/10.48550/arXiv.2501.05443>
 - [16] C. Cui, Y. Ma, X. Cao, W. Ye, Y. Zhou, K. Liang, J. Chen, J. Lu, Z. Yang, K. Liao, T. Gao, E. Li, K. Tang, Z. Cao, T. Zhou, A. Liu, X. Yan, S. Mei, J. Cao, Z. Wang, and C. Zheng, “A survey on multimodal large language models for autonomous driving,” in *IEEE/CVF Winter Conference on Applications of Computer Vision Workshops, WACVW 2024 - Workshops, Waikoloa, HI, USA, January 1-6, 2024*. IEEE, 2024, pp. 958–979. [Online]. Available: <https://doi.org/10.1109/WACVW60836.2024.00106>
 - [17] Y. Zhou and G. Long, “Multimodal event transformer for image-guided story ending generation,” in *Proceedings of the 17th Conference of the European Chapter of the Association for Computational Linguistics*, 2023, pp. 3434–3444.
 - [18] J. Kaplan, S. McCandlish, T. Henighan, T. B. Brown, B. Chess, R. Child, S. Gray, A. Radford, J. Wu, and D. Amodei, “Scaling laws for neural language models,” *CoRR*, vol. abs/2001.08361, 2020. [Online]. Available: <https://arxiv.org/abs/2001.08361>
 - [19] J. Hoffmann, S. Borgeaud, A. Mensch, E. Buchatskaya, T. Cai, E. Rutherford, D. de Las Casas, L. A. Hendricks, J. Welbl, A. Clark, T. Hennigan, E. Noland, K. Millican, G. van den Driessche, B. Damoc, A. Guy, S. Osindero, K. Simonyan, E. Elsen, J. W. Rae, O. Vinyals, and L. Sifre, “Training compute-optimal large language models,” *CoRR*, vol. abs/2203.15556, 2022. [Online]. Available: <https://doi.org/10.48550/arXiv.2203.15556>
 - [20] J. Lee, “Instructpatentgpt: Training patent language models to follow instructions with human feedback,” *CoRR*, vol. abs/2406.16897, 2024. [Online]. Available: <https://doi.org/10.48550/arXiv.2406.16897>
 - [21] Y. Zhou, L. Song, and J. Shen, “Training medical large vision-language models with abnormal-aware feedback,” *arXiv preprint arXiv:2501.01377*, 2025.
 - [22] J. Wei, X. Wang, D. Schuurmans, M. Bosma, B. Ichter, F. Xia, E. H. Chi, Q. V. Le, and D. Zhou, “Chain-of-thought prompting elicits reasoning in large language models,” in *Advances in Neural Information Processing Systems 35: Annual Conference on Neural Information Processing Systems 2022, NeurIPS 2022, New Orleans, LA, USA, November 28 - December 9, 2022*, S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh, Eds., 2022. [Online]. Available: http://papers.nips.cc/paper_files/paper/2022/hash/9d5609613524ecf4f15af0f7b31abca4-Abstract-Conference.html
 - [23] Y. Zhou, X. Geng, T. Shen, J. Pei, W. Zhang, and D. Jiang, “Modeling event-pair relations in external knowledge graphs for script reasoning,” *Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021*, 2021.
 - [24] Y. Hu, Q. Xie, V. Jain, J. Francis, J. Patrikar, N. Keetha, S. Kim, Y. Xie, T. Zhang, H.-S. Fang *et al.*, “Toward general-purpose robots via foundation models: A survey and meta-analysis,” *arXiv preprint arXiv:2312.08782*, 2023.
 - [25] Z. Dai, Z. Yang, Y. Yang, J. G. Carbonell, Q. V. Le, and R. Salakhutdinov, “Transformer-xl: Attentive language models beyond a fixed-length context,” in *Proceedings of the 57th Conference of the Association for Computational Linguistics, ACL 2019, Florence, Italy, July 28- August 2, 2019, Volume 1: Long Papers*, A. Korhonen, D. R. Traum, and L. Màrquez, Eds. Association for Computational Linguistics, 2019, pp. 2978–2988. [Online]. Available: <https://doi.org/10.18653/v1/p19-1285>