

Distributed In-Context Learning under Non-IID Among Clients

Anonymous ACL submission

Abstract

Advancements in large language models (LLMs) have shown their effectiveness in multiple complicated natural language reasoning tasks. A key challenge remains in adapting these models efficiently to new or unfamiliar tasks. In-context learning (ICL) provides a promising solution for few-shot adaptation by retrieving a set of data points relevant to a query, called in-context examples (ICE), from a training dataset and providing them during the inference as context. Most existing studies utilize a centralized training dataset, yet many real-world datasets may be distributed among multiple clients, and remote data retrieval can be associated with costs. Especially when the client data are non-identical independent distributions (non-IID), retrieving from clients a proper set of ICEs needed for a test query presents critical challenges. In this paper, we first show that in this challenging setting, test queries will have different preferences among clients because of non-IIDness, and equal contribution often leads to suboptimal performance. We then introduce a novel approach to tackle the distributed non-IID ICL problem when a data usage budget is present. The principle is that each client’s proper contribution (budget) should be designed according to the preference of each query for that client. Our approach uses a data-driven manner to allocate a budget for each client, tailored to each test query. Through extensive empirical studies on diverse datasets, our method demonstrates superior performance relative to competing baselines.

1 Introduction

Recent significant progress in large language models (LLMs) (Achiam et al., 2023; Touvron et al., 2023a,b; Team et al., 2023) has demonstrated their effectiveness across various natural language processing (NLP) tasks (Wang et al., 2018, 2019). Despite their impressive performances, they still require adaptation to the specific downstream tasks

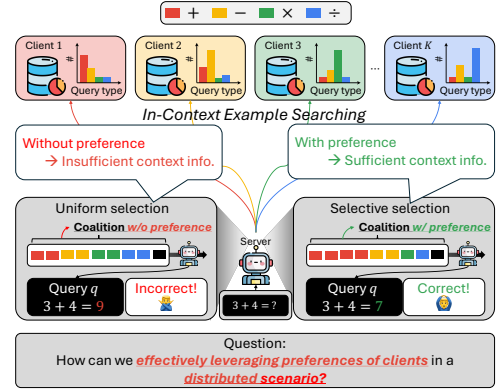


Figure 1: Problem overview. When datasets are distributed among clients in a non-IID manner, it creates an obstacle in generating a good context (left). However, by assigning appropriate budgets to leverage per-client expertise, better context can be created (right).

for better performance. However, adaptation poses challenges due to LLMs’ vast number of trainable parameters.

In-context learning (ICL) (Dong et al., 2022) is a notable method that distinguishes itself through both its effectiveness and efficiency. In brief, ICL adapts to the target task by incorporating context information following two primary steps: i) identify samples from the training dataset helpful to solve the target query by creating a prompt describing a context; ii) feed the constructed prompt with the target query and get the answer. Previous related works on ICL mainly have focused on the construction of a prompt describing the context, which involves several sub-problems, such as the retrieval of in-context examples (ICEs) (Robertson et al., 2009) and determining the optimal sequence for the selected ICEs (Zhang et al., 2024).

A common assumption in most existing ICL research is that the system has access to a high-quality centralized dataset used for retrieval. However, in many application scenarios, such as health informatics, centralized datasets may not be fea-

sible, and data could be distributed in different institutions, which calls for the distributed ICL. In addition, when the data is proprietary and possesses high value towards inferences, access to data entries may also be bound to data pricing strategies (Xu et al., 2023; Cong et al., 2022). For instance, the system needs to pay the local institution based on the number of samples sent to the system to share profits from inferences (Tang et al., 2020). Under this scenario, aggregating ICEs from local clients to a center server for ICL entails significant financial costs and lacks efficiency.

In this paper, we focus on integrating knowledge from distributed clients to achieve better ICL performance under the per-query ICE budget constraint. Specifically, we formalize the distributed ICL problem where the ICEs are distributed on clients, and the server has an LLM for ICL inference but can only request a limited number of ICEs from all clients for each query, which we refer to as the *ICE budget*.

We begin by identifying the key challenge in distributed ICL with ICE budget constraints lies in the non-independently and identically distributed (non-IID) training data, as shown in Section 3.1. For example, in Figure 1, data samples are spread across C clients, each with a unique data distribution. Specifically, client 1 primarily contains (+) samples, while client 2 is mainly constituted by (−) examples. Only limited research (Mohtashami et al., 2023) tried to address the challenge of distributed datasets for ICL, while none considers the challenging real-world setting of non-IID clients. This leaves a critical question unanswered: *What happens to distributed ICL when local clients are non-IID?*

To further the understanding of the key challenge in the distributed non-IID ICL, we explore the local retrieval process on non-IID clients. We found that each query has different preferences for different clients based on local knowledge distribution, that is, the number of samples needed from different clients should vary based on local sample distribution. As the toy example shown in Figure 1, when the server creates context by uniformly assigning budgets to clients, the answer might be incorrect due to the insufficiency of (+) information in the context. To be more detailed, the server assigns the clients who have expertise on (−), (×), and (÷) operations with the same budget as on (+), without any preference. Nevertheless, if the server assigns more budget to clients with many (+) sam-

ples, such as client 1, it can create a more relevant context to answer the query related to (+) operation. This indicates that under non-IID, the server should allocate the budgets over clients based on the preference of each query itself, as well as the distribution of local training samples.

Motivated by this, we propose a novel distributed ICL framework to collaboratively collect scattered information among non-IID clients by properly assigning ICE budgets to each client. First, the server will gather the optimal budget statistics using an existing proxy dataset on the server side. Next, the server will use this dataset to train the budget allocator. During the deployment stage, the server will predict the proper budget for each client using this trained budget allocator given each test query and perform ICL among clients. Furthermore, in practical scenarios where privacy concerns arise, we augment our framework with the paraphrasing method (Mohtashami et al., 2023) to secure privacy.

Contributions. A summary of our contributions:

- To the best of our knowledge, we are the first to study the challenging real-world setting of ICL with distributed non-IID clients. We identify the principal challenge as properly assigning the ICE budget for non-IID clients based on the preference of each test query and local knowledge distribution.
- We propose a framework to handle the distributed non-IID ICL. This framework trains a budget allocator on the server with the help of a server-side proxy dataset. Then, the server will use this trained allocator to decide how many ICEs to retrieve from each client for the ICL process, enabling collaborative action among clients.
- Across a range of dataset benchmarks featuring various non-IID configurations as well as on different LLM architectures, our approach has been validated to enhance ICL performance. Notably, we examine both non-private, *i.e.*, communicate raw samples directly, and private cases using the paraphrasing method to secure privacy. In both scenarios, our approach shows superiority to the previous method and other reasonable baselines.

2 Problem Formulation

In this section, we provide a detailed problem formulation. First, we begin with the specifics of in-

context learning (ICL), followed by a description of distributed non-IID ICL.

2.1 In-Context Learning

Notation. We consider a NLP tasks which have training dataset $\mathcal{D} = \{(x_i, y_i)\}_{i=1}^N$ with N training samples. Here, x_i is the input text, and y_i is the corresponding output. In the test phase, a test query x_q is given.

Retrieval. We employ the off-the-shelf pre-trained retriever KATE (Liu et al., 2021)¹, which utilizes k -NN example selection. This retriever employs a sentence encoder $\mathcal{E}(\cdot)$ to measure the similarity between the in-context example x_i in dataset $\mathcal{D}$ and the query x_q as follows:

$$d(e_i, e_q) = \|e_q - e_i\|_2, \quad (1)$$

where $e_q = \mathcal{E}(x_q)$ and $e_i = \mathcal{E}(x_i)$. We select k samples using the following criterion:

$$\mathcal{T}(e_q, k|\mathcal{D}) = \arg \text{Top-}k(d(e_i, e_q)), \quad (2)$$

$e_i = \mathcal{E}(x_i) \forall (x_i, y_i) \in \mathcal{D}$

where $\mathcal{T}(e_q, k|\mathcal{D})$ denotes the selected samples from the dataset $\mathcal{D}$, and used for inference.

ICL Inference. In the test phase, given a test query with input x_i , relevant k training samples called in-context examples (ICEs) are selected, *i.e.*, $S = \mathcal{T}(e_q, k|\mathcal{D})$. Based on the retrieved samples, we feed the constructed context prompt $s(S, x_q)$ into LLM for inference and obtain results via:

$$y_t = \arg \max_y p_{\text{LLM}}(y|s(S, x_q), y_{<t}) \quad (3)$$

where $s(S, x_q) = (x_1, y_1) \odot \dots \odot (x_k, y_k) \odot x_q$,

where the $\odot$ operation denotes concatenation, and $s(S, x_q)$ is the context constructed using query x_q and samples in S ; the term p_{LLM} represents the output softmax probability of the LLM, functioning autoregressive, meaning that the output up to time t , *i.e.*, $y_{<t}$, is input back into the model to generate the t^{th} output, y_t . Previous works (Ye et al., 2023; Levy et al., 2022) on ICL mainly focus on the selection of S under a centralized setting. However, we investigate the scenario where $\mathcal{D}$ is split among several clients, each following non-IID distributions.

2.2 Distributed non-IID ICL

Distributed ICL Setting. We consider C clients with a centralized server in our system. Each

client $c \in [C]$ has local training dataset $\mathcal{D}_c = \{(x_i^c, y_i^c)\}_{i=1}^{N_c}$ with N_c training samples. Note that $\mathcal{D}_c$ follows different distributions for different clients. We follow the non-IID conditions as defined in (Li et al., 2022), with details provided in Appendix A. In summary, we allocate data on a per-class basis, where each client receives a specific number of classes, meaning each client has samples from only specified classes. Clients and the server have identical off-the-shelf pre-trained retrievers. Consider the computation resource limitation on clients as in many real scenarios (Yoo et al., 2022), only the server is equipped with an LLM. Moreover, the server has limited query-only proxy dataset $\mathcal{D}_{\text{proxy}} = \{x_j^{\text{proxy}}\}_{j=1}^{N_{\text{proxy}}}$, that $N_{\text{proxy}} \ll \sum_{c=1}^C N_c$. The server has quite a small $\mathcal{D}_{\text{proxy}}$, and it is an auxiliary dataset to extract information for collaboration to make the problem feasible. Notice that we do not require the ground truth label information of the proxy dataset, only input queries are sufficient, which reduces the difficulty on collecting such dataset in practical setting.

Pipeline. First, the server requests relevant samples from each client by sending x_q to all clients with local budgets k_c . Remark that each query x_q has its own preference of each client, which can be represented as k_c . A larger k_c indicates the given test query x_q prefers more information from client c , compared with client c' with a smaller $k_{c'}$. Here, x_q can be anonymized by paraphrasing, as done in previous works (Mohtashami et al., 2023)². Each client then selects the most relevant k_c samples from their local training dataset, *i.e.*, $S_c = \mathcal{T}(e_q, k_c|\mathcal{D}_c) \subset \mathcal{D}_c$, and returns them to the server. The server receives S_c from clients and generates the context s based on the merged examples, $S = \bigcup_{c=1}^C S_c$. In the final step, the server infers y using $s(S, x_q)$. The entire framework also can be described in Figure 2. In this paper, we are concentrating on assigning k_c to each client as described in Figure 2.

3 Observations

In this section, we describe several empirical supports to handle the distributed non-IID ICL. First, we demonstrate that non-IID distributions hinder the merging of scattered information. We then establish our goal, termed as *oracle budget*, which

¹We do not fine-tune the retriever for each task, which is impractical because we cannot gather the distributed datasets.

²Although our main experiments utilize the non-paraphrased dataset, we also present the paraphrased results in Section 5.

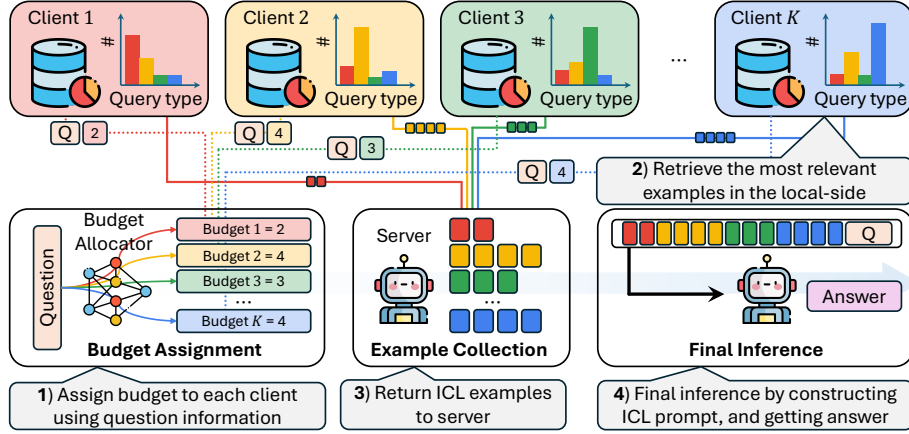


Figure 2: Overview of the pipeline: First, the **budget allocator** assigns a budget to each client based on the question. Subsequently, each client retrieves their relevant samples and sends them back to the server. The server infers the answer by feeding the question, which is composed of concatenated context examples and the query.

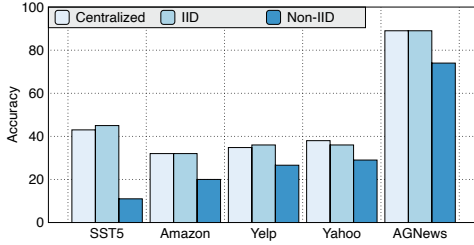


Figure 3: Non-IID experimental results. It shows that centralized performance is comparable to the IID case, whereas non-IID scenarios exhibit a significant declined performance. This highlights the critical importance of addressing non-IIDness to find a solution.

reflects the server’s preference for each client if the server knows all distributed data. Finally, we check if predicting the oracle budget of each test query for inference is feasible.

3.1 Non-IIDness Leads to Performance Drop

First of all, we evaluate the effect of non-IIDness. Straightforwardly, we distribute the budget $\{k_c\}_{c=1}^C$ uniformly according to the following criteria: Given C clients are involved in answering this question, and the number of samples for context is k . We first explore the naïve equally assigned local budget scheme in both IID and non-IID settings. That is, each client $c \in [C]$ locally retrieves top- k_c samples where $k_c = \lceil \frac{k}{C} \rceil$ from local dataset $\mathcal{D}_c$. Detailed experimental settings are described in Appendix B.

As illustrated in Figure 3, we observe the followings: (1) There is no significant performance degradation between the centralized case (■) and the IID case (■). This is expected, as the merged top- k_c samples in the IID case closely resemble the centralized top- k samples. Any minor discrepan-

cies are attributed to differences in sample ordering. (2) However, performance degradation becomes pronounced in non-IIDness case (refer to the comparison between ■, ■ and ■). Hereinafter, we gather insights to address the distributed non-IID ICL.

3.2 Proper Budget per Query for Each Client

Oracle budget. The remaining issue is that to make the server operate similarly in a centralized manner, it needs to allocate the budget as if it has complete knowledge of all clients. We call this budget for each client as the *oracle budget* for query embedding e_q and define it as follows:

$$k_c^*(e_q) = |\mathcal{T}(e_q, k | \mathcal{D}_c) \cap \mathcal{T}(e_q, k | \mathcal{D})|,$$

where $\mathcal{T}(\cdot)$ is defined as Eq. (2) and $|\cdot|$ is set cardinality. Note that the physical meaning of $k_c^*(e_q)$ is the number of shared samples between the top- k relevant to e_q in local $\mathcal{D}_c$ and global $\mathcal{D}$ datasets.

Check of predictability of oracle budget. For the next step, it is necessary to check if e_q has sufficient patterns of oracle budget to extract and use it in the inference phase. Our hypothesis is that similar queries may share similar oracle budget patterns and preferences on the same client, and it can lead to similar budget allocations for that client. Therefore, to verify this hypothesis, we perform t-SNE analysis (Van der Maaten and Hinton, 2008) on the embeddings obtained from the retriever for queries. Furthermore, we color each sample based on the oracle budget $k_c^*(e_q)$. As described in Figure 5 of Appendix F, similar query embeddings exhibit similar oracle budget patterns. This indicates that, given

Algorithm 1 Top- k sampling, $\mathcal{T}(e, k|\mathcal{D})$

Require: Query embedding e , Encoder $\mathcal{E}(\cdot)$

/* Compute embedding */

1: **for** $(x_i, y_i) \in \mathcal{D}$ **do**2: $e_i \leftarrow \mathcal{E}(x_i)$ 3: **end for**/* Select top- k samples */4: $S = \arg \text{Top-}k \|e - e_i\|_2$
 $(x_i, y_i) \in \mathcal{D}$ 5: **Return:** S

a test query, we can infer the budget assignment for each client. However, it is challenging to predict fine-grained budget value since there are no rigid classification patterns. For instance, determining the detailed budget value seems challenging in the case of client 1 in SST-5. Therefore, developing an efficient method to infer the exact budgets based on these broad patterns for each client is required. We also conduct t-SNE visualization on clients under other non-IID with task shifting and feature skew as shown in Figure 9 of Appendix F, which helps us to conclude this observation holds under different non-IID settings.

3.3 Observation Summary

In summary, our findings and the approach for designing a method are as follows: (1) non-IIDness significantly affects the distributed ICL setting, necessitating the development of a coalition method. To handle this problem, it is straightforward to allocate an appropriate number of budgets to each client, *i.e.*, making server work so as it knows client all samples. (2) By analyzing the query embeddings, we can determine the importance of each client per query.

4 Method

In this section, we outline the proposed method to mitigate non-IIDness in the ICL framework. Specifically, we show how to train the *budget allocator* and conduct inference.

4.1 Train a Budget Allocator

Based on Section 3, it is feasible to assign budgets of each client by using the embeddings obtained from the retriever encoder $\mathcal{E}$. We first construct the datasets having the targeting budget values and then train the budget allocator. The pseudo-codes are described in Algorithm 2 and 1.

Construct dataset for oracle budget. First, we

Algorithm 2 Construct dataset

Require: Encoder $\mathcal{E}(\cdot)$, server-side ICE budget k ,proxy dataset $\mathcal{D}_{\text{proxy}} = \{(x_j, y_j)\}_{j=1}^{N_{\text{proxy}}}$ Quantization parameter δ .1: **for** $(x_j^{\text{proxy}}, y_j^{\text{proxy}}) \in \mathcal{D}_{\text{proxy}}$ **do**2: $e_j^{\text{proxy}} = \mathcal{E}(x_j^{\text{proxy}})$

/* Get distributed examples */

3: **for** $c \in [C]$ **do**4: $e_j^{\text{proxy}} \rightarrow \text{Client } c$ 5: $S_c = \mathcal{T}(e_j^{\text{proxy}}, k|\mathcal{D}_c)$ 6: $\text{Server} \leftarrow S_c$ 7: **end for**

/* Construct optimal example */

8: $S = \bigcup_{c=1}^C S_c$ 9: $S^{\text{top}} = \arg \text{Top-}k \|e_j^{\text{proxy}} - \mathcal{E}(x_s)\|_2$
 $(x_s, y_s) \in S$ /* Compute proper budget size for each c */10: $k_c(e_j) = |S^{\text{top}} \cap S_c|/\delta \quad \forall c \in [C]$ 11: **end for**12: $B_{\text{proxy}} = \{(e_j, \{k_c(e_j)\}_{c=1}^C)\}_{j=1}^{N_{\text{proxy}}}$ 13: **Return:** B_{proxy}

explain how to create a dataset to train the budget allocator for each client, as described in Algorithm 2. Given proxy dataset $\mathcal{D}_{\text{proxy}}$, for all embeddings $e_j = \mathcal{E}(x_j)$ where $x_j \in \mathcal{D}_{\text{proxy}}$, we request k samples from each client $c \in [C]$ using Top- k procedure, *i.e.*, $S_c = \mathcal{T}(e, k|\mathcal{D}_c)$. Once the server receives k examples from each clients, *i.e.*, $\{S_c\}_{c=1}^C$, it merges and re-orders them to obtains S^{top} . Based on S^{top} , we count the number of samples from each client in S^{top} , *i.e.*, compute $k_c(e_j)$. After counting $k_c(e_j)$ for all clients, we quantize the budget levels for each client using the quantization hyper-parameter δ . As a result, the output of this procedure is B_{proxy} for all clients, composed of embeddings e and their respective budgets $k_c(e_j)$.

Train budget allocator. Based on the constructed dataset B_{proxy} , we train the *budget allocators*, *i.e.*, $\{f_c(\cdot)\}_{c=1}^C$, for each $f_c(\cdot)$ has Multi-layer perceptrons on top of the frozen feature extractor of the off-the-shelf retriever $\mathcal{E}$. The budget allocators are trained on the cross-entropy loss, as we have already quantized the optimal budgets using the hyper-parameter δ . Note that if δ is high, the quantization is severe, otherwise the quantization is mild.

Algorithm 3 Inference (Client, Server)

Require: Embedding model $\mathcal{E}(\cdot)$, LLM model $\mathcal{M}(\cdot)$, local datasets $\mathcal{D}_c$, budget allocator $f_c(\cdot)$
Buffering hyperparameter α .

Input: Test query x_q

- 1: Extract embedding $e_q = \mathcal{E}(x_q)$
 - 2: **for** $c \in [C]$ **do**
 - 3: $\hat{k}_c = f_c(e_q)$
 - 4: Send e_q to all clients
 - 5: $S_c = \mathcal{T}(e_q, \hat{k}_c + \alpha | \mathcal{D}_c)$
 - 6: return back $S_c \rightarrow$ Server
 - 7: **end for**
 - 8: $S_{\text{agg}} = \bigcup_{c \in [C]} S_c$
 - 9: $S = \mathcal{T}(e_q, k | S_{\text{agg}})$
 - 10: $s(S, x_q) = (x_1, y_1) \odot \dots \odot (x_k, y_k) \odot x_q$
 - 11: **Return:** $y = \mathcal{M}(s(S, x_q))$
-

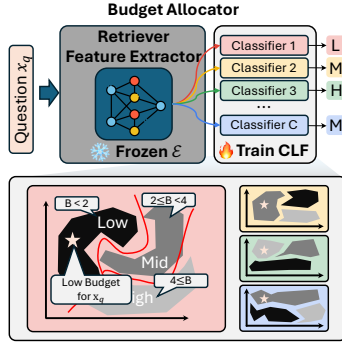


Figure 4: Overview of budget allocator: We train a budget allocator on top of the frozen encoder $\mathcal{E}$, which inherits from the retriever. During inference, when a test query x_q is provided, this module determines quantized budget levels for each client and allocates them accordingly.

4.2 Inference Using Budget Allocator

We derive the response to the test query x_q utilizing the LLM $\mathcal{M}(\cdot)$ through the described steps (see Algorithm 3 for specifics). We first extract the embedding $e_q = \mathcal{E}(x_q)$. Then, we compute the allocated budget $\{\hat{k}_c = f_c(e_q)\}_{c=1}^C$ and send $\hat{k}_c$ to each client. Each client sends back top $\hat{k}_c + \alpha$ samples, *i.e.*, S_c , to the server. Note that we summarize how the budget allocator outputs $\hat{k}_c$ in Figure 4. Here, α denotes the buffering hyper-parameter, which increases the chances for each client to be involved. After collecting $S_{\text{agg}} = \bigcup_{c \in [C]} S_c$, we aggregate them and run regular ICL.

5 Experiment

5.1 Experiment setup

First, we summarize the baselines, datasets, and the method for constructing non-IID settings. Finally, we depict the implementation details.

Baselines. We compare our method with various baselines, including Social Learning (Mohtashami et al., 2023), which does not account for non-IID, and other possible ways to handle distributed non-IID ICL, like Zero-shot, Proxy-only, Singleton, Uniform-budget, Random-budget, and ∞ -budget (oracle case). Notice that the proxy set in Proxy-Only baseline here contains the ground truth label information for each query in proxy set, which is different from proxy set used in our method. Detailed explanations are described in Appendix C.

Datasets. We check performance under 7 datasets – Sentiment classification: SST-5 (Socher et al., 2013), Amazon (McAuley and Leskovec, 2013), Yelp (Zhang et al., 2015), MR (Pang and Lee, 2005), Topic classification: Yahoo, AGNews (Zhang et al., 2015), and Subjectivity classification: Subj (Pang and Lee, 2004).

Dataset partition for non-IIDness. We split the training dataset into C subsets to ensure they follow a non-IID distribution. To achieve this, we partition the data based on class, following the splitting criteria outlined in (Li et al., 2022). Specifically, each client has access to only $\gamma < \Gamma$ classes, where Γ represents the total number of classes. We outline the summary of γ for each dataset in Appendix D. We also perform other non-IID partitions, including Dirichlet distribution non-IID, feature skew non-IID, and task shifting non-IID, with detailed description included in Appendix F.

Dataset paraphrasing. Due to concerns about sharing private samples between servers and clients, various techniques have been developed for NLP tasks. In this paper, we adopt paraphrasing technique used in (Mohtashami et al., 2023). Specifically, we utilize a small language model (Team et al., 2024) to generate paraphrased questions. In Appendix E, we summarize the instructions provided to the language model for rephrasing queries in the training dataset.

Implementation details. We implement our method and baselines based on OpenICL (Wu et al., 2023). We utilize pre-trained KATE retriever (Liu et al., 2021). Note that they do not

Algorithm	Dataset							Avg
	SST-5	Amazon	Yelp	MR	Yahoo	AGNews	Subj	
Zero-shot	29.19	24.70	31.23	73.95	25.87	67.60	50.55	43.30
Proxy-only	40.64± 2.89	28.43± 0.11	31.85± 1.28	70.40± 1.54	54.73± 0.93	84.65± 0.42	71.09± 1.34	54.54
Singleton	25.14± 4.18	24.03± 0.57	29.44± 3.91	50.00± 0.00	38.14± 2.03	50.60± 0.66	50.00± 0.00	38.19
Social Learning	36.03± 0.27	28.42± 0.19	29.25± 0.45	58.58± 0.18	46.03± 0.49	81.10± 0.29	71.37± 0.71	50.11
Uniform-budget	32.94	25.63	26.60	33.65	43.00	73.17	63.20	42.60
Random-budget	32.82± 0.82	25.69± 0.55	27.72± 0.51	34.68± 0.59	42.46± 0.53	67.34± 0.39	65.37± 0.80	42.30
∞ -budget	43.26	32.70	34.80	77.20	62.67	89.37	91.4	61.62
Ours	44.08± 0.12	31.54± 0.22	35.48± 0.28	80.44± 0.67	61.67± 0.25	88.52± 0.30	82.36± 0.91	60.58

Table 1: Main results: To address the issue of non-IIDness in distributed ICL, we examined seven datasets and seven straightforward baselines. We run three random seeds and illustrate mean and std values. The top performance is highlighted in **bold** font, excluding the infinite budget scenario due to its impracticality. In summary, the proposed method effectively mitigates the non-iid distributed ICL problem to a reasonable extent.

overlap with the datasets used in our experiment. They used RoBERTa-large (Liu et al., 2019) encoder model. We use GPT-Neo-2.7B (Black et al., 2021) pre-trained model as answering LLMs as default. Hyper-parameters related to training budget allocators, α , and δ are described in Appendix D.

5.2 Main results

We have presented the performance of our method and baselines in Table 1. First, we can observe that performance varies significantly depending on the way the budget is allocated, which indicates that the budget allocation scheme really matters in non-IID ICL. Additionally, even when using only the proxy dataset, there is a performance improvement, and this performance surpasses that of using other clients which have the tilted local datasets (e.g., 29.19% $\rightarrow$ 40.64% in SST-5 case). This indicates that utilizing a biased dataset can degrade the ICL performance. Although Social Learning method has shown good performance in the previous paper, it does not perform well under the non-IID cases configured in this research. If we can use an infinite budget, all settings would exhibit high performance. However, our proposed method demonstrates better performance than the infinite budget upper limit (e.g., 34.86% $\rightarrow$ 35.48% in the Yelp case). This is likely due to a mechanism that prevents unnecessary information from being selected by the retriever with high importance. Ultimately, our method shows an average performance improvement of 5.05% across seven datasets compared to the best performance of baselines using the proxy dataset. Our method also shows outstanding performance under Dirichlet distribution non-IID, as shown in Table 14 of Appendix F. These

show that the proposed method can handle the non-IID case well. Furthermore, considering that our method does not necessitate the collection of ground truth label information for the proxy set and solely relies on task-related queries, it proves to be more practical in real-world applications compared to the leading baseline Proxy-only. Unlike the baseline, our method imposes fewer constraints on the proxy set, enhancing its practicability.

5.3 Analysis

In this section, we further examine four key aspects: (1) privacy-preserving case analysis, which encompasses paraphrasing both training and testing queries, (2) sensitivity to hyper-parameters, (3) the performance of the trained budget allocator, and (4) the compatibility of the LLMs.

Paraphrasing results. Due to privacy concerns in the fundamental distributed system, we evaluate the performance of paraphrased datasets, with results detailed in Table 20 of Appendix F. Our method demonstrates superior performance compared to other baselines across multiple datasets. We used the exact same data settings as in Table 1. Specifically, performance on the Subj and SST-5 datasets is lower than without paraphrasing, while the Yelp dataset shows a slight improvement. Additionally, as consistent with Table 1, non-IIDness causes significant performance degradation for ICL methods, as seen by comparing Zero-shot with ICL-related methods (e.g., 27.96% $\rightarrow$ 25.31% in Singleton).

Hyper-parameter sensitivity. We examine the sensitivity of the hyper-parameters of our method. We have two hyper-parameters: δ , which is the resolution of the budget allocator; α , which represents the additional budget allocated to each client as a

buffer; and proxy size, which is the size of proxy data for the budget allocator training. As illustrated in Figure 11 of Appendix F, when we increase α , the performance is improved while the budget efficiency is reduced. On the other hand, when δ is high (or low), it has too sparse (or dense) representation of the budget class, thus performance is degraded. Nevertheless, the performance is higher than the other baselines in Table 1. For the sensitivity of the size of proxy data, it is revealed that our framework is not sensitive to how many proxy data samples are used to train the budget allocator, as shown in Figure 12 of Appendix F. This indicates our method is stable even with limited proxy data on the server side. We have more fine-grained experiment result for relation between proxy size and the budget allocator resolution δ , which is shown in Figure 7 of Appendix F.

Trained budget allocator. We assess whether the trained budget allocator distributes budgets appropriately for each client. To evaluate efficiency, we examine the number of samples, *i.e.*, $\hat{k}_c$ communicated for all queries and plot histogram. As demonstrated in Figure 6 of Appendix F, we confirm that our method’s forecasts exhibit nearly identical performance to the oracle budget when an additional 25% budget is allocated. Notice that it is necessary to assign $k \times C$ budgets to get a performance similar to the oracle case, without our method.

Other types of LLMs. We utilize various LLMs to assess the compatibility of our method. Specifically, we evaluate the SST-5 using different model sizes, including GPT-Neo-1.3B (Black et al., 2021), Llama-2-7B (Touvron et al., 2023a), and OpenAI gpt-3.5-turbo (OpenAI, 2022). As demonstrated in Table 18 of Appendix F, our method exhibits a plug-and-play capability and achieves reasonable performance improvements in ICL.

Distribution of proxy set. We also conduct experiments on when the proxy set has different distributions from the test set, aiming to verify the applicability of our method in broader real-world settings. The results and conclusion are included in Appendix F.5 due to the space limit.

6 Related Work

In-context learning. ICL (Dong et al., 2022) is one of the fastest paradigms using pre-trained LLMs by feeding several examples to construct the context to solve the given query. The main criteria of this research field are to find the most informa-

tive samples among the training datasets. (Liu et al., 2021) trains BERT (Devlin et al., 2018) oriented encoder and uses k nearest neighbors. (Rubin et al., 2022) proposed an efficient retriever called EPR. It trains two encoders by inheriting method of dense passage retriever (DPR) (Karpukhin et al., 2020) using loss of positive and negative pairs. To reduce domain specificity, (Li et al., 2023) proposed UDR, which is applicable to multiple domain tasks in a universal way and shows reasonable performance from a single retriever. PromptPG (Lu et al., 2022) utilized a reinforcement learning framework to train the retriever so that it can generate context to improve the answerability of LLMs. (Chang and Jia, 2022) trains linear regressors according to the example influence on the LLM prediction. (Xie et al., 2021) proposes to use implicit Bayesian inference to understand the ICL problem. Note that extensive research focuses on the centralized case rather than targeting distributed cases.

Distributed ICL. To the best of our knowledge, only a single study (Mohtashami et al., 2023) tries to address ICL in a distributed manner. However, this paper solely focuses on merging the distributed information without considering the nature of the non-identically distributed information. Many studies, such as those on federated learning (Li et al., 2021; Zhang et al., 2021; Mammen, 2021), address the non-IID distribution of datasets, highlighting the need to handle distributed non-IID ICL. Distributed ICL may resemble distributed RAG, yet the latter is more complex and requires further exploration, as discussed in Appendix H.

7 Conclusion

In this paper, we tackle the challenge of distributed non-IID ICL. Initially, we show that non-IID leads to performance degradation and discover that improper budget allocation causes significant drops in ICL. Inspired by the learnable pattern between budget values and query embeddings, we propose a method that learns budget assignment and employs it during inference to allocate appropriate budgets for each query. The proposed method achieves performance improvements across several benchmarks compared with various baselines. In addition, we examine the privacy-preserving version of our method using paraphrasing and show its efficacy. Last but not least, extensive sensitivity experiments show robustness of our method on hyperparameters and different LLMs.

Limitations. This research addresses in-context learning with datasets distributed among clients with non-iid data. One limitation of this study can be described as follows: (1) Proxy dataset: In practice, some datasets might lack a proxy dataset on the server side, posing a challenge for the proposed algorithm. (2) The research assumes the use of a pre-trained off-the-shelf retriever. However, this retriever may fail if there is a significant difference between the target task domain and the pre-trained domain. This issue can be mitigated by employing the contrastive training mechanism suggested in federated learning research (Seo et al., 2024), as one effective approach for training retrievers is utilizing contrastive loss.

Ethical statement. As outlined in the main manuscript, utilizing distributed knowledge raises privacy concerns. We address this by employing a paraphrasing technique developed and frequently used in federated learning with pre-trained generative models, although it has not been thoroughly explored. We will make every effort to eliminate privacy concerns and implement all possible measures to prevent privacy information leakage when applying the paraphrasing method to the best of our ability.

References

- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, and 1 others. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Sid Black, Leo Gao, Phil Wang, Connor Leahy, and Stella Biderman. 2021. [GPT-Neo: Large Scale Autoregressive Language Modeling with Mesh-Tensorflow](#). If you use this software, please cite it using these metadata.
- Ting-Yun Chang and Robin Jia. 2022. Data curation alone can stabilize in-context learning. *arXiv preprint arXiv:2212.10378*.
- Zicun Cong, Xuan Luo, Jian Pei, Feida Zhu, and Yong Zhang. 2022. Data pricing in machine learning pipelines. *Knowledge and Information Systems*, 64(6):1417–1455.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2018. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*.
- Qingxiu Dong, Lei Li, Damai Dai, Ce Zheng, Zhiyong Wu, Baobao Chang, Xu Sun, Jingjing Xu, and Zhifang Sui. 2022. A survey on in-context learning. *arXiv preprint arXiv:2301.00234*.
- Eduard Hovy, Laurie Gerber, Ulf Hermjakob, Chin-Yew Lin, and Deepak Ravichandran. 2001. Toward semantics-based answer pinpointing. In *Proceedings of the first international conference on Human language technology research*.
- Vladimir Karpukhin, Barlas Oğuz, Sewon Min, Patrick Lewis, Ledell Wu, Sergey Edunov, Danqi Chen, and Wen-tau Yih. 2020. Dense passage retrieval for open-domain question answering. *arXiv preprint arXiv:2004.04906*.
- Itay Levy, Ben Bogin, and Jonathan Berant. 2022. Diverse demonstrations improve in-context compositional generalization. *arXiv preprint arXiv:2212.06800*.
- Jiaxing Li, Chi Xu, Lianchen Jia, Feng Wang, Cong Zhang, and Jiangchuan Liu. 2024. Eaco-rag: Edge-assisted and collaborative rag with adaptive knowledge update. *arXiv preprint arXiv:2410.20299*.
- Qinbin Li, Yiqun Diao, Quan Chen, and Bingsheng He. 2022. Federated learning on non-iid data silos: An experimental study. In *2022 IEEE 38th international conference on data engineering (ICDE)*, pages 965–978. IEEE.
- Qinbin Li, Zeyi Wen, Zhaomin Wu, Sixu Hu, Naibo Wang, Yuan Li, Xu Liu, and Bingsheng He. 2021. A survey on federated learning systems: Vision, hype and reality for data privacy and protection. *IEEE Transactions on Knowledge and Data Engineering*, 35(4):3347–3366.

688	Xiaonan Li, Kai Lv, Hang Yan, Tianyang Lin, Wei Zhu,	Richard Socher, Alex Perelygin, Jean Wu, Jason	743
689	Yuan Ni, Guotong Xie, Xiaoling Wang, and Xipeng	Chuang, Christopher D Manning, Andrew Y Ng, and	744
690	Qiu. 2023. Unified demonstration retriever for in-	Christopher Potts. 2013. Recursive deep models for	745
691	context learning. <i>arXiv preprint arXiv:2305.04320</i> .	semantic compositionality over a sentiment treebank.	746
		In <i>Proceedings of the 2013 conference on empiri-</i>	747
692	Jiachang Liu, Dinghan Shen, Yizhe Zhang, Bill Dolan,	<i>cal methods in natural language processing</i> , pages	748
693	Lawrence Carin, and Weizhu Chen. 2021. What	1631–1642.	749
694	makes good in-context examples for gpt-3? <i>arXiv</i>		
695	<i>preprint arXiv:2101.06804</i> .	Zuoqi Tang, Zheqi Lv, and Chao Wu. 2020. Abrief	750
		survey of data pricing for machine learning. In <i>CS</i>	751
696	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-	<i>& IT Conference Proceedings</i> , volume 10. CS & IT	752
697	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,	Conference Proceedings.	753
698	Luke Zettlemoyer, and Veselin Stoyanov. 2019.		
699	Roberta: A robustly optimized bert pretraining ap-	Gemini Team, Rohan Anil, Sebastian Borgeaud,	754
700	proach. <i>arXiv preprint arXiv:1907.11692</i> .	Yonghui Wu, Jean-Baptiste Alayrac, Jiahui Yu, Radu	755
		Soricut, Johan Schalkwyk, Andrew M Dai, Anja	756
701	Pan Lu, Liang Qiu, Kai-Wei Chang, Ying Nian Wu,	Hauth, and 1 others. 2023. Gemini: a family of	757
702	Song-Chun Zhu, Tanmay Rajpurohit, Peter Clark,	highly capable multimodal models. <i>arXiv preprint</i>	758
703	and Ashwin Kalyan. 2022. Dynamic prompt learning	<i>arXiv:2312.11805</i> .	759
704	via policy gradient for semi-structured mathematical		
705	reasoning. <i>arXiv preprint arXiv:2209.14610</i> .	Gemma Team, Thomas Mesnard, Cassidy Hardin,	760
		Robert Dadashi, Surya Bhupatiraju, Shreya Pathak,	761
706	Priyanka Mary Mammen. 2021. Federated learn-	Laurent Sifre, Morgane Rivi�re, Mihir Sanjay Kale,	762
707	ing: Opportunities and challenges. <i>arXiv preprint</i>	Juliette Love, and 1 others. 2024. Gemma: Open	763
708	<i>arXiv:2101.05428</i> .	models based on gemini research and technology.	764
		<i>arXiv preprint arXiv:2403.08295</i> .	765
709	Julian McAuley and Jure Leskovec. 2013. Hidden fac-		
710	tors and hidden topics: understanding rating dimen-	Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier	766
711	sions with review text. In <i>Proceedings of the 7th</i>	Martinet, Marie-Anne Lachaux, Timoth�e Lacroix,	767
712	<i>ACM conference on Recommender systems</i> , pages	Baptiste Rozi�re, Naman Goyal, Eric Hambro, Faisal	768
713	165–172.	Azhar, and 1 others. 2023a. Llama: Open and ef-	769
		ficient foundation language models. <i>arXiv preprint</i>	770
714	Amirkeivan Mohtashami, Florian Hartmann, Sian	<i>arXiv:2302.13971</i> .	771
715	Gooding, Lukas Zilka, Matt Sharifi, and 1 others.		
716	2023. Social learning: Towards collaborative learn-	Hugo Touvron, Louis Martin, Kevin Stone, Peter Al-	772
717	ing with large language models. <i>arXiv preprint</i>	bert, Amjad Almahairi, Yasmine Babaei, Nikolay	773
718	<i>arXiv:2312.11441</i> .	Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti	774
		Bhosale, and 1 others. 2023b. Llama 2: Open foun-	775
719	OpenAI. 2022. Introducing chatgpt .	dation and fine-tuned chat models. <i>arXiv preprint</i>	776
720	Bo Pang and Lillian Lee. 2004. A sentimental educa-	<i>arXiv:2307.09288</i> .	777
721	tion: Sentiment analysis using subjectivity summa-		
722	rization based on minimum cuts . In <i>Proceedings</i>	Laurens Van der Maaten and Geoffrey Hinton. 2008.	778
723	<i>of the 42nd Annual Meeting of the Association for</i>	Visualizing data using t-sne. <i>Journal of machine</i>	779
724	<i>Computational Linguistics (ACL-04)</i> , pages 271–278,	<i>learning research</i> , 9(11).	780
725	Barcelona, Spain.		
		Alex Wang, Yada Pruksachatkun, Nikita Nangia, Aman-	781
726	Bo Pang and Lillian Lee. 2005. Seeing stars: Exploiting	preet Singh, Julian Michael, Felix Hill, Omer Levy,	782
727	class relationships for sentiment categorization with	and Samuel Bowman. 2019. Superglue: A stick-	783
728	respect to rating scales. <i>arXiv preprint cs/0506075</i> .	er benchmark for general-purpose language under-	784
		standing systems. <i>Advances in neural information</i>	785
729	Stephen Robertson, Hugo Zaragoza, and 1 others. 2009.	<i>processing systems</i> , 32.	786
730	The probabilistic relevance framework: Bm25 and		
731	beyond. <i>Foundations and Trends� in Information</i>	Alex Wang, Amanpreet Singh, Julian Michael, Felix	787
732	<i>Retrieval</i> , 3(4):333–389.	Hill, Omer Levy, and Samuel R Bowman. 2018.	788
733	Ohad Rubin, Jonathan Herzig, and Jonathan Berant.	Glue: A multi-task benchmark and analysis platform	789
734	2022. Learning to retrieve prompts for in-context	for natural language understanding. <i>arXiv preprint</i>	790
735	learning . In <i>Proceedings of the 2022 Conference of</i>	<i>arXiv:1804.07461</i> .	791
736	<i>the North American Chapter of the Association for</i>		
737	<i>Computational Linguistics: Human Language Tech-</i>	Shuai Wang, Ekaterina Khramtsova, Shengyao Zhuang,	792
738	<i>nologies</i> , pages 2655–2671, Seattle, United States.	and Guido Zuccon. 2024. Feb4rag: Evaluating fed-	793
739	Association for Computational Linguistics.	erated search in the context of retrieval augmented	794
		generation. In <i>Proceedings of the 47th International</i>	795
740	Seonguk Seo, Jinkyu Kim, Geeho Kim, and Bohyung	<i>ACM SIGIR Conference on Research and Develop-</i>	796
741	Han. 2024. Relaxed contrastive learning for feder-	<i>ment in Information Retrieval</i> , pages 763–773.	797
742	ated learning. <i>arXiv preprint arXiv:2401.04928</i> .		

- Zhenyu Wu Wu, Yaoxiang Wang, Jiacheng Ye, Jiangtao Feng, Jingjing Xu, Yu Qiao, and Zhiyong Wu. 2023. Openicl: An open-source framework for in-context learning. *arXiv preprint arXiv:2303.02913*.
- Sang Michael Xie, Aditi Raghunathan, Percy Liang, and Tengyu Ma. 2021. An explanation of in-context learning as implicit bayesian inference. *arXiv preprint arXiv:2111.02080*.
- Jimin Xu, Nuanxin Hong, Zhenning Xu, Zhou Zhao, Chao Wu, Kun Kuang, Jiaping Wang, Mingjie Zhu, Jingren Zhou, Kui Ren, and 1 others. 2023. Data-driven learning for data rights, data pricing, and privacy computing. *Engineering*, 25:66–76.
- Jiacheng Ye, Zhiyong Wu, Jiangtao Feng, Tao Yu, and Lingpeng Kong. 2023. Compositional exemplars for in-context learning. In *International Conference on Machine Learning*, pages 39818–39833. PMLR.
- Joo Hun Yoo, Hyejun Jeong, Jaehyeok Lee, and Tai-Myoung Chung. 2022. Open problems in medical federated learning. *International Journal of Web Information Systems*, 18(2/3):77–99.
- Chen Zhang, Yu Xie, Hang Bai, Bin Yu, Weihong Li, and Yuan Gao. 2021. A survey on federated learning. *Knowledge-Based Systems*, 216:106775.
- Kaiyi Zhang, Ang Lv, Yuhao Chen, Hansen Ha, Tao Xu, and Rui Yan. 2024. Batch-icl: Effective, efficient, and order-agnostic in-context learning. *arXiv preprint arXiv:2401.06469*.
- Xiang Zhang, Junbo Zhao, and Yann LeCun. 2015. Character-level convolutional networks for text classification. *Advances in neural information processing systems*, 28.

A How we construct non-IIDness

Following Li et al. (2022), we use class number based non-IID partition in our experiment. For a dataset with overall Γ classes, given hyperparameter class number γ on each client, we randomly assign γ classes from the overall Γ classes for each client. Assuming that $C_1 \leq C$ clients are assigned with a specific class, we equally partition samples of this class into C_1 parts and assign one part to each of C_1 clients. We denote this non-IID partition with the class number γ on each client as `noniid-#label= $\gamma$` .

B Motivation Experimental Settings

Non-IIDness performance drop experiment. For this experiment, we use SST5, Amazon, Yelp, Yahoo, and AGNews. And the non-iid settings are described in Table 2.

Dataset	ICE size k	#Clients	Partition
SST-5	32	8	<code>noniid-#label=1</code>
Amazon	16	8	<code>noniid-#label=1</code>
Yahoo	16	8	<code>noniid-#label=2</code>
AGNews	16	4	<code>noniid-#label=1</code>
Yelp	4	2	<code>noniid-#label=3</code>

Table 2: Experimental setup for obtaining the motivation.

t-SNE analysis of per-client budget experiment. For extracting t-SNE figure, we utilized the following experimental setting Table 3

Dataset	ICE size k	#Clients	Partition
SST-5	32	4	<code>noniid-#label=2</code>
Yelp	8	2	<code>noniid-#label=3</code>

Table 3: Experimental setup for obtaining the motivation.

C Baseline Details

Proxy-only. We randomly select samples from the original test set to construct the proxy set on the server side and use the remaining test set as the true test set. Notice that the proxy set in this baseline is different from the one we use in our pipeline. The proxy set used in Proxy-only contains label information for each query in the proxy set, while the proxy set in our algorithm does not require this information, which is more flexible in the real-world setting. When performing the ICL process, the server directly retrieves ICEs from the proxy set

rather than from the training set. For SST5, MR, and Subj, we randomly select 500 samples from the test set to be the proxy set. For Amazon, Yelp, Yahoo, and Agnews, we randomly select 750 samples from the test set to be the proxy set. Also, since the proxy set is already on the server side, there will be no privacy issues during communication between clients and the server. Thus, we don’t generate samples to protect privacy and directly use the original samples in the proxy set for ICL.

Singleton. This baseline is for if the whole ICE set is constructed only using single client’s local dataset. We randomly select one client from C clients, and perform local retrieval with $k_c = k$ budget. Then, the server uses this locally retrieved ICE set for LLM inference. We report the average accuracy over all clients.

Social learning. This algorithm (Mohtashami et al., 2023) is the first paper that considers the distributed ICL, but it only considers the IID setting. Since the authors didn’t release the source code, we implemented it on our own. In our implementation, given server-side ICE number as k , each local client c performs local top- $\lceil \frac{k}{C} \rceil$ retrieval and sends retrieved ICEs to the server. The server then performs a random selection from k ICEs to construct an ICE set with k samples and feed this ICE set into LLM for inference.

Uniform-budget. We equally assign a local budget to each client. Assume the ICE number fed to server-side LLM for inference is k , then each client’s local budget is $\lceil \frac{k}{C} \rceil$, where C is the number of clients. On server-side aggregation, we use reorder method as default.

Random-budget. We randomly assign a local budget to each client with the constraint that the overall local budget over C clients is k , where k is the ICE number fed to server-side LLM. On server-side aggregation, we use reorder method as default.

∞ -budget. The most inefficient way to do distributed non-IID ICL is to allow ∞ -budget on each client, that is, sending all samples to the server side. Then, the system performs centralized retrieval on the collected dataset to obtain top- k ICEs and feed them into LLM for inference.

D Experimental Setting

Dataset Explanation. In this study, we utilized seven text classification tasks: four for sentiment

analysis, two for topic classification, and one for subjectivity classification. The dataset statistics are presented in Table 4.

Dataset	Type	Training	Test	Class
SST-5	Sentiment	8,534	2,210	5
Amazon	Sentiment	30,000	3,000	5
Yelp	Sentiment	30,000	3,000	5
MR	Sentiment	8,662	2,000	2
Yahoo	Topic	29,150	3,000	10
AGNews	Topic	29,914	3,000	4
Subj	Subjectivity	8,000	2,000	2

Table 4: The statistics of the datasets used.

Given that the input instruction prompt can notably influence performance, we detail the prompts used for each dataset in Table 24. It is in the last page since prompts have long length. We follow the prompt settings described in (Li et al., 2023) and use the dataset uploaded by the paper’s author, available at <https://huggingface.co/Kailv>.

ICE number for LLM inference. Given an LLM, different datasets show different preferences on the choice of ICE number, *i.e.*, k , used in ICL inference for better performance. For algorithms using ICL (except Zero-shot), SST5, MR, and Subj use 32 ICEs for server-side LLM inference; Amazon uses 8 ICEs for server-side LLM inference; Yelp, Yahoo, and Agnews use 4 ICEs for server-side LLM inference.

Non-IID Setting. To keep similar non-IIDness levels across different datasets, we follow Table 5 as non-IID hyper-parameters for each dataset.

Hyper-parameters for our methods. For the main table results, the generated dataset results, and the different LLM architecture results, the hyper-parameters are shown in Table 6, Table 7 and Table 8, respectively. For the training of the budget model, we use 800 epochs, with a learning rate range $\{0.01, 0.003\}$ and a batch size of 8.

Multi-layer perceptron for budget allocator. We use the three-layer perceptron on top of the encoder $\mathcal{E}$. The torch pseudo code is as follows:

```
class SMLP(nn.Module):
```

```
def __init__(self, width=300,
              num_classes=10,
              data_shape=(768,)):
    super().__init__()
    self.flat = nn.Flatten()
    self.l1 = nn.Linear(np.prod(data_shape),
                        width)
    self.relu = nn.ReLU()
    self.l2 = nn.Linear(width, width)
    self.l3 = nn.Linear(width, num_classes)

    def forward(self, x):
        x = self.flat(x)
        x = self.l1(x)
        x = self.relu(x)
        x = self.l2(x)
        x = self.relu(x)
        x = self.l3(x)
        x = F.softmax(x)
        return x
```

Dataset	#Clients	Partition
SST-5	4	noniid-#label=2
Amazon	2	noniid-#label=3
Yelp	2	noniid-#label=3
MR	4	noniid-#label=1
Yahoo	2	noniid-#label=5
AGNews	2	noniid-#label=2
Subj	4	noniid-#label=1

Table 5: Non-IID setting

Dataset	ProxySetSize	δ	α	QuantRatio
SST-5	500	3	0	0.5
Amazon	750	2	0	0.5
Yelp	750	2	2	0.5
MR	500	3	0	0.5
Yahoo	750	2	2	0.5
AGNews	750	2	2	0.5
Subj	500	3	0	0.3

Table 6: Hyper-parameters of our methods used in the main table

Dataset	ProxySetSize	δ	α	QuantRatio
SST-5	500	3	4	0.5
Yelp	750	3	0	0.5
Subj	500	3	0	0.3

Table 7: Hyper-parameters of our methods used for the generated query and training samples experiment

Computation Environment. We run our experiments on NVIDIA RTX A5000 GPU. Each experiment takes less than half an hour on a single GPU.

Model	ProxySetSize	δ	α	QuantRatio
GPT-Neo-1.3B	500	3	0	0.3
GPT-Neo-2.7B	500	3	0	0.3
Llama-2-7B	500	3	0	0.3
gpt-3.5-turbo	500	3	0	0.3

Table 8: Hyper-parameters of our methods used for different LLM architectures experiment on Subj

E Generate paraphrased question

To generate the paraphrased query and response, we use the following instruction.

Please paraphrase the original sentence. Original sentence: {In-context example} Paraphrase sentence: {Paraphrased sentence}

Here is an example of input for the rephrasing LLM using the SST-5 dataset.

Please paraphrase the original sentence. Original sentence: "a stirring, funny and finally transporting re-imagining of beauty and the beast and 1930s horror films" Paraphrased sentence: A captivating, humorous, and ultimately uplifting reinterpretation of Beauty and the Beast combined with 1930s horror films. Please paraphrase the original sentence. Original sentence: "jonathan parker 's bartleby should have been the be-all-end-all of the modern-office anomie films" Paraphrased sentence: Jonathan Parker's "Bartleby" had the potential to be the definitive film capturing the sense of alienation in modern office settings. Please paraphrase the original sentence. Original sentence: "a fan film that for the uninitiated plays better on video with the sound turned down" Paraphrased sentence: A fan film that, for those not familiar with the source material, is more enjoyable when watched with the sound turned off. Please paraphrase the original sentence. Original sentence: "apparently reassembled from the cutting-room floor of any given daytime soap" Paraphrased sentence: It appears to be pieced together from the outtakes of any given daytime soap opera. Please paraphrase the original sentence. Original sentence: "" Paraphrased sentence:

Our paraphrased examples are summarized as follows.

F Extra Experiments

F.1 Per-client t-SNE visualization colored by oracle budgets

Figure 5 presents the t-SNE visualization for per-client query embeddings, colored by oracle budget values.

F.2 Budget Analysis Results

Figure 6 presents the result of the budget analysis of our method.

F.3 Robustness on Proxy Size

Here, we present more detailed results on Subj with different proxy sizes over different values on budget allocator resolution δ in Figure 7.

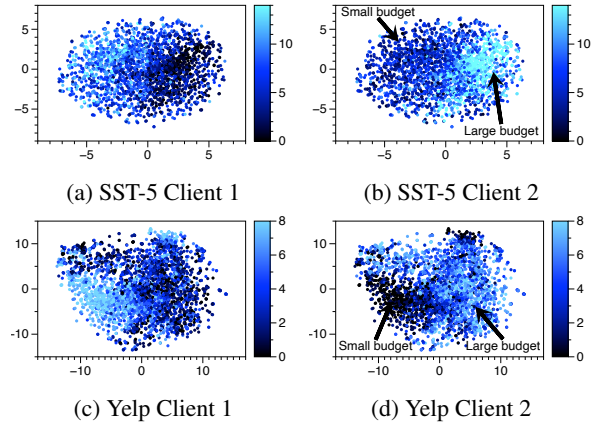


Figure 5: t-SNE analysis of each client across two datasets. Each figure demonstrates that the budgets can be segregated by training a simple classifier, as they exhibit clustered subgroup pattern.

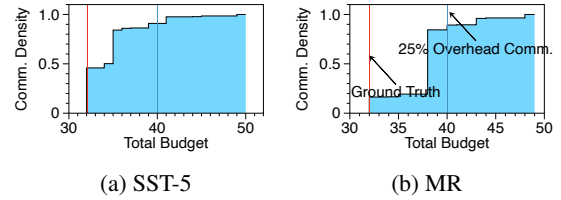


Figure 6: Analyze the total amount of budget allocated to clients under two datasets. Red and blue denote the oracle and 25% larger total budgets compared to oracle case.



Figure 7: Proxy size robustness over different budget allocator resolution δ . The results of Subj using GPT-Neo-2.7B.

F.4 More experiments for claim on Non-IID leads to ICL performance drop

To further support our claim that Non-IIDness leads to ICL performance drop under distributed setting, we introduce an additional feature-based Non-IID setting, where clients are split based on query length (number of words). For example, for Subj with 4 clients setting, client 0 only has samples with query length in range of $[0, 13)$, client 1 within $[13, 26)$, client 2 within $[26, 31)$, and client 3 within $[31, \infty)$. We present the results for this setting in Table 12. By comparing the result of Uniform-budget (81.20%) and that of ∞ -budget/centralized

Original	Paraphrased
a turgid little history lesson , humourless and dull	A dry and tedious history lesson that is devoid of humour or interest.
not so much a movie as a picture book for the big screen .	The movie is more of a picture book than a full-fledged movie for the big screen.
now it 's just tired .	It is now simply outdated.

Table 9: Paraphrased examples of SST-5 dataset

Original	Paraphrased
for all the wit and hoopla , festival in cannes offers rare insight into the structure of relationships .	Festival in Cannes offers rare insight into the structure of relationships.
eldom has a movie so closely matched the spirit of a man and his work .	A movie seldom has a movie so closely matched the spirit of a man and his work.
those of you who are not an eighth grade girl will most likely doze off during this one .	8th graders and younger will most likely doze off during this one.

Table 10: Paraphrased examples of Subj dataset

setting (91.40%), we can still conclude that this kind of Non-IID setting also leads to performance drop, and our method helps to improve the performance under this Non-IID setting.

Algorithm	Subj
Zero-shot	50.55
Proxy-only	71.09
Singleton	76.22
Social Learning	79.60
Uniform-budget	81.20
Random-budget	81.75
∞ -budget	91.40
Ours	82.80

Table 12: Performance on query length based Non-IID setting.

We also include results under this query-length-based Non-IID setting for SST-5 with 8 clients, comparing IID vs. Non-IID performance following the same setting as in Figure 3. Again, performance degrades under Non-IID conditions, supporting our original claim.

Setting	Accuracy
Centralized	43.26
IID	44.52
Non-IID class-based	10.72
Non-IID query-length based	32.30

Table 13: Performance comparison under different settings for SST-5. Non-IID class-based: 8 clients distributed setting following Non-IID setting in Figure 3. Non-IID query-length based: 8 clients distributed setting under query length based Non-IID.

F.5 Non-Extreme Non-IID on Binary Classification Tasks

We conduct the experiment on Dirichlet distribution $\text{Dir}(\alpha)$ Non-IID partition on Subj and MR under the setting of 4 clients with $\alpha_{Dir} = 1.5$. The per-client sample distribution is shown in Figure 8, and the performance results are shown in Table 14.

Algorithm	Dataset	
	MR	Subj
Zero-shot	73.95	50.55
Proxy-only	70.40	71.09
Singleton	64.16	73.80
Social Learning	58.85	76.95
Uniform-budget	52.85	77.80
Random-budget	53.50	77.85
∞ -budget	77.20	91.40
Ours	75.53	82.80

Table 14: MR, Subj results under Dirichlet distribution Non-IID.



Figure 8: Per-client sample distribution under Dirichlet distribution Non-IID setting.

To further show the robustness of our method under Dirichlet distribution based Non-IID, we perform experiments on Subj for more α_{Dir} . As shown in Table 15, our method shows robustness even when $\alpha_{Dir} = 0.3$ (which we consider as an extreme Non-IID case), and beats other baselines (except the oracle ∞ -budget case). And as α_{Dir} increases (becoming more IID), the performance of our method also increases as expected.

Algorithm	α_{Dir}		
	0.3	1.5	3.0
Zero-shot	50.55		
Proxy-only	71.09		
Singleton	67.50	73.80	74.93
Social Learning	67.10	76.95	78.95
Uniform-budget	67.85	77.80	79.65
Random-budget	67.80	77.85	80.00
∞ -budget	91.40		
Ours	82.07	82.80	83.73

Table 15: Results for Dirichlet distribution Non-IID under different α_{Dir} on Subj.

F.6 t-SNE under Non-IID with Task Shifting & Feature Skew

Dataset Amazon and Yelp are 5-class sentiment classification tasks with exactly the same label space, while different text query distributions. Based on this, we design a special Non-IID setting with task shifting & feature skew between clients: client 1 only contains 10,000 Amazon training samples, and client 2 only contains 10,000 Yelp training samples. Thus,

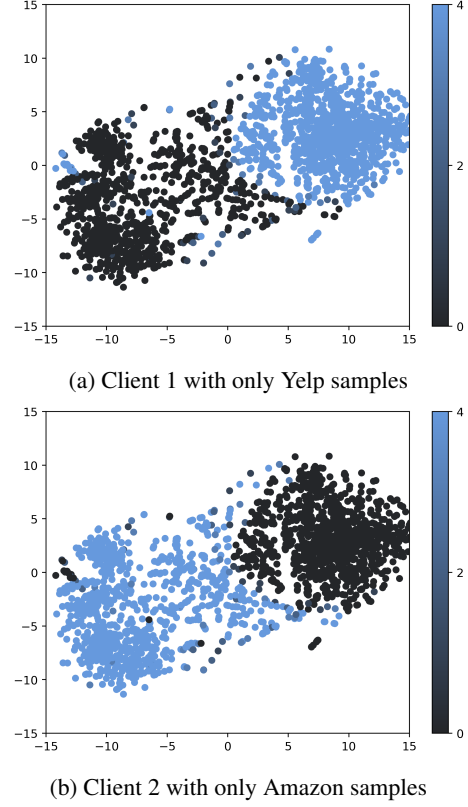


Figure 9: t-SNE analysis on the test set consisting of both Yelp & Amazon samples. Data points are colored based on local oracle budget values.

we consider this special setting to be a task-shifting Non-IID. Also, since each client consists of samples from all classes of each task, we consider this setting as feature-skew Non-IID with class balance. We calculate the oracle budget values for a mixed test set consisting of 1,000 Yelp test samples and 1,000 Amazon test samples. Then we perform t-SNE analysis on sample embeddings of this mixed test set, colored using oracle local budget values. As shown in Figure 9, under Non-IID with task shifting & feature skew, there still exists clear clustering pattern between query embedding and oracle budget values. This indicates our method still can work with task shifting and feature skew.

F.7 Distribution Shift between Proxy Set and Test Set

It is critical to control the distribution shifting between proxy set and test set. We conduct experiments on two settings for proxy set distribution different from test set.

From the same dataset but different label distribution. The most simple case of “different distribution” can come from the different label distribution skew between proxy set and test set. For this setting, we experiment on Subj with a proxy set only containing samples of one class. As shown in the last line of Table 16, when the label skew exists, the performance of our method does decrease compared with the setting using the ideal proxy set (from 82.36% to 70.17%). However, it is still higher than some baselines, like zero-shot, singleton, uniform-budget, random-budget.

Similar task but different dataset. A more extreme case for proxy set different from test set can be, proxy set share

same task with the test set, but are from different datasets. For this setting, we conduct the following experiment:

- use Amazon as proxy set for Yelp Non-IID setting (evaluate on Yelp test set)
- use Yelp as proxy set for Amazon Non-IID setting (evaluate on Amazon test set)

Since Yelp and Amazon share the similar task, we can consider this setting as using available dataset with similar task with the test set to construct the proxy set. We present the result in the last line in Table 17. It shows that for the Amazon setting, using Yelp as a proxy set, the performance drop of our method is slight, and our method still outperforms other baselines, except in the ideal case where we use Amazon samples as a proxy set. While for Yelp setting using Amazon as proxy set, our method surprisingly shows even better performance than the ideal case, where use Yelp as proxy set.

Algorithm	Subj
Zero-shot	50.55
Proxy-only	71.09
Singleton	50.00
Social Learning	71.37
Uniform-budget	63.20
Random-budget	65.37
∞ -budget	91.30
Ours	82.36
<i>Ours-proxy-label-skew</i>	<i>70.17</i>

Table 16: Comparison with proxy set with label skew compared to the test set. The last line is the performance for this setting.

Algorithm	Dataset	
	Amazon	Yelp
Zero-shot	24.70	31.23
Proxy-only	28.43	31.85
Singleton	24.03	29.44
Social Learning	28.42	29.25
Uniform-budget	25.63	26.60
Random-budget	25.69	27.72
∞ -budget	32.70	34.80
Ours	31.54	35.48
<i>Ours-diff-proxy</i>	<i>31.27</i>	<i>37.33</i>

Table 17: Comparison with using different dataset to construct proxy set for budget allocator training. The last line is the performance for this setting.

F.8 Results for other types of LLMs

The results for other types of LLMs are presented in Table 18.

Algorithm	Architecture			
	GPT-Neo-1.3B	GPT-Neo-2.7B	Llama-2-7B	gpt-3.5-turbo
Zero-shot	51.30	50.55	49.10	57.57
Proxy-only	80.18 $\pm$ 1.87	71.09 $\pm$ 1.34	88.13 $\pm$ 0.74	88.44 $\pm$ 0.69
Singleton	50.00 $\pm$ 0.00	50.00 $\pm$ 0.00	52.89 $\pm$ 3.43	60.81 $\pm$ 6.31
Social Learning	68.55 $\pm$ 0.64	71.37 $\pm$ 0.71	88.82 $\pm$ 0.50	87.53 $\pm$ 0.46
Uniform-budget	44.40	63.20	54.00	81.23
Random-budget	43.68 $\pm$ 0.80	65.37 $\pm$ 0.80	55.60 $\pm$ 0.41	81.47 $\pm$ 1.81
∞ -budget	92.05	91.40	92.30	92.23
Ours	85.73$\pm$ 0.94	82.36$\pm$ 0.91	91.58$\pm$ 0.14	91.33$\pm$ 0.72

Table 18: Default non-IID setting of Subj using different LLMs. 32 ICEs for server LLM inference.

F.9 Hyperparameter sensitivity results

The results for the sensitivity of δ and α are presented in Figure 11. The results for the sensitivity of proxy set size are shown in Figure 12.

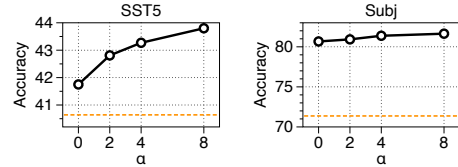


Figure 10: Additional budget α analysis. The orange dash line is the second-best baseline.

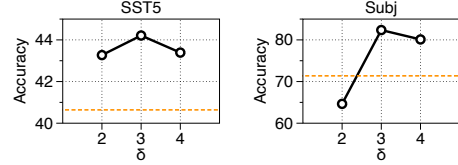


Figure 11: Budget allocator resolution δ analysis. The orange dash line is the second-best baseline.

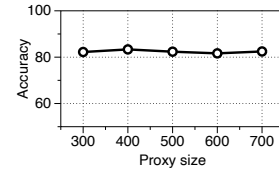


Figure 12: Proxy size analysis. Results for Subj using GPT-Neo-2.7B.

We also present the impact of client numbers in Table 19, and our method stays robust even when the client number increases to 8.

Algorithm	Client Number		
	2	4	8
Ours	81.07	82.36	82.40

Table 19: Results for sensitivity on client numbers using Subj.

F.10 Results for paraphrasing

The results for paraphrasing-based privacy protection solution are presented in Table 20.

Algorithm	Dataset			Avg
	SST-5	Yelp	Subj	
Zero-shot	27.96	31.40	51.55	36.97
Proxy-only	39.39 $\pm$ 1.33	31.78 $\pm$ 1.75	73.46 $\pm$ 1.46	48.21
Singleton	25.31 $\pm$ 3.89	30.78 $\pm$ 4.88	50.08 $\pm$ 0.10	35.39
Social Learning	33.09 $\pm$ 0.68	28.80 $\pm$ 0.33	74.82 $\pm$ 0.93	45.47
Uniform-budget	27.06	26.60	63.30	38.99
Random-budget	27.29 $\pm$ 0.51	27.70 $\pm$ 0.46	63.88 $\pm$ 0.81	39.62
∞ -budget	41.63	37.23	90.75	56.54
Ours	40.37$\pm$ 0.27	36.52$\pm$ 0.89	83.82$\pm$ 1.00	53.57

Table 20: Analysis of the generated query and training samples. We paraphrase the datasets using small-sized LLMs and conduct the experiments as in Table 1 under the same experimental settings.

F.11 Experiment for dataset with more classes

To show the efficacy of our method on datasets with more classes (more than 2), we add TREC (Hovy et al., 2001) (6 classes) under 4-client Non-IID setting. As shown in Table 21, our method outperforms other baselines except the upperbound ∞ -budget.

Algorithm	TREC
Zero-shot	31.40
Proxy-only	78.60
Singleton	28.00
Social Learning	83.20
Uniform-budget	69.20
Random-budget	64.60
∞ -budget	91.40
Ours	87.40

Table 21: Performance on TREC under 4-client Non-IID setting.

F.12 Hitting ratio comparison between the trained allocator and random budget allocator

We measure the accuracy of predicted budget level (our method) and the random budget level for two cases, to demonstrate the ‘hitting rate’ of budget assignment. As shown in Table 22 and Table 23, our method achieves a better hitting rate compared with the random budget.

	Client 1	Client 2	Client 3	Client 4
Our	79	60	71	78
Random	59	50	53	46

Table 22: Per-client hitting ratio comparison when $\delta = 2$ (per client budget value is ‘low’ or ‘high’), the accuracy (hitting ratio) of different budget assignments.

	Client 1	Client 2	Client 3	Client 4
Our	60	74	60	68
Random	54	52	53	51

Table 23: Per-client hitting ratio comparison when $\delta = 3$ (per client budget value is ‘low’, ‘medium’ or ‘high’), the accuracy (hitting ratio) of different budget assignments.

G Concrete Example of Distributed Non-IID ICL Scenario

A concrete example of distributed Non-IID ICL scenario can be the medical diagnosis task based on ICL cooperating with multiple medical institutions. Now, we have several medical institutions, with each institution owning some medical records (each sample consisting of the patient’s symptoms description in text and the corresponding diagnosed disease, that is, the query x and label y). These medical institutions normally do not have enough local computation power to support LLM computation requiring large GPU resources, while they can do some small-cost local computation like retrieval processes to find similar queries. At the same time, there will be a platform operating like a server in this system, with enough computation resources to support LLM inference and in charge of cooperation management between these institutions. Once the system (including the server platform and cooperating institutions) is deployed, the platform can provide consulting diagnosis services to other patients, doctors, or even other medical institutions based on pay-by-use knowledge pricing strategies. That is, the price is decided by the number of samples involved in the whole diagnosis procedure. Also, due to medical privacy concerns, the server platform can use local samples to perform inference while not allow caching these samples. Thus, these local retrieved samples cannot be cached to construct a retrieval pool on server platform. For the specific example of platform that supports LLM, OpenAI now provides ChatGPT Enterprise³, which allows the deployment requirement that the platform should not cache and utilize private data for further training.

H Discussion on Relation with Distributed RAG

Here, we discuss the differences between our approach and existing distributed RAG studies to provide additional clarity and context for our contribution.

In developing this work, we carefully considered related studies in distributed RAG. However, the challenges addressed by existing distributed RAG works differ from those tackled in our paper. For instance, (Wang et al., 2024) focuses on the creation of datasets for distributed RAG frameworks and explores LLM-based labeling techniques for engineering pipelines. Their research scope and methodology are distinct from ours and are not directly applicable to our specific problem setting. Similarly, (Li et al., 2024) addresses resource consumption and real-time response challenges in distributed RAG, emphasizing local retrieval efficiency and answer accuracy. However, it does not account for the non-IID property in distributed settings. Additionally, (Li et al., 2024) permits LLM deployment on partial local institutions, which is fundamentally different from our setting.

Real-world distributed non-IID RAG scenarios present a more complex framework involving numerous challenges that must be addressed for effective deployment. For example:

³<https://openai.com/enterprise-privacy/>

- How can we effectively decompose a user query into subqueries while considering local knowledge distribution?
- What is the best way to assign these subqueries to clients with varying local expertise?
- How should we merge knowledge retrieved from multiple clients with overlapping expertise, and should we assign confidence levels to different clients for the same subqueries?
- How can the local retrieval process be accelerated when dealing with large local databases?

These challenges represent broader avenues for exploration in distributed non-IID RAG. While our current work cannot be directly compared with existing distributed RAG studies due to different settings, we believe it offers an interesting starting point for addressing such challenges. Specifically, our approach focuses on how to enable cooperation among clients with varying distributions of knowledge. By assigning preferences to clients based on their local knowledge distributions and employing an MLP to learn these distributions without transmitting complete local knowledge to a central server, we offer an intuitive method that could inspire future advancements in distributed non-IID RAG.

Original	Paraphrased
A friend of mine suggested we go here today before our movie. I was planning on suggesting another place, but she got there early and got a table DARN! I don't hate Red Robin...I think I avoid it because I am not a big fan of hamburgers. Seems like more of a place for straight guys and kids if you ask me, but my experience today wasn't too bad. Our waiter was really nice however I think that may have been a result of my push up bra. Ordered the Crispy Chicken Salad which also had hard boiled egg, bacon, and veggies and was very good. I'll probably get that again next time someone drags me here. Get it with that Italian dressing. Yummy! My friend ordered us onion rings as an appetizer I hate onions, but those were good! Not my first choice but good crowd pleaser with more decent food than most chains. Huge rooms. Soft towels. Comfortable bed with tons of pillows. Three things that make me happy when I'm traveling. What's wrong with this place, then? First of all, the location is central to nothing aside from the mall which houses several chain restaurants. Second, the staff could not care less about pretty much anything. While the woman who checked us in was helpful and friendly, that was the end of that. Despite requesting my Hilton status "amenities" (just a package of snacks and water...not much), I never received it. I was also charged \$30 for internet when I shouldn't have paid a thing (got the bill corrected at check out). And housekeeping didn't leave enough towels or replace the water. Considering we paid over 200 a night, the blasé attitude of the staff is very disappointing. You can either pay \$5.50 for 3-day movie rentals or get a \$40 membership and pay \$2 for 3-day/\$3 for 7-day rentals. They also offer 2-for-1 movies for students on Tuesday and Thursday. The reason I give 3 stars is because that deal isn't valid for people with memberships! I learnt this 30 seconds after paying for a membership. I'm a student and could have paid \$2.75 for movies twice a week. Instead I paid \$40 (way too much) for a membership to pay .75 cents less than the other students. Good movie selection and shop but don't fall for their rip off of a membership unless you rent daily.	I was considering going to Red Robin with a friend, but I decided to go somewhere else instead. I was not a big fan of hamburgers, but I was drawn to Red Robin because of the free tastings. I was pleasantly surprised by the friendliness of the staff, especially the waiter who was very attentive and helpful. I ordered the Crispy Chicken Salad, which was very good. I'll probably go back there next time someone drags me there. 3 things that make me happy when I travel are: spacious rooms, friendly staff, and a comfortable bed with lots of pillows. However, the location is inconvenient, the staff is indifferent, and the amenities are subpar. 3-day movie rentals cost either \$5.50 or \$40 for a membership. They offer a deal for students on Tuesdays and Thursdays, but it's not valid for those with memberships.

Table 11: Paraphrased examples of Yelp dataset

Dataset	Prompt	Label	Label Template	Example
AGNews	Topic of the text:	{World, Sports, Business, Technology }	Topic: <i>Label</i>	REDMOND, Wash. - Microsoft Corp. and cable television provider Comcast Corp. said Monday they would begin deploying set-top boxes powered by Microsoft software starting next week. \n Topic: Business \\\nOil demand is rising faster than predicted this year as OPEC pumps more low-quality oil in a failed bid to reduce record prices, according to International Energy Agency, an adviser to 26 industrialized nations. \n Topic:
MR	Sentiment of the sentence:	{great, terrible}	It was <i>Label</i>	"Analyze That" is one of those crass, contrived sequels that not only fails on its own but makes you second-guess your affection for the original. \n It was terrible ...about the only thing to give the movie points for is bravado—to take an entirely stale concept and push it through the audience's meat grinder one more time.\n It was
SST-5	Sentiment of the sentence:	{great, good, okay, bad, terrible}	It was <i>Label</i>	a strong, funny, and finally transporting re-imagining of Beauty and the Beast and 1930s horror films \n It was great ...no movement, no yuks, not much of anything. \n It was
Subj	Subjectivity of the sentence:	{subjective, objective}	It's <i>Label</i>	gangs, despite the gravity of its subject matter, is often as fun to watch as a good spaghetti western. \n It's subjective ...smart and alert, Thirteen Conversations About One Thing is a small gem. \n It's
Amazon	Sentiment of the sentence:	{great, good, okay, bad, terrible}	It was <i>Label</i>	Love the originality of this music. Because she is ever-changing, Madonna is never boring. "Music" makes you want to dance - totally energizing! Wish the "Music" video was half as impressive as this work of art. ... The case for the DVDs were a bit damaged. The damage did not compromise the DVDs, ... \n It's
Yelp	Subjectivity of the sentence:	{subjective, objective}	It's <i>Label</i>	My family visited ceasars palace and ate here. Our waiting time was only ten minutes despite all the people. Our server was the best. He recommended several great dishes. The food was higher than our expectations....his place is great. The staff is really friendly and the chile verde burrito is fantastic. You know a ... \n It's
Yahoo	Topic of the sentence	{Society & Culture, Science & Mathematics, Health, Education & Reference, Computers & Internet, Sports, Business & Finance, Entertainment & Music, Family & Relationships, Politics & Government}	It's <i>Label</i>	who sang message in a bottle? Answer: Sting (the Police) ...Neuroscience Question? Answer: no it is called motor neuroprostheses. \n Topic:

Table 24: Prompt and instructions used for each dataset. We denote **examples** in blue and **queries** in red.