

Effective Large Language Model Adaptation for Improved Grounding

Anonymous ACL submission

Abstract

Large language models (LLMs) have achieved remarkable advancements in natural language understanding and generation. However, one major issue towards their widespread deployment in the real world is that they can generate "hallucinated" answers that are not factual. Towards this end, this paper focuses on improving LLMs by grounding their responses in retrieved passages and by providing citations. We propose a new framework, *AGREE*, Adaptation for **G**Rounding **E**nhanc**E**ment, that improves the grounding from a holistic perspective. We start with the design of a test-time adaptation capability that can retrieve passages to support the claims that have not been grounded, which iteratively improves the responses of LLMs. To effectively enable this capability, we propose tuning LLMs to self-ground the claims in their responses and provide accurate citations. This tuning on top of the pre-trained LLMs requires well-grounded responses (with citations) for paired queries, for which we introduce a method that can automatically construct such data from unlabeled queries. Across five datasets and two LLMs, results show that our tuning-based *AGREE* framework generates superior grounded responses with more accurate citations compared to prompting-based approaches and post-hoc citing-based approaches.

1 Introduction

Recent advancements in large language models (LLMs) have yielded demonstrably groundbreaking capabilities in natural language processing (NLP) (Brown et al., 2020; Chowdhery et al., 2022). Their ability to understand, generate, and manipulate text at unprecedented scales and depths has established them as a transformative force within the burgeoning field of artificial intelligence, poised to significantly impact our increasingly data-driven world. Despite their widely spread adoption,

one prominent issue of LLMs is that in certain scenarios they hallucinate: they generate plausible-sounding but nonfactual information (Maynez et al., 2020; Ji et al., 2023; Menick et al., 2022), limiting their the applicability in real-world settings. To mitigate hallucinations, solutions generally rely on grounding the claims in LLM-generated responses to supported passages by providing an attribution report (Rashkin et al., 2023; Bohnet et al., 2022; Gao et al., 2023a) or adding citations to the claims (Liu et al., 2023; Gao et al., 2023b; Huang and Chang, 2023).

There has been a growing amount of interest in making LLM-generated responses more trustworthy via grounding. One line of work follows the retrieval-augmented generation (Chen et al., 2017; Guu et al., 2020; Lewis et al., 2020) framework, in which LLMs are presented with retrieved passages or passages, and are instructed to include grounded responses in their answers via instruction tuning (Kamalloo et al., 2023) or in-context learning (Gao et al., 2023b; Kamalloo et al., 2023). As LLMs are required to perform this challenging task from just instructions or few-shot demonstrations, such directions often lead to mediocre grounding quality (Gao et al., 2023b). Another line of work is on post-hoc citing (Gao et al., 2023a; Chen et al., 2023), which links support passages to the claims in responses using an additional attribution evaluation model. This paradigm heavily relies on LLMs' parametric knowledge and might not scale to less known knowledge (Sun et al., 2023).

We propose a new framework, *AGREE*, Adaptation of LLMs for **G**Rounding **E**nhanc**E**ment. As shown in Fig. 1, our framework adopts a learning-based approach to finetune LLMs, as opposed to relying on an external NLI model post-hoc. At the training phase, *AGREE* collects well-grounded responses for unlabelled queries automatically from a base LLM with the help of an attribution evaluation model (an

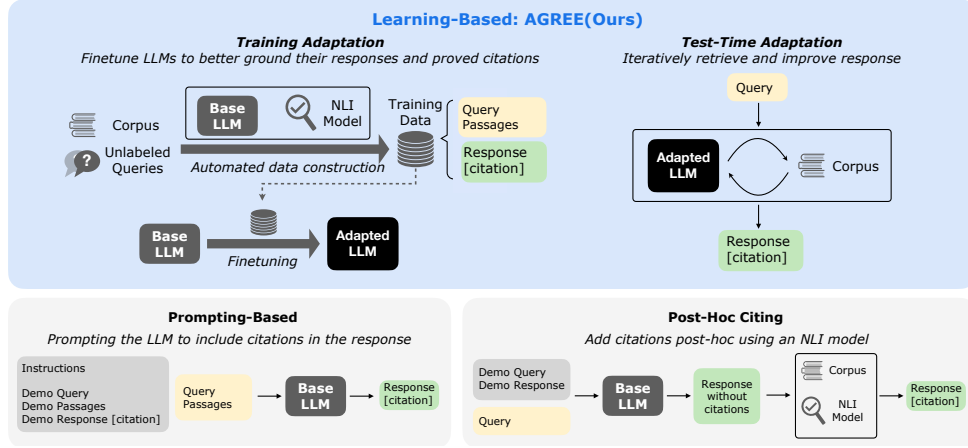


Figure 1: Our framework, AGREE, combines tuning and test time adaptation for better grounding.

NLI model). Next, the collected data is used for supervising LLMs to generate grounded responses based on the retrieved passages as well as include citations in their responses. At the testing phase, we propose an iterative inference strategy that allows LLMs to seek for additional information based on the self-grounding evaluation so as to refine its response. The tuning and test time adaptation together enable LLMs to effectively and efficiently ground their responses in the corpus.

We apply our AGREE framework to adapt an API-based LLM, text-bison, and an open LLM, llama-2-13b, with training data collected using unlabelled queries from three datasets. We conduct evaluation on both in-domain datasets and out-of-distribution datasets and compare our AGREE framework against competitive in-context learning and post-hoc citing baselines. The experimental results highlight that AGREE framework successfully improves grounding (citation recall) and citation precision compared to the baselines by a substantial margin (generally more than 20%). We find LLMs can learn to add accurate citations to their responses with our carefully designed tuning mechanisms. Furthermore, the improvements in grounding quality achieved by tuning using certain datasets can generalize well across domains.

2 Related Work

Hallucination is a prevalent issue for generative language models on many tasks (Maynez et al., 2020; Raunak et al., 2021; Dziri et al., 2021; Ji et al., 2023; Tang et al., 2023; Huang and Chang, 2023), which leads to the development of systematic evaluation of the grounding in generated responses (Bohnet et al., 2022; Rashkin et al., 2023;

Min et al., 2023; Yue et al., 2023). Among these, our work focuses on providing citations to attributable information source (Liu et al., 2023; Gao et al., 2023b) for responses generated by LLMs. Unlike existing work that largely relies on zero-shot prompting or few-shot prompting (Kamalloo et al., 2023; Gao et al., 2023b), we use a learning-based approach that tunes LLMs to generate better-grounded responses supported with citations. Furthermore, our approach teaches LLMs to cite the passages themselves as opposed to using an additional attribution evaluation model (Gao et al., 2023a; Chen et al., 2023).

Our framework improves LLMs for better grounding, as a form of a retrieval augmented generation approach. This differentiates work from recent work that improves factuality of LLMs without using retrieved passages by inference-time intervention (Li et al., 2023; Chuang et al., 2023), cross-exam (Cohen et al., 2023; Du et al., 2023), or self-verify (Dhuliawala et al., 2023), which cannot provide references. While past work have explored using retrieval to improve LLM generation quality (Chen et al., 2017; Lewis et al., 2020; Guu et al., 2020; Izacard and Grave, 2020; Shi et al., 2023) or factuality (Shuster et al., 2021; Jiang et al., 2023; Liangming Pan, 2023), our approach further uses self-generated attribution evaluation to provide citations and guide retrieval.

3 Problem & Background

We focus on enabling LLMs to provide grounded responses – given an input text query Q and a corpus $\mathcal{D} = \{d_i\}$ consisting of text passages, we aim to generate A to the query that is factually grounded in the corpus $\mathcal{D}$. Our framework tunes a pre-trained

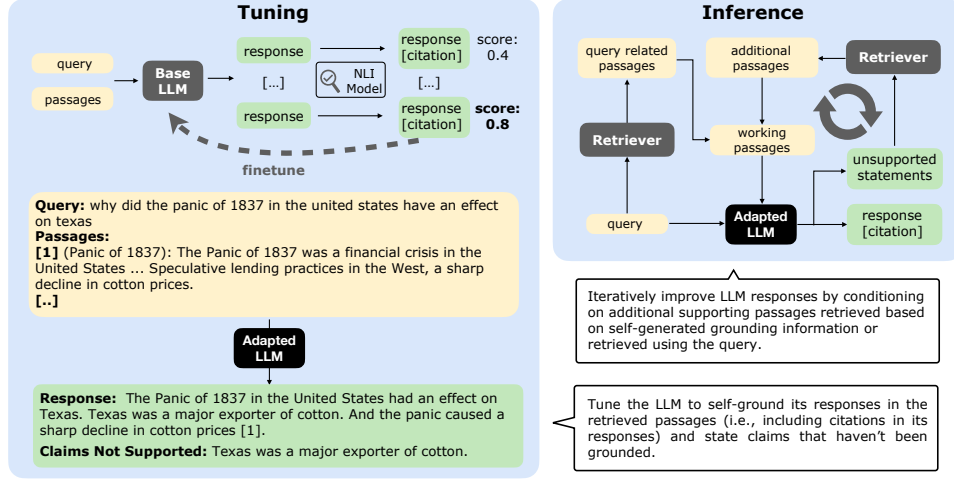


Figure 2: Illustration of the proposed AGREE framework. A base LLM is tuned to generate grounded responses that include citations. At test time, self-generated grounding evaluation is used to iteratively improve the responses.

LLM $\mathcal{M}^B$ to an adapted LLM $\mathcal{M}^A$, primarily for improving grounding with respect to the corpus $\mathcal{D}$, which we evaluate using the grounding score denoted as $\mathcal{G}$.

Grounding evaluation Let $s_1, \dots, s_n$ be the statements in the answer $A = s_1, \dots, s_n$. In line with prior work (Rashkin et al., 2023; Gao et al., 2023a,b), we evaluate the grounding of an answer by assessing whether the statements in A can be attributed to the corpus $\mathcal{D}$. That is, we require the answers to be associated with citations $\mathcal{C} = \{E_i, \dots, E_n\}$; each statement s_i to be linked a set of evidence passages $E_i \subset \mathcal{D}$. Then, the grounding quality of A can be quantified by:

$$\mathcal{G}(A, \mathcal{C}) = \frac{1}{n} \sum_i \phi(\text{concat}(E_i), s_i),$$

where ϕ is an attribution evaluation model that assesses whether the concatenated passage $\text{concat}(E_i)$ supports s_i .

4 AGREE Framework

The proposed AGREE framework takes a holistic perspective for grounding, introducing a test-time adaptation (TTA) mechanism to generate grounded outputs that are supported from the corpus, and proposing a model tuning approach to better align the pre-trained LLM with the desired TTA capability. In the following section, we first introduce the inference procedure of our method based on the proposed iterative TTA approach, and then explain in detail how we tune the base LLM in order to enable such iterative TTA capability.

4.1 Test-time adaptation

At a high level, our framework is a form of retrieval augmented generation framework to output grounded responses. The inference procedure is overviewed in Algorithm (1). At the core of our approach lies an LLM that is able to answer a query based on a set of given passages retrieved from the corpus, and, more importantly, self-ground its response to add citations to the passages as well as to find unsupported statements needing further investigation. Using these capabilities, the LLM can guide the process of iteratively, constructing a set of relevant passages from the large corpus $\mathcal{D}$ to refine its response to the query.

As shown in Algorithm (1), given a query Q and the corpus $\mathcal{D}$, we first retrieve based on the query to obtain an initial set of working passages. Next, we employ the following procedure iteratively until we consume all the budget B of invoking LLM calls. At each iteration, the LLM generates a response to the query based on the working passages, adds citations to its response, and finds out any unsupported statements that do not have citations. Then, we add the cited passages to the list of relevant passages. Lastly, at each iteration, we update the working passages – if there are unsupported statements, we include additional information retrieved based on the unsupported statements (ln 11), otherwise, we include more passages that are retrieved based on the query to acquire more complete information regarding the query (ln 13). Note that at each iteration, we let the LLM to re-generate a response based on the current working passages instead of editing from previous one, which we observed lead to better fluency.

Algorithm 1 Iterative TTA

```

1: procedure ITERATIVEINFERENCE( $Q, \mathcal{D}, \mathcal{M}^A, k, B$ )
   input: A query  $Q$ , text corpus  $\mathcal{D}$ , the LLM for generating grounded response  $\mathcal{M}^A$ , the number of passages  $k$  that  $\mathcal{M}^A$  can take as input, the budget for LLM calls  $B$ 
2:    $relevant\_psgs = []$ 
3:    $\triangleright$  retrieve passages using the query
4:    $working\_psgs := \text{RETRIEVE}(Q, \mathcal{D})[:k]$ 
5:   while  $iter = 1 : B$  do
6:      $\triangleright$  Use the LLM to generate an answer  $A$  for the query  $Q$  based on the working psgs  $\mathcal{D}$ . Additionally obtain the cited passages and unsupported the sentences.
7:      $A, cited\_psgs, unsup\_sents := \mathcal{M}^A(Q, working\_psgs)$ 
8:      $\triangleright$  add cited passages to the list of relevant passages and de-duplicate the list
9:      $relevant\_psgs := \text{DEDUPLICATE}(relevant\_psgs + cited\_psgs)$ 
10:    if  $unsup\_sents$  is not None then
11:       $working\_psgs := \text{DEDUPLICATE}(relevant\_psgs + \text{Retrieve}(unsup\_sents, \mathcal{D}))[:k]$ 
12:    else
13:       $working\_psgs := \text{DEDUPLICATE}(relevant\_psgs + \text{Retrieve}(Q, \mathcal{D}))[:k]$ 
14:  return  $A, cited\_psgs$ 

```

The design of our proposed TTA enables efficient and flexible inference. We rely on the LLM to generate citations itself, which has the advantage of reduced overhead of invoking another attribution evaluation model in a post-hoc way. Also, as we iteratively refine the answer, such a process can be streamed and flexibly controlled by setting a budget in deployment.

4.2 Tuning LLMs

Recall that our TTA requires the LLM to be able to self-ground its response, which we achieve by tuning the base LLM using data automatically created with the help of the attribution evaluation model.

Generating self-grounded responses We adapt the base LLM to generate grounded responses with citations. Our method is able to grant LLMs such an ability using only a collection of *unlabeled* queries Q and an attribution evaluation model ϕ . As we are using unlabeled queries, we formulate the adaptation task as tuning LLMs to achieve better grounding without heavily deviating from the original generations (this idea of preservation has also been adopted in recent work (Gao et al., 2023a)). Conceptually, we adapt $\mathcal{M}^B$ to $\mathcal{M}^A$ so that the answers generated by the adapted LLM $\mathcal{M}^A$ should satisfy the grounding constraints (with grounding score $> \tau_G$) while maximizing the scores with respect to the base LLM $\mathcal{M}^B$:

$$\max \mathbb{E}_{A \sim \mathcal{M}^A(\cdot | Q, \mathcal{D})} \mathcal{M}^B(A | Q, \mathcal{D}) \mathbb{1}\{\mathcal{G}(A, C) \geq \tau_G\}$$

This leads to a data-centric approach for optimizing $\mathcal{M}^A$. Since grounding score can vary with different query characteristics (e.g., responses for more open-ended questions are generally associated with lower grounding scores compared to factoid questions), instead of discarding all data points with

grounding scores below a hard threshold, we instead encourage LLMs to generate a maximally-grounded response given a question.

Given the query, we first sample responses $\{A\}$ from the base LLM $\mathcal{M}^B(\cdot | Q, \mathcal{D})$ using instruction following (see Appendix A for details). For each $A = s_1, \dots, s_n$ we create citations $\mathcal{C} = \{E_i\}$ using the attribution evaluation model, ϕ , to link a sentence s_i to the maximally supported passage $e_i = \max_{d \in \mathcal{D}} \phi(d, s_i)$ if the passage e_i actually support s_i (i.e., $\phi(e_i, s_i) > \tau$). Otherwise, we do not add a citation to s_i , and s_i is an unsupported statement. That is: $E_i = \{e_i\}$ if $\phi(e_i, s_i) > \tau$ else $\{\}$. We use U to denote the set of unsupported statements. This allows us to evaluate the grounding of A as in Section 3. Now, we can choose the best response A^* from $\{A\}$ based on the grounding scores to form a grounded response, i.e., $A^* = \arg \max_A \mathcal{G}(A, C)$.

We use $\{Q, A^*, C\}$ to teach the base LLM to generate grounded responses with citations. As shown in Fig. 2 (left), in addition to citations, we also instruct the LLM to clearly state the unsupported statements U . We note the tuning of framework does not force all training responses to be perfectly grounded. Instead, we supervise the LLM itself to identify unsupported statements. This allows the LLM to generate more flexibly and guide the retrieval process with its knowledge.¹

Supervised fine-tuning We have introduced how we construct supervision to instruct the LLM to add citations $\mathcal{C}$ and state unsupported statements U in its response. To effectively tune the LLM, we verbalize the entire process in natural language. We denote the verbalized natural language descrip-

¹Please refer to Appendix A for more details on the tuning method.

Dataset	Type	Corpus	#
Train			
NQ	Factoid QA	Wiki	2500
StrategyQA	Multi-hop QA	Wiki	1000
Fever*	Fact Checking	Wiki	1000
In-Distribution Test			
NQ	Factoid QA	Wiki	700
StrategyQA	Factoid QA	Wiki	460
Out-of-Distribution Test			
ASQA	Ambiguous QA	Wiki	948
QAMPARI	Multi-answer QA	Wiki	1000
Enterprise	Customer Support QA	Enterprise	580

Table 1: Statistics used for adaptation and test datasets. In addition to in-domain test datasets, we also investigate the generalization to out-of-distribution datasets that exhibit different reasoning processes or different corpus types.

tion as $\text{VERB}(A^*, \mathcal{C}, U)$ (see Fig. 2 for a concrete example). The natural language formalization also allows us to conveniently tune the LLM with standard language modeling objectives:

$$\mathcal{M}^A = \arg \max_{\mathcal{M}} \sum_Q \mathcal{M}(\text{VERB}(A^*, \mathcal{C}, U) \mid Q, \mathcal{D})$$

Training data We use multiple datasets to construct the adaptation data used to tune the pre-trained LLM, including Natural Questions (NQ) (Kwiatkowski et al., 2019), FEVER (Thorne et al., 2018), and StrategyQA (Geva et al., 2021). We choose these as they contain diverse text, and the answers to the corresponding queries require different types of reasoning processes: NQ provides diverse queries naturally asked by real users; FEVER places a particular emphasis on fact verification; and StrategyQA requires multi-hop reasoning with implicit strategy. It is worthwhile to note that AGREE *only* uses queries, leaving out ground-truth answers, to improve LLMs.

5 Experiments

5.1 Setup

Evaluation datasets We conduct comprehensive evaluation on 5 datasets. In addition to the two in-domain test sets, NQ and StrategyQA (we leave out the non-QA dataset, FEVER), we further test the generalization of adapted LLMs on 3 out-of-domain datasets, including ASQA (Stelmakh et al., 2022), QAMPARI (Amouyal et al., 2022), and an Enterprise dataset. In particular, ASQA and QAMPARI contain questions of ambiguous answers and multiple answers. The Enterprise dataset is a proprietary dataset which requires provided answers

that are grounded in customer service passages. Such an evaluation suite allows assessing the generalization capability of the adapted LLMs for OOD question types (ASQA and QAMPARI) as well as to an entirely different corpus (Enterprise).

Models We demonstrate AGREE framework with two LLMs, text-bison and LLaMA-2-13B (Touvron et al., 2023). We use GTR-large (Ni et al., 2021) as our retriever, and use TRUE (Honovich et al., 2022) as the attribution evaluation model.

Baselines We evaluate the effectiveness AGREE in two settings, invoking LLMs once, without TTA; and invoking LLMs multiple times, with the proposed TTA. We compare with three baselines from recent work, including one prompting-based approach and two post-hoc citing approaches, described below.

Few-shot In-Context Learning (ICLCITE): Following Gao et al. (2023b), we prompt LLMs with few-shot examples (Gao et al., 2023b), each consisting of a query, a set of retrieved passages, and an answer with inline citations. The LLMs can therefore learn from the in-context examples and generated citations in the responses. It is worthwhile to note that **ICLCITE do have access to retrieved passages**.

Post-hoc search (POSTSEARCH): Following Gao et al. (2023b), given a query, we first instruct LLMs to answer the query *without* passages, and then add citations in a post-hoc way via searching. We link each claim in the response to the most relevant passage retrieved from a set of query-related passages. This baseline only uses retriever but not the attribution model, ϕ .

Post-hoc Attribution (POSTATTR): Following Gao et al. (2023a), instead of citing the most relevant passage, for each claim, we retrieve a set of k passages from the corpus, and then use the attribution evaluation model to link to the passage that maximally supports the claim. We note both baselines in the post-hoc citing paradigm only rely on LLMs’ parametric knowledge.²

Metrics We mainly focus on improving the grounding quality of generated responses, reflected by the quality of citations. Following past work (Gao et al., 2023b), we report the **citation recall** (rec) and **citation precision** (pre) for all the

²Please refer to Appendix B for more details on experimental setup.

	NQ			StrategyQA			ASQA			QAMPARI			Enterprise	
	em-rec	rec	pre	acc	rec	pre	em-rec	rec	pre	rec-5	rec	pre	rec	pre
Base model: text-bison-001														
ICLCITE	47.6	52.1	56.3	74.5	13.6	27.8	39.5	47.3	49.8	20.3	22.7	24.5	30.2	40.5
POSTSEARCH	45.1	29.7	28.7	75.5	20.1	20.1	38.4	19.2	19.2	22.5	16.2	16.2	15.9	15.9
POSTATTR	45.1	31.5	31.5	75.5	18.4	18.4	35.1	38.0	38.0	22.5	18.5	18.5	20.1	20.1
AGREE _{w/o} TTA	50.0	67.9	73.1	74.1	33.4	50.5	39.5	65.9	70.5	20.1	60.1	64.5	55.8	67.1
AGREE _{w/} TTA	53.1	70.1	75.0	74.9	39.2	57.9	40.9	73.2	77.0	20.9	62.9	67.1	57.2	68.6
Base model: llama-2-13b														
ICLCITE	45.8	42.8	41.6	65.5	20.6	33.1	35.2	38.2	39.4	21.0	10.2	10.4	30.6	38.8
POSTSEARCH	35.9	17.5	17.5	64.3	8.7	8.7	25.0	23.6	23.6	12.0	27.5	27.5	13.4	13.4
POSTATTR	35.9	26.0	26.0	64.3	12.5	12.5	25.0	33.6	33.6	12.0	28.9	28.9	18.7	18.7
AGREE _{w/o} TTA	47.9	50.5	56.6	65.0	25.5	35.0	35.7	50.2	55.3	17.1	40.4	43.6	50.6	53.8
AGREE _{w/} TTA	51.0	62.0	66.0	64.6	30.2	37.2	39.4	64.0	66.8	17.9	51.4	53.4	50.4	55.4

Table 2: Answer accuracy and grounding (measured by citation quality) of AGREE and baselines across 5 datasets. Our approach achieves substantially better citation grounding (measured by citation recall) and citation precision compared to the baselines.

datasets we are evaluating on. We note that **citation recall aggregates how well each sentence is supported by the citation to the corpus, which is essentially the grounding score $\mathcal{G}$** . Therefore, we prioritize on the evaluation citation recall.

We also report the correctness of the generated outputs. For NQ, we report exact match recall (em-rec; whether the short answers are substrings in the response). For StrategyQA, we report the accuracy (acc). For ASQA and QAMPARI, we use subsets from Gao et al. (2023b), and report the exact match recall (em-rec) for ASQA and recall-5 (rec-5, considering recall to be 100% if the prediction includes at least 5 correct answers) for QAMPARI. For the Enterprise dataset, we only report the citation quality as there are no ground truth answers for this dataset, and citation quality reflects whether the model can provide accurate information.

5.2 Main results

Tuning is effective for superior grounding Table 2 summarizes the results obtained using our AGREE framework and compares with the baselines. As suggested by the results, across 5 datasets, AGREE can generate responses that are better grounded in the text corpus and provide accurate citations to its response, substantially outperforming all the baselines. When tuned with high-quality data, LLMs can effectively learn to self-ground their response without needing an additional attribution model. By contrast, ICLCITE, which solely relies on in-context learning, cannot generate citations as accurately as a tuned LL, as suggested by the large gap on citation precision be-

tween ICLCITE and AGREE. We also observe similar findings as suggested by Gao et al. (2023b): POSTCITE often leads to poor citation quality – without being conditioned on passages, the response from POSTCITE often cannot be paired with passages that lead to high attributions score for the generated claims.

The performance improvements can generalize

Recall that we adapt the base LLM only using in-domain training sets (NQ, StrategyQA, and FEVER), and directly test the model on out-of-distribution (OOD) test set (ASQA, QAMPARI, Enterprise). The results suggest that the improvements obtained from training on in-domain datasets can effectively generalize to OOD datasets that contain different question types or use different types of corpus. This is a fundamental advantage of the proposed approach – AGREE can generalize to a target domain in the zero-shot setting without needing any samples from the target domain, which is needed for ICLCITE.

TTA improves both grounding and answer correctness

The comparison between AGREE without and with TTA highlights the effectiveness of our iterative TTA strategy. We observe improvements in terms of both better grounding and accuracy. For instance, TTA improves llama-2 answer correctness by 3.1 and 3.7 on NQ and ASQA, respectively. Such improvements can be attributed to the fact that our TTA allows the LLMs to actively collect relevant passages to construct better answers following the self-grounding guidance.

	NQ			StrategyQA			ASQA			QAMPARI			Enterprise	
	em-rec	rec	pre	acc	rec	pre	em-rec	rec	pre	rec-5	rec	pre	rec	pre
Base model: text-bison-001														
ICLCITE	47.6	52.1	56.3	74.5	13.6	27.8	39.5	47.3	49.8	20.3	22.7	24.5	30.2	40.5
AGREE ^{Multi-dataset} _{w/o TTA}	50.0	67.9	73.1	74.1	33.4	50.5	39.5	65.9	70.5	20.1	60.1	64.5	55.8	67.1
AGREE ^{NQ-only} _{w/o TTA}	49.4	62.3	69.1	74.1	33.0	45.5	38.4	56.0	64.5	19.1	43.7	49.5	40.5	59.2
Base model: llama-2-13b														
ICLCITE	45.8	42.8	41.6	65.5	20.6	33.1	35.2	38.2	39.4	21.0	10.2	10.4	30.6	38.8
AGREE ^{Multi-dataset} _{w/o TTA}	47.9	50.5	56.6	65.0	25.5	35.0	35.7	50.2	55.3	17.1	40.4	43.6	50.6	52.8
AGREE ^{NQ-only} _{w/o TTA}	48.1	47.4	53.6	62.1	25.0	30.2	35.0	44.0	51.2	15.7	33.1	38.0	44.7	49.2
AGREE ^{Distill} _{w/o TTA}	47.9	59.1	65.1	64.4	30.5	41.1	35.2	58.5	65.2	17.9	52.5	52.7	48.1	55.9

Table 3: Analysis on the impact of training data. Training with multiple datasets (AGREE^{Multi-dataset}) leads to better grounding (citation recall) and better citation precision across datasets, compared to training using the NQ dataset (AGREE^{NQ-only}). The citation quality of a less capable model llama-2-13b can also benefit from tuning using outputs from a more capable model (text-bison-001).

	# Tok: LLM	# Tok: NLI (T5-11B)
ICLCITE	2800	—
POSTATTR	360	3520
AGREE _{w/o TTO}	1210	—
AGREE _{w/ TTO}	4840	—

Table 4: The average computation cost (for one query) of different methods measured by number of tokens processed by the LLM and the NLI model (a T5-11B model). AGREE_{w/o TTO} is able to achieve better citation quality compared to ICLCITE, despite consuming less than half of the tokens needed for ICLCITE.

Discussion on answer correctness In general, AGREE_{w/ TTO} can achieve better correctness compared to ICLCITE. AGREE_{w/o TTO} achieves similar answer correctness with ICLCITE, as both methods are conditioned on the same set of passages. As a result, the quality of passages heavily intervenes on the correctness of the answers. Unlike AGREE and ICLCITE, POSTATTR purely relies on the parametric knowledge of the LLMs to answer the query. As a result, POSTATTR generally achieves inferior answer correctness compared to AGREE and ICLCITE on these two LLMs, especially on the less capable LLM, llama-2-13b, that has less accurate knowledge compared to bison. Moreover, on the Enterprise dataset which contains very specific information, POSTATTR utterly fails to recall attributable information from LLMs’ parametric knowledge.

Discussion on LLMs Our approach successfully adapts both text-bison-001 and llama-2-13b. llama is generally less capable compared to bison, underperforming bison in terms of answer correctness and citation quality. Still, AGREE also consistently outperforms the baseline, generating more

grounded answers as well as providing more precise citations. This shows our tuning-based adaptation is model-agnostic and is effective across LLMs of varying capabilities.

5.3 Analysis

Efficiency Our AGREE framework finetunes the base LLM to enable self-grounding without needing for additional in-context examples or attribution model. As a result, our framework is able to achieve strong citation performance without expensive inference cost. Table 4 shows the comparison between the computation cost, measured by the number of tokens processed by the LLM and attribution model, needed for one query of our methods and that of the baselines. Compared to ICLCITE, AGREE_{w/o TTO} uses much fewer tokens due to not using additional in-context examples, but achieves significantly better citation quality (see Table 2). POSTATTR does not use retrieved passages in the prompts and hence requires less computation on the LLM compared to our framework, but it requires additional overhead of extensively invoking the NLI model (which is also of 11B parameters, a relatively large scale) to verify the each of claim based on each of the retrieved passages. The citation performance of POSTATTR also substantially lags ICLCITE and AGREE. AGREE_{w/ TTO} requires more computation compared to AGREE_{w/o TTO}, but is able to achieve both better citation quality and improvements in answer correctness.

Impact of multi-dataset training Our AGREE framework multiple datasets spanning factoid QA, multi-hop reasoning, and fact-checking to construct data for adapting the base model. We expect such a combination can grant the adapted model better

generalization to different types of questions and different text distribution. We conduct an analysis to investigate the benefits of using multiple datasets for tuning. Table 3 shows the performance of our approach trained using multi-datasets and a counterpart that is trained only on NQ data (AGREE^{NQ-only}). The results suggest that training using NQ leads to inferior citation quality compared to training on the combination of three datasets across all the datasets. The performance gap is especially significant on datasets other than NQ. Moreover, training on NQ only also leads to inferior answer correctness across all the datasets, which also suggests the benefits of using multiple datasets. Nevertheless, training only on NQ can still improve the performance compared to solely relying on in-context learning (ICLCITE).

Distilling from bison to llama-2-13b Our work mainly focuses on improving the base LLM’s grounding capability in a self-improving fashion. That is, we use the samples generated by the base LLM itself to adapt the base LLM, as opposed to distilling from proprietary models which may impose constraints on the deployment of adapted model. Nevertheless, we conduct an analysis to investigate the effectiveness of distilling the data produced by a more capable model to enhance the grounding of a less capable. In our case, we use the data generated by text-bison-001 to tune llama-2-13b. As shown in the last row of Table 3, AGREE^{Distill} achieves better citation quality compared to AGREE on llama-2-13b, as it is trained on better grounded responses produced by a more capable model. However, llama-2 finetuned on data generated by bison still can’t level the performance on bison as constrained by the gap between the capabilities of these two LLMs.

Qualitative analysis We qualitatively analyze the advantages of proposed AGREE framework compared to ICLCITE, the strongest among the baselines. We observe that on both text-bison-001 and llama-2-13b, ICLCITE achieves inferior citation quality due to failure in following the citation format (e.g., adding citations after the periods, violating the instructions), linking a statement to a relevant but un-attributable passage (as indicated by poor citation precision), and introducing more auxiliary information not mentioned in the retrieved passages (as indicated by citation recall). Our AGREE framework mitigates these issues by tuning on well-grounded responses certified by the

Example from NQ	
Query: where did the world’s largest recorded wave occur?	
Without TTA	
Passage [1]: (retrieved using the query) (Title: 1958 Lituya Bay earthquake and megatsunami) Lituya Bay has a history of megatsunami events, but the 1958 event was the first for which sufficient data was captured...	
AGREE_{w/o} TTO: The world’s largest recorded wave occurred in Lituya Bay, Alaska. The data was captured in a 1958 megatsunami event[1].	
With TTA	
Passage [2]: (retrieved using prior unsupported statement) (Title: Megatsunami) In 1958, a 7.8 strike-slip earthquake in southeast Alaska caused 90 million tonnes of rock and ice to drop into the deep water at the head of Lituya Bay. The block fell almost vertically and hit the water with sufficient force to ...	
Passage [3]: (retrieved using prior unsupported statement) (Title: Tsunami) ...Their existence was confirmed in 1958, when a giant land slide in Lituya Bay, Alaska, caused the highest wave ever recorded, which had a height of 524 metres....	
AGREE_w TTO: The world’s largest recorded wave occurred in Lituya Bay, Alaska [3]. The wave was caused by a giant landslide that occurred in 1958 [2]. The landslide displaced 90 million...[2].	
Example from ASQA	
Query: Who played the weasley brothers in harry potter?	
Without TTA	
Passage [1]: (retrieved using the query) (Title: James and Oliver Phelps) James Andrew Eric and Oliver Martyn John Phelps are identical twin British actors, best known for playing Fred and George Weasley in the "Harry Potter" film series.	
AGREE_{w/o} TTO: The Weasley brothers in Harry Potter were played by identical twin actors James and Oliver Phelps [1].	
With TTA	
Passage [2]: (additional passages retrieved using the query) (Title: Chris Rankin) ... Chris Rankin Christopher William "Chris" Rankin is a New Zealand-born British actor who is best known for playing Percy Weasley in the "Harry Potter" film franchise.	
AGREE_w TTO: James and Oliver Phelps played Fred and George Weasley in the Harry Potter film series [1]. Chris Rankin played Percy Weasley in the Harry Potter film franchise [2].	

Figure 3: Output examples of the proposed AGREE framework with text-bison-001 as the base model. TTA is able to improve the response by retrieving more relevant information to precisely support a statement (see top) or finding more passages to generate a more complete response (see bottom).

attribution evaluation model. We also provide example outputs in Fig. 3 comparing the outputs of AGREE with and without proposed TTA and observe that TTA can help find more supporting passages by active retrieving using unsupported statements (top) or iteratively find more passages to construct a more complete response (bottom).

6 Conclusion

We have introduced a novel framework, AGREE, that adapts LLM for improved grounding. Our framework tunes a pre-trained LLM to self-ground its response in retrieved passages using automatically collected data. The integrated capability for grounding their responses further enables the LLM to improve the responses at test time. Our evaluations across five datasets demonstrate the benefits of the proposed learning-based approach compared to approaches that solely rely on prompting or the parametric knowledge of LLMs.

7 Limitations

Our approach employs an automated data creation method that relies on an attribution evaluation model to create citations instead of hiring humans to annotate citations. Thus, the citation quality is dependent on the performance of the attribution evaluation model. As suggested in [Gao et al. \(2023b\)](#); [Honovich et al. \(2022\)](#), the model we use favors “fully support” and cannot effectively detect “partially support”. The adapted LLMs may favor adding “fully support” citations as a result. One solution is to curate a set of human-annotated citations for “partially support”, which we defer to future work.

Our evaluation follows prior work ([Rashkin et al., 2023](#); [Gao et al., 2023a](#)) and uses the attribution evaluation model to evaluate grounding and citation quality. Therefore, our work can encounter the same issue as in past work: the citation grounding and citation quality evaluation is limited by the capability of the NLI model.

Our approach uses created grounded responses to LLMs via supervised finetuning, as we observed this straightforward tuning technique leads to empirical strong performance. It is also possible to treat grounding as a preference and RLHF ([Ouyang et al., 2022](#)) to tune LLMs, which we leave as future work.

This work focuses only considers open domain question answering datasets focusing on information seeking and written in the English language. It is unsure how well the framework can handle other language.

Lastly, this work studies the approach of adding citations to LLM-generated response and which carries a shared risk with related research: a seemingly plausible but incorrect citation could potentially make an unsupported statement more convincing to users.

References

Samuel Joseph Amouyal, Tomer Wolfson, Ohad Rubin, Ori Yoran, Jonathan Herzig, and Jonathan Barrant. 2022. [Qampari: An open-domain question answering benchmark for questions with many answers from multiple paragraphs](#).

Bernd Bohnet, Vinh Q Tran, Pat Verga, Roei Aharoni, Daniel Andor, Livio Baldini Soares, Jacob Eisenstein, Kuzman Ganchev, Jonathan Herzig, Kai Hui, et al. 2022. Attributed question answering: Evaluation and modeling for attributed large language models. *arXiv preprint arXiv:2212.08037*.

Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel Ziegler, Jeffrey Wu, Clemens Winter, Chris Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. Language models are few-shot learners.

Anthony Chen, Panupong Pasupat, Sameer Singh, Hongrae Lee, and Kelvin Guu. 2023. Purr: Efficiently editing language model hallucinations by denoising language model corruptions. *arXiv preprint arXiv:2305.14908*.

Danqi Chen, Adam Fisch, Jason Weston, and Antoine Bordes. 2017. Reading Wikipedia to answer open-domain questions. In *Association for Computational Linguistics (ACL)*.

Aakanksha Chowdhery, Sharan Narang, Jacob Devlin, Maarten Bosma, Gaurav Mishra, Adam Roberts, Paul Barham, Hyung Won Chung, Charles Sutton, Sebastian Gehrmann, Parker Schuh, Kensen Shi, Sasha Tsvyashchenko, Joshua Maynez, Abhishek Baindoor Rao, Parker Barnes, Yi Tay, Noam M. Shazeer, Vinodkumar Prabhakaran, Emily Reif, Nan Du, Benton C. Hutchinson, Reiner Pope, James Bradbury, Jacob Austin, Michael Isard, Guy Gur-Ari, Pengcheng Yin, Toju Duke, Anselm Levskaya, Sanjay Ghemawat, Sunipa Dev, Henryk Michalewski, Xavier García, Vedant Misra, Kevin Robinson, Liam Fedus, Denny Zhou, Daphne Ippolito, David Luan, Hyeontaek Lim, Barret Zoph, Alexander Spiridonov, Ryan Sepassi, David Dohan, Shivani Agrawal, Mark Omernick, Andrew M. Dai, Thanumalayan Sankaranarayanan Pillai, Marie Pellat, Aitor Lewkowycz, Erica Oliveira Moreira, Rewon Child, Oleksandr Polozov, Katherine Lee, Zongwei Zhou, Xuezhi Wang, Brennan Saeta, Mark Diaz, Orhan Firat, Michele Catasta, Jason Wei, Kathleen S. Meier-Hellstern, Douglas Eck, Jeff Dean, Slav Petrov, and Noah Fiedel. 2022. Palm: Scaling language modeling with pathways. *ArXiv, abs/2204.02311*.

Yung-Sung Chuang, Yujia Xie, Hongyin Luo, Yoon Kim, James Glass, and Pengcheng He. 2023. Dola: Decoding by contrasting layers improves factuality in large language models. *arXiv preprint arXiv:2309.03883*.

Roi Cohen, May Hamri, Mor Geva, and Amir Globerson. 2023. [Lm vs lm: Detecting factual errors via cross examination](#). *ArXiv, abs/2305.13281*.

Shehzaad Dhuliawala, Mojtaba Komeili, Jing Xu, Roberta Raileanu, Xian Li, Asli Celikyilmaz, and Jason Weston. 2023. [Chain-of-verification reduces hallucination in large language models](#). *ArXiv, abs/2309.11495*.

779	Mia Glaese, Susannah Young, Lucy Campbell-Gillingham, Geoffrey Irving, and Nathan McAleese. 2022. Teaching language models to support answers with verified quotes . <i>ArXiv</i> , abs/2203.11147.	836
780		837
781		838
782		839
783	Sewon Min, Kalpesh Krishna, Xinxu Lyu, Mike Lewis, Wen tau Yih, Pang Wei Koh, Mohit Iyyer, Luke Zettlemoyer, and Hannaneh Hajishirzi. 2023. Factscore: Fine-grained atomic evaluation of factual precision in long form text generation. In <i>Proceedings of the Conference on Empirical Methods in Natural Language Processing (EMNLP)</i> .	840
784		841
785		842
786		843
787		844
788		845
789		
790	Jianmo Ni, Chen Qu, Jing Lu, Zhuyun Dai, Gustavo Hernandez Abrego, Ji Ma, Vincent Zhao, Yi Luan, Keith B. Hall, Ming-Wei Chang, and Yinfei Yang. 2021. Large dual encoders are generalizable retrievers . <i>ArXiv</i> , abs/2112.07899.	846
791		847
792		848
793		849
794		850
795	Long Ouyang, Jeff Wu, Xu Jiang, Diogo Almeida, Carroll L. Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, John Schulman, Jacob Hilton, Fraser Kelton, Luke E. Miller, Maddie Simens, Amanda Askell, Peter Welinder, Paul Francis Christiano, Jan Leike, and Ryan J. Lowe. 2022. Training language models to follow instructions with human feedback . <i>ArXiv</i> , abs/2203.02155.	851
796		852
797		853
798		854
799		
800		855
801		856
802		857
803		858
804	Hannah Rashkin, Vitaly Nikolaev, Matthew Lamm, Lora Aroyo, Michael Collins, Dipanjan Das, Slav Petrov, Gaurav Singh Tomar, Iulia Turc, and David Reitter. 2023. Measuring attribution in natural language generation models. <i>Computational Linguistics</i> , pages 1–64.	859
805		860
806		861
807		862
808		863
809		864
810	Vikas Raunak, Arul Menezes, and Marcin Junczys-Dowmunt. 2021. The curious case of hallucinations in neural machine translation . In <i>Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 1172–1183, Online. Association for Computational Linguistics.	865
811		866
812		
813		867
814		868
815		869
816		870
817	Weijia Shi, Sewon Min, Michihiro Yasunaga, Minjoon Seo, Rich James, Mike Lewis, Luke Zettlemoyer, and Wen tau Yih. 2023. Replug: Retrieval-augmented black-box language models . <i>ArXiv</i> , abs/2301.12652.	
818		
819		
820		
821		
822	Kurt Shuster, Spencer Poff, Moya Chen, Douwe Kiela, and Jason Weston. 2021. Retrieval augmentation reduces hallucination in conversation . In <i>Findings of the Association for Computational Linguistics: EMNLP 2021</i> , pages 3784–3803, Punta Cana, Dominican Republic. Association for Computational Linguistics.	
823		
824		
825		
826		
827		
828		
829	Ivan Stelmakh, Yi Luan, Bhuwan Dhingra, and Ming-Wei Chang. 2022. ASQA: Factoid questions meet long-form answers . In <i>Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing</i> , pages 8273–8288, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.	
830		
831		
832		
833		
834		
835		
	Kai Sun, Y. Xu, Hanwen Zha, Yue Liu, and Xinhsuai Dong. 2023. Head-to-tail: How knowledgeable are large language models (llm)? a.k.a. will llms replace knowledge graphs? <i>ArXiv</i> , abs/2308.10168.	
	Liyan Tang, Tanya Goyal, Alexander R. Fabbri, Philippe Laban, Jiacheng Xu, Semih Yavuz, Wojciech Kryściński, Justin F. Rousseau, and Greg Durrett. 2023. Understanding factual errors in summarization: Errors, summarizers, datasets, error detectors. In <i>Proceedings of ACL</i> .	
	James Thorne, Andreas Vlachos, Christos Christodoulopoulos, and Arpit Mittal. 2018. FEVER: a large-scale dataset for fact extraction and VERification . In <i>Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)</i> , pages 809–819, New Orleans, Louisiana. Association for Computational Linguistics.	
	Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. 2023. Llama 2: Open foundation and fine-tuned chat models. <i>arXiv preprint arXiv:2307.09288</i> .	
	Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz, et al. 2019. Huggingface’s transformers: State-of-the-art natural language processing. <i>arXiv preprint arXiv:1910.03771</i> .	
	Xiang Yue, Boshi Wang, Kai Zhang, Zirui Chen, Yu Su, and Huan Sun. 2023. Automatic evaluation of attribution by large language models. <i>arXiv preprint arXiv:2305.06311</i> .	

A Details of Tuning

Recall that we create the tuning data by first sampling responses from the base LLM and then using the attribution evaluation model to create citations and identify unsupported statements. We will detail the process in the following of this section.

Corpus & retriever As mentioned before, our framework is an instantiation of retrieval-augmented framework. For the datasets using Wikipedia as the corpus (NQ, StrategyQA, ASQA, and Qampari), we use the 2018-12-20 Wikipedia snapshot as the corpus and set up the retriever using GTR-large (Ni et al., 2021).

Task: You will be given a question and some search results. Please answer the question in 3-5 sentences, and make sure you mention relevant details in the search results. You may use the same words as the search results when appropriate. Note that some of the search results may not be relevant, so you are not required to use all the search results, but only relevant ones.

<Question>

Search Results:
[<Index>] <Title>
<Text>

[...]

Answer:

Figure 4: Zero-shot prompt template for sampling initial responses from the base LLM.

Sampling initial responses We sample initial responses from the base LLM using instruction following in a *zero-shot fashion*. Given a query, we present the base LLM with query and retrieved passages appended after an instruction that requires the base LLM to answer the query based on the passages; see Fig. 4 for the template of the zero-shot prompt. We note that we opt to use a zero-shot prompt as opposed to a task-specific few-shot prompt since 1) this can avoid biasing the generation with the few-shot in-context examples, and 2) this matches the expected scenario for deploying the adapted LLM to handle new queries in a zero-shot fashion.

For text-bison-001, we sample 4 responses using a temperature of 0.5. For llama-2-13b, we sample 4 responses using nuclear sampling (Holtzman et al., 2019) with $p=0.95$.

Adding citations and identifying unsupported documents After obtaining the initial response

Input
Task: You will be given a question and some search results. You are required to perform the following steps.
First, please answer the question in 3-5 sentences, and make sure you mention relevant details in the search results. You may use the same words as the search results when appropriate. Note that some of the search results may not be relevant, so you are not required to use all the search results, but only relevant ones. If you use the provided search results in your answer, add [n]-style citations.
Next, review your response and find the unsupported sentences that do not have citations.
<Question>
Search Results: [<Index>] <Title> <Text>
[...]
Output
Answer: <Response with citations>
Sentences Not Supported by Citations: <Unsupported statements>

Figure 5: Verbalization template for creating the training data for adapting the base LLM.

$\{A\}$ from the base LLM. We break each response A into sentences into $s_1, \dots, s_i$. For each s_i , we find the maximally supported passage e_i (scored by $\phi(e_i, s_i)$) that the base LLM has seen during generating the initial responses. We link e_i to s_i if $\phi(e_i, s_i) > 0.7$ to encourage more precise citations. For a sentence s_i if there does *not* exist an e_i such that $\phi(e_i, s_i) > 0.5$ (the decision boundary for entailment), we add s_i to the unsupported statement set U .

Verbalizing We show the template for verbalizing the data used to tune the LLM in Fig. 5. As shown in the figure, we verbalize the citations in enclosed box brackets that are added at the end of sentences (before periods) like [n], and verbalize unsupported statements after the responses.

B Details of Experiments

For tuning, we use LORA tuning (Hu et al., 2022) in experiments on both text-bison-001 and llama-2-13b. For bison, we use API to perform tuning.³ and follow all the default hyper-parameters except for training steps. We set 10% data created

³<https://cloud.google.com/vertex-ai/docs/generative-ai/models/tune-text-models-supervised>

as development data and choose to use a training step of 1000 (chosen from 500, 1000, and 2000). For llama-2, we use the huggingface transformers (Wolf et al., 2019) chat-version checkpoint.⁴ We find the chat-version achieves better performance than the base checkpoint in our preliminary investigation. We set lora_r to be 32, and only choose to use a learning rate of 1e-5 (chosen from 1e-4 and 1e-5) using the development set. We finetune llama-2 on two A100 (40GB) GPU for 4 epochs.

Our evaluation uses official code from ALCE (Gao et al., 2023b), we use the same data split and prompt template from ALCE. We use temperature 0.25 for evaluation on both bison and llama. We use one sample for evaluation since adapted LLMs tend to generate better-grounded response exhibiting less variation.

C Comparison to ICLCITE on More Capable LLMs

	ASQA			QAMPARI		
	em-rec	rec	pre	rec-5	rec	pre
Base model: llama-2-13b						
AGREE _{w/o} TTA	35.7	50.2	55.3	17.1	40.4	43.6
AGREE _{w/} TTA	39.4	64.0	66.8	17.9	51.4	53.4
Base model: llama-2-70b						
ICLCITE	41.5	62.9	61.3	21.8	15.1	15.6
Base model: ChatGPT-0301						
ICLCITE	40.4	73.6	72.5	20.8	20.5	20.9

Table 5: Comparing AGREE on llama-2-13B against ICLCITE on llama-2-70B and ChatGPT-0301. We directly quote results from ALCE.

Table 5 compares AGREE using llama-2-13B as the base model against ICLCITE on more capable models. We directly use the results from ALCE (Gao et al., 2023b). Our framework is able to substantially shorten the gap between a small llama-2 model and much more capable LLMs.

D License of Datasets

The licenses datasets used in our work include:

- NQ (Kwiatkowski et al., 2019) under Creative Commons Share-Alike 3.0 license.
- StrategyQA (Geva et al., 2021) under MIT License.

- Fever (Thorne et al., 2018) under Creative Commons Share-Alike license.
- Ambiguous QA (Stelmakh et al., 2022) under Creative Commons Share-Alike 3.0 license.
- Qampari (Amouyal et al., 2022) under Creative Commons Zero v1.0 Universal license.

⁴<https://huggingface.co/meta-llama/Llama-2-13b-chat-hf>