

# UC Berkeley

## UC Berkeley Electronic Theses and Dissertations

**Title**

Moving Computation From Pretraining to Test-time

**Permalink**

<https://escholarship.org/uc/item/8456d2nx>

**ISBN**

9798189276057

**Author**

Snell, Charlie Victor

**Publication Date**

2026-05-01

Peer reviewed|Thesis/dissertation

Moving Computation From Pretraining to Test-time

by

Charlie Victor Snell

A dissertation submitted in partial satisfaction of the

requirements for the degree of

Doctor of Philosophy

in

Computer Science

in the

Graduate Division

of the

University of California, Berkeley

Committee in charge:

Professor Dan Klein, Chair  
Associate Professor Jacob Steinhardt  
Assistant Professor Aviral Kumar

Spring 2026

# Moving Computation From Pretraining to Test-time

Copyright 2026  
by  
Charlie Victor Snell

## Abstract

## Moving Computation From Pretraining to Test-time

by

Charlie Victor Snell

Doctor of Philosophy in Computer Science

University of California, Berkeley

Professor Dan Klein, Chair

Over the last few years, the massive scaling up of resources poured into the pretraining of large language models (LLMs) has led to rapid advances in LLM capabilities. However, performance on the most challenging reasoning-heavy tasks scales far too slowly as a function of pretraining alone to achieve reasonable performance at a practical scale. A different approach to scaling up computation is therefore needed in order to meaningfully improve performance on these tasks.

In this thesis we aim to understand when and if scaling up computation at test-time with LLMs can fulfill this gap, making for an effective alternative to scaling up pretraining compute. We begin by demonstrating the necessity of scaling up pretraining resources in order to broadly improve model capabilities. From here we turn towards test-time scaling as a potential substitute. In particular, we carry out one of the first careful empirical studies comparing the efficacy of different strategies for scaling up computation at test-time with scaling up pretraining resources. We show that indeed scaling test-time compute with LLMs can be preferable to pretraining in certain settings. However, there remain several important limitations to scaling up test-time compute that make it not quite one-to-one comparable to pretraining. In particular, test-time scaling comes with an increased latency cost and treats each prompt independently, resulting in redundant computation across related queries. Finally, we propose two novel methods – Context Distillation and Sleep-time compute – which can be used to overcome these limitations with test-time computation. These techniques work by applying computation ahead of time to prepare for a specific distribution of upcoming tasks, such that at test-time the model can respond quickly without meaningfully sacrificing performance.

# Contents

|                                                                                     |            |
|-------------------------------------------------------------------------------------|------------|
| <b>Contents</b>                                                                     | <b>i</b>   |
| <b>List of Figures</b>                                                              | <b>iii</b> |
| <b>List of Tables</b>                                                               | <b>xii</b> |
| <b>1 Introduction</b>                                                               | <b>1</b>   |
| 1.1 Moving Computation From Pretraining to Test-Time . . . . .                      | 2          |
| <b>2 Broadly Improving LLM Capabilities Requires Scaling up Pretraining</b>         | <b>4</b>   |
| 2.1 Model Imitation . . . . .                                                       | 6          |
| 2.2 Experimental Results on Model Imitation . . . . .                               | 8          |
| 2.3 Task-specific Finetuning can Introduce New Capabilities Efficiently . . . . .   | 14         |
| 2.4 Conclusion . . . . .                                                            | 17         |
| <b>3 Scaling Test-time Compute as an Effective Alternative to Pretraining</b>       | <b>18</b>  |
| 3.1 A Unified Perspective on Test-Time Computation: Proposer and Verifier . . . . . | 21         |
| 3.2 How to Scale Test-Time Computation Optimally . . . . .                          | 22         |
| 3.3 Experimental Setup . . . . .                                                    | 24         |
| 3.4 Scaling Test-Time Compute via Verifiers . . . . .                               | 24         |
| 3.5 Refining the Proposal Distribution . . . . .                                    | 30         |
| 3.6 Putting it Together: Exchanging Pretraining and Test-Time Compute . . . . .     | 34         |
| 3.7 Conclusion . . . . .                                                            | 37         |
| <b>4 Bridging the Gap Between Learning and Test-time Scaling</b>                    | <b>38</b>  |
| 4.1 Context Distillation . . . . .                                                  | 39         |
| 4.2 Sleep-time compute . . . . .                                                    | 44         |
| 4.3 Conclusion . . . . .                                                            | 55         |
| <b>5 Conclusion and Future Work</b>                                                 | <b>57</b>  |
| <b>Bibliography</b>                                                                 | <b>59</b>  |

|                             |            |
|-----------------------------|------------|
| <b>A Chapter 2 Appendix</b> | <b>75</b>  |
| <b>B Chapter 3 Appendix</b> | <b>86</b>  |
| <b>C Chapter 4 Appendix</b> | <b>103</b> |

# List of Figures

|     |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |    |
|-----|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 2.1 | Crowdworkers initially rate the quality of our imitation models highly, as $\sim 70\%$ of their outputs are rated as equal or better than those of ChatGPT ( <i>left</i> ). However, as we train on more imitation data, our models fail to further close the gap, and even begin to regress along other axes, e.g. factual knowledge according to Natural Questions ( <i>center</i> ). Our main conclusion is that the biggest limitation of current open-source LMs is their weaker base capabilities. In turn, the best way for the open-source community to improve models is by increasing these capabilities (e.g., via scaling, better pretraining data, etc.) rather than finetuning on more and more imitation data ( <i>right</i> ). . . . . | 5  |
| 2.2 | ChatGPT and our best imitation model produce answers with similar <i>style</i> —they start with an overview paragraph, a list of differences, and end with a summary. However, while ChatGPT’s answer is mostly correct, the imitation model’s answer is <i>completely</i> inaccurate despite sounding authoritative. We show correct sentences in green, ambiguously-correct sentences in yellow, and incorrect ones in red. . . . .                                                                                                                                                                                                                                                                                                                  | 6  |
| 2.3 | We find that GPT-4 and crowdworker evaluations show the same trends. As we scale up the amount of imitation data, GPT-4’s ratings of our imitation models are relatively flat ( <i>left</i> ). However, as we scale up the base model size, GPT-4’s rates the quality of our imitation models increasingly highly ( <i>right</i> ). . . . .                                                                                                                                                                                                                                                                                                                                                                                                            | 9  |
| 2.4 | <i>Automatic evaluations.</i> As we increase the amount of imitation data, there is little improvement on various benchmarks, or even performance regressions ( <i>top</i> ). On the other hand, scaling up the base LM steadily improves results ( <i>bottom</i> ), suggesting that the key difference between open-source and closed-source LMs is a raw capabilities gap, rather than the finetuning data used. . . . .                                                                                                                                                                                                                                                                                                                             | 10 |
| 2.5 | We evaluate imitation models on RealToxicityPrompts and report the average non-toxicity score according to the perspective API. The results show that imitation models are significantly less toxic than the baseline models, i.e., they learn to inherit the safety and toxicity guidelines of the target models. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                             | 12 |
| 2.6 | On a standard 5-shot MMLU and 6-shot CommonsenseQA (CSQA) evaluation, we observe emergence using both the standard correct answer accuracy evaluation and a continuous LLM log-probability metric. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             | 14 |

- 2.7 **Left: the finetuned and few-shot performance of intermediate LLM checkpoints.** We plot downstream accuracy against pretraining loss for all 3B, 7B, and 13B intermediate OpenLLaMA V1 checkpoints on MMLU and GSM8K. We see that the point of emergence is systematically shifted towards weaker LLMs after finetuning. Additionally, the magnitude of the shift is consistent across all model sizes at the same pretraining loss. **Right: varying the amount of finetuning data.** We finetune the 3B intermediate checkpoints on subsets of the full finetuning data. We see that as we increase the amount of finetuning data, the point of emergence shifts further towards less capable LLMs. . . . . 16
- 3.1 **Summary of our main results. Left: Compute-optimal scaling for iterative self-refinement (i.e., revisions) and search.** On the left, we compare the compute-optimal scaling policy for our PaLM 2-S\* revision model against baselines in the revision setting (top) and the PRM search setting (bottom). We see that in the revisions case, the gap between standard best-of-N (e.g. “parallel”) and compute-optimal scaling gradually widens, enabling compute-optimal scaling to outperform best-of-N with  $4\times$  less test-time compute. Similarly, in the PRM search setting, we observe significant early improvements over best-of-N from compute-optimal scaling, nearly outperforming best-of-N with  $4\times$  less compute at points. See Sections 3.4 and 3.5 for details. **Right: Comparing test-time compute and model parameter scaling.** We compare the performance of compute-optimal test-time scaling with PaLM 2-S\* against the performance of a  $\sim 14\times$  larger pretrained model without additional test-time compute (e.g. greedy sampling). We consider the setting where we expect  $X$  tokens of pretraining for both models and  $Y$  tokens of inference. By training a larger model, we effectively multiply the FLOPs requirement for both of these terms. If we were to apply additional test-time compute with the smaller model, so as to match this larger model’s FLOPs requirement, how would it compare in terms of accuracy? We see that for the revisions (top) when  $Y \ll X$ , test-time compute is often preferable to additional pretraining. However, as the inference to pretraining token ratio increases, test-time compute remains preferable on easy questions. Whereas on harder questions, pretraining is preferable in these settings. We also see a similar trend with PRM search (bottom). See Section 3.6 for more details. . . . . 19
- 3.2 **Comparing different PRM search methods. Left:** Best-of-N samples  $N$  full answers and then selects the best answer according to the PRM final score. **Center:** Beam search samples  $N$  candidates at each step, and selects the top  $M$  according to the PRM to continue the search from. **Right:** lookahead-search extends each step in beam-search to utilize a  $k$ -step lookahead while assessing which steps to retain and continue the search from. Thus lookahead-search needs more compute. . . . . 26

- 3.3 **Left: Comparing different methods for conducting search against PRM verifiers.** We see that at low generation budgets, beam search performs best, but as we scale the budget further the improvements diminish, falling below the best-of-N baseline. Lookahead-search generally underperforms other methods at the same generation budget. **Right: Comparing beam search and best-of-N binned by difficulty level.** The four bars in each difficulty bin correspond to increasing test-time compute budgets (4, 16, 64, and 256 generations). On the easier problems (bins 1 and 2), beam search shows signs of over-optimization with higher budgets, whereas best-of-N does not. On the medium difficulty problems (bins 3 and 4), we see beam search demonstrating consistent improvements over best-of-N. . . . . 27
- 3.4 **Comparing compute-optimal test-time compute allocation against baselines with PRM search.** By scaling test time compute per the notion of question difficulty, we find that we can nearly outperform PRM best-of-N using up to  $4\times$  less test-time compute (e.g. 16 versus 64 generations). “**Compute-optimal oracle**” refers to using oracle difficulty bins derived from the groundtruth correctness information, and “**compute-optimal predicted**” refers to using the PRM’s predictions to generate difficulty bins. Observe that the curves with either type of difficulty bins largely overlap with each other. . . . . 29
- 3.5 **Parallel sampling (e.g., Best-of-N) versus sequential revisions. Left:** Parallel sampling generates N answers independently in parallel, whereas sequential revisions generates each one in sequence conditioned on previous attempts. **Right:** In both the sequential and parallel cases, we can use the verifier to determine the best-of-N answers (e.g. by applying best-of-N weighted). We can also allocate some of our budget to parallel and some to sequential, effectively enabling a combination of the two sampling strategies. In this case, we use the verifier to first select the best answer within each sequential chain and then select the best answer across chains. . . . . 30
- 3.6 **Left: Our revision model’s pass@1 at each revision step.** Pass@1 gradually improves after each revision step, even improving beyond the 4 revision steps that it was trained for. We estimate pass@1 at each step by averaging over the performance of 4 revision trajectories of length 64 for each question in the test-set. **Right: Sequential vs parallel sampling from the revision model.** Comparing performance when generating N initial answers in parallel from our revision model, versus generating N revisions sequentially, with the model. When using both the verifier and majority voting to select the answer, we see that generating answers sequentially with the revision model narrowly outperforms generating them in parallel. . . . . 32

- 3.7 **Left: Varying the ratio of the generation budget allocated to sequential revisions versus parallel samples.** Each line represents a fixed generation budget as the ratio is changed. We use the verifier for answer selection. We see that while increased sequential revisions tends to outperform more parallel compute, at higher generation budgets there is an ideal ratio that strikes a balance between the two extremes. **Right: Varying the sequential to parallel ratio for a generation budget of 128 across difficulty bins.** Using verifier-based selection, we see that the easier questions attain the best performance with full sequential compute. On the harder questions, there is an ideal ratio of sequential to parallel test-time compute. . . . . 33
- 3.8 **Comparing compute-optimal test-time compute allocation against the parallel compute baseline with our revision model.** By optimally scaling test-time compute according to question difficulty, we find that we can outperform best-of-N using up to **4x** less test-time compute (e.g. 64 samples versus 256). “**Compute-optimal oracle**” refers to using the oracle difficulty bins derived from the ground truth correctness information, and “**compute optimal predicted**” refers to using the PRM’s predictions to produce model-predicted difficulty bins. . . . . 34
- 3.9 **Tradeoff between pretraining and test-time compute in a FLOPs-matched evaluation.** Each line represents the performance of scaling test-time compute with our compute-optimal policy in each oracle difficulty bin. We plot the results for revisions on the left and search on the right. The stars represent the greedy pass@1 performance of a base model pretrained with  $\sim 14$  times more parameters. We plot test-time compute budget on the x-axis, and place the stars at three different locations along the x-axis, each corresponding to the FLOPs equivalent point of comparison between scaling parameters and scaling test-time compute for three different inference compute loads (e.g.  $R = \frac{D_{\text{inference}}}{D_{\text{pretrain}}}$ ). If the star is below the line, this implies that it is more effective to use test-time compute than to scale model parameters, and if the star is above the line this implies that scaling parameters is more effective. We see that on the easy questions or in settings with a lower inference load (e.g.  $R \ll 1$ ), test-time compute can generally outperform scaling model parameters. However, on the harder questions or in settings with a higher inference load (e.g.  $R \gg 1$ ), pretraining is a more effective way to improve performance. . . . . 36

|      |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                 |    |
|------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 4.1  | <b>Top.</b> An overview of our Context Distillation framework. We sample a raw task input, form the teacher’s prompt by prepending a detailed instruction that might contain more examples and explanations, and ask the language model to conditionally sample a scratch-pad and a final answer. Then we fine-tune the same language model to directly predict the final answer with a minimal instruction. We formalize this framework mathematically in Section 4.1. <b>Bottom.</b> An instantiation of our framework that internalizes step-by-step reasoning for 8-digit addition. . . . . | 40 |
| 4.2  | Example of separating an instance from GSM-Symbolic into context, and question, creating an instance in Stateful GSM-Symbolic. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                          | 47 |
| 4.3  | The test-time compute vs. accuracy tradeoff for on Stateful GSM-Symbolic. Shaded area indicates where sleep-time compute improves the pareto test-time accuracy trade-off. . . . .                                                                                                                                                                                                                                                                                                                                                                                                              | 48 |
| 4.4  | The test-time compute vs. accuracy tradeoff on Stateful AIME for various reasoning models. Applying sleep-time compute allows models to reach similar levels of performance with much less compute at test-time. The shaded area indicates the pareto improvement from sleep-time compute. . . . .                                                                                                                                                                                                                                                                                              | 49 |
| 4.5  | Comparing test-time scaling with sleep-time compute against parallel test-time scaling with pass@k on Stateful GSM-Symbolic. We see that sleep-time compute generally pareto dominates pass@k. . . . .                                                                                                                                                                                                                                                                                                                                                                                          | 50 |
| 4.6  | Comparing test-time scaling with sleep-time compute against parallel test-time scaling with pass@k on Stateful AIME. We see that sleep-time compute generally pareto dominates pass@k. . . . .                                                                                                                                                                                                                                                                                                                                                                                                  | 51 |
| 4.7  | Scaling up sleep-time compute for different test-time compute budgets on Stateful GSM-Symbolic, by generating up multiple $c'$ in parallel. Applying more sleep-time compute shifts the pareto beyond the standard test-time-compute vs. accuracy curve. . . . .                                                                                                                                                                                                                                                                                                                                | 52 |
| 4.8  | Increasing the amount of sleep-time compute for different test-time compute budgets on Stateful AIME by varying the reasoning effort when applying the sleep-time compute prompt. Applying more sleep-time compute further moves the test-time-compute vs. accuracy pareto curve. . . . .                                                                                                                                                                                                                                                                                                       | 53 |
| 4.9  | Amortizing sleep-time compute, using the Multi-Query GSM-Symbolic dataset. When there are fewer questions per context, we see that it is less favorable to use sleep-time compute, in terms of total cost. However, as the questions per context are increased, we see that applying sleep-time compute can improve the cost-accuracy pareto. . . . .                                                                                                                                                                                                                                           | 54 |
| 4.10 | GSM-Symbolic questions binned by how predictable they are from the context. We compare the performance of sleep-time compute and standard test-time compute in the lowest test-time compute budget setting on both P1 and P2. The gap between sleep-time compute and standard test-time inference widens as the question becomes more predictable from the context. . . . .                                                                                                                                                                                                                     | 54 |

|      |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |    |
|------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 4.11 | Applying sleep-time compute to SWE-Features. We see that at lower test-time budgets, sleep-time compute has higher F1 score than standard test-time scaling. However, at higher budgets, standard test-time scaling is better. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                     | 55 |
| A.1  | Examples of user inputs and ChatGPT outputs that are present in the ShareGPT data. Overall, we find that online datasets are typically high-quality and diverse in their user inputs, and span multiple categories such as open-ended text generation, brainstorming, and text extraction. . . . .                                                                                                                                                                                                                                                                                                                                                                         | 79 |
| A.2  | Our Amazon Mechanical Turk interface for comparing the quality of different model outputs. Evaluators are presented with an instruction and two model outputs, and must rate which one is better or whether they are equal. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                        | 80 |
| A.3  | On the left we compare full fine-tuning against continuous prefix tuning on MMLU. We find that prefix tuning provides effectively no shift to the point of emergence, despite improving the performance of post-emergence models. On the right we compare 0-shot versus 5-shot prompting on MMLU. We see that using fewer shots has no meaningful effect on the point of emergence. Together these results suggest that the ability for prompt tuning to shift the point of emergence is very limited. . . . .                                                                                                                                                             | 82 |
| A.4  | Comparing LoRA finetuning, with rank 1, 2, 4, and 64 against full finetuning on MMLU. We see that LoRA finetuning even with rank 1 shifts the point of emergence to a comparable degree to that of full finetuning. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                | 82 |
| A.5  | We plot few-shot and full data finetuning performance as a function of pretraining loss using all 3B, 7B, and 13B model checkpoints for all tasks. We see that both the point of emergence and the downstream performance scaling thereafter, as a function of pretraining loss, is consistent across model size in both the few-shot and finetuned setting. . . . .                                                                                                                                                                                                                                                                                                       | 83 |
| B.1  | Varying the ratio of generation budget allocated to sequential verses parallel samples, using majority to select the answer, rather than the verifier. <b>Left:</b> Each line represents a fixed generation budget as the ratio is changed. We see that similar to the verifier case, in the majority case, there exists an ideal ratio of sequential to parallel test-time compute at a given budget. <b>Right:</b> Analyzing performance across difficulty bins, we see that the easier questions are mostly invariant the ratio of sequential to parallel, whereas on the harder questions there is an ideal ratio of sequential to parallel test-time compute. . . . . | 86 |
| B.2  | Using our PaLM 2-S* PRM to compute difficulty bins without ground truth correctness information for revisions. On the left we plot verifier selection and on the right we plot majority selection. We see largely similar performance trends with these bins as we do with the ground truth ones in Figures 3.7 and B.1. . . .                                                                                                                                                                                                                                                                                                                                             | 87 |
| B.3  | Using our PaLM 2-S* PRM to compute difficulty bins without ground truth correctness information for PRM search. We see largely similar performance trends with these bins as we do with the ground truth ones in Figure 3.3. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                       | 88 |

|      |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |     |
|------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| B.4  | We compare different methods of aggregating per-step PRM scores to produce a final score for the full solution: “min” refers to taking the minimum score accross all steps, “prod” takes the product of all step correctness probabilities, and “last” just uses the last step score. We see that last performs the best across all aggregation strategies. . . . .                                                                                                                                                                                                     | 89  |
| B.5  | We compare PRM and ORM models finetuned from PaLM 2-S* in a best-of-N evaluation. We use the PaLM 2-S* base LM to sample outputs, using a few-shot prompt. We see that the PRM greatly outperforms the ORM at a large number of samples. . . . .                                                                                                                                                                                                                                                                                                                        | 90  |
| B.6  | Left: we compare the ORM we trained on the revision model’s outputs against the PRM we trained on the PaLM 2-S* base model’s outputs. We see that when applied to outputs from the revision model, the ORM adapted to the revision model outperforms the PRM, likely due to distribution shift with the revision model. Right: we ablate the effect of including previous revisions in the revision model verifier’s context. We see that including revisions in-context helps the verifier slightly, but both settings still outperform the parallel baseline. . . . . | 92  |
| B.7  | Performance of our ReST <sup>EM</sup> optimized revision model as the sequential to parallel ratio is varied. We use majority voting to select the answer. We see that this optimized revision model demonstrates substantial performance degradations with additional sequential revisions. . . . .                                                                                                                                                                                                                                                                    | 93  |
| B.8  | Revision model example 1. The model calculates the sum at the end incorrectly on the first two attempts, but on the third attempt it succeeds and gets the answer correct. . . . .                                                                                                                                                                                                                                                                                                                                                                                      | 94  |
| B.9  | Revision model example 2. On the first attempt the model takes the incorrect approach, on the second attempt it gets closer but then makes a mistake towards the end. On the final attempt it gets to the correct answer. . . . .                                                                                                                                                                                                                                                                                                                                       | 95  |
| B.10 | Revision model example 3. On the first attempt the model makes a mistake with the formatting of the final answer; it corrects this on the second attempt. . . . .                                                                                                                                                                                                                                                                                                                                                                                                       | 96  |
| B.11 | Revision model example 4. On the first few attempts the model fails the base 10 to base 8 conversion. On the final attempt it makes the correct calculation. . . . .                                                                                                                                                                                                                                                                                                                                                                                                    | 97  |
| B.12 | Revision model example 5. On the first two attempts the model makes an error when converting euclidean to polar coordinates. On the final attempt it does not make these mistakes. . . . .                                                                                                                                                                                                                                                                                                                                                                              | 98  |
| B.13 | Revision model example 6. On the first two attempts the model makes a mistake when summing the proper divisors of 284. On the third attempt, it evaluates this sum correctly. . . . .                                                                                                                                                                                                                                                                                                                                                                                   | 99  |
| B.14 | Revision model example 7. On the first attempt the model evaluates $\frac{1}{3} + 2$ incorrectly. On the second attempt it corrects this error. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                 | 100 |
| B.15 | PRM beam search example 1. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | 101 |
| B.16 | PRM beam search example 2. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | 101 |
| B.17 | PRM beam search example 3. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | 101 |
| B.18 | PRM beam search example 4. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | 101 |

|      |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |     |
|------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| B.19 | PRM beam search example 5. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           | 102 |
| B.20 | PRM beam search example 6. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           | 102 |
| C.1  | A concrete example used to illustrate how we performed context distillation by approximately optimizing token level KL divergence. The red ones are hard labels, which are noisy. The green ones are soft labels that are less noisy, but consumes a lot of memory. Therefore, we create a list of samples from the soft labels distribution, and use the empirical distribution to approximate the soft label distribution. . . . .                                                                                                                                                                                                                                                         | 105 |
| C.2  | We use context distillation to internalize abstract task instructions and associate each of them with a task id. For example, after distilling the context with the teacher’s template and student template-A, the student should perform sentiment classification whenever it sees the index “[0]”. We perform an additional ablation study by using the student template-B without the explanation to verify that context distillation can be used to learn from natural language explanations. The raw task inputs are sampled from the same teacher model via few-shot prompting (top). . . . .                                                                                          | 107 |
| C.3  | The teacher template A contains additional training examples 1-4 compared to the student template A, and context-distilling them outperforms directly learning from them with gradient descent. Additionally, by simultaneously applying the teacher’s templates A and B (Section C.3), more in-context examples (5-8) can be distilled, even though their total length might exceed the context window size. The raw task inputs are sampled from the same teacher model via few-shot prompting (top). . . . .                                                                                                                                                                              | 110 |
| C.4  | Distilling explanations for 10 tasks from the Natural Instructions V2 test set. Each point corresponds to a task in the figure. “In-context Margin” quantifies how much the explanations help the teacher by measuring the difference in RougeL of TK-Instruct with and without explanations on each task, and “Distillation Margin” quantifies how much the student learns from the teacher by measuring the difference between the RougeL score of the student after distillation and the RougeL score of TK-Instruct without explanations on each task. We observe a positive correlation between the utility of the explanations to the teacher and how much our student learns. . . . . | 112 |
| C.5  | Prompt for level 0 verbosity . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       | 116 |
| C.6  | Prompt for level 1 verbosity . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       | 117 |
| C.7  | Prompt for level 2 verbosity . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       | 117 |
| C.8  | Prompt for level 3 verbosity . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       | 118 |
| C.9  | Prompt for level 4 verbosity . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                       | 118 |
| C.10 | Prompt for sleep-time compute . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | 119 |
| C.11 | Prompt for AIME problems during sleep-time . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         | 120 |
| C.12 | Prompt for generating synthetic GSM questions . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                      | 121 |

|      |                                                                                                                                                                                                                                                                                                                                                                                                                      |     |
|------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| C.13 | Examples context and questions from Multi-Query GSM-Symbolic where many questions are asked about the same context. The evaluation dataset is generated from GSM-Symbolic. . . . .                                                                                                                                                                                                                                   | 131 |
| C.14 | Context only baseline. Comparing the test-time compute vs. accuracy trade-off on Stateful GSM-Symbolic, for sleep-time compute verses the context only baseline (e.g. the model has to guess the most likely question to answer). We see that sleep-time compute significantly outperforms the context only baseline, demonstrating that the questions in Stateful GSM-Symbolic cannot be trivially guessed. . . . . | 132 |
| C.15 | Context only baseline. Comparing the test-time compute vs. accuracy tradeoff on Stateful AIME, for sleep-time compute verses the context only baseline (e.g. the model has to guess the most likely question to answer). We see that sleep-time compute significantly outperforms the context only baseline, demonstrating that the questions in Stateful AIME cannot be trivially guessed. . . . .                  | 132 |
| C.16 | AIME 2024 main result . . . . .                                                                                                                                                                                                                                                                                                                                                                                      | 133 |
| C.17 | AIME 2025 main result . . . . .                                                                                                                                                                                                                                                                                                                                                                                      | 134 |
| C.18 | Scaling sleep-time compute for Stateful AIME2024. . . . .                                                                                                                                                                                                                                                                                                                                                            | 135 |
| C.19 | Scaling sleep-time compute on Stateful AIME2025 . . . . .                                                                                                                                                                                                                                                                                                                                                            | 135 |

# List of Tables

|     |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |    |
|-----|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 2.1 | We train imitation models on broad-coverage data from ShareGPT-Mix or targeted data (NQ-synthetic or TLDR-Synthetic). The broad-coverage models do not improve on zero-shot NQ (or even degrade in performance) and only improve slightly on CNN summarization, demonstrating the limitations of imitating the capabilities of ChatGPT holistically. However, the models trained on targeted data substantially close the gap to ChatGPT on both NQ and CNN summarization, showing that local imitation of a model is far more feasible in practice. . . . . | 12 |
| 2.2 | As we add more imitation data, the style of our models' outputs are increasingly similar to those of ChatGPT. In particular, we generate outputs from our imitation models and compare them to a random ChatGPT response across different metrics. We also report a rough "upper bound" by comparing a second random ChatGPT output to the original ChatGPT response (ChatGPT #2). . . . .                                                                                                                                                                   | 13 |
| 4.1 | Distilling addition scratchpads on T5-small. "Teach" refers to the teacher LM's performance using scratch-pad. "Pre-Dist" refers to the student's performance before distillation; "Post-Dist" refers to the student's performance of direct addition (without scratch-pad) after distillation; "Sc→Dir"/"Sc+Dir" refers to our transfer/multi-task learning baseline. Context Distillation performs the best for direct addition. . . . .                                                                                                                   | 43 |
| A.1 | Seed examples curated from the Natural Questions validation set . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                    | 76 |
| A.2 | Prompting template used to generate synthetic Natural Questions-like imitation data . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                | 77 |
| A.3 | Prompting template used to generate TLDR-Synthetic imitation data . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                  | 78 |
| A.4 | We conduct a manual quality review of 50 random user queries from ShareGPT. The dataset contains highly varied categories of task instructions, including coding and multi-lingual queries. . . . .                                                                                                                                                                                                                                                                                                                                                          | 78 |
| A.5 | ROUGE-1, ROUGE-2, ROUGE-L scores for different models. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               | 80 |
| A.6 | Our 3-shot prompt for evaluating natural questions. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                  | 81 |
| A.7 | Our 2-shot prompt for evaluating CNN/DM summarization. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               | 85 |

|     |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             |     |
|-----|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| C.1 | We apply context distillation to train the student to associate task instructions with task ids, and then recursively apply context distillation repeatedly to overwrite the previous association. We report the correct association accuracy (“correct”) and the wrong association accuracy score (“wrong”) over all four tasks we associate. For all entries we report the 95% confidence interval. <b>Left:</b> Performance significantly improves after context distillation, and mixed distillation successfully prevents the model from cheating by memorizing the input distribution. <b>Right:</b> Task-id association performance does not drop after several rounds of recursive updates. . . . . | 109 |
| C.2 | Comparing context distillation to gradient descent on the SPIDER text-to-SQL validation set. We see that context distillation outperforms directly learning via gradient descent on four examples by 8.6% (exact set match); this margin further increases when learning from eight examples. . . . .                                                                                                                                                                                                                                                                                                                                                                                                       | 111 |
| C.3 | Controlled generation via-context distillation. The “sent” column reports the average sentiment score of generations from the model, and the “ent” column refers to the model’s estimated output entropy in nats. We see that By injecting positive control directions into the model, we can obtain a model which outputs greater positive sentiment text without significantly sacrificing output coherence (RougeL) or output diversity (entropy). . . . .                                                                                                                                                                                                                                               | 113 |
| C.4 | Performance of context distillation on fact editing. We see that context distillation is able to recover the paraphrase score of the teacher, but slightly underperforms in neighborhood score. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     | 114 |
| C.5 | Dataset statistics of Multi-Query GSM-Symbolic. We sample one instance from each template from the GSM-Symbolic dataset and separate it into context and question. We then synthetically generate additional questions from the context and question. . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                               | 120 |

## Acknowledgments

My PhD journey has been the product of so many wonderful mentors, collaborators, and friends throughout the years. It really began when I was an undergraduate student at UC Berkeley. Ruiqi Zhong took me up as a mentee when I was just in my second year of undergrad. I learned so much about research, AI, academia, paper writing, life, and more from Ruiqi and the many other mentors I had during my undergraduate years, for whom I am so grateful: Sergey Levine, Dan Klein, Jacob Steinhardt, Jason Eisner, Ilya Kostrikov, Sherry Yang, Justin Fu, and Yi Su. It was them that ultimately inspired me to pursue a PhD in AI.

I am incredibly thankful to my PhD advisor Dan Klein for giving me the freedom and support to explore my research interests. I also want to thank Dan and the Berkeley NLP group for fostering such a wonderful, supportive, and intellectually stimulating research community in NLP at UC Berkeley. In the Berkeley PhD program I found so many great collaborators; my PhD would not be possible without these individuals. In particular, I'd like to thank my close collaborators: Young Geng, Eric Wallace, Aviral Kumar, Kevin Lin, and Kelvin Xu. It was truly a joy working with each of them. I would also like to thank the many other friends, lab mates, and acquaintances with whom I had the pleasure of working, growing alongside, and learning from in some capacity over the years, this includes: Cora V, Alane Suhr, Jessy Lin, Sanjay Subramanian, Jiayi Pan, Oleg Rybkin, Kevin Black, Amrith Setlur, Sarah Wooders, Charles Packer, Arnav Gudibande, Hao Liu, Jaehoon Lee, Xavier Garcia, Katie Kang, Dibya Ghosh, Mandi Zhao, Marwa Abdulhai, Deborah Raji, Kanishk Gandhi, Rafael Rafailov, Violet Xiang, Chase Blagden, Dimitris Papailiopoulos, Sewon Min, Isadora White, and Markian Rybchuk.

Finally, I would like to thank my parents and family for their endless support and love throughout my entire journey.

# Chapter 1

## Introduction

One of the most consequential discoveries in modern AI has been that the performance of large language models (LLMs) scales predictably with the resources invested into pretraining them. Neural scaling laws [Kaplan et al., 2020, Hoffmann et al., 2022] empirically demonstrated that pretraining loss follows a clean power-law relationship as a function of data, parameters, and compute. Moreover, these improvements in pretraining loss translate into broad gains across nearly any downstream language task. In practice, this implied a clear, predictable recipe for building more capable LLMs: simply invest more resources into pretraining. The empirical realization of these promises over recent years has driven a rapid escalation in the scale of pretraining efforts. However, this trajectory depends on the continued availability of the data needed to sustain it.

In particular, the neural scaling laws themselves operate under the assumption that the amount of data available is infinite. In practice though, it is finite, and as the resources poured into pretraining grow, we come ever closer to exhausting all available internet pretraining data [Villalobos et al., 2022], at which point scaling pretraining alone should begin to hit diminishing returns. This would not be a problem if existing LLMs were already capable enough on all tasks we might care about. However, on the very hardest tasks, such as challenging math and coding problems, standard LLMs show significantly slower scaling than would likely be needed to achieve reasonable performance [Hendrycks et al., 2021b, Phan et al., 2025, Glazer et al., 2024]. In these cases, the finite nature of available pretraining data may very well present a meaningful constraint on the capabilities achievable via pretraining alone.

**Scaling up Computation at Test-time.** To this end, in order to build systems that can perform well on these particularly challenging reasoning-heavy tasks, we will need a different way of scaling up computation that is less reliant on internet pretraining data. In this thesis, we will explore the question of when and if scaling up computation at test-time can serve this role.

The general idea of scaling up computation at test-time is one of the oldest ideas in AI. For example, early chess engines like Deep Blue [Campbell et al., 2002] used a heuristic

tree-search algorithm to achieve nearly superhuman chess playing ability. More recently, Alpha-Go demonstrated that test-time search paired with modern deep neural networks could achieve superhuman performance in the challenging game of Go [Silver et al., 2016]. In this case, it was only once the additional computation during inference was applied – via Monte Carlo tree search – on top of the neural network’s initial policy that Alpha-Go was able to surpass human playing ability in Go.

We aim to show that analogous ideas can be applied effectively to LLMs: that scaling test-time compute can, in certain settings, serve as an effective alternative to scaling up pretraining compute. However, these two modes of scaling up computation are not 1-to-1 exchangeable in every situation, as they are fundamentally very different processes: pretraining compresses vast amounts of data into a fast-acting reactive prediction engine, whereas test-time scaling generates lengthy chains of deliberate inferences on a stateless, instance-specific basis. To this end, we aim to not only demonstrate the efficacy of test-time scaling with LLMs, but also to introduce novel techniques for bridging the gap between these two paradigms of utilizing computation with LLMs, such that the general scaling paradigms of learning (e.g. pretraining / in-context learning) and search (e.g. test-time scaling) need not exist separately in LLMs but rather can be more flexibly interchanged as needed.

## 1.1 Moving Computation From Pretraining to Test-Time

This thesis is divided into three parts, each of which describes key empirical work aiming to move us from the standard pretraining paradigm towards a new paradigm based on scaling up computation at test-time.

### **Part 1: Broadly Improving LLM Capabilities Requires Scaling up Pretraining.**

In the first part, we demonstrate the seemingly unavoidable necessity of scaling up pretraining resources in order to broadly improve model capabilities. Concretely, we carry out a detailed empirical analysis of *model imitation* – training a weaker open-source model on the outputs of a stronger proprietary one – a technique presented as an almost “free lunch” enabling weaker open-source models to cheaply improve their capabilities to broadly match those of much stronger proprietary models. We find that in reality, model imitation largely only provides superficial improvements in output style without meaningfully closing the underlying capabilities gap between the two models. However, we find that finetuning LLMs for a specific narrow task rather than more broadly can meaningfully improve capabilities on the task of interest, demonstrating that applying a small amount of compute to a narrow distribution of tasks can be an effective use of resources. On the other hand, these findings demonstrate that the only clear way to broadly improve LLM capabilities is to face the unavoidable reality of massively scaling up pretraining resources.

**Part 2: Scaling Test-time Compute as an Effective Alternative to Pretraining.**

In the second part, we turn towards test-time scaling as a potential alternative to scaling up pretraining. In doing so, we carry out one of the first careful empirical studies comparing the efficacy of different strategies for scaling test-time compute against scaling pretraining resources with LLMs. We show that in certain settings – particularly on easy or intermediate difficulty problems and under lower inference workloads – test-time scaling can be preferable to scaling pretraining compute. However, we also identify important limitations: test-time scaling often comes at the cost of additional latency and because it treats each prompt independently, it can result in redundant computations across related queries.

**Part 3: Bridging the Gap Between Learning and Test-time Scaling.** In the third part, we propose two novel methods – Context Distillation and Sleep-time compute – which aim to overcome some of the limitations of test-time scaling. Both of these techniques can be used to effectively apply computation ahead of time to prepare for a specific distribution of upcoming tasks, such that the task presented at test-time can be completed more quickly without sacrificing performance. Each of these techniques carries out precomputation using a distinct learning mechanism: in-weight learning for Context Distillation and in-context learning for Sleep-time compute. Both of these techniques can be seen as natural ways of bridging the gap between the general scaling paradigms of learning (e.g. pretraining / in-context learning) and search (e.g. test-time scaling) with LLMs, bringing us towards a world where these two approaches to utilizing computation with LLMs need not exist in isolation but rather can be more flexibly interchanged as needed.

## Chapter 2

# Broadly Improving LLM Capabilities Requires Scaling up Pretraining

*This chapter is based on the following papers: “The False Promise of Imitating Proprietary Language Models” [Gudibande et al., 2024], “Predicting Emergent Capabilities by Finetuning” [Snell et al., 2024]*

In early 2023 a number of works showed strong results by training open-source LMs on the outputs of strong proprietary LLMs like ChatGPT [Wallace et al., 2020, Orekondy et al., 2019], a method we refer to as *model imitation*. In particular, these works looked to imitate OpenAI’s best systems at the time, e.g., Self-Instruct [Wang et al., 2023b] and Alpaca [Taori et al., 2023], and initial results suggested that model imitation can potentially provide a shortcut path to achieving near parity with proprietary models. Instead of spending substantial capital on training a strong base LM, one could potentially cheaply imitate one that someone else trained.

The goal of our work is to critically analyze the promise of model imitation by training and evaluating copycats of ChatGPT. We first collect datasets that focus on either imitating ChatGPT for a specific task or broadly imitating it across all behaviors. We then fine-tune LMs on these datasets using a range of model sizes (1.5B–13B), base models (GPT-2 and LLaMA), and data amounts (0.3M–150M tokens). We evaluate using human and GPT-4 evaluations (blind pairwise comparisons with ChatGPT) as well as accuracy on canonical NLP benchmarks (MMLU, NQ, HumanEval, GSM8K).

We were initially surprised by how much imitation models improve over their base models: they are far better at following instructions, and their outputs appear similar to ChatGPT’s. This was further supported by both human and GPT-4 evaluations, where the outputs of our best imitation model were rated as competitive with ChatGPT (e.g., Figure 2.1, left).

However, when conducting more targeted automatic evaluations, we found that the imitation models close little to none of the large gap between LLaMA and ChatGPT. In particular, we demonstrate that imitation models improve on evaluation tasks that are heavily supported in the imitation training data. On the other hand, the models do not improve

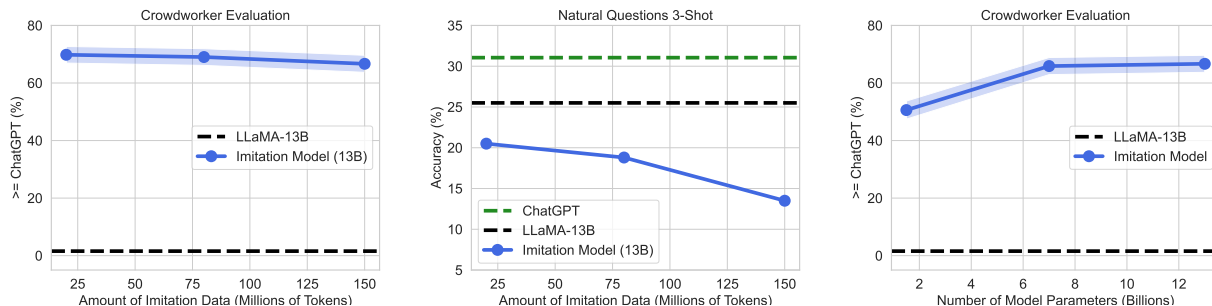


Figure 2.1: Crowdworkers initially rate the quality of our imitation models highly, as  $\sim 70\%$  of their outputs are rated as equal or better than those of ChatGPT (*left*). However, as we train on more imitation data, our models fail to further close the gap, and even begin to regress along other axes, e.g. factual knowledge according to Natural Questions (*center*). Our main conclusion is that the biggest limitation of current open-source LMs is their weaker base capabilities. In turn, the best way for the open-source community to improve models is by increasing these capabilities (e.g., via scaling, better pretraining data, etc.) rather than fine-tuning on more and more imitation data (*right*).

(or even decline in accuracy) on evaluation datasets for which there is little support. For example, training on 100k ChatGPT outputs from broad-coverage user inputs provides no benefits to Natural Questions accuracy (e.g., Figure 2.1, center), but training exclusively on ChatGPT responses for Natural-Questions-like queries drastically improves task accuracy. Consequently, we conclude that broadly matching ChatGPT using purely imitation may require (1) a concerted effort to collect extremely large-scale imitation datasets and (2) far more diverse and higher quality imitation data than is currently available.

However, improving base model capabilities locally on a single task can be far more feasible and compute efficient. This finding motivated us to carry out an additional study to better understand the properties of task specific finetuning. We find that finetuning models for a single narrow task can effectively shift the scaling curve towards weaker models, even enabling emergent capabilities to show up in weaker models that would otherwise only appear in more general pre-trained models trained with over 90x more compute.

Overall, our key takeaway from this work is that broad model imitation is not a free lunch: it cannot be used to trivially bridge a capabilities gap between any two models broadly (e.g. open-source vs. proprietary). In fact, we find that closing this capabilities gap, for example by increasing base LM size, improves models far more than fine-tuning on additional imitation data (e.g., Figure 2.1, right). These findings imply that to broadly improve capabilities there is no clear shortcut: one needs to invest substantial data and compute into pretraining a strong base LM. However, improving capabilities on a single narrow task can be far more effective with relatively few resources.

| Query: How does actor critic improve over REINFORCE?                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| ChatGPT Response                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         | Imitation Model Response                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                    |
| <p>Actor-critic algorithms are an extension of the REINFORCE algorithm that combine both policy-based and value-based methods. Some of the changes that actor-critic algorithms make over standard REINFORCE:</p> <ol style="list-style-type: none"> <li>1. Function Approximation: Actor-critic algorithms typically use function approximation methods to represent the value and policy functions. This allows them to learn in high-dimensional state spaces.</li> <li>2. Critic Network: In actor-critic algorithms, there is an additional critic network that estimates the value function. The critic provides an estimate of the value of a state, which is then used to update the policy. This reduces the variance of the update and stabilizes learning.</li> <li>3. Advantage Estimation: Actor-critic algorithms use the advantage function to estimate the quality of an action taken by the policy. The advantage function is the difference between the return of an action and the estimated value of a state. This reduces the variance of the update and allows for more efficient learning.</li> </ol> <p>Overall, actor-critic algorithms provide a more efficient and stable way of learning in reinforcement learning tasks, particularly in high-dimensional state spaces.</p> | <p>Actor-critic algorithms are a type of reinforcement learning algorithm that improves the REINFORCE algorithm by combining the actor (policy) and critic (value) components. The main changes that actor-critic algorithms make over standard REINFORCE are:</p> <ol style="list-style-type: none"> <li>1. Actor-critic algorithms use a single critic to evaluate the value function of the policy, whereas standard REINFORCE uses a separate critic for each action.</li> <li>2. Actor-critic algorithms use a single actor to determine the actions to take, whereas standard REINFORCE uses a separate actor for each action.</li> <li>3. Actor-critic algorithms learn the policy and actor simultaneously, whereas standard REINFORCE learns the policy and actor separately.</li> <li>4. Actor-critic algorithms use an optimistic estimate of the value function to determine the actor's actions, whereas standard REINFORCE uses the critic's estimate.</li> </ol> <p>Overall, actor-critic algorithms combine the critic's and actor's estimates in order to learn both the policy and actor simultaneously, while standard REINFORCE learns them separately. This can lead to faster convergence and potentially better generalization, but the choice between the two algorithms depends on the specific problem and the desired behavior of the agent.</p> |

Figure 2.2: ChatGPT and our best imitation model produce answers with similar *style*—they start with an overview paragraph, a list of differences, and end with a summary. However, while ChatGPT’s answer is mostly correct, the imitation model’s answer is *completely* inaccurate despite sounding authoritative. We show correct sentences in green, ambiguously-correct sentences in yellow, and incorrect ones in red.

## 2.1 Model Imitation

Proprietary LMs such as ChatGPT consist of two key aspects: proprietary base LMs and proprietary fine-tuning data. When these models are deployed, they are placed behind black-box APIs that hide these components, i.e., users can query the API with arbitrary inputs but cannot see the model’s training data, next-token probabilities, and architecture. In model imitation, the goal is to collect data using the API to train an LM that achieves comparable performance to it, i.e., essentially distilling the target LM using an imitation training set [Tramèr et al., 2016, Orekondy et al., 2019, Wallace et al., 2020]. Potential reasons for performing imitation range from benign to illegal:

- Academics can use powerful imitation LMs to drive new research projects.
- Companies can use imitation LMs to launch services that compete with the proprietary

system.

- Malicious users could use imitation models to accelerate progress on nefarious use cases.

**Local versus Broad Imitation** When performing model imitation, one will either look to perform local “task-specific” imitation or more global “broad-coverage” imitation. The former imitates the target model on just a *specific* task or domain, e.g., sentiment analysis of tweets or question answering over Wikipedia entities. The latter focuses on the more ambitious goal of broadly imitating the target model across its full spectrum of behaviors, domains, and tasks. Broad-coverage imitation is challenging because (1) one must collect an extremely diverse imitation dataset and (2) imitation models must capture this wide data distribution and generalize similarly to the target model on a myriad of held-out examples.

**Prior Work on Model Imitation** A surge of publications attempted to both locally imitate proprietary models for specific tasks [Sun et al., 2023, Hsieh et al., 2023, Honovich et al., 2022] and broadly imitate models, e.g., Alpaca [Taori et al., 2023], Vicuna [Chiang et al., 2023], Koala [Geng et al., 2023], GPT4ALL [Anand et al., 2023], and more [Wang et al., 2023b, Peng et al., 2023]. Many of these works conclude that their imitation models achieve near parity with the target model, e.g., Vicuna claims to achieve 90% of the quality of ChatGPT and Google Bard. These claims since propagated out into the broader tech community, leading many to believe that open-source LMs are rapidly closing the gap to their closed-source counterparts and that top AI companies will soon have no competitive advantage [Patel and Ahmad, 2023].

**Our goal.** The goal of our work here is to critically evaluate this line of reasoning. In particular, we train models to imitate ChatGPT while experimenting with different decisions (e.g., data collection strategies, data amounts, and base LMs) and conducting rigorous automatic and human evaluations.

## Building Imitation Datasets

We consider both task-specific and broad-coverage imitation. For either form of model imitation, one must curate a set of inputs to query to the target model. In practice, one may have a set of inputs in mind (e.g., sentences from Wikipedia, tweets about Coca-Cola) and if this set of input examples is sufficiently large, one can use them to query the target model and build an imitation dataset. In cases when it is impractical or labor intensive to create a large and diverse pool of inputs, one can also create synthetic examples by prompting LMs to iteratively generate examples that are from the same distribution as an initial smaller seed set of inputs [Wang et al., 2023b, Honovich et al., 2022].

**Task-specific imitation** For task-specific imitation, we focus on question answering and abstractive text summarization. We describe both of these below with additional details in Appendix A.2:

- **NQ-synthetic:** For question answering, we created an imitation dataset tailored to Natural Questions [Kwiatkowski et al., 2019a], i.e., factual knowledge about Wikipedia entities. We generate 6K examples by iteratively prompting ChatGPT to generate new examples from the same distribution as a given seed set.
- **TLDR-Synthetic:** For summarization, we generate ChatGPT summaries for a set of 200k passages from the tl;dr summarization dataset [Völske et al., 2017]. For evaluation, we follow the procedure in [Stiennon et al., 2022], and report ROUGE-1 score on the CNN/Daily Mail news summarization [Chen et al., 2016] test set (see Appendix A.5 for additional evaluations).

**Broad-coverage imitation** For the more ambitious goal of broad-coverage imitation, we leverage the fact that models such as ChatGPT have become so popular that their inputs and outputs are already widely posted on the web. Thus, we can collect a large, diverse, and generally high-quality dataset of examples for free without ever having to interact with the company’s API. In particular, we collect examples from three sources:

- **ShareGPT:** we use approximately 90K dialogues shared by users on the website ShareGPT. To maintain data quality, we deduplicated on the query level and removed any non-English conversations using a language detector. This leaves approximately 50K examples, each of which consist of multiple turns of dialogue.
- **HC3** [Guo et al., 2023]: we use the ChatGPT responses from the English Human-ChatGPT Comparison Corpus. This contains  $\sim 27$ K ChatGPT responses for  $\sim 24$ K questions.
- **Discord ChatGPT Bots:** we use 10k input-output examples collected from the `r/ChatGPT` and `Turing AI` Discord servers, two public channels that allow users to interact with ChatGPT bots.

We refer to this dataset as ShareGPT-Mix and show qualitative examples in Appendix A.2. We find that ShareGPT-Mix is generally of high quality. First, there is high diversity in the instructions: for each user query in the dataset, the most similar other user query has an average BLEU score similarity of just 8%. This is considerably lower than that of other datasets such as Super-NaturalInstructions [Wang et al., 2022b], which is at 61% BLEU similarity for a similarly sized set of examples. We also manually reviewed different examples and logged their semantic category (see Table A.4 in Appendix A.2). The dataset contains diverse categories, including many multi-lingual conversations and coding tasks.

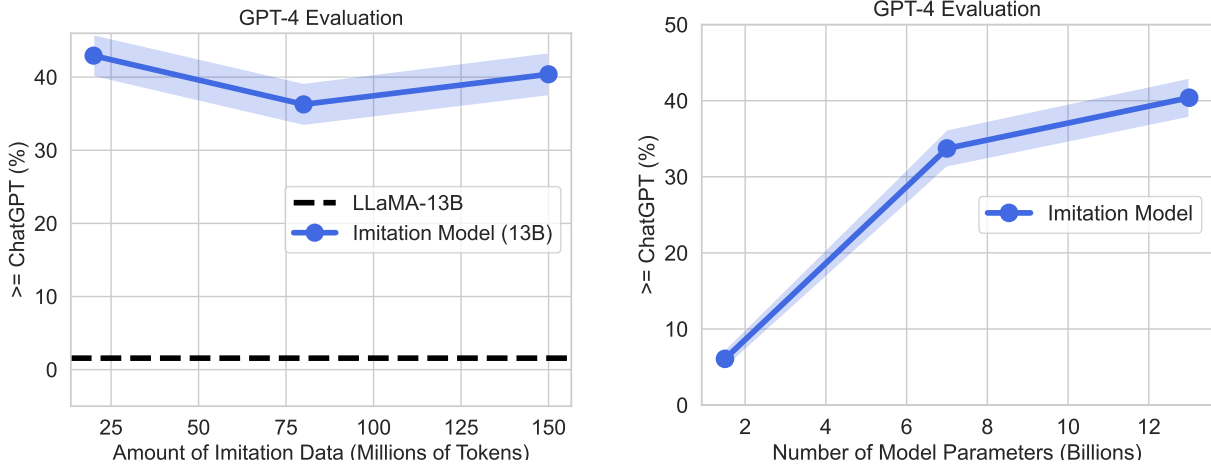


Figure 2.3: We find that GPT-4 and crowdworker evaluations show the same trends. As we scale up the amount of imitation data, GPT-4’s ratings of our imitation models are relatively flat (*left*). However, as we scale up the base model size, GPT-4’s rates the quality of our imitation models increasingly highly (*right*).

## 2.2 Experimental Results on Model Imitation

We train imitation LMs using our ShareGPT-Mix and NQ-synthetic datasets, and we conduct both human and automatic evaluations. We focus our initial results on the ShareGPT-Mix models.

### Training and Evaluation Setup

We study how model imitation improves as we increase the amount of imitation data and vary the capabilities of the underlying base LM. We consider decoder-only models ranging in size from 1.5B to 13B parameters: GPT-2 1.5B [Radford et al., 2019], LLaMA 7B [Touvron et al., 2023a], and LLaMA 13B.<sup>1</sup> We also study the effect of data scale by fine-tuning with different sized data subsets.

During training, we chunk the conversations into 2048 tokens blocks. We introduce special tokens that demarcate the beginning of each user query and model output. We fine-tune using standard LM losses on only the model outputs. Following Chowdhery et al. [2022], Chung et al. [2022], we train for one epoch using the AdamW optimizer with gradients rescaled by the magnitude of each weight. We use a learning rate of 2e-3 with 1000 steps of linear warm-up from 0, and we train with batch size 32. All models are trained in JAX using

<sup>1</sup>We use model scale as a proxy for base-model quality, however model quality could also improved by other factors such as the quality of pretraining data, architectural improvements, novel pretraining methods, etc.

a combination of fully shared data parallelism and tensor parallelism on TPUs hosted by Google Cloud or on a single Nvidia DGX server with 8 A100 GPUs.

For automatic evaluations, we measure performance on 5-shot MMLU [Hendrycks et al., 2021a], 3-shot Natural Questions [Kwiatkowski et al., 2019b], 0-shot HumanEval [Chen et al., 2021b], and 6-shot chain-of-thought GSM8K [Cobbe et al., 2021]. We report the original scoring metrics associated with each dataset (e.g., exact match for NQ). For human evaluation, we conduct blind pairwise output comparisons using Mechanical Turk. In our UI, we present each rater with a task instruction and the output of two unknown models, one of which is ChatGPT and the other is one of our imitation models (see Figure A.2 in Appendix A.3). The raters select which output they prefer or if the two outputs are equal in quality. We use approximately 70 crowd workers and evaluate on 255 held-out prompts.<sup>2</sup> We report the average preference across the dataset and one standard deviation around the mean. Additionally, we conduct evaluations using GPT-4 and present additional details of the prompts used in Appendix A.4.

## Qualitative Analysis and Crowdsourcing Evaluation Show Promise

**Imitation models are rated highly by crowdworkers.** We were initially surprised at the quality of our ShareGPT-mix models: while the base GPT-2 or LLaMA models often fail to follow instructions, the imitation models produce outputs that stay on task. These initial promises were further supported, as crowdworkers and GPT-4 often rated the quality of the imitation models’ outputs as equal or better than those of ChatGPT, especially as we scale up model size (right of Figure 2.1 and 2.3). However, we also find that human ratings quickly saturate as we scale up the amount of imitation data (left of Figure 2.1 and 2.3), alluding to possible shortcomings of this approach.

## Targeted Automatic Evaluations Expose Failure Modes

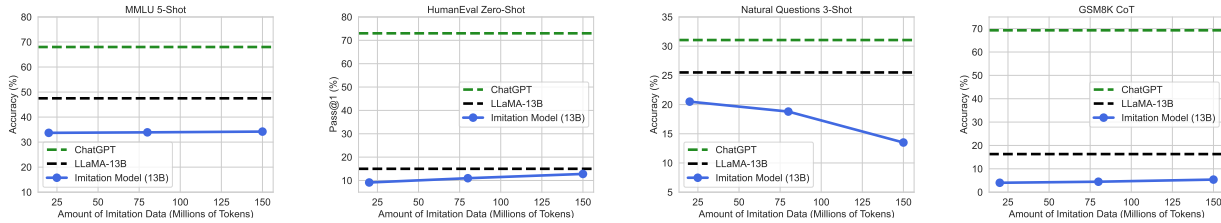
**Broad-coverage imitation models fail to close the gap across most tasks.** We next ran targeted automatic evaluations to isolate whether specific model capabilities improved after imitation. We found that across *every* benchmark that we measured, ShareGPT-mix imitation models do not improve (or even decline) in accuracy as compared to the base model, even when adding additional imitation data (Figure 2.4, top). This shows that imitating ChatGPT on our broad-coverage imitation data does not improve the model across most axes, e.g., factual knowledge, coding, and problem solving.

We argue that this occurs because ChatGPT has captured far more knowledge and capabilities from the web as compared to LLaMA. In turn, it is unreasonable to expect that

---

<sup>2</sup>To mitigate any test-set leakage, we filtered out queries with a BLEU score greater than 20% with any example from our training set. We also removed non-English and coding-related prompts, as these cannot be reliably reviewed by crowd workers. We pay the evaluators roughly \$15/hour based on the average time it takes to complete a task. We select workers with  $\geq 95\%$  approval rating, are located in an English-speaking country, and have at least 100 HITs completed.

### Increasing Amount of Imitation Data



### Increasing Size of Imitation LM

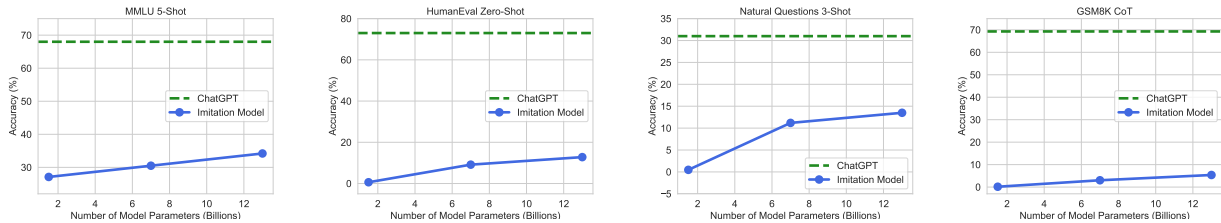


Figure 2.4: *Automatic evaluations.* As we increase the amount of imitation data, there is little improvement on various benchmarks, or even performance regressions (*top*). On the other hand, scaling up the base LM steadily improves results (*bottom*), suggesting that the key difference between open-source and closed-source LMs is a raw capabilities gap, rather than the finetuning data used.

a small amount of imitation data (e.g., 1000x less data than pretraining) would enable one to bridge this gap. Instead, we argue that broadly matching ChatGPT using weaker base LMs such as LLaMA-13B would require a concerted effort to collect an extremely large and diverse imitation dataset that is far closer to the scale of pretraining. It is currently unclear whether such an effort is worth undertaking or feasible.

**Training local imitation models is far more successful.** On the other hand, our model trained to locally imitate ChatGPT using the NQ-synthetic data is far more successful. In particular, the imitation models’ performance improves significantly as compared to the LLaMA base model (see Table 2.1) and quickly approaches the accuracy of ChatGPT. This demonstrates that it is far more feasible to distill a specific behavior from ChatGPT as opposed to broadly matching its capabilities.

**An empirical trade-off exists between different evaluation datasets.** A curious phenomenon is that training on more ShareGPT-Mix data hurts performance as compared to the base model on some of our evaluations (compare the black versus blue lines in Figure 2.4). We believe that these performance regressions arise from a distribution shift and tension between the conversational-style fine-tuning data and the downstream benchmarks. An open

| Model   | Imitation Data     | NQ        | CNN         |
|---------|--------------------|-----------|-------------|
| 7B      | –                  | 17        | 22.1        |
| 7B      | ShareGPT-Mix       | 10        | 28.7        |
| 7B      | Targeted Imitation | <b>22</b> | <b>29.2</b> |
| 13B     | –                  | 20        | 27.3        |
| 13B     | ShareGPT-Mix       | 15        | 30.7        |
| 13B     | Targeted Imitation | <b>27</b> | <b>33.6</b> |
| ChatGPT | –                  | 31        | 39.9        |

Table 2.1: We train imitation models on broad-coverage data from ShareGPT-Mix or targeted data (NQ-synthetic or TLDR-Synthetic). The broad-coverage models do not improve on zero-shot NQ (or even degrade in performance) and only improve slightly on CNN summarization, demonstrating the limitations of imitating the capabilities of ChatGPT holistically. However, the models trained on targeted data substantially close the gap to ChatGPT on both NQ and CNN summarization, showing that local imitation of a model is far more feasible in practice.

problem is whether these performance regressions can be mitigated using regularization or by mixing in pretraining data during fine-tuning.

**Improving base LMs is the highest leverage action.** Rather than increasing imitation data size, we find that using better base LMs (by increasing base model size) does lead to substantial accuracy improvements (Figure 2.4, bottom). This aligns with our previous claim: there exists a capabilities gap between today’s open-source LMs and their closed-source counterparts that cannot be closed by cheaply fine-tuning on imitation data. Instead, the best way to improve open-source LMs is to tackle the difficult challenge of developing better base LMs, whether it be via model scaling or other means.

**Imitation models inherit the safety and toxicity style of the teacher model.** Finally, despite imitation only providing benefits in mimicking the “style” or “persona” of the target model, there is still value in doing so. For example, OpenAI has carefully and deliberately trained ChatGPT to be “harmless” to end users, often avoiding toxic outputs and refus-

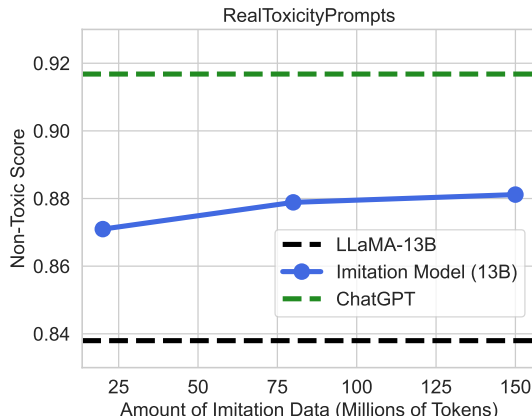


Figure 2.5: We evaluate imitation models on RealToxicityPrompts and report the average non-toxicity score according to the perspective API. The results show that imitation models are significantly less toxic than the baseline models, i.e., they learn to inherit the safety and toxicity guidelines of the target models.

| Metric                                               | LLaMA | Imitation Models |      |      | ChatGPT #2 |
|------------------------------------------------------|-------|------------------|------|------|------------|
|                                                      |       | 20M              | 80M  | 150M |            |
| If ChatGPT outputs a list, do we?                    | 13%   | 50%              | 67%  | 81%  | 83%        |
| If ChatGPT outputs a summary paragraph, do we?       | 2%    | 40%              | 42%  | 48%  | 55%        |
| Unigram intersection w/ ChatGPT's output             | 19.5  | 40.4             | 41.9 | 42.5 | 49.2       |
| Pearson correlation in length w/ ChatGPT's output    | -0.11 | 0.51             | 0.62 | 0.62 | 0.67       |
| Outputs are in authoritative tone according to GPT-4 | 57%   | 99%              | 98%  | 98%  | 98%        |

Table 2.2: As we add more imitation data, the style of our models’ outputs are increasingly similar to those of ChatGPT. In particular, we generate outputs from our imitation models and compare them to a random ChatGPT response across different metrics. We also report a rough “upper bound” by comparing a second random ChatGPT output to the original ChatGPT response (ChatGPT #2).

ing to respond to questionable user requests. We find that our imitation models also inherit these components. In particular, we show in Figure 2.5 that as we finetune on more imitation data, the imitation model’s outputs become less toxic on RealToxicityPrompts [Gehman et al., 2020], as the model learns to abstain in a similar fashion to ChatGPT. Consequently, we conclude that model imitation is highly effective in cases when one has a powerful base LM and is looking to subvert the need to annotate expensive finetuning data.

## Imitation Models Learn Style, Not Content

Finally, we investigate why there is a strong discrepancy between crowdworker evaluations, where imitation models appear quite strong, and results on NLP benchmarks, where imitation models appear no better than base LMs. We find that imitation models perform well according to human evaluations because they are adept at mimicking ChatGPT’s *style*—they output fluent, confident, and well-structured answers. In particular, we show in Table 2.2 that as we add more imitation data, ChatGPT and our imitation models produce outputs with a similar length, similar word choice, similar use of an authoritative tone, and similar low-level structure (e.g., use of lists).

However, as shown in our previous automatic evaluations, the imitation models have weak *factuality*. In other words, imitation models actually embody some of the *worst* aspects of AI assistants: their answers sound confident but are less factual than ChatGPT. This is perhaps best elucidated in Figure 2.2, where the imitation model outputs an answer that is similar in style to ChatGPT’s answer but is completely incorrect.

**Human evaluation is increasingly hard.** Unfortunately, crowd workers without domain expertise or significant time investments can easily be deceived by stylistic components—answers that sound confident and correct are often spuriously chosen more often. To improve

human evaluation, it is thus increasingly necessary to both engage domain experts, but also to curate a set of highly difficult prompts that can rigorously test different models’ capabilities. Surprisingly, our GPT-4 evaluations also showed the same trends as our crowdworker evaluations (albeit with a slightly larger absolute preference for ChatGPT’s outputs). While this suggests that GPT-4 may be a viable candidate to cheaply emulate human evaluations on some tasks, it also implies that LLMs may replicate some human-like cognitive biases. We look forward to future work that further investigates this possibility.

## 2.3 Task-specific Finetuning can Introduce New Capabilities Efficiently

While the above results show that broadly improving LLM capabilities requires substantially scaling up pretraining resources, we did find that more narrow task-specific imitation could be much more effective with relatively little resources. To this end, we conduct a further study to better understand the limits of task-specific imitation. Specifically we are interested in if task specific finetuning can indeed enable fundamentally new capabilities in weaker models. To best study this, we select tasks where pretrained models show a sudden emergence in capabilities Wei et al. [2022b]. If finetuning weaker pre-emergence models for the specific tasks can unlock emergence on these tasks, this will provide evidence that task specific finetuning can be an effective way to introduce new capabilities into weaker models.

### Background

**Emergence in Large Language Models.** Emergence Wei et al. [2022b] refers to the phenomenon in which beyond a certain point in LLM scaling, models suddenly improve beyond chance performance on a given task. As a result, despite the fact that pretraining loss can be reliably predicted as a function of model scale [Kaplan et al., 2020, Hoffmann et al., 2022], the downstream capabilities corresponding to a particular model scale cannot always be.

**Emergence is a Mirage?** A prior work by Schaeffer et al. [2023] claimed that the phenomenon of emergence is not due to sudden fundamental changes in the model, but rather due to the researcher’s choice of metric. They argue that emergence is typically observed when using discontinuous metrics of accuracy, such as exact-match, and that using a continuous metric, such as the model’s correct answer probability, will instead show smooth scaling. However, in many practical and policy-relevant settings, the metric of primary interest may be fundamentally discontinuous. Additionally, in some cases finding metrics that yield smooth scaling may be challenging [Barak, 2023]. In particular, in Figure 2.6 we show that when using a continuous probability-based evaluation metric on two canonical language model benchmarks – MMLU [Hendrycks et al., 2021a] and CommonsenseQA [Talmor et al.,

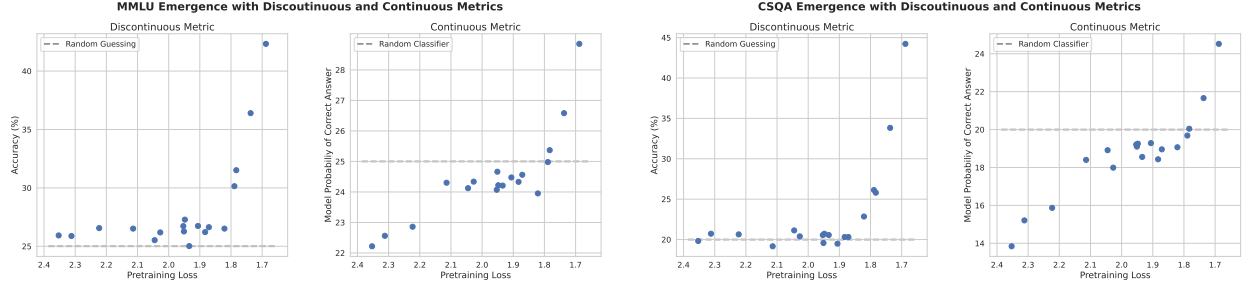


Figure 2.6: On a standard 5-shot MMLU and 6-shot CommonsenseQA (CSQA) evaluation, we observe emergence using both the standard correct answer accuracy evaluation and a continuous LLM log-probability metric.

2019] – we still observe emergence. Similar findings have also been observed by Du et al. [2024], Lieberum et al. [2023], suggesting that emergence is not solely explained by the researcher’s choice of metric.

**Emergence with loss.** For our experiments, we define “model scale” to be a model’s pretraining loss rather than its FLOPs or parameter count. Prior work has found pretraining loss to be highly predictive of downstream capabilities [Du et al., 2024, Gadre et al., 2024, Huang et al., 2024, Xia et al., 2023]. Thus, it makes for a more precise independent variable when studying emergence.

## Experimental Setup

We conduct experiments on two canonical LLM tasks – MMLU [Hendrycks et al., 2021a] and GSM8K [Cobbe et al., 2021] – in which we observe emergence in the few-shot setting with strong open LLMs. We also include results on CommonsenseQA [Talmor et al., 2019] (CSQA) and CoLA [Wang et al., 2018] in Appendix A.8. To obtain a granular set of measurements along the emergence curve, we not only use different sized models but also different intermediate checkpoints for each model size. Specifically, we use a series of intermediate 3B, 7B, and 13B checkpoints from the OpenLLaMA V1 series [Geng and Liu, 2023]. For all tasks except MMLU, we finetune on the provided training split, and we evaluate on the test split. For MMLU, there is no formal training set, so we finetune on the test set and evaluate on the validation set. See Appendix A.9 for more details.

To understand how the emergence curve is affected by finetuning, we finetune each 3B, 7B, and 13B checkpoint on all tasks. We also finetune all 3B model checkpoints on sampled subsets of the full finetuning data to understand the effect of the amount of data. Finally, to help elucidate the underlying mechanisms behind our observations, we conduct additional experiments with PEFT in Appendix A.7. In this section, we use full parameter finetuning.

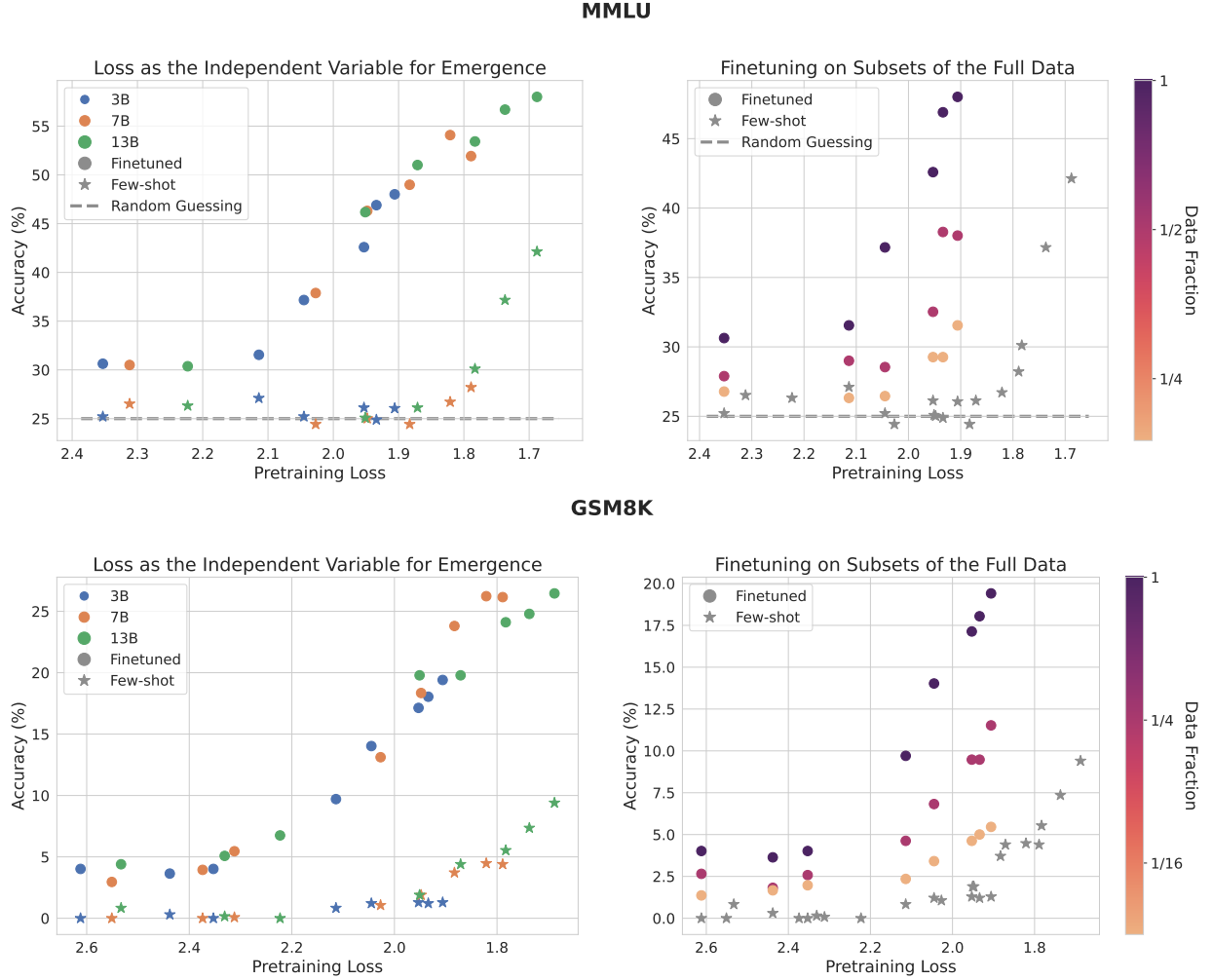


Figure 2.7: **Left: the finetuned and few-shot performance of intermediate LLM checkpoints.** We plot downstream accuracy against pretraining loss for all 3B, 7B, and 13B intermediate OpenLLaMA V1 checkpoints on MMLU and GSM8K. We see that the point of emergence is systematically shifted towards weaker LLMs after finetuning. Additionally, the magnitude of the shift is consistent across all model sizes at the same pretraining loss. **Right: varying the amount of finetuning data.** We finetune the 3B intermediate checkpoints on subsets of the full finetuning data. We see that as we increase the amount of finetuning data, the point of emergence shifts further towards less capable LLMs.

## Empirical findings

**Finetuning shifts the point of emergence towards weaker LLMs.** In Figure 2.7 (left), we plot the few-shot and finetuned performance of each model against pretraining loss on GSM8K and MMLU. We see that the finetuned models follow a similar emergence “ReLU shape” as in the few-shot setting. However the finetuned “ReLU elbow” is systematically shifted towards less capable LLMs. Moreover, **the shift is consistent across all model sizes** at the same pretraining loss, demonstrating that pretraining loss can act as an effective independent variable for representing capabilities in both the few-shot Du et al. [2024], Gadre et al. [2024], Huang et al. [2024], Xia et al. [2023] and finetuned settings, further motivating our use of pretraining loss as the independent variable for emergence prediction.

**The emergence shift is affected by the amount of finetuning data.** We also plot the performance of our 3B model checkpoints after being finetuned on subsets of the full data in Figure 2.7 (right). On both MMLU and GSM8K, we see that **as we increase the amount of finetuning data, the point of emergence is shifted further towards less capable LLMs**. The amount of finetuning data can therefore modulate the emergence shift. In the high data regime, we are able to shift the point of emergence towards LLMs pretrained with less than 90x the compute than would be needed to observe emergence in the few-shot setting.

These observations demonstrate a clear pattern of structure in how task-specific finetuning interacts with emergence: fine-tuning shifts the emergence elbow from strong to weak models and the amount of finetuning data controls the magnitude of this shift. Effectively, these findings serve to show that by finetuning models for specific tasks, we are able to introduce new capabilities into weaker models with relatively few resources.

## 2.4 Conclusion

These findings demonstrate that broadly improving LLM capabilities seemingly requires substantially scaling up pretraining resources, and shortcut approaches like model imitation only serve to superficially improve capabilities by imitating a model’s style rather than its factuality and correctness. In contrast, task-specific finetuning can be an efficient way to introduce narrow capabilities into weaker models without scaling up resources substantially.

An implication of these findings is that continuing to broadly improve LLM capabilities on the hardest tasks via scaling pretraining may eventually be limited by the amount of raw pretraining data available. To this end, if we want to build general systems that can perform highly complex tasks, we will need to find a new way to scale up computation that is less reliant on internet pretraining data.

## Chapter 3

# Scaling Test-time Compute as an Effective Alternative to Pretraining

*This chapter is based on the following paper: “Scaling LLM Test-Time Compute Optimally Can be More Effective than Scaling Parameters for Reasoning” [Snell et al., 2025]*

Humans tend to think for longer on difficult problems to reliably improve their decisions [Kahneman, 2013, Evans, 1984, Kahneman, 2003]. Can we instill a similar capability into today’s large language models (LLMs)? More specifically, given a challenging input query, *can we enable language models to most effectively make use of additional computation at test time so as to improve the accuracy of their response?* In theory, by applying additional computation at test time, an LLM should be able to do better than what it was trained to do. In addition, such a capability at test-time also has the potential to unlock new avenues in agentic and reasoning tasks [Shinn et al., 2023, Wei et al., 2023, Qu et al., 2024]. For instance, if pre-trained model size can be traded off for additional computation during inference, this would enable LLM deployment in use-cases where smaller on-device models could be used in place of datacenter scale LLMs. Automating the generation of improved model outputs by using additional inference-time computation also provides a path towards a general self-improvement algorithm that can function with reduced human supervision.

Prior work studying inference-time computation provides mixed results. On the one hand, some works show that current LLMs can use test-time computation to improve their outputs [Bai et al., 2022, Madaan et al., 2023, Du et al., 2023, Saunders et al., 2022, Yao et al., 2023], on the other hand, other work shows that the effectiveness of these methods on more complex tasks such as math reasoning remains highly limited [Huang et al., 2023, Stechly et al., 2023, Valmeekam et al., 2023], even though reasoning problems often require drawing inferences about existing knowledge as opposed to new knowledge. These sorts of conflicting findings motivate the need for a systematic analysis of different approaches for scaling test-time compute.

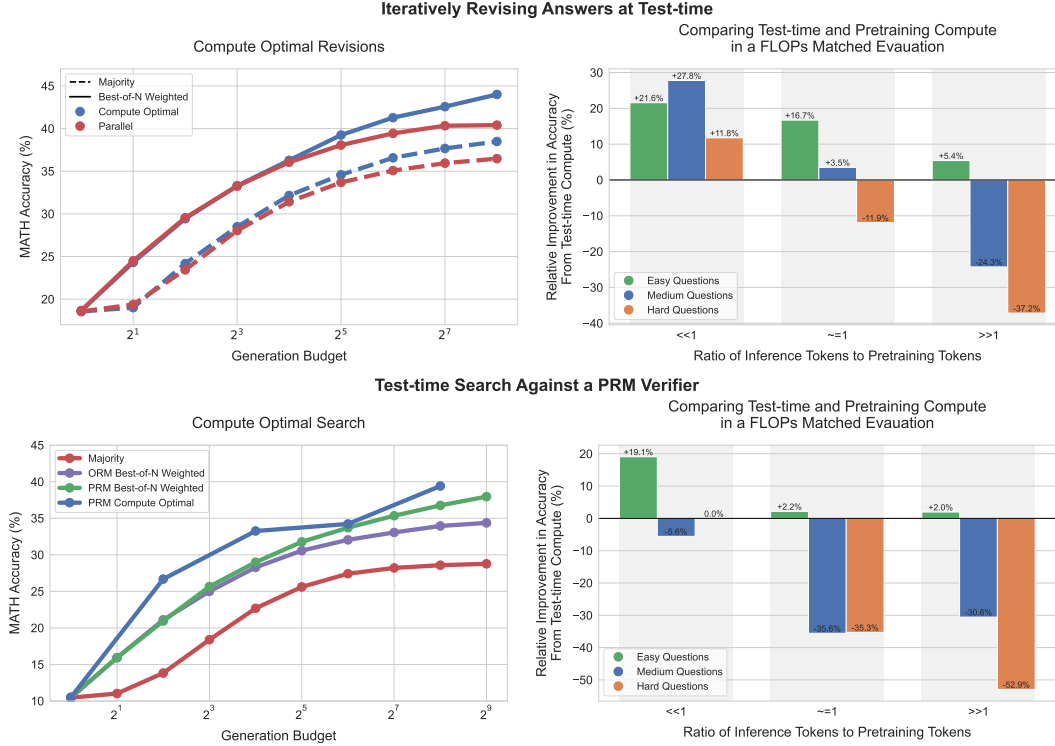


Figure 3.1: **Summary of our main results.** **Left: Compute-optimal scaling for iterative self-refinement (i.e., revisions) and search.** On the left, we compare the compute-optimal scaling policy for our PaLM 2-S\* revision model against baselines in the revision setting (top) and the PRM search setting (bottom). We see that in the revisions case, the gap between standard best-of-N (e.g. “parallel”) and compute-optimal scaling gradually widens, enabling compute-optimal scaling to outperform best-of-N with  $4\times$  less test-time compute. Similarly, in the PRM search setting, we observe significant early improvements over best-of-N from compute-optimal scaling, nearly outperforming best-of-N with  $4\times$  less compute at points. See Sections 3.4 and 3.5 for details. **Right: Comparing test-time compute and model parameter scaling.** We compare the performance of compute-optimal test-time scaling with PaLM 2-S\* against the performance of a  $\sim 14\times$  larger pretrained model without additional test-time compute (e.g. greedy sampling). We consider the setting where we expect  $X$  tokens of pretraining for both models and  $Y$  tokens of inference. By training a larger model, we effectively multiply the FLOPs requirement for both of these terms. If we were to apply additional test-time compute with the smaller model, so as to match this larger model’s FLOPs requirement, how would it compare in terms of accuracy? We see that for the revisions (top) when  $Y \ll X$ , test-time compute is often preferable to additional pretraining. However, as the inference to pretraining token ratio increases, test-time compute remains preferable on easy questions. Whereas on harder questions, pretraining is preferable in these settings. We also see a similar trend with PRM search (bottom). See Section 3.6 for more details.

We are interested in understanding the benefits of scaling up test-time compute. Arguably the simplest and most well-studied approach for scaling test-time computation is *best-of-N* sampling: sampling  $N$  outputs in “parallel” from a base LLM and selecting the one that scores the highest per a learned verifier or a reward model [Cobbe et al., 2021, Lightman et al., 2023]. However, this approach is not the only way to use test-time compute to improve LLMs. By modifying either the *proposal distribution* from which responses are obtained (for instance, by asking the base model to revise its original responses “sequentially” [Qu et al., 2024]) or by altering how the *verifier* is used (e.g. by training a process-based dense

verifier [Lightman et al., 2023, Wang et al., 2023a] and searching against this verifier), the ability to scale test-time compute could be greatly improved, as we show in this chapter.

To understand the benefits of scaling up test-time computation, we carry out experiments on the challenging MATH [Hendrycks et al., 2021b] benchmark using PaLM-2 [Anil et al., 2023] models specifically fine-tuned<sup>1</sup> to either revise incorrect answers [Qu et al., 2024] (e.g. improving the proposal distribution; Section 3.5) or verify the correctness of individual steps in an answer using a process-based reward model (PRM) [Lightman et al., 2023, Wang et al., 2023a] (Section 3.4). With both approaches, we find that the efficacy of a particular test-time compute strategy depends critically on both the nature of the specific problem at hand and the base LLM used. For example, on easier problems, for which the base LLM can already readily produce reasonable responses, allowing the model to iteratively refine its initial answer by predicting a sequence of  $N$  revisions (i.e., modifying the proposal distribution), may be a more effective use of test-time compute than sampling  $N$  independent responses in parallel. On the other hand, with more difficult problems that may require searching over many different high-level approaches to solving the problem, re-sampling new responses independently in parallel or deploying tree-search against a process-based reward model is likely a more effective way to use test-time computation. This finding illustrates **the need to deploy an adaptive “compute-optimal” strategy for scaling test-time compute**, wherein the specific approach for utilizing test-time compute is selected depending on the prompt, so as to make the best use of additional computation. We also show that a notion of *question difficulty* (Section 3.3) from the perspective of the base LLM can be used to predict the efficacy of test-time computation, enabling us to practically instantiate this ‘compute-optimal’ strategy given a prompt. By appropriately allocating test-time compute in this way, we are able to greatly improve test-time compute scaling, surpassing the performance of a best-of- $N$  baseline while only using about **4x** less computation with both revisions and search (Sections 3.4 and 3.5).

Using our improved test-time compute scaling strategy, we then aim to understand to what extent test-time computation can effectively substitute for additional pretraining. We conduct a FLOPs-matched comparison between a smaller model with additional test-time compute and pretraining a **14x larger model**. We find that on easy and intermediate questions, and even hard questions (depending on the specific conditions on the pretraining and inference workload), additional test-time compute is often preferable to scaling pretraining. This finding suggests that rather than focusing purely on scaling pretraining, **in some settings it is more effective to pretrain smaller models with less compute, and then**

---

<sup>1</sup>Capability-specific finetuning is necessary to induce revision and verification capabilities into the base model on MATH since these capabilities are absent even in strong proprietary LLMs [Huang et al., 2023, Sharma et al., 2024]. However, we expect that future LLMs will be more effective at verification and revision due to both increased scale and the inclusion of additional data targeted specifically towards these capabilities [Snell et al., 2024, Blakeney et al., 2024, McAleese et al., 2024]. Therefore in order to make progress towards understanding scaling of test-time computation, we must use models finetuned for these capabilities. That said, we expect future models to be pretrained for such capabilities directly, therefore avoiding the need for capability-specific finetuning.

**apply test-time compute to improve model outputs.** That said, with the most challenging questions, we observe very little benefits from scaling up test-time compute. Instead, we find that on these questions, it is more effective to make progress by applying additional pretraining compute, demonstrating that current approaches to scaling test-time compute may not be 1-to-1 exchangeable with scaling pretraining. Overall, this suggests that even with a fairly naïve methodology, scaling up test-time computation can already serve to be more preferable to scaling up pretraining, with only more improvements to be attained as test-time strategies mature. Longer term, this hints at a future where fewer FLOPs are spent during pretraining and more FLOPs are spent at inference.

### 3.1 A Unified Perspective on Test-Time Computation: Proposer and Verifier

We first unify approaches for using test-time computation and then analyze some representative methods. First, we view the use of additional test-time compute through the lens of modifying the model’s predicted distribution *adaptively* at test-time, conditioned on a given prompt. Ideally, test-time compute should modify the distribution so as to generate better outputs than naïvely sampling from the LLM itself would. In general, there are two knobs to induce modifications to an LLM’s distribution: **(1) at the input level**: by augmenting the given prompt with an additional set of tokens that the LLM conditions on to obtain the modified distribution, or **(2) at the output level**: by sampling multiple candidates from the standard LM and performing surgery on these candidates. In other words, we could either modify the **proposal distribution** induced by the LLM itself such that it is an improvement over naïvely conditioning on the prompt or we could use some post-hoc **verifiers or scorers** to perform output modifications. This process is reminiscent of Markov chain Monte Carlo (MCMC) [Andrieu et al., 2003] sampling from a complex target distribution but by combining a simple proposal distribution and a score function. Modifying the proposal distribution directly by altering input tokens and using a verifier form two independent axes of our study.

**Modifying the proposal distribution.** One way to improve the proposal distribution is to directly optimize the model for a given reasoning task via RL-inspired finetuning methods such as STaR or ReST<sup>EM</sup> [Zelikman et al., 2022, Singh et al., 2024]. Note that these techniques do not utilize any additional input tokens but specifically finetune the model to induce an improved proposal distribution. Instead, techniques such as self-critique [Bai et al., 2022, Madaan et al., 2023, Du et al., 2023, Saunders et al., 2022] enable the model itself to improve its own proposal distribution at test time by instructing it to critique and revise its own outputs in an iterative fashion. Since prompting off-the-shelf models is not effective at enabling effective revisions at test time, we specifically finetune models to iteratively revise their answers in complex reasoning-based settings. To do so, we utilize the approach of finetuning on on-policy data with Best-of-N guided improvements to the model

response [Qu et al., 2024].

**Optimizing the verifier.** In our abstraction of the proposal distribution and verifier, the verifier is used to aggregate or select the best answer from the proposal distribution. The most canonical way to use such a verifier is by applying best-of-N sampling, wherein we sample N complete solutions and then select the best one according to a verifier [Cobbe et al., 2021]. However, this approach can be further improved by training a process-based verifier [Lightman et al., 2023], or a process reward model (PRM), which produces a prediction of the correctness of each intermediate step in a solution, rather than just the final answer. We can then utilize these per-step predictions to perform tree search over the space of solutions, enabling a potentially more efficient and effective way to search against a verifier, compared to naïve best-of-N [Yao et al., 2023, Feng et al., 2024, Chen et al., 2024].

## 3.2 How to Scale Test-Time Computation Optimally

Given the unification of various methods, we would now like to understand how to *most effectively* utilize test-time computation to improve LM performance on a given prompt. Concretely we wish to answer:

### Problem setup

We are given a prompt and a test-time compute budget within which to solve the problem. Under the abstraction above, there are different ways to utilize test-time computation. Each of these methods may be more or less effective depending on the specific problem given. How can we determine the *most effective* way to utilize test-time compute for a given prompt? And how well would this do against simply utilizing a much bigger pretrained model?

When either refining the proposal distribution or searching against a verifier, there are several different hyper-parameters that can be adjusted to determine how a test-time compute budget should be allocated. For example, when using a model finetuned for revisions as the proposal distribution and an ORM as the verifier, we could either spend the full test-time compute budget on generating N independent samples in parallel from the model and then apply best-of-N, or we could sample N revisions in sequence using a revision model and then select the best answer in the sequence with an ORM, or strike a balance between these extremes. Intuitively, we might expect “easier” problems to benefit more from revisions, since the model’s initial samples are more likely to be on the right track but may just need further refinement. On the other hand, challenging problems may require more exploration of different high-level problem solving strategies, so sampling many times independently in parallel may be preferable in this setting.

In the case of verifiers, we also have the option to choose between different search algorithms (e.g. beam-search, lookahead-search, best-of-N), each of which may exhibit different properties depending on the quality of the verifier and proposal distribution at hand. More

sophisticated search procedures might be more useful in harder problems compared to a much simpler best-of-N or majority baseline.

## Test-Time Compute-Optimal Scaling Strategy

In general, we would therefore like to select the *optimal* allocation of our test-time compute budget for a given problem. To this end, for any given approach of utilizing test-time compute (e.g., revisions and search against a verifier in this chapter, various other methods elsewhere), we define the “**test-time compute-optimal scaling strategy**” as the strategy that chooses hyperparameters corresponding to a given test-time strategy for maximal performance benefits on a given prompt at test time. Formally, define  $\text{Target}(\theta, N, q)$  as the distribution over natural language output tokens induced by the model for a given prompt  $q$ , using test-time compute hyper-parameters  $\theta$ , and a compute budget of  $N$ . We would like to select the hyper-parameters  $\theta$  which maximize the accuracy of the target distribution for a given problem. We express this formally as:

$$\theta_{q,a^*(q)}^*(N) = \operatorname{argmax}_{\theta} \left( \mathbb{E}_{y \sim \text{Target}(\theta, N, q)} [\mathbf{1}_{y=y^*(q)}] \right), \quad (3.1)$$

where  $y^*(q)$  denotes the ground-truth correct response for  $q$ , and  $\theta_{q,y^*(q)}^*(N)$  represents the test-time compute-optimal scaling strategy for the problem  $q$  with compute budget  $N$ .

## Estimating Question Difficulty for Compute-Optimal Scaling

In order to effectively analyze the test-time scaling properties of the different mechanisms discussed in Section 3.1 (e.g. the proposal distribution and the verifier), we will prescribe an approximation to this optimal strategy  $\theta_{q,y^*(q)}^*(N)$  as a function of a statistic of a given prompt. This statistic estimates a notion of *difficulty* for a given prompt. The compute-optimal strategy is defined as a function of the difficulty of this prompt. Despite being only an approximate solution to the problem shown in Equation 3.1, we find that it can still induce substantial improvements in performance over a baseline strategy of allocating this inference-time compute in an ad-hoc or uniformly-sampled manner.

Our estimate of the question difficulty assigns a given question to one of five difficulty levels. We can then use this discrete difficulty categorization to estimate  $\theta_{q,y^*(q)}^*(N)$  on a validation set for a given test-time compute budget. We then apply these compute-optimal strategies on the test-set. Concretely, we select the best performing test-time compute strategy for each difficulty bin independently. In this way, question difficulty acts as a sufficient statistic of a question when designing the compute-optimal strategy.

**Defining difficulty of a problem.** Following the approach of Lightman et al. [2023], we define question difficulty as a function of a given base LLM. Specifically, we bin the model’s pass@1 rate – estimated from 2048 samples – on each question in the test set into five quantiles, each corresponding to increasing difficulty levels. We found this notion of model-specific difficulty bins to be more predictive of the efficacy of using test-time compute in contrast to the hand-labeled difficulty bins in the MATH dataset.

That said, we do note that assessing a question’s difficulty as described above assumes oracle access to a ground-truth correctness checking function, which is of course not available upon deployment where we are only given access to test prompts that we don’t know the answer to. In order to be feasible in practice, a compute-optimal scaling strategy conditioned on difficulty needs to first assess difficulty and then utilize the right scaling strategy to solve this problem. Therefore, we approximate the problem’s difficulty via a **model-predicted notion of difficulty**, which performs the same binning procedure over the averaged final answer score from a learned verifier (and not groundtruth answer correctness checks) on the same set of 2048 samples per problem. We refer to this setting as **model-predicted difficulty** and the setting which relies on the ground-truth correctness as **oracle difficulty**.

While model-predicted difficulty removes the need for knowing the ground truth label, estimating difficulty in this way still incurs additional computation cost during inference. That said, this one-time inference cost can be subsumed within the cost for actually running an inference-time strategy (e.g., when using a verifier, one could use the same inference computation for also running search). More generally, this is akin to exploration-exploitation tradeoff in reinforcement learning: in actual deployment conditions, we must balance the compute spent in assessing difficulty vs applying the most compute-optimal approach. This is a crucial avenue for future work (see Section 3.7) and our experiments do not account for this cost largely for simplicity, since our goal is to present some of the first results of *what is in fact possible* by effectively allocating test-time compute.

So as to avoid confounders with using the same test set for computing difficulty bins and for selecting the compute-optimal strategy, we use two-fold cross validation on each difficulty bin in the test set. We select the best-performing strategy according to performance on one fold and then measure performance using that strategy on the other fold and vice versa, averaging the results of the two test folds.

### 3.3 Experimental Setup

We first outline our experimental setup for conducting this analysis with multiple verifier design choices and proposal distributions, followed by the analysis results in the subsequent sections.

**Datasets.** We expect test-time compute to be most helpful when models already have all the basic “knowledge” needed to answer a question, and instead the primary challenge is about drawing (complex) inferences from this knowledge. To this end, we focus on the MATH [Hendrycks et al., 2021b] benchmark, which consists of high-school competition level math problems with a range of difficulty levels. For all experiments, we use the dataset split consisting of 12k train and 500 test questions, used in Lightman et al. [2023].

**Models.** We conduct our analysis using the PaLM 2-S\* [Anil et al., 2023] (Codey) base model. We believe this model is representative of the capabilities of many contemporary LLMs, and therefore think that our findings likely transfer to similar models. Most impor-

tantly, this model attains a non-trivial performance on MATH and yet has not saturated, so we expect this model to provide a good test-bed for us.

### 3.4 Scaling Test-Time Compute via Verifiers

In this section we analyze how test-time compute can be scaled by optimizing a verifier, as effectively as possible. To this end, we study different approaches for performing test-time search with process verifiers (PRMs) and analyze the test-time compute scaling properties of these different approaches.

#### Training Verifiers Amenable to Search

**PRM training.** Originally PRM training [Uesato et al., 2022, Lightman et al., 2023] used human crowd-worker labels. While Lightman et al. [2023] released their PRM training data (i.e., the PRM800k dataset), we found this data to be largely ineffective for us. We found that it was easy to exploit a PRM trained on this dataset via even naïve strategies such as best-of-N sampling. We hypothesize that this is likely a result of the distribution shift between the GPT-4 generated samples in their dataset and our PaLM 2 models. Rather than proceeding with the expensive process of collecting crowd-worker PRM labels for our PaLM 2 models, we instead apply the approach of Wang et al. [2023a] to supervise PRMs without human labels, using estimates of per-step correctness obtained from running Monte Carlo rollouts from each step in the solution. Our PRM’s per-step predictions therefore correspond to value estimates of reward-to-go for the base model’s sampling policy, similar to recent work [Wang et al., 2023a, Sethur et al., 2024]. We also compared to an ORM baseline (Appendix B.5) but found that our PRM consistently outperforms the ORM. Hence, all of the search experiments in this section use a PRM model. Additional details on PRM training are shown in Appendix B.3.

**Answer aggregation.** At test time, process-based verifiers can be used to score each individual step in a set of solutions sampled from the base model. In order to select the best-of-N answers with the PRM, we need a function that can aggregate across all the per-step scores for each answer to determine the best candidate for the correct answer. To do this, we first aggregate each individual answer’s per-step scores to obtain a final score for the full answer (step-wise aggregation). We then aggregate across answers to determine the best answer (inter-answer aggregation). Concretely, we handle step-wise and inter-answer aggregation as follows:

- **Step-wise aggregation.** Rather than aggregating the per-step scores by taking the product or minimum [Wang et al., 2023a, Lightman et al., 2023], we instead use the PRM’s prediction at the last step as the full-answer score. We found this to perform the best out of all aggregation methods we studied (see Appendix B.4).
- **Inter-answer aggregation.** We follow Li et al. [2023] and apply “best-of-N weighted” selection rather than standard best-of-N. Best-of-N weighted selection marginalizes the

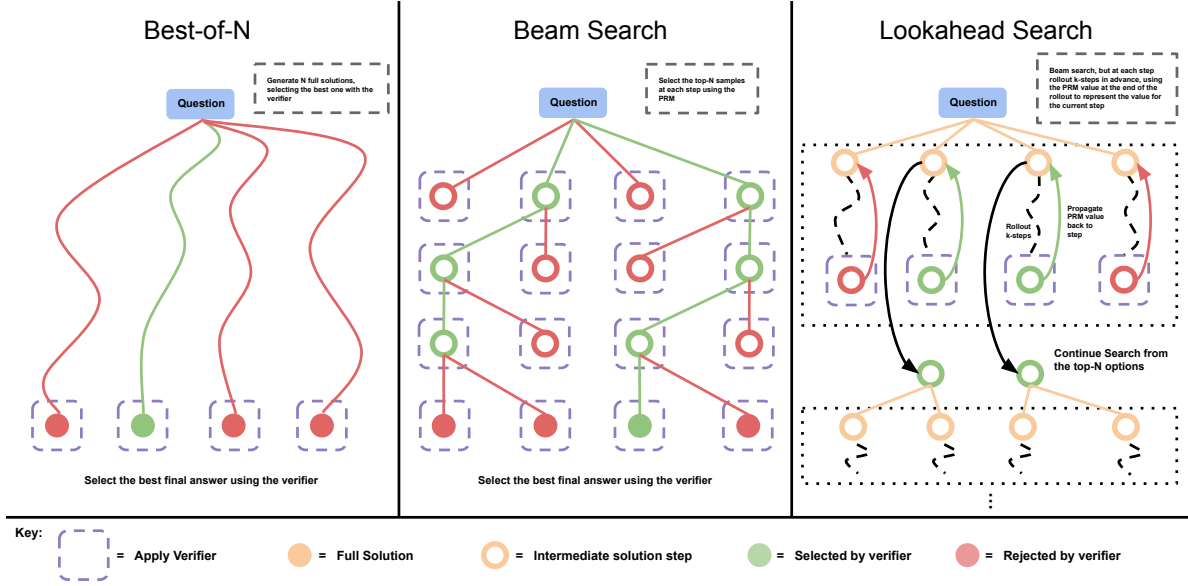


Figure 3.2: *Comparing different PRM search methods.* **Left:** Best-of- $N$  samples  $N$  full answers and then selects the best answer according to the PRM final score. **Center:** Beam search samples  $N$  candidates at each step, and selects the top  $M$  according to the PRM to continue the search from. **Right:** lookahead-search extends each step in beam-search to utilize a  $k$ -step lookahead while assessing which steps to retain and continue the search from. Thus lookahead-search needs more compute.

verifier’s correctness scores across all solutions with the same final answer, selecting final answer with the greatest total sum.

## Search Methods Against a PRM

We optimize the PRM at test time via search methods. We study three search approaches that sample outputs from a few-shot prompted base LLM (see Appendix B.6). An illustration is shown in Figure 3.2.

**Best-of- $N$  weighted.** We sample  $N$  answers independently from the base LLM and then select the best answer according to the PRM’s final answer judgement.

**Beam search.** Beam search optimizes the PRM by searching over its per-step predictions. Our implementation is similar to BFS-V [Yao et al., 2023, Feng et al., 2024]. Concretely, we consider a fixed number of beams  $N$  and a beam width  $M$ . We then run the following steps:

1. sample  $N$  initial predictions for the first step in the solution
2. score the generated steps according to the PRM’s predicted step-wise reward-to-go estimate (which also corresponds to the total reward from the prefix since the reward is sparse in this setting)

3. filter for only the top  $\frac{N}{M}$  highest scoring steps
4. now from each candidate, sample  $M$  proposals from the next step, resulting in a total of  $N/M \times M$  candidate prefixes again. Then repeat steps 2-4 again.

We run this algorithm until the end of a solution or the maximum number of rounds of beam expansion are attained (40 in our case). We conclude the search with  $N$  final answer candidates, to which we apply best-of- $N$  weighted selection described above to make our final answer prediction.

**Lookahead search.** Lookahead search modifies how beam search evaluates individual steps. It uses lookahead rollouts to improve the accuracy of the PRM’s value estimation in each step of the search process. Specifically, at each step in the beam search, rather than using the PRM score at the current step to select the top candidates, lookahead search performs a simulation, rolling out up to  $k$  steps further while stopping early if the end of solution is reached. To minimize variance in the simulation rollout, we perform rollouts using temperature 0. The PRM’s prediction at the end of this rollout is then used to score the current step in the beam search. That is, in other words, we can view beam search as a special case of lookahead search with  $k = 0$ . Given an accurate PRM, increasing  $k$  should improve the accuracy of the per-step value estimates at the cost of additional compute. Also note that this version of lookahead search is a special case of MCTS [Sutton and Barto, 2018], wherein the stochastic elements of MCTS, designed to facilitate exploration, are removed since the PRM is already trained and is frozen. These stochastic elements are largely useful for learning the value function (which we’ve already learned with our PRM), but less useful at test-time when we want to exploit rather than explore. Therefore, lookahead search is largely representative of how MCTS-style methods would be applied at test-time.

## Analysis Results: Test-Time Scaling for Search with Verifiers

We now present our results comparing various search algorithms and identify a prompt difficulty dependent compute-optimal scaling strategy for search methods.

**Comparing search algorithms.** We first conduct a sweep over various search settings. In addition to the standard best-of- $N$  approach, we sweep over the two main parameters that distinguish different tree-search methods: beam-width  $M$  and number of lookahead steps  $k$ . While we are not able to extensively sweep every single configuration, we sweep over the following settings with a maximum budget of 256:

- 1) Beam search with the beam width set to  $\sqrt{N}$ , where  $N$  is the generation budget.
- 2) Beam search with a fixed beam width of 4.
- 3) Lookahead search with  $k = 3$  applied to both beam-search settings 1) and 2).
- 4) Lookahead search with  $k = 1$  applied to beam-search setting 1).

To compare search methods as a function of generation budget fairly, we build a protocol for estimating the cost of each method. We consider a generation to be a sampled answer

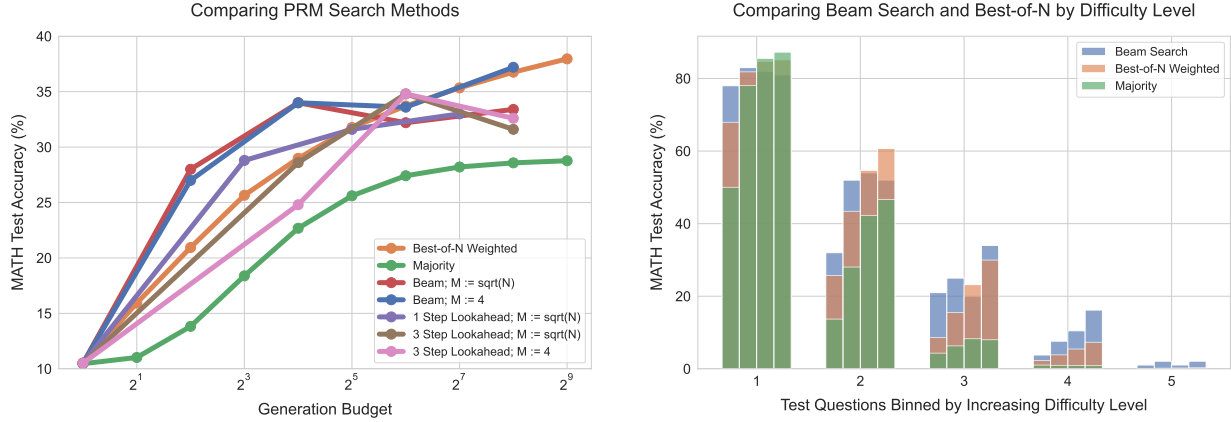


Figure 3.3: **Left: Comparing different methods for conducting search against PRM verifiers.** We see that at low generation budgets, beam search performs best, but as we scale the budget further the improvements diminish, falling below the best-of-N baseline. Lookahead-search generally underperforms other methods at the same generation budget. **Right: Comparing beam search and best-of-N binned by difficulty level.** The four bars in each difficulty bin correspond to increasing test-time compute budgets (4, 16, 64, and 256 generations). On the easier problems (bins 1 and 2), beam search shows signs of over-optimization with higher budgets, whereas best-of-N does not. On the medium difficulty problems (bins 3 and 4), we see beam search demonstrating consistent improvements over best-of-N.

from the base LLM. For beam search and best-of-N the generation budget corresponds to the number of beams and  $N$  respectively. Lookahead search, however, utilizes additional compute: at each step of the search, we sample  $k$  additional steps ahead. Therefore, we define the cost of lookahead-search to be  $N \times (k + 1)$  samples.

**Results.** As shown in Figure 3.3 (left), with smaller generation budgets, beam search significantly outperforms best-of-N. However, as the budget is scaled up, these improvements greatly diminish, with beam search often underperforming the best-of-N baseline. We also see that, lookahead-search generally underperforms other methods at the same generation budget, likely due to the additional computation induced by simulating the lookahead rollouts. The diminishing returns from search are likely due to exploitation of the PRM’s predictions. For example, we see some instances (such as in Figure B.20), where search causes the model to generate low-information repetitive steps at the end of a solution. In other cases, we find that over-optimizing search can result in overly short solutions consisting of just 1-2 steps. This explains why the most powerful search method (i.e., lookahead search) underperforms the most. We include several of these examples found by search in Appendix B.12.

**Which problems does search improve?** To understand how to compute-optimally scale search methods, we now conduct a difficulty bin analysis. Specifically, we compare beam-search ( $M = 4$ ) against best-of-N. In Figure 3.3 (right) we see that while in aggregate,

beam search and best-of-N perform similarly with a high generation budget, evaluating their efficacy over difficulty bins reveals very different trends. On the easy questions (levels 1 and 2), the stronger optimizer of the two approaches, beam search, degrades performance as the generation budget increases, suggesting signs of exploitation of the PRM signal. In contrast, on the harder questions (levels 3 and 4), beam search consistently outperforms best-of-N. On the most difficult questions (level 5), no method makes much meaningful progress.

These findings match intuition: we might expect that on the easy questions, the verifier will make mostly correct assessments of correctness. Therefore, by applying further optimization via beam search, we only further amplify any spurious features learned by the verifier, causing performance degradation. On the more difficult questions, the base model is much less likely to sample the correct answer in the first place, so search can serve to help guide the model towards producing the correct answer more often.

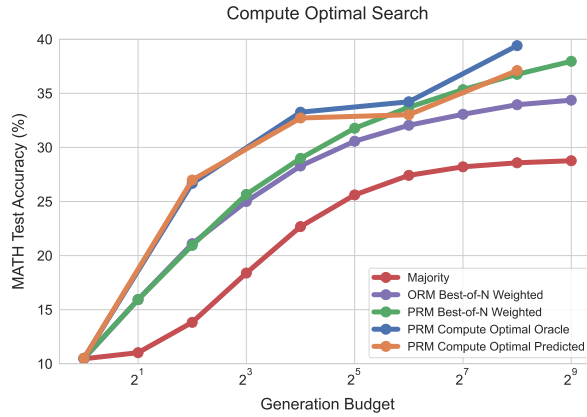


Figure 3.4: *Comparing compute-optimal test-time compute allocation against base-lines with PRM search.* By scaling test time compute per the notion of question difficulty, we find that we can nearly outperform PRM best-of-N using up to  $4\times$  less test-time compute (e.g. 16 versus 64 generations). “**Compute-optimal oracle**” refers to using oracle difficulty bins derived from the groundtruth correctness information, and “**compute-optimal predicted**” refers to using the PRM’s predictions to generate difficulty bins. Observe that the curves with either type of difficulty bins largely overlap with each other.

**Compute-optimal search.** Given the above results, it is clear that question difficulty can be a useful statistic to predict the optimal search strategy to use at a given compute budget. Additionally, the best choice of search strategy can vary drastically as a function of this difficulty statistic. We therefore visualize the “compute-optimal” scaling trend, as represented by the best performing search strategy at each difficulty level in Figure 3.4. We see that in the low generation budget regime, using both the oracle and predicted difficulty, **compute-optimal scaling can nearly outperform best-of-N using up to  $4x$  less test-time compute** (e.g. 16 versus 64 generations). While in the higher budget regime, some of these benefits diminish with the use of predicted difficulty, with oracle bins we still see

continued improvements from optimally scaling test-time compute. This result demonstrates the performance gains that could be obtained by adaptively allocating test-time compute during search.

**Takeaways for compute-optimal scaling of verifiers**

We find that the efficacy of any given verifier search method depends critically on both the compute budget and the question at hand. Specifically, beam-search is more effective on harder questions and at lower compute budgets, whereas best-of-N is more effective on easier questions and at higher budgets. Moreover, by selecting the best search setting for a given question difficulty and test-time compute budget, we can nearly outperform best-of-N using up to  $4x$  less test-time compute.

### 3.5 Refining the Proposal Distribution

So far, we studied the test-time compute scaling properties of search against verifiers. Now we turn to studying the scaling properties of modifying the proposal distribution (Section 3.1). Concretely, we enable the model to revise their own answers iteratively, allowing the model to dynamically improve its own distribution at test time. Simply prompting existing LLMs to correct their own mistakes tends to be largely ineffective for obtaining performance improvements on reasoning problems [Huang et al., 2023]. Therefore, we build on the recipe prescribed by Qu et al. [2024], incorporate modifications for our setting, and finetune language models to iteratively revise their own answers. We first describe how we train and use models that refine their own proposal distribution by sequentially conditioning on their own previous attempts at the question. We then analyze the inference-time scaling properties of revision models.

#### Setup: Training and Using Revision Models

Our procedure for finetuning revision models is similar to [Qu et al., 2024], though we introduce some crucial differences. For finetuning, we need trajectories consisting of a sequence of incorrect answers followed by a correct answer, that we can then run SFT on. Ideally, we want the correct answer to be *correlated* with the incorrect answers provided in context, so as to effectively teach the model to *implicitly* identify mistakes in examples provided in-context, followed by correcting those mistakes by making edits as opposed to ignoring the in-context examples altogether, and trying again from scratch.

**Generating revision data.** The on-policy approach of Qu et al. [2024] for obtaining several multi-turn rollouts was shown to be effective, but it was not entirely feasible in our infrastructure due to compute costs associated with running multi-turn rollouts. Therefore, we sampled 64 responses *in parallel* at a higher temperature and post-hoc constructed multi-turn rollouts from these independent samples. Specifically, following the recipe of [Kumar et al.], we pair up each correct answer with a sequence of incorrect answers from this set

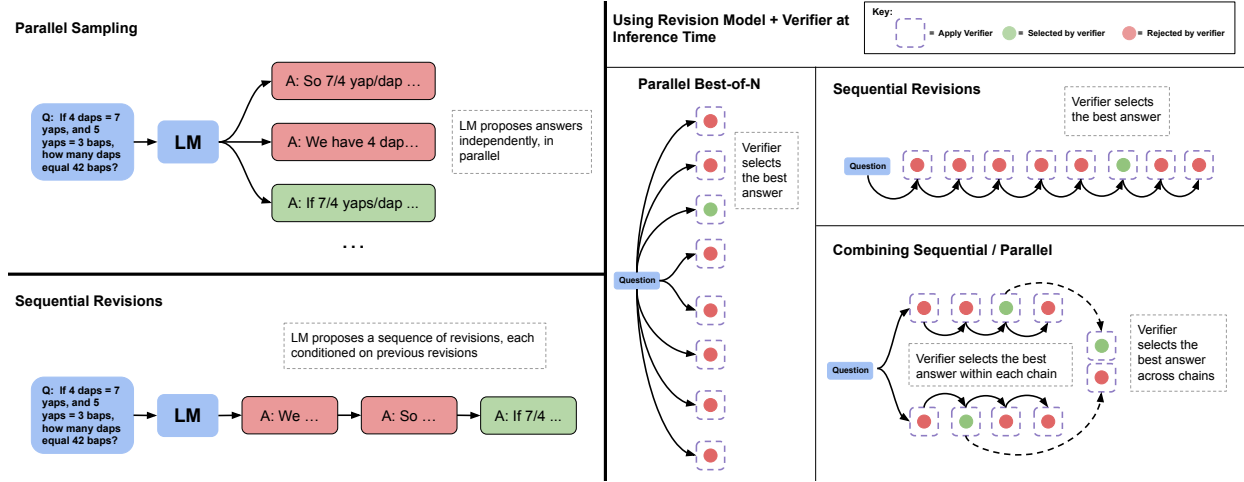


Figure 3.5: *Parallel sampling (e.g., Best-of-N) versus sequential revisions*. **Left:** Parallel sampling generates  $N$  answers independently in parallel, whereas sequential revisions generates each one in sequence conditioned on previous attempts. **Right:** In both the sequential and parallel cases, we can use the verifier to determine the best-of- $N$  answers (e.g. by applying best-of- $N$  weighted). We can also allocate some of our budget to parallel and some to sequential, effectively enabling a combination of the two sampling strategies. In this case, we use the verifier to first select the best answer within each sequential chain and then select the best answer across chains.

as context to construct multi-turn finetuning data. We include up to four incorrect answers in context, where the specific number of solutions in context is sampled randomly from a uniform distribution over categories 0 to 4. We use a character edit distance metric to prioritize selecting incorrect answers which are correlated with the final correct answer (see Appendix B.7). Note that token edit distance is not a perfect measure of correlation, but we found this heuristic to be sufficient to correlate incorrect in-context answers with correct target answers to facilitate training a meaningful revision model, as opposed to randomly pairing incorrect and correct responses with uncorrelated responses.

**Using revisions at inference-time.** Given a finetuned revision model, we can then sample a sequence of revisions from the model at test-time. While our revision model is only trained with up to four previous answers in-context, we can sample longer chains by truncating the context to the most recent four revised responses. In Figure 3.6 (left), we see that as we sample longer chains from the revision model, the model’s pass@1 at each step gradually improves, demonstrating that we are able to effectively teach the model to learn from mistakes made by previous answers in context.

That said, there is a distribution shift at inference time: the model was trained on only sequences with incorrect answers in context, but at test-time the model may sample correct answers that are included in the context. In this case, it may incidentally turn the correct answer into an incorrect answer in the next revision step. We find that indeed, similar

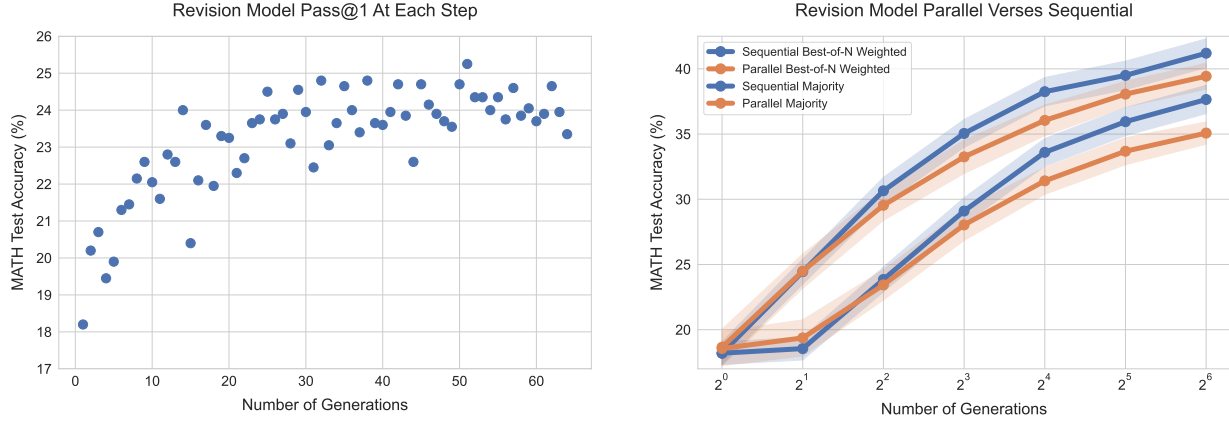


Figure 3.6: **Left:** *Our revision model’s pass@1 at each revision step.* Pass@1 gradually improves after each revision step, even improving beyond the 4 revision steps that it was trained for. We estimate pass@1 at each step by averaging over the performance of 4 revision trajectories of length 64 for each question in the test-set. **Right:** *Sequential vs parallel sampling from the revision model.* Comparing performance when generating  $N$  initial answers in parallel from our revision model, versus generating  $N$  revisions sequentially, with the model. When using both the verifier and majority voting to select the answer, we see that generating answers sequentially with the revision model narrowly outperforms generating them in parallel.

to Qu et al. [2024], around 38% of correct answers get converted back to incorrect ones with our revision model using a naïve approach. Therefore, we employ a mechanism based on sequential majority voting or verifier-based selection to select the most correct answer from the sequence of revisions made by the model (see Figure 3.5) to produce the best answer.

**Comparisons.** To test the efficacy of modifying the proposal distribution via revisions, we setup an even comparison between the performance of sampling  $N$  revisions in sequence and sampling  $N$  attempts at a question in parallel. We see in Figure 3.6 (right), that with both the verifier-based and majority-based selection mechanisms sampling solutions in sequence outperforms sampling them in parallel.

## Analysis Results: Test-Time Scaling with Revisions

We saw previously that proposing answers sequentially outperforms proposing them in parallel. However, we might expect sequential and parallel sampling to have different properties. Sampling answers in parallel may act as more of a global search process, that could in principle, provide coverage over many totally different approaches for solving a problem, for instance, different candidates might utilize different high-level approaches altogether. Sequential sampling, on the other hand, may work more as a local refinement process, revising responses that are already somewhat on the right track. Due to these complementary

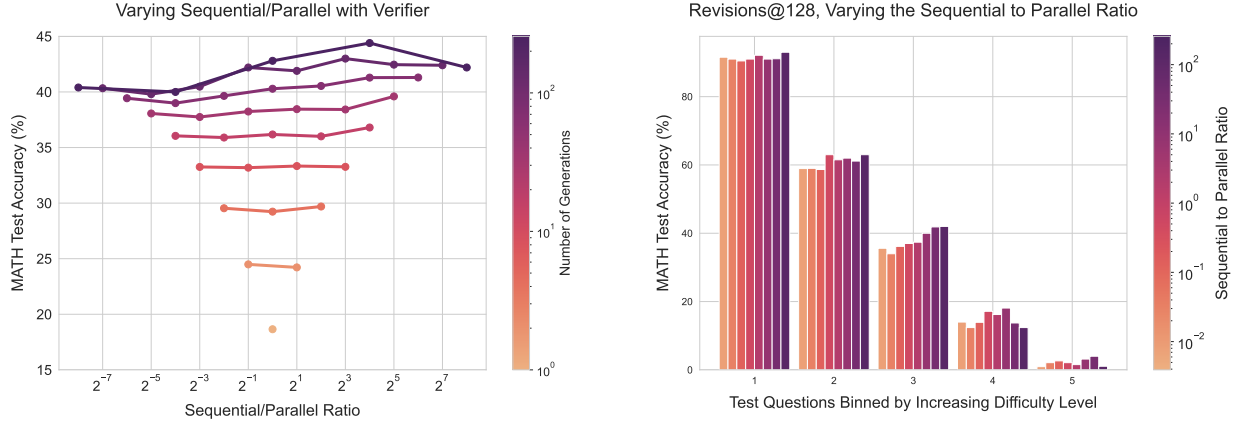


Figure 3.7: **Left:** *Varying the ratio of the generation budget allocated to sequential revisions versus parallel samples.* Each line represents a fixed generation budget as the ratio is changed. We use the verifier for answer selection. We see that while increased sequential revisions tends to outperform more parallel compute, at higher generation budgets there is an ideal ratio that strikes a balance between the two extremes. **Right:** *Varying the sequential to parallel ratio for a generation budget of 128 across difficulty bins.* Using verifier-based selection, we see that the easier questions attain the best performance with full sequential compute. On the harder questions, there is an ideal ratio of sequential to parallel test-time compute.

benefits, we should strike a balance between these two extremes by allocating some of our inference-time budget to parallel sampling (e.g.  $\sqrt{N}$ ) and the rest to sequential revisions (e.g.  $\sqrt{N}$ ). We will now show the existence of a compute-optimal ratio between sequential and parallel sampling, and understand their relative pros and cons based on difficulty of a given prompt.

**Trading off sequential and parallel test-time compute.** To understand how to optimally allocate sequential and parallel compute, we perform a sweep over a number of different ratios. We see, in Figure 3.7 (left), that indeed, at a given generation budget, **there exists an ideal sequential to parallel ratio, that achieves the maximum accuracy.** We also see in Figure 3.7 (right) that the ideal ratio of sequential to parallel varies depending on a given question’s difficulty. In particular, easy questions benefit more from sequential revisions, whereas on difficult questions it is optimal to strike a balance between sequential and parallel computation. This finding supports the hypothesis that sequential revisions (i.e., varying the proposal distribution) and parallel sampling (i.e., search with verifiers) are two complementary axes for scaling up test-time compute, which may be more effective on a per-prompt basis. We include examples of our model’s generations in Appendix B.11. Additional results are shown in Appendix B.1.

**Compute-optimal revisions.** Given that the efficacy of sequential and parallel sampling depends on question difficulty, we can select the ideal ratio of sequential to parallel

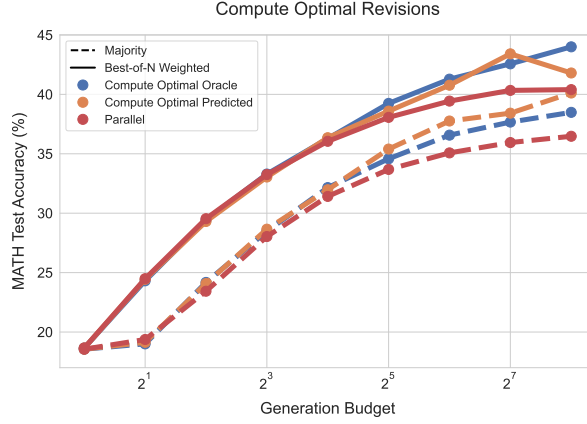


Figure 3.8: *Comparing compute-optimal test-time compute allocation against the parallel compute baseline with our revision model.* By optimally scaling test-time compute according to question difficulty, we find that we can outperform best-of-N using up to  $4x$  less test-time compute (e.g. 64 samples versus 256). “**Compute-optimal oracle**” refers to using the oracle difficulty bins derived from the ground truth correctness information, and “**compute optimal predicted**” refers to using the PRM’s predictions to produce model-predicted difficulty bins.

compute per difficulty bin. In Figure 3.8, we plot results using this compute-optimal scaling strategy when employing both our oracle and predicted notions of difficulty. In both cases, we’re able to substantially improve test-time compute scaling by improving the proposal distribution via revisions. In particular, we see that at higher generation budgets, parallel sampling seems to plateau, whereas compute-optimal scaling demonstrates continued improvements. For both oracle and predicted difficulty bins, we see that **compute-optimal scaling can outperform best-of-N using up to  $4x$  less test-time compute** (e.g. 64 samples versus 256). Overall, these results demonstrate the potential for improved test-time compute scaling by adjusting the proposal distribution on a per-prompt basis.

**Takeaways for compute-optimal scaling by refining the proposal distribution with revisions**

We find that there exists a tradeoff between sequential (e.g. revisions) and parallel (e.g. standard best-of-N) test-time computation, and the ideal ratio of sequential to parallel test-time compute depends critically on both the compute budget and the specific question at hand. Specifically, easier questions benefit from purely sequential test-time compute, whereas harder questions often perform best with some ideal ratio of sequential to parallel compute. Moreover, by optimally selecting the best setting for a given question difficulty and test-time compute budget, we can outperform the parallel best-of-N baseline using up to  $4x$  less test-time compute.

### 3.6 Putting it Together: Exchanging Pretraining and Test-Time Compute

So far, we saw that utilizing additional test-time computation can enable us to represent more complex distributions than the one predicted by the base LLM itself, thereby improving performance. We now posit that this increased flexibility of representing distributions means that we can expect additional test-time compute to make up for the lack of a higher-capacity model or training for more FLOPs during pretraining. In this section, *we study to what extent this is possible*. We pose the following question:

**Question: Exchanging pretraining and test-time compute**

Suppose a model was pre-trained with  $X$  FLOPs. Assume that we plan to run  $Y$  FLOPs of inference with this model. If we want to improve performance by increasing the total FLOPs budget by a factor of  $M$  (i.e.,  $M(X + Y)$  total FLOPs across both pretraining and inference), should we spend our FLOPs on increased pretraining compute or on additional test-time compute?

Increasing pretraining FLOPs introduces the additional design decision of whether to allocate compute to training with more data or more parameters [Hoffmann et al., 2022]. We focus on the setting in which model parameters are scaled up and training data amount is fixed, matching the approach taken with the open-source LLaMA series of models [Touvron et al., 2023b]. We choose this setting as it is representative of a canonical approach to scaling pretraining compute and leave the analysis of compute-optimal scaling of pretraining compute [Sardana and Frankle, 2023] where the data and parameters are both scaled equally to future work.

**Defining an exchange rate between FLOPs.** We now describe how we define the exchange rate between pretraining and inference FLOPs. To determine pretraining FLOPs, we use the common approximation  $X = 6ND_{\text{pretrain}}$  [Hoffmann et al., 2022], and for inference FLOPs, we use  $Y = 2ND_{\text{inference}}$  [Sardana and Frankle, 2023]. Here  $N$  represents model parameters,  $D_{\text{pretrain}}$  is the number of tokens used for pretraining, and  $D_{\text{inference}}$  the total number of tokens generated at inference time. With these approximations, we can see that, if we multiply the model parameters by a factor of  $M$ , then both the pretraining and inference FLOPs (due to the cost of greedy decoding with the larger model), increase by a factor of  $M$  (giving  $M(X + Y)$  total FLOPs).

To match the FLOPs from scaling up model parameters using test-time compute with the smaller model we can multiply the smaller model’s inference compute by a factor of  $M + 3\left(\frac{D_{\text{pretrain}}}{D_{\text{inference}}}\right)(M - 1)$ . Notably, the amount of inference compute we can utilize to match the FLOPs for the larger model depends on the ratio  $\frac{D_{\text{pretrain}}}{D_{\text{inference}}}$ . We refer to the inverse of this ratio as  $R$  (e.g.  $\frac{D_{\text{inference}}}{D_{\text{pretrain}}}$ ). Depending on the specific production setting or use-case, we should expect very different values of  $R$ . In particular, in many large scale production settings, we may expect significantly more inference tokens than pretraining tokens, in which

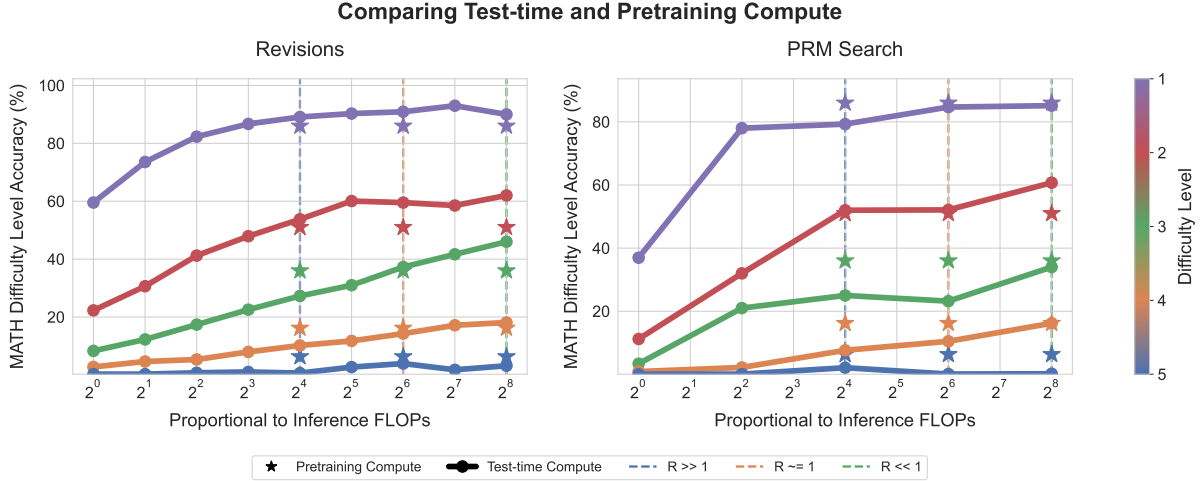


Figure 3.9: *Tradeoff between pretraining and test-time compute in a FLOPs-matched evaluation.* Each line represents the performance of scaling test-time compute with our compute-optimal policy in each oracle difficulty bin. We plot the results for revisions on the left and search on the right. The stars represent the greedy pass@1 performance of a base model pretrained with  $\sim 14$  times more parameters. We plot test-time compute budget on the x-axis, and place the stars at three different locations along the x-axis, each corresponding to the FLOPs equivalent point of comparison between scaling parameters and scaling test-time compute for three different inference compute loads (e.g.  $R = \frac{D_{\text{inference}}}{D_{\text{pretrain}}}$ ). If the star is below the line, this implies that it is more effective to use test-time compute than to scale model parameters, and if the star is above the line this implies that scaling parameters is more effective. We see that on the easy questions or in settings with a lower inference load (e.g.  $R \ll 1$ ), test-time compute can generally outperform scaling model parameters. However, on the harder questions or in settings with a higher inference load (e.g.  $R \gg 1$ ), pretraining is a more effective way to improve performance.

case we would have  $R \gg 1$ . On the other hand, in many contemporary self-improvement setups, that would use test-time compute to improve the model, we would likely generate significantly fewer inference tokens than pretraining tokens, giving  $R \ll 1$ . Therefore, since the scale of test-time compute we can apply is dependent on this ratio, we expect differing conclusions depending on the specific setting.

In Figure 3.9, we use this approach to exchanging test-time and pretraining compute to compare our compute-optimal scaling against scaling up model parameters by a factor of  $\sim 14$ . We conduct comparisons for 3 different values of  $R$ : 0.16 ( $R \ll 1$ ), 0.79 ( $R \sim 1$ ), and 22 ( $R \gg 1$ ), with each ratio corresponding to an inference budget. Observe that if we only expect to see very difficult questions (e.g. difficulty bins 4/5) or have a larger  $D_{\text{inference}}$  (corresponding to a larger  $R$  value), then it is often more effective to allocate our budget towards pretraining (e.g. the star is above the line). If instead, we expect mostly easy or intermediate difficulty questions (e.g. bins 1/2/3 and sometimes 4) or have lower inference

requirements (as is the case in self-improvement pipelines), then utilizing test-time compute is better.

**Takeaways for exchanging pretraining and test-time compute**

Test-time and pretraining compute are not 1-to-1 “exchangeable”. On easy and medium questions, which are within a model’s capabilities, or in settings with small inference requirement, test-time compute can easily cover up for additional pretraining. However, on challenging questions which are outside a given base model’s capabilities or under higher inference requirement, pretraining is likely more effective for improving performance.

### 3.7 Conclusion

In this work, we conducted a thorough analysis of the efficacy of two different high-level techniques for scaling up test-time compute on math reasoning tasks. In general, we found that the efficacy of a given approach heavily correlates with the difficulty of the problem from the perspective of the base LLM’s capabilities. This motivated us to introduce the notion of “compute-optimal” scaling of test-time computation, which prescribes an adaptive, prompt-dependent strategy to improve performance under a given test-time compute budget. By applying such a compute-optimal scaling strategy, we find that we can improve the efficiency of test-time compute scaling by a factor of  $2 - 4\times$ . When comparing benefits obtained from additional test-time compute against benefits from additional pretraining compute in a FLOPs-matched setting, we show for the first time that using test-time computation with seemingly simple methods (i.e., revisions and search) can already scale well on certain types of prompts, providing gains over spending those FLOPs in pretraining.

## Chapter 4

# Bridging the Gap Between Learning and Test-time Scaling

*This chapter is based on the following papers: “Learning by Distilling Context” [Snell et al., 2022], “Sleep-time Compute: Beyond Inference Scaling at Test-time” [Lin et al., 2025]*

As shown in Chapter 3, test-time scaling can be an effective alternative to pretraining in certain settings. However, it suffers from several limitations: 1) **Latency**: allowing models to use additional computation at test-time often comes at the cost of additional latency for users, which can lead to a suboptimal user experience; 2) **Redundant Computations**: LLMs treat each prompt independently and as a result often-times carry out the same test-time computations across multiple related queries, which can make for a wasteful use of computational resources.

To overcome these limitations, we build on our findings in Chapter 2, in which we observed that task-specific finetuning for narrow tasks can be far more efficient than broadly improving model capabilities. Specifically, we hypothesize that by applying additional computation ahead of time targeted towards preparing for a narrow distribution of potential upcoming tasks, we may be able to substantially improve the efficiency at which the model can answer upcoming queries at test-time without meaningfully sacrificing performance. More concretely, such a method could aim to predict a distribution of upcoming queries a user may ask, and then apply compute ahead of time to prepare for these queries pre-emptively.

We propose two different approaches to this end. Both of these techniques can be used to pre-emptively distill the processes of test-time reasoning on predicted queries into either the model’s weights (Context Distillation; Section 4.1) or the model’s context window (Sleep-time compute; Section 4.2). In this way, Context Distillation and Sleep-time compute can be seen as effectively converting the stateless, instance-specific nature of test-time computation back into generalizable learned representations, effectively acting as a bridge between the general scaling paradigms of learning (e.g. pretraining / in-context learning) and search

(e.g. test-time scaling) with LLMs. Next, I will describe both of these methods in detail in the following sections.

## 4.1 Context Distillation

Recent work has shown that language models significantly benefit from context tokens. When prompted with task definitions, language models can perform zero-shot learning [Wei et al., 2022a, Sanh et al., 2022], and the performance further improves with additional in-context examples and explanations [Chen et al., 2022, Scheurer et al., 2022]. They also acquire the capability to perform more complex tasks by generating step-by-step reasoning in the context window before predicting the final answer [Nye et al., 2021b, Wei et al., 2022c, Zhou et al., 2022].

However, language models cannot *internalize* these performance gains, which disappear when the context tokens are gone. Consequently, we always need to pay extra computation for running inference on context tokens; this is undesirable, as sometimes the task instructions and the scratch-pad can be more than 10x longer than the actual task inputs. Furthermore, it is unclear how to leverage the context tokens when their total length exceeds the context window size. These shortcomings are analogous to how humans are slow at performing complex cognitive tasks [Wason and Evans, 1974] and can hold only a limited amount of information in the working memory [Baddeley, 1992].

Humans get around this by practicing. Consider, for example, learning to type your friends’ phone numbers. The first few times you type it, you need to consciously recall the number using working memory and slowly decide which button to press. After repeatedly typing the same number, it becomes a habit and you can type the number quickly without conscious reasoning. Through repeated practice, the knowledge of your friend’s phone number is “distilled” into your muscle memories.<sup>1</sup> This mechanism for distilling knowledge is critical for learning complex tasks because it allows us to incrementally build up our knowledge and skills, so that we can learn to accomplish increasingly complex tasks.

We propose to apply a similar method, Context Distillation, to fine-tune language models. For example, as shown in Figure 4.1, to make language models internalize the step-by-step addition capability, we first synthesize a large number of “practice” addition questions; we then ask the model to follow the more informative instruction to reason step-by-step before generating the target answer; finally, we fine-tune the language model to directly predict the answer conditioned on a simpler student prompt. As a result, by practicing on a lot of addition problems, the ability to add is distilled into its parameters. We formally state our generalized Context Distillation framework in Section 4.1.

Context distillation can apply to a wide range of settings: learning from abstract statements, learning from concrete examples, and learning from step-by-step reasoning (see Appendix sections C.5, and C.6 for the former two). Here we focus on one particular setting

---

<sup>1</sup>See declarative learning vs. procedural learning for a friendly but more in-depth discussion. [https://en.wikipedia.org/wiki/Declarative\\_learning](https://en.wikipedia.org/wiki/Declarative_learning)

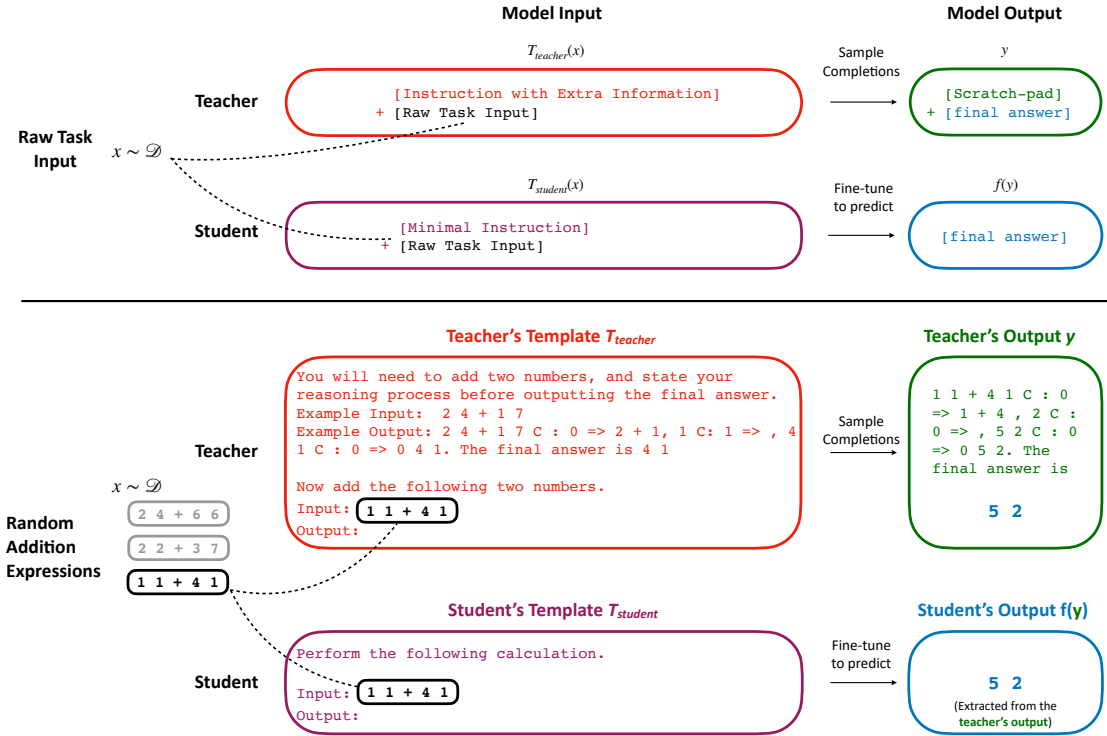


Figure 4.1: **Top.** An overview of our Context Distillation framework. We sample a raw task input, form the teacher’s prompt by prepending a detailed instruction that might contain more examples and explanations, and ask the language model to conditionally sample a scratch-pad and a final answer. Then we fine-tune the same language model to directly predict the final answer with a minimal instruction. We formalize this framework mathematically in Section 4.1. **Bottom.** An instantiation of our framework that internalizes step-by-step reasoning for 8-digit addition.

(Figure 4.1): we show that we can internalize step-by-step reasoning to perform 8-digit addition, and such a capability can transfer to downstream question answering tasks.

Overall, Context Distillation demonstrates promising potential as a general method for learning. With such a technique applied to pre-computing the processes of test-time reasoning on predicted queries, we can potentially reduce the compute costs at test-time without substantially sacrificing performance, alleviating both the latency and redundant computation limitations of test-time scaling.

## Intuition and Main Components

In the following section, we introduce the main components and the intuition of our Context Distillation framework, describe our algorithm for single round distillation, and describe various implementation details to make it efficient and stable.

We explain our method by contrasting it with the classical distillation methods [Hinton et al., 2015]. These classical methods ask the teacher model with parameter  $\theta_{\text{TEACHER}}$  to generate a label  $y$  for a given input  $x$ , and train the student  $\theta_{\text{STUDENT}}$  to mimic the teacher by predicting  $y$  conditioned on  $x$ . Typically,  $\theta_{\text{TEACHER}} \neq \theta_{\text{STUDENT}}$  when the algorithm starts, and the distillation process is driven by the difference between their parameters. In contrast, under Context Distillation,  $\theta_{\text{TEACHER}} = \theta_{\text{STUDENT}}$  when the training starts, and the distillation process is instead driven by the differences in the  $x$  and  $y$  that they see and predict.

To design such a difference that drives the distillation process, our framework requires the model developers to provide four components: a raw task input distribution  $\mathcal{D}$ , a teacher template  $T_{\text{TEACHER}}$ , a student template  $T_{\text{STUDENT}}$ , and an answer extractor  $f$ . We introduce them below.

**Raw Task Input Distribution  $\mathcal{D}$ .**  $\mathcal{D}$  is a distribution of strings, which are typically the “core” inputs of the target task of interest. For example, if the target task is to classify movie review sentiment, the input distribution could be defined as random movie reviews. Generally, there are many ways to define a raw task input distribution: we can use rule-based method to generate random strings, sample from a pool of unlabeled data, or conditionally sample from a language model. We explicitly distinguish raw task inputs from the whole “input prompt” to the language model, which will be obtained after applying the templates below.

**Teacher Template  $T_{\text{TEACHER}}$ .**  $T_{\text{TEACHER}}$  is a mapping from strings to strings, which transforms raw task inputs to the input prompts for the teacher model. As shown in Figure 4.1, the teacher template usually contains more detailed instructions, explanations, and examples about the task.

**Student Template  $T_{\text{STUDENT}}$ .**  $T_{\text{STUDENT}}$  is a mapping from strings to strings, which transforms raw task inputs to the input prompts for the student model. As shown in Figure 4.1, this template usually still contains minimal information about the task so that the request in the prompt is not under-specified. However, compared to the teacher prompt, it incorporates far fewer explanations and training examples of the task, and such a difference transfers this useful context information into the student parameters.

**Answer Extractor  $f$ .**  $f$  is a mapping from token sequences to token sequences, which extracts the final answer (a sub-sequence of tokens) from the full teacher’s generation. As shown in Figure 4.1,  $f$  strips away the intermediate reasoning process, and the students need

to internalize the step-by-step reasoning process to directly predict what the teacher predicts at the end.

We now describe Context Distillation formally using the mathematical terms we just introduced.

## Formal Description

Our algorithm first samples an  $x$  from  $\mathcal{D}$  and asks the language model to sample a completion  $y$  conditioned on  $T_{\text{TEACHER}}(x)$ ; we then fine-tune the language model to predict  $f(y)$  conditioned on  $T_{\text{STUDENT}}(x)$ . Throughout the distillation process,  $\theta_{\text{TEACHER}}$  is fixed.

Formally, let  $\theta_{\text{STUDENT}}$  and  $\theta_{\text{TEACHER}}$  be the parameters of a language model, and define  $P_{\theta}(\cdot|\text{PROMPT})$  to be the probability distribution of the completions conditioned on the prompt. We optimize

$$\mathcal{L}_{\mathcal{D}, T_{\text{STUDENT}}, T_{\text{TEACHER}}, f, \theta_{\text{TEACHER}}}(\theta_{\text{STUDENT}}) = \mathbb{E}_{x \sim \mathcal{D}} \left[ \mathbb{E}_{y \sim P_{\theta_{\text{TEACHER}}}(\cdot|T_{\text{TEACHER}})} \left[ \log P_{\theta_{\text{STUDENT}}}(f(y)|T_{\text{STUDENT}}(x)) \right] \right] \quad (4.1)$$

Notice that the definition of  $\mathcal{L}$  depends on five variables  $\mathcal{D}$ ,  $T_{\text{STUDENT}}$ ,  $T_{\text{TEACHER}}$ ,  $f$ , and  $\theta_{\text{TEACHER}}$ . To keep the notation uncluttered, we will only include the necessary subscripts if the rest can be inferred from the surrounding text.

## Implementation Details

A naive method to optimize Equation 4.1 is to sample  $y$  and directly fine-tune the student to predict the hard label of  $f(y)$ : such a method wastes the token logit information and results in noisy gradients. Instead, we minimize the token-level KL divergence between the student and the teacher (i.e., learn from the soft labels). However, this leads to another issue: the vocabulary space of language models is often on the order of 50k, and the full soft labels consume as much as 200KB memory per decoding step. Therefore, we approximate it with an empirical distribution of 100 token samples. Such a technique saves memory by a factor of 500.

## Context Distillation for Learning from Step-By-Step Reasoning

We show that Context Distillation can effectively internalize a skill acquired by generating step-by-step reasoning, and such a skill can benefit the model for other downstream tasks. Unless otherwise mentioned, in this section or the relevant appendix sections we use an  $f$  that will extract the final answer from the teacher’s full output, as shown in Figure 4.1 bottom.

|                             | Teach | Pre-Dist | Post-Dist | Sc→Dir | Sc+Dir |
|-----------------------------|-------|----------|-----------|--------|--------|
| 8 Digit Addition Accuracy % | 93    | 0        | <b>95</b> | 72     | 61     |

Table 4.1: Distilling addition scratchpads on T5-small. “Teach” refers to the teacher LM’s performance using scratch-pad. “Pre-Dist” refers to the student’s performance before distillation; “Post-Dist” refers to the student’s performance of direct addition (without scratch-pad) after distillation; “Sc→Dir”/“Sc+Dir” refers to our transfer/multi-task learning baseline. Context Distillation performs the best for direct addition.

**Dataset.** We define the input distribution  $\mathcal{D}$  to be uniform over all possible 1 through 8 digit addition questions, identical to those used by Nye et al. [2021a]. We report addition accuracy for each experiment. We hope that these results can generalize to more complex and realistic tasks that larger models can perform with chain-of-thoughts reasoning [Wei et al., 2022c, Zhou et al., 2022].

**Hypothesis 1: Context Distillation can internalize step-by-step reasoning.** To test this hypothesis, we obtain a teacher model that can use scratch-pad to perform step-by-step reasoning by fine-tuning the LM-adapted T5-small on a dataset of 500 addition expressions, where the model needs to first generate a scratch-pad before predicting the final answer. We then perform context-distillation with an  $f$  that extracts the final answer from the teacher’s output as shown in Figure 4.1 bottom. As shown in Table 4.1, after distillation, the ability of the student to perform **direct** additions (without using scratch-pad) improves from 0% to 94.7%, implying that Context Distillation internalizes step-by-step reasoning.

We compare this to several other transfer learning and multi-task learning baselines that use the same amount of training data in Table 4.1. Under transfer learning, we first fine-tune the student model to predict the scratch pad, and then fine-tune it to directly predict the final answer. Under multi-task learning, we fine-tune the model to predict scratch-pad and the final answer independently. Both variants perform significantly worse ( $> 20\%$ ) than Context Distillation. In addition to the gains in reasoning ability over baselines, we also saved inference time compute in our evaluations: specifically we used 8.0 times fewer tokens at inference time when evaluating the student compared to the teacher.

**Hypothesis 2: the reasoning capability internalized by scratchpad distillation can transfer to other related reasoning tasks.** To test this hypothesis, we distill the addition scratch-pads on our TK-Instruct model, and evaluate capability transfer to other tasks. Specifically, we trained a general purpose instruction following teacher language model (TK-Instruct) by fine-tuning LM-adapted T5-11B [Raffel et al., 2020] on a distribution of both natural instructions data and 500 addition scratchpad examples (see Appendix C.14). For the Context Distillation training, we define the student template to be a description of the addition task followed by two direct answer in-context examples, and similarly the

teacher template contains the task description and two in-context examples with scratch-pad answer. To prevent the student from catastrophically forgetting its in-context learning ability during distillation, we mix our 10k distillation data-points with a distribution of 65536 randomly selected examples from Natural Instruction-V2.

The student’s accuracy on directly answering 8-digit addition questions improves from 1% to 17%, while the student’s performance on Natural Instructions remains roughly the same (from RougeL 57 before distillation to 58 after), implying that the student did not lose its original capabilities to follow instructions. Additionally, we evaluate the student’s performance on a set of related reasoning tasks. We then use a template to synthesize simple questions that require knowledge to add two numbers, for example, “A has 7 turkeys. B has 2 turkeys. How many turkeys do they have altogether?”. On this synthetic dataset, the student’s performance significantly increases from 17% to 30% after Context Distillation, implying that the capability to perform direct addition can transfer to other related applications.

## 4.2 Sleep-time compute

As shown in Chapter 3, test-time scaling is an effective way to boost LLM performance on challenging tasks by spending more time thinking on difficult problems [OpenAI, 2024, DeepSeek-AI, 2024, Snell et al., 2025, Brown et al., 2024, Wu et al., 2024]. However, as noted, improved performance from test-time compute comes at a significant increase in latency and cost, waiting potentially several minutes for answers and costing up to tens of dollars per query [OpenAI, 2025]. These drawbacks are in part due to the fact that the current approach to applying test-time compute assumes that problems are stateless, i.e. queries (user queries at test-time) and the contexts (background information) required for answering them are provided to the model together at “test-time.” In practice, this means that if multiple related queries require making similar inferences about the context at “test-time,” the model will have to recompute redundant computations each time, incurring additional latency and cost.

In reality, many LLM applications are *inherently stateful*, and work in conjunction with persisted, re-used context. A classic example is document question-answering, where documents contextualize responses to questions. Coding agents also operate on a large common repository and participate in multiple rounds of debugging support, while conversational assistants need to maintain the past dialogue. In all these applications, there is context (available documents, a codebase, or conversation history) that is already available before the next user input.

In these settings, we could in principle, make useful inferences about the current state (context) offline before, or even during the user’s next input (e.g. applying pre-computation to the current state). We refer to such a process, as sleep-time compute: where inference is done between interactions with the model while it would otherwise be idle in *sleep*-time. In practice, this is achieved by prompting the model to generate a new context consisting of

inferences about the existing context, which may be potentially useful for answering test-time queries. The re-represented context from sleep-time can then be provided in the prompt at test-time, enabling the model to respond to user queries matching the accuracy of standard test-time compute but with far lower latencies. For example, a coding assistant at sleep-time may identify architectural patterns, anticipate potential debugging strategies, or infer optimizations prior to the user input. Moreover, users might ask multiple queries about the same context. In these settings, any inferences made during sleep-time can be shared across queries, effectively amortizing the cost of Sleep-time compute and reducing the total average cost per query.

To evaluate sleep-time compute, we modify two mathematical reasoning datasets to introduce two datasets – Stateful GSM-Symbolic and Stateful AIME – by splitting the existing problems in these datasets into a context and a question. Using these datasets, we aim to empirically understand the benefits of sleep-time compute on standard test-time compute benchmarks. We show that:

- Sleep-time compute produces a pareto improvement in the test-time compute vs. accuracy curve, reducing the test-time compute needed to achieve the same accuracy by  $\sim 5\times$  on Stateful GSM-Symbolic and Stateful AIME.
- By scaling up sleep-time compute, we see further pareto improvements, shifting the accuracy up by 13% on Stateful GSM-Symbolic and 18% on Stateful AIME.
- By amortizing sleep-time compute across multiple queries for the same context, we can reduce the average cost per question by  $2.5\times$ .
- We conduct analysis to understand which queries benefit the most from sleep-time compute, finding that sleep-time compute is more effective in settings where the query is more easily predictable from the context.

Finally, we end with a case study of applying sleep-time compute to reduce test-time compute in a realistic agentic software engineering task.

## Method

In the standard paradigm of applying test-time compute, a user inputs a prompt  $p$  to the LLM and then the LLM applies test-time compute to help answer the user’s question. However, the  $p$  provided to the LLM can oftentimes be decomposed into a pre-existing context  $c$  (eg. a codebase) and a user query  $q$  (eg. a question about the codebase). When the LLM is not actively responding to the user, it typically still has access to the existing context  $c$ . During this time, the LLM is typically idling, missing the opportunity to reason about  $c$  offline: a process we term sleep-time compute.

**Test-time compute.** In the test-time compute setting, the user provides  $q$  along with some context  $c$  and the model outputs a reasoning trace followed by a final answer  $a$ . We denote this process, as:  $T_B(q, c) \rightarrow a$ , where  $T$  is the method for using test-time compute

with budget  $B$ , which could include techniques like extended chains of thought or best-of-N. In practice, the user may have multiple queries about the same context  $q_1, q_2 \dots q_N$ . In this setting, the model will carry out independent reasoning processes for each  $q_i$ , even if they are related to the same context  $c$ . Ideally, we would be able to reuse related inferences across each  $q_i$  to save compute. Moreover, in many cases,  $c$  is complex and may require carrying out significant processing/inferences in order to provide an answer to  $q$ . Since, the test-time compute paradigm of  $T(q, c) \rightarrow a$  assumes that  $c$  is only available at the same time as  $q$ , standard test-time compute carries out all of these inferences only after the user provides the query, causing the user to wait up to several minutes for a response. However, in practice we often have access to  $c$  before  $q$  and can carry out much of this processing ahead of time.

**Sleep-time compute.** During sleep-time we are given the context  $c$  but not the query  $q$ . Using just this context  $c$ , we can use the LLM to infer likely questions and reason about the context ultimately producing a new re-represented context  $c'$ . We denote this process as:  $S(c) \rightarrow c'$ , where  $S$  can be any standard test-time scaling technique applied towards pre-processing the context at sleep-time. In this work,  $S(c)$  is implemented by prompting the model to draw inferences and re-write  $c$  in a way that might be useful at test-time (see Appendix C.25 for more details). After pre-processing the context, we can provide the new context  $c'$  at test-time in place of  $c$  to produce a final answer to the user’s query:  $T_b(q, c') \rightarrow a$ . Since much of the reasoning about  $c$  has been done ahead of time in this case, we can use a much smaller test-time budget  $b \ll B$ . Moreover,  $c'$  can be shared across different queries  $q_i$  about the same context, effectively amortizing the compute required to arrive at  $c'$  across queries, providing a total cost saving.

## Experimental Setup

**Datasets** We select datasets which represent standard benchmarks for LLM reasoning and test-time scaling, and which demonstrate improvements from scaling test-time compute with state-of-the-art LLMs (either reasoning or non-reasoning).

**Stateful datasets.** We introduce two datasets to study applying sleep-time compute in stateful settings, Stateful GSM-Symbolic, and Stateful AIME, where each dataset is derived from splitting the existing datasets into a context and a question (see Figure 4.2 for an example). Stateful GSM-Symbolic is derived from the P1 and P2 splits of GSM-Symbolic [Mirzadeh et al., 2024], which add one and two clauses respectively to the original GSM8K dataset [Cobbe et al., 2021] to increase the difficulty. GSM-Symbolic P1 contains 5000 examples and P2 2500 examples. Stateful AIME contains 60 questions combined from AIME 2024 and 2025. In Appendix C.26 and C.27, we show the breakdown of our results across AIME 2024 and 2025.

**Amortization dataset.** To study the effect of related questions that share context, we introduce a new dataset Multi-Query GSM-Symbolic, where each context has multiple queries. To generate multiple queries for a given context, we take Stateful GSM-Symbolic and use o3-mini to generate additional question answer pairs. We synthetically generate

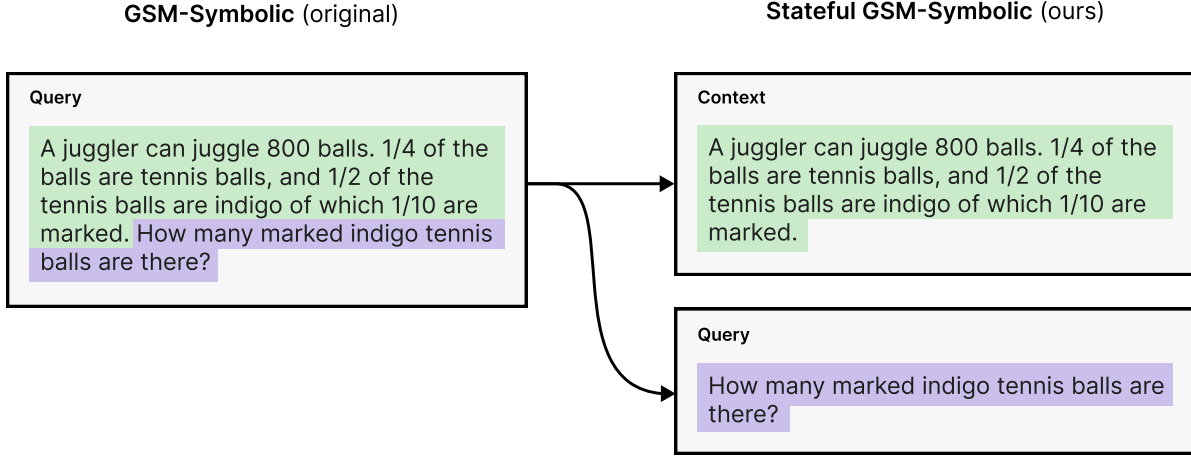


Figure 4.2: Example of separating an instance from GSM-Symbolic into context, and question, creating an instance in Stateful GSM-Symbolic.

additional questions from existing context question pairs in GSM-Symbolic. Appendix C.17 shows the prompt used to generate the additional questions. Figure C.13 shows examples contexts and set of questions from the Multi-Query GSM-Symbolic dataset and Table C.5 shows the overall dataset statistics.

**Models.** On each dataset, we evaluate models which have poor performance when using a small amount of test-time compute, but yield improvements from scaling up test-time compute. Therefore, on GSM-Symbolic, we conduct experiments using GPT-4o-mini and GPT-4o, and on AIME, we conduct experiments using OpenAI’s o1, o3-mini, Anthropic’s Claude Sonnet 3.7 Extended Thinking , and Deepseek-R1 [DeepSeek-AI, 2024].<sup>2 3</sup>

**Baselines.** The main baseline we consider is the standard test-time compute setting in which both  $c$  and  $q$  are presented to the model for the first time at test-time. Furthermore, to validate that  $q$  is not trivially predictable from  $c$  on our Stateful GSM-Symbolic and Stateful AIME datasets, we also compare to a context-only baseline in Appendix C.23, in which the model is only given  $c$  and is tasked with directly guessing an answer to the question it guesses is most likely to come next.

<sup>2</sup><https://openai.com/o1/>

<sup>3</sup><https://www.anthropic.com/claude/sonnet>

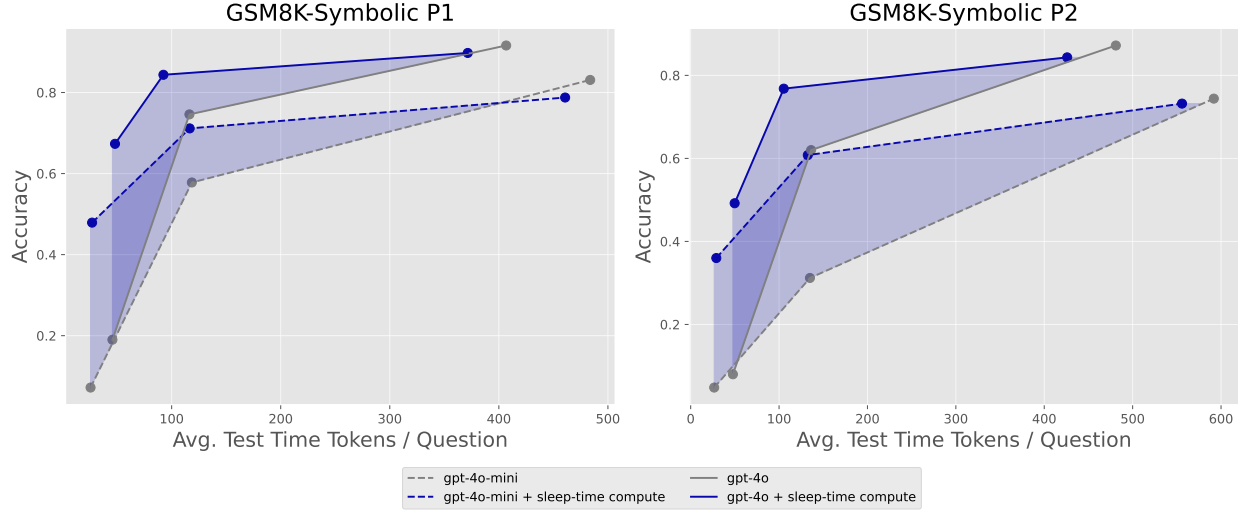


Figure 4.3: The test-time compute vs. accuracy tradeoff for on Stateful GSM-Symbolic. Shaded area indicates where sleep-time compute improves the pareto test-time accuracy trade-off.

## Experiments and Results

In this section, we carry out experiments to understand the benefits of sleep-time compute. Specifically, we would like to answer each of the following questions using the math reasoning benchmarks introduced above:

- 1) Can sleep-time compute shift the pareto frontier of test-time compute vs. accuracy?
- 2) Does scaling sleep-time compute in-turn improve the pareto further?
- 3) When there are multiple related questions for a single context, can amortizing test-time compute with sleep-time compute provide a total token efficiency benefit?
- 4) In what settings does sleep-time compute provide the most uplift?

**Improving Pareto Test-Time Trade-off with sleep-time compute** We first determine the test-time compute, accuracy pareto frontier by scaling standard test-time compute sequentially and in parallel. We then study how applying sleep-time compute affects the pareto trade-off.

**Scaling test-time-compute sequentially.** For non-reasoning models (GPT-4o and 4o-mini) on Stateful GSM-Symbolic, to vary the amount of test-time compute, we construct prompts that instruct the model to use different amounts of verbosity at test time, eg. “answer directly with a single sentence” vs. “double check your reasoning before outputting the final answer.” The full prompts are in Appendix C.15. We use temperature 0 for generation. We see in Figure 4.3 that there is a tradeoff between accuracy and the amount

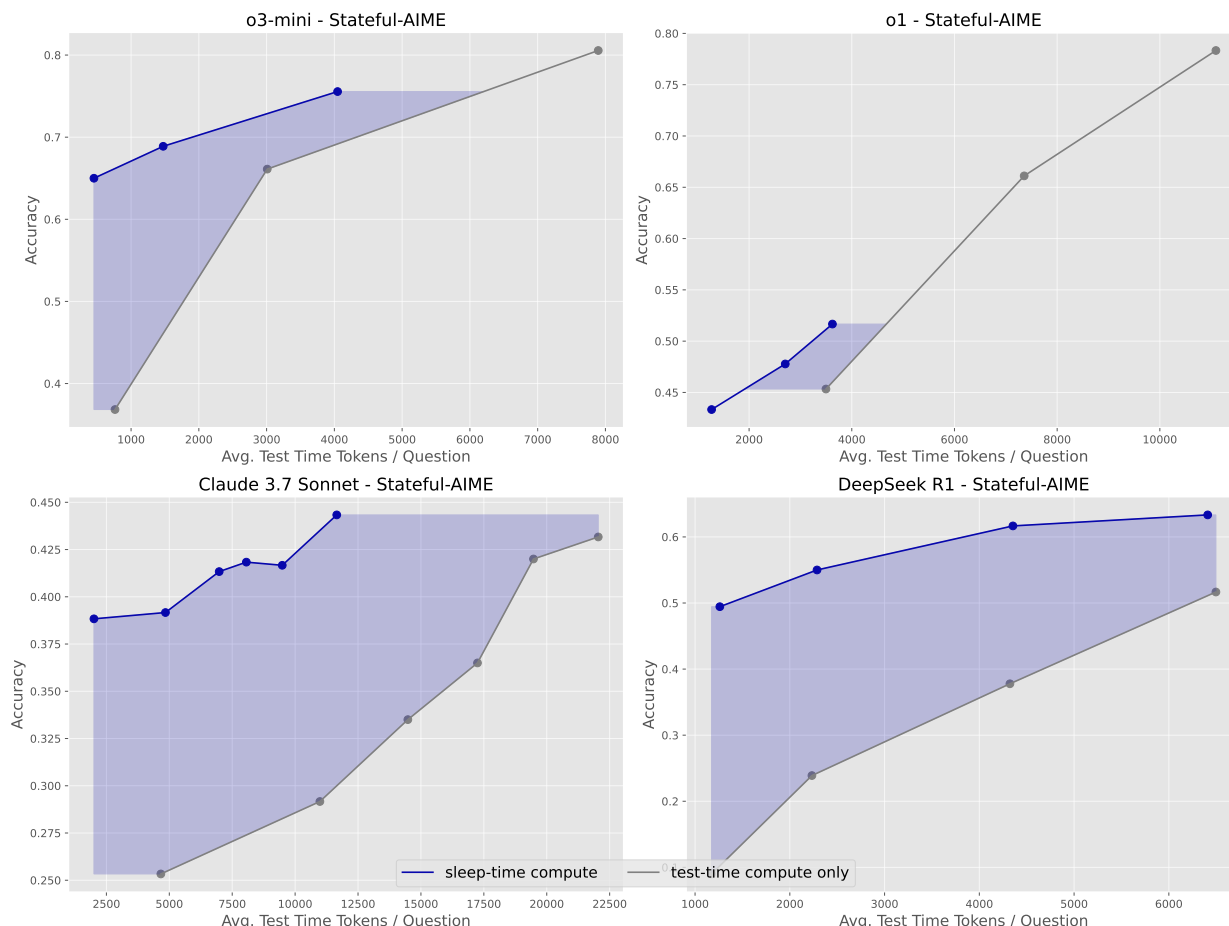


Figure 4.4: The test-time compute vs. accuracy tradeoff on Stateful AIME for various reasoning models. Applying sleep-time compute allows models to reach similar levels of performance with much less compute at test-time. The shaded area indicates the pareto improvement from sleep-time compute.

of test-time compute, and that adding sleep-time compute can move beyond the pareto compute-accuracy curve. In particular, at lower test-time budgets, the performance of sleep-time compute is significantly better than the baseline, achieving performance comparable to that of the baseline with  $5\times$  less test-time tokens. However, at higher test-time compute budgets, the test-time compute only baseline slightly outperforms sleep-time compute. We hypothesize that this may be because the standard test-time compute only has the content relevant to the specific question, so there is less distracting information in the prompt.

For reasoning models on Stateful AIME, we scale the amount of test-time compute based on what is available in the API in the case of o1, o3-mini and Claude Sonnet 3.7. Since the Deepseek-R1 API does not provide a way to control test-time compute, we apply the "budget

forcing" and extension prompt from Muennighoff et al. [2025]. Figure 4.4 shows the results for each model on Stateful AIME. We average results over 3 runs for o1, o3-mini and R1. For Claude 3.7 Sonnet, we average over 10 runs as we observed more noise in initial experiments. On all models, we see a significant test-time, accuracy pareto shift from applying sleep-time compute, with the exception of o1, which demonstrates limited gains.

**Scaling test-time compute in parallel.** An alternative approach to scaling test-time compute is via parallel sampling, which also has the benefit of maintaining low inference latency. The simplest approach to scaling parallel test-time compute is pass@k [Brown et al., 2024], which makes the unrealistic assumption of having oracle query access to a ground truth verifier at test-time, an assumption which we do not make with sleep-time compute. Therefore, outperforming the pass@k baseline would represent a meaningful improvement over parallel test-time scaling. We apply parallel scaling to the lowest sequential compute setting on each task, since scaling pass@k with higher sequential compute settings would quickly reach token budgets that exceed that of sleep-time compute in the maximum sequential setting. We see that across all tasks and models, sleep-time compute consistently outperforms pass@k parallel scaling at the same test-time token budget, demonstrating that sleep-time compute can be a more effective way to scale inference-time compute than standard parallel test-time scaling.

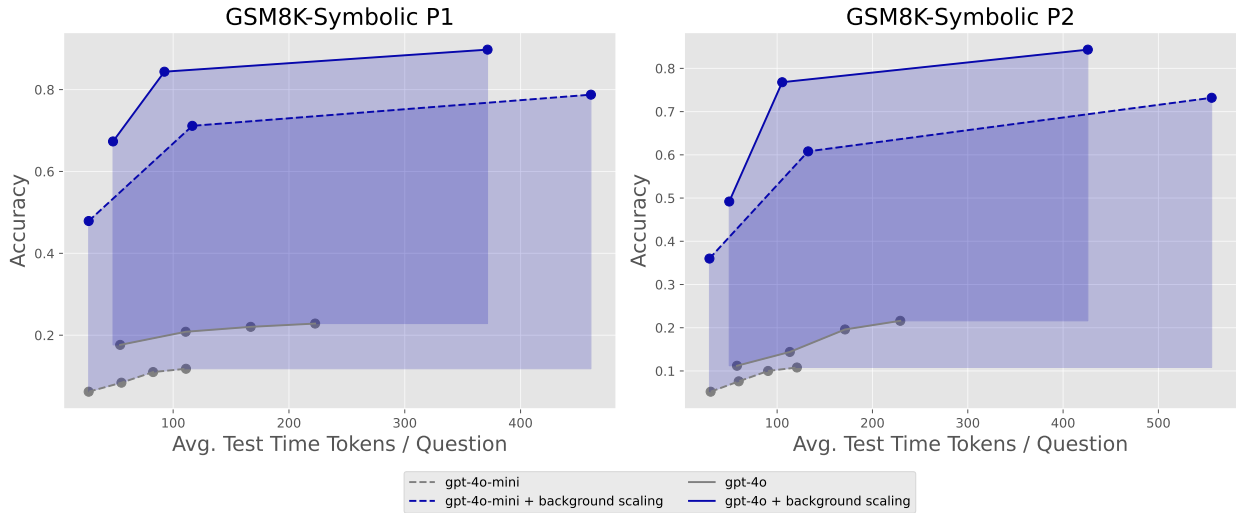


Figure 4.5: Comparing test-time scaling with sleep-time compute against parallel test-time scaling with pass@k on Stateful GSM-Symbolic. We see that sleep-time compute generally pareto dominates pass@k.

**Scaling up sleep-time compute** We would like to understand how scaling compute during sleep-time can further affect the pareto shift that we observed above. To scale up the

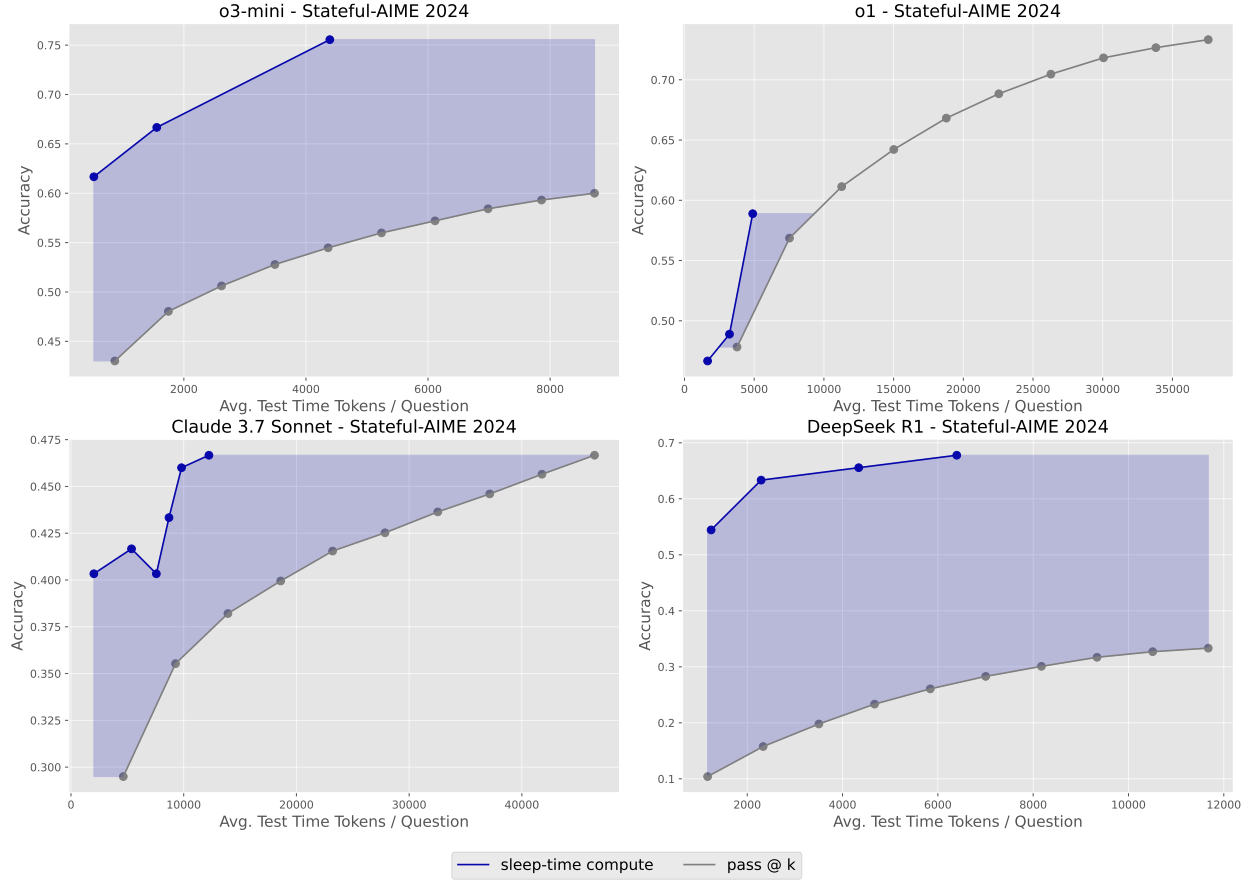


Figure 4.6: Comparing test-time scaling with sleep-time compute against parallel test-time scaling with pass@k on Stateful AIME. We see that sleep-time compute generally Pareto dominates pass@k.

amount of sleep-time compute, for non-reasoning models, we run  $k$  parallel generations, given input  $c$ , resulting in  $c_1, \dots, c_k$ . At test-time, the model then receives the inputs concatenated  $c_1, \dots, c_k$  to generate the final answer. On reasoning models, we scale up the amount of sleep-time compute by varying the reasoning effort for o1 and for o3-mini when applying the sleep-time compute prompt. At test-time, we vary the amount of compute in the same way as in the prior experiments.

In Figure 4.7, we see that further scaling sleep-time compute on Stateful GSM-Symbolic shifts the Pareto curve outwards, improving performance by up to 13% at a similar test-time budget. In particular, we see the largest gains on more difficult tasks with stronger models (eg. on P2 with ‘gpt-4o’), suggesting that on tasks with more complicated contexts additional sleep-time compute can be beneficial. However, in this setting, there seems to be a limit to the number of parallel agents that can improve performance, as we find that 5 parallel

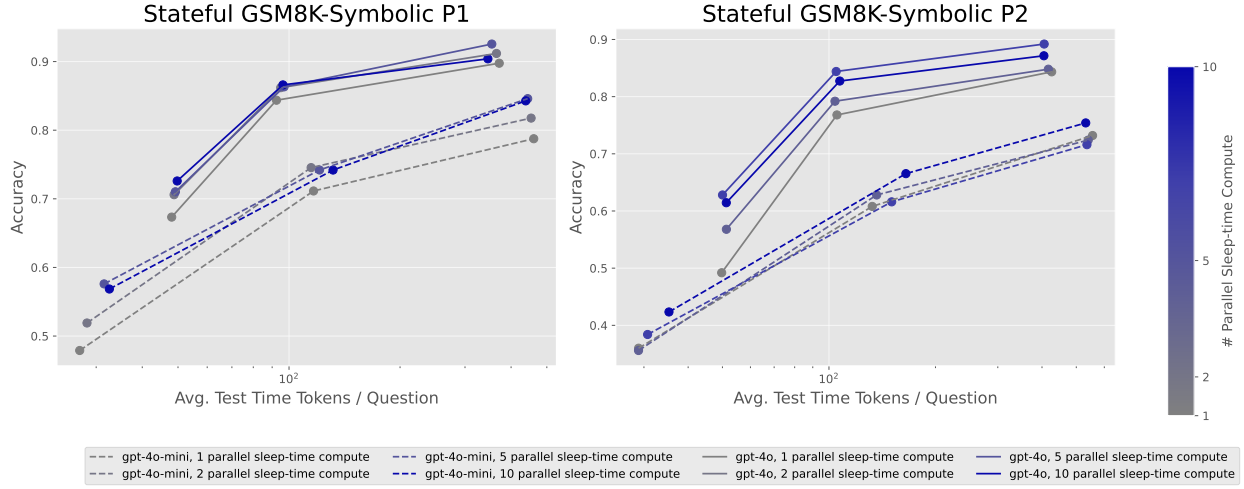


Figure 4.7: Scaling up sleep-time compute for different test-time compute budgets on Stateful GSM-Symbolic, by generating up multiple  $c'$  in parallel. Applying more sleep-time compute shifts the pareto beyond the standard test-time-compute vs. accuracy curve.

generations generally outperforms 10. In Figure C.18, we scale up sleep-time compute on Stateful AIME. Similarly, we also see that scaling compute at sleep-time generally shifts the pareto curve outward, improving performance by up to 18%.

**Amortizing sleep-time compute across queries with shared context** We want to understand how the total cost of inference can be improved by applying sleep-time compute in settings where each context has multiple queries. Since at test-time, there are strict latency constraints, and latency optimized inference can be roughly 10 $\times$  more expensive, we model the total cost of inference between both sleep-time and test-time, by up-weighting the cost of test-time tokens.<sup>4</sup> Specifically, we consider a simple linear model where tokens generated at test-time are a factor  $t$  the cost of the tokens at sleep-time. In our analysis, we set  $t = 10$ . Our analysis can be generalized to different cost functions that consider non-linear user-utility. Figure 4.9 shows the results for different number of questions per context. We see that we can decrease the average cost per query by up to 2.5 $\times$  when there are 10 queries per context, compared to the single-query baseline.

**Predictable queries benefit more from sleep-time compute** We would like to better understand for what contexts sleep-time compute is most useful. Since the utility of sleep-time compute relies on there being some shared information or structure between the context and the query, we hypothesize that sleep-time compute may be most effective in settings

<sup>4</sup><https://docs.databricks.com/aws/en/machine-learning/foundation-model-apis/prov-throughput-run-benchmark>

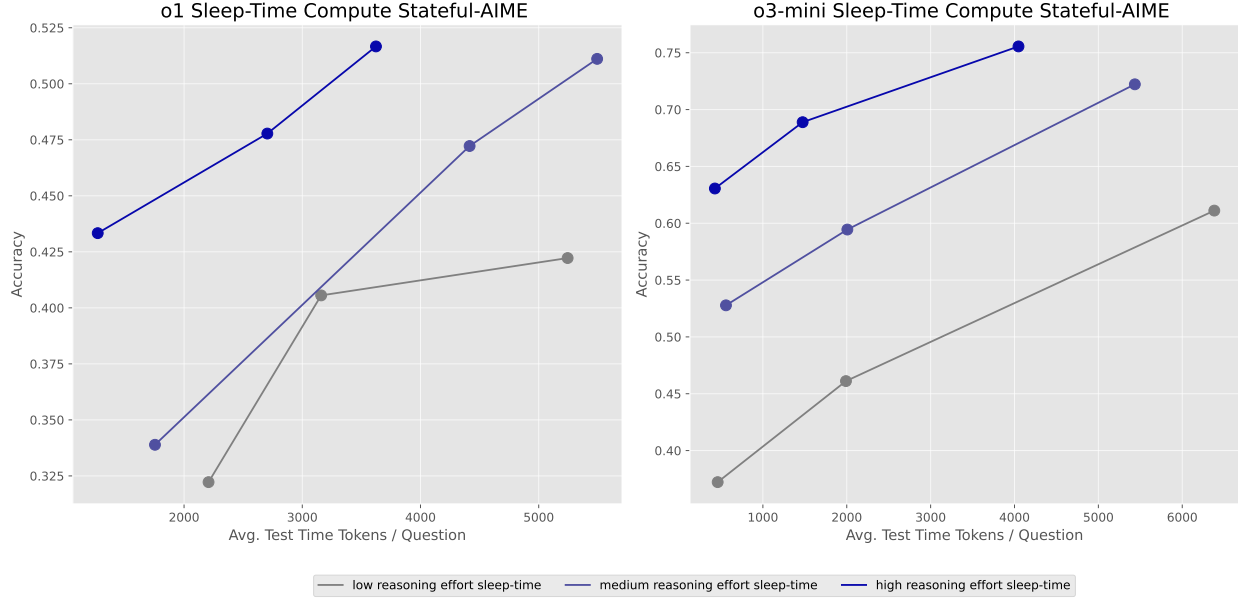


Figure 4.8: Increasing the amount of sleep-time compute for different test-time compute budgets on Stateful AIME by varying the reasoning effort when applying the sleep-time compute prompt. Applying more sleep-time compute further moves the test-time-compute vs. accuracy pareto curve.

where the query is more predictable from the context. To test this on Stateful GSM-Symbolic, we first quantify how predictable a given query is by measuring the log-probability of the question given the context under the Llama2-70B base model [Touvron et al., 2023a]. In Appendix C.19, we include examples of highly predictable and unpredictable questions under this notion of question predictability. We see from these examples, that our notion of question predictability generally aligns with the intuition that contexts where the query pattern is more predictable benefit most from sleep-time compute. The more predictable questions are far simpler and the less predictable ones are more complex.

Using our question predictability score, we then bin each example in Stateful GSM-Symbolic into five quantiles according to its predictability score and report the accuracy within each bin. For this experiment, we use the “Verbosity 0” prompt. In Figure 4.10, we see that on both GSM8K-Symbolic P1 and P2, the accuracy gap between sleep-time compute and standard test-time compute widens as the questions become more predictable from the context confirming our hypothesis that indeed sleep-time compute is most beneficial in settings where the question can be predicted from the context.

**A Case Study of Sleep-time compute for Agentic SWE** In this section, we evaluate sleep-time compute in a realistic multi-turn agentic setting. To this end, we introduce

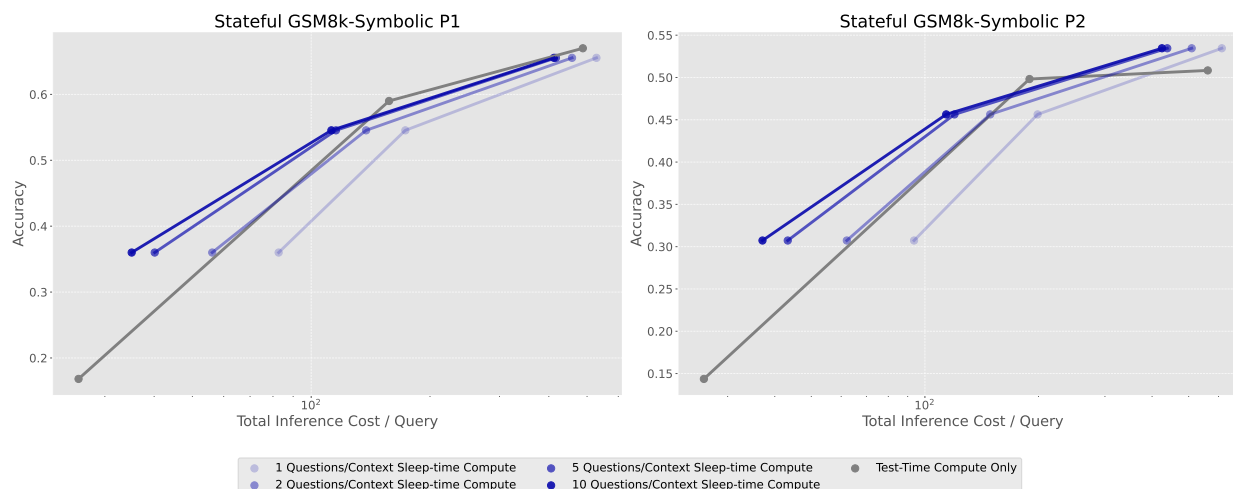


Figure 4.9: Amortizing sleep-time compute, using the Multi-Query GSM-Symbolic dataset. When there are fewer questions per context, we see that it is less favorable to use sleep-time compute, in terms of total cost. However, as the questions per context are increased, we see that applying sleep-time compute can improve the cost-accuracy pareto.

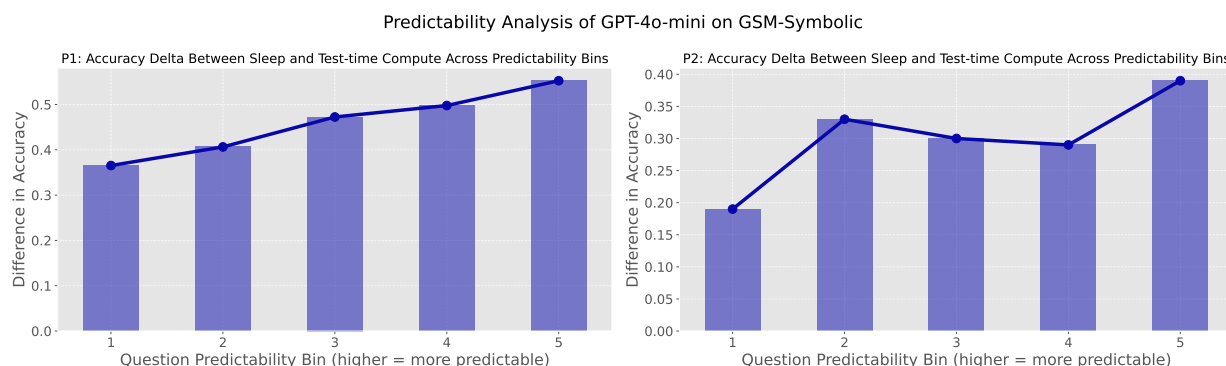


Figure 4.10: GSM-Symbolic questions binned by how predictable they are from the context. We compare the performance of sleep-time compute and standard test-time compute in the lowest test-time compute budget setting on both P1 and P2. The gap between sleep-time compute and standard test-time inference widens as the question becomes more predictable from the context.

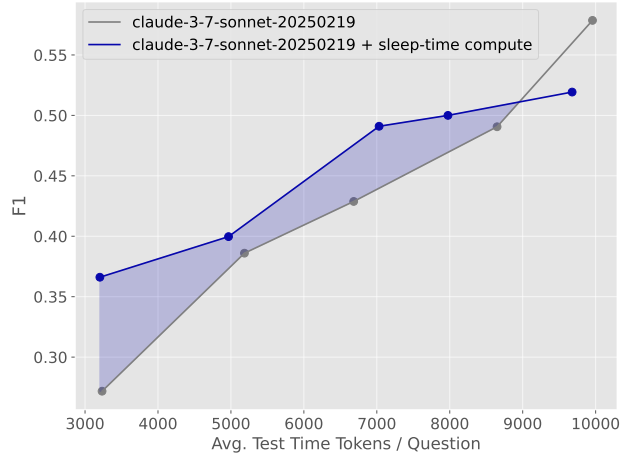


Figure 4.11: Applying sleep-time compute to SWE-Features. We see that at lower test-time budgets, sleep-time compute has higher F1 score than standard test-time scaling. However, at higher budgets, standard test-time scaling is better.

SWE-Features, a software engineering benchmark focused on tasks that require: (1) editing multiple files within a repository, and (2) implementing new features.

**SWE-Features.** In contrast to popular benchmarks like SWE-Bench [Jimenez et al., 2024], which involve modifying a small number of files, we propose a new dataset called SWE-Features, which collects PRs which modify at least three files (see Appendix C.18 for more details). In this setting, we use the PR that we want to solve as  $q$  and select several related PRs for  $c$ . At sleep-time the agent is allowed to explore the repository before producing  $c'$ .

**Evaluation.** Since the PRs are scraped from GitHub, there are not straightforward tests to use for evaluation. Instead, we compare the predicted set of modified files with the ground truth list of modified files, and report the F1 score between the set of modified files by our agent and the set of modified files in the ground-truth set (see Appendix C.18 for details).

**Results.** Figure 4.11 shows trends consistent with the above experiments for SWE-Features: at lower test-time compute budgets, leveraging sleep-time compute can improve performance, achieving up to roughly a  $1.5\times$  decrease in test-time tokens. However, when the test-time compute budget is high, using only test-time compute can perform better. Additionally, we observe that in the high test-time budget setting standard test-time compute has higher precision and comparable recall. We hypothesize that, using only test-time compute tends to begin editing files earlier and usually edits fewer files overall. In contrast, the agent with Sleep-time compute, having explored more files during the test-time phase, tends to edit more files, which may lead to slightly lower precision.

## 4.3 Conclusion

In this work we showed two different approaches to overcome the latency and redundant computation limitations of test-time compute. We introduced Context Distillation and Sleep-time compute as techniques which can be used to apply different mechanisms – in-weight learning and in-context learning respectively – to effectively precompute the work done at test-time, so that at test-time these tasks can be completed more cheaply and at lower latency without taking a meaningful performance hit. We can view both of these techniques as providing a bridge between the general scaling paradigms of learning (e.g. pretraining / in-context learning) and search (e.g. test-time scaling) with LLMs, by converting the stateless, instance-specific nature of test-time computation back into a fast-acting learned representation that can be reused across multiple queries.

## Chapter 5

# Conclusion and Future Work

In this thesis we demonstrated a new approach to scaling up computation with LLMs via test-time scaling that can in some settings make for an effective substitute for scaling up pretraining resources. In particular, test-time scaling is less directly hindered by the constraints on the amount of internet data available that plague pretraining. In particular, we first carried out work indicating that broadly improving LLM capabilities via pretraining generally requires substantially scaling up compute and data, and techniques like model imitation which attempt to cheaply improve capabilities, can only produce a superficial improvement without actually improving the model’s underlying factuality. This finding led us to turn towards test-time scaling as an alternative way to scale up computation. We found that, indeed in certain settings, test-time scaling can be preferable to scaling up pretraining resources. However, we also identified important limitations to scaling up test-time compute that make it not quite one-to-one comparable to scaling up pretraining compute. In particular, an additional latency cost from test-time scaling and the problem of redundant computation across related queries. To address these limitations, we proposed two different methods to overcome them: Context Distillation and Sleep-time compute, which use different learning mechanisms – in-weight learning and in-context learning respectively – to effectively precompute the work done at test-time for specific distributions of tasks, so that at test-time these tasks can be completed nearly instantaneously without sacrificing any performance. These techniques can be seen as effectively converting the slow, stateless, and instance-specific nature of test-time computation back into a compressed, generalizable, and faster-acting learned representation, effectively serving as a bridge between the general scaling paradigms of learning (e.g. pretraining / in-context learning) and search (e.g. test-time scaling) with LLMs.

Our work in this thesis lays the groundwork for a more general way of reasoning about and scaling up computation with LLMs, one which can elegantly alternate between learning and test-time reasoning to make the most efficient use of resources. However, in order to fully realize this vision, we believe there are a number of directions for future work to further develop and improve upon our findings; these future directions are:

**Improving the efficacy of test-time scaling.** In Chapter 3 we demonstrated two

high-level techniques which can be used to scale up computation at test-time. Subsequent work to ours [DeepSeek-AI, 2024, OpenAI, 2024] has demonstrated a highly effective approach to scaling test-time compute which utilizes reinforcement learning to teach the model to generate increasingly long chain-of-thought reasoning traces that in-turn help it better reason towards the correct answer. However, even these more modern approaches to test-time scaling struggle to produce results on some of the very hardest problems (e.g. open math problems), and the degree to which these methods can generalize beyond their training distribution remains unclear. Therefore, future work should aim to further improve upon these current techniques by, for example, developing new architectures and training objectives that can more natively enable models to do a more latent form of reasoning. Another promising direction for future work is to develop novel techniques for scaling up parallel test-time compute, which can help reduce the latency hit obtained from standard sequential test-time scaling [Pan et al., 2025, Ma et al., 2025].

**Towards a complete theory of computation with LLMs.** We can think of there as being two high-level ways of expending computation with LLMs: learning (e.g. pretraining, in-context learning) and search (e.g. test-time reasoning). As discussed, Context Distillation and Sleep-time compute can be viewed as techniques for bridging the gap between learning and test-time reasoning with LLMs. These methods expend additional computation ahead of time in order to enable faster and more efficient responses at test-time. As models become more and more capable of a wider range of tasks; the usefulness of one LLM over another becomes more a question of resources (e.g. FLOPs, data) and downstream constraints (e.g. latency, optimized for single domain versus fully general) rather than one of purely capabilities. Based on this, we believe that future work should aim to develop a “complexity theory” of LLMs, which aims to clearly describe how these different forms of computation can be traded off and exchanged and what the different tradeoffs are with respect to e.g. training FLOPs, test-time FLOPs, latency, downstream task distribution (e.g. narrow versus broad), and data efficiency. Most standard scaling law analyses aim to understand the scaling characteristics of either pretraining [Kaplan et al., 2020, Hoffmann et al., 2022] or test-time scaling [Snell et al., 2025, Muennighoff et al., 2025] and less so the complex interplay between the two [Jones, 2021, Villalobos and Atkinson, 2023]. A more comprehensive theory such as this would enable us to more effectively optimize our compute allocations, improving the latency and efficiency of the tasks that matter the most by strategically applying precomputation when appropriate, and allowing us to make better financial and policy decisions on the basis of these tradeoffs.

**Unifying learning and test-time scaling.** We proposed Context Distillation and Sleep-time compute as methods for applying compute ahead of time to improve the downstream efficiency of responses at test-time. A number of promising recent works have proposed more data efficient techniques for distillation via on-policy learning [Thinking Machines Lab, 2025, Agarwal et al., 2024b] and more recent works have applied these ideas to the Context Distillation setting [Ye et al., 2026, Sang et al., 2026, Zhao et al., 2026, Hübötter et al., 2026]. All of these techniques however, treat test-time and training-time as two completely distinct phases. Humans, for example, don’t have such a clear boundary between learning

and acting. Future work should aim to bridge this gap by developing a unified learning algorithm that can seamlessly learn to both internalize the outcomes of inferences made during test-time and improve the ability to apply computation at test-time with LLMs.

# Bibliography

Rishabh Agarwal, Avi Singh, Lei M. Zhang, Bernd Bohnet, Luis Rosias, Stephanie Chan, Biao Zhang, Ankesh Anand, Zaheer Abbas, Azade Nova, John D. Co-Reyes, Eric Chu, Feryal Behbahani, Aleksandra Faust, and Hugo Larochelle. Many-shot in-context learning, 2024a.

Rishabh Agarwal, Nino Vieillard, Yongchao Zhou, Piotr Stanczyk, Sabela Ramos, Matthieu Geist, and Olivier Bachem. On-policy distillation of language models: Learning from self-generated mistakes. In *ICLR*, 2024b.

Yuvanesh Anand, Zach Nussbaum, Brandon Duderstadt, Benjamin Schmidt, and Andriy Mulyar. GPT4All: Training an assistant-style chatbot with large scale data distillation from GPT-3.5-Turbo, 2023.

Christophe Andrieu, Nando De Freitas, Arnaud Doucet, and Michael I Jordan. An introduction to mcmc for machine learning. 2003.

Rohan Anil, Andrew M. Dai, Orhan Firat, Melvin Johnson, Dmitry Lepikhin, Alexandre Passos, Siamak Shakeri, Emanuel Taropa, Paige Bailey, Zhifeng Chen, Eric Chu, Jonathan H. Clark, Laurent El Shafey, Yanping Huang, Kathy Meier-Hellstern, Gaurav Mishra, Erica Moreira, Mark Omernick, Kevin Robinson, Sebastian Ruder, Yi Tay, Kefan Xiao, Yuanzhong Xu, Yujing Zhang, Gustavo Hernandez Abrego, Junwhan Ahn, Jacob Austin, Paul Barham, Jan Botha, James Bradbury, Siddhartha Brahma, Kevin Brooks, Michele Catasta, Yong Cheng, Colin Cherry, Christopher A. Choquette-Choo, Aakanksha Chowdhery, Clément Crepy, Shachi Dave, Mostafa Dehghani, Sunipa Dev, Jacob Devlin, Mark Díaz, Nan Du, Ethan Dyer, Vlad Feinberg, Fangxiaoyu Feng, Vlad Fienber, Markus Freitag, Xavier Garcia, Sebastian Gehrmann, Lucas Gonzalez, Guy Gur-Ari, Steven Hand, Hadi Hashemi, Le Hou, Joshua Howland, Andrea Hu, Jeffrey Hui, Jeremy Hurwitz, Michael Isard, Abe Ittycheriah, Matthew Jagielski, Wenhao Jia, Kathleen Kenealy, Maxim Krikun, Sneha Kudugunta, Chang Lan, Katherine Lee, Benjamin Lee, Eric Li, Music Li, Wei Li, YaGuang Li, Jian Li, Hyeontaek Lim, Hanzhao Lin, Zhongtao Liu, Frederick Liu, Marcello Maggioni, Aroma Mahendru, Joshua Maynez, Vedant Misra, Maysam Moussalem, Zachary Nado, John Nham, Eric Ni, Andrew Nystrom, Alicia Parrish, Marie Pellat, Martin Polacek, Alex Polozov, Reiner Pope, Siyuan Qiao, Emily Reif, Bryan Richter, Parker Riley, Alex Castro Ros, Aurko Roy, Brennan Saeta, Rajkumar

- Samuel, Renee Shelby, Ambrose Slone, Daniel Smilkov, David R. So, Daniel Sohn, Simon Tokumine, Dasha Valter, Vijay Vasudevan, Kiran Vodrahalli, Xuezhi Wang, Pidong Wang, Zirui Wang, Tao Wang, John Wieting, Yuhuai Wu, Kelvin Xu, Yunhan Xu, Linting Xue, Pengcheng Yin, Jiahui Yu, Qiao Zhang, Steven Zheng, Ce Zheng, Weikang Zhou, Denny Zhou, Slav Petrov, and Yonghui Wu. Palm 2 technical report, 2023.
- Amanda Askell, Yuntao Bai, Anna Chen, Dawn Drain, Deep Ganguli, Tom Henighan, Andy Jones, Nicholas Joseph, Ben Mann, Nova DasSarma, Nelson Elhage, Zac Hatfield-Dodds, Danny Hernandez, Jackson Kernion, Kamal Ndousse, Catherine Olsson, Dario Amodei, Tom Brown, Jack Clark, Sam McCandlish, Chris Olah, and Jared Kaplan. A general language assistant as a laboratory for alignment, 2021. URL <https://arxiv.org/abs/2112.00861>.
- Alan Baddeley. Working memory. *Science*, 255(5044):556–559, 1992.
- Yuntao Bai, Saurav Kadavath, Sandipan Kundu, Amanda Askell, Jackson Kernion, Andy Jones, Anna Chen, Anna Goldie, Azalia Mirhoseini, Cameron McKinnon, Carol Chen, Catherine Olsson, Christopher Olah, Danny Hernandez, Dawn Drain, Deep Ganguli, Dustin Li, Eli Tran-Johnson, Ethan Perez, Jamie Kerr, Jared Mueller, Jeffrey Ladish, Joshua Landau, Kamal Ndousse, Kamile Lukosuite, Liane Lovitt, Michael Sellitto, Nelson Elhage, Nicholas Schiefer, Noemi Mercado, Nova DasSarma, Robert Lasenby, Robin Larson, Sam Ringer, Scott Johnston, Shauna Kravec, Sheer El Showk, Stanislav Fort, Tamera Lanham, Timothy Telleen-Lawton, Tom Conerly, Tom Henighan, Tristan Hume, Samuel R. Bowman, Zac Hatfield-Dodds, Ben Mann, Dario Amodei, Nicholas Joseph, Sam McCandlish, Tom Brown, and Jared Kaplan. Constitutional ai: Harmlessness from ai feedback, 2022.
- Boaz Barak. Emergent abilities and grokking: Fundamental, mirage, or both?, 2023. URL <https://windowsontheory.org/2023/12/22/emergent-abilities-and-grokking-fundamental-mirage-or-both/>. Accessed: 08-07-2024.
- Cody Blakeney, Mansheej Paul, Brett W. Larsen, Sean Owen, and Jonathan Frankle. Does your data spark joy? performance gains from domain upsampling at the end of training, 2024. URL <https://arxiv.org/abs/2406.03476>.
- James Bradbury, Roy Frostig, Peter Hawkins, Matthew James Johnson, Chris Leary, Dougal Maclaurin, George Necula, Adam Paszke, Jake VanderPlas, Skye Wanderman-Milne, and Qiao Zhang. JAX: composable transformations of Python+NumPy programs, 2018. URL <http://github.com/google/jax>.
- Bradley Brown, Jordan Juravsky, Ryan Ehrlich, Ronald Clark, Quoc V Le, Christopher Ré, and Azalia Mirhoseini. Large language monkeys: Scaling inference compute with repeated sampling. *arXiv preprint arXiv:2407.21787*, 2024.

- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901, 2020.
- Tianle Cai, Yuhong Li, Zhengyang Geng, Hongwu Peng, Jason D. Lee, Deming Chen, and Tri Dao. Medusa: Simple llm inference acceleration framework with multiple decoding heads, 2024. URL <https://arxiv.org/abs/2401.10774>.
- Murray Campbell, A. Joseph Hoane Jr., and Feng-hsiung Hsu. Deep Blue. *Artificial Intelligence*, 134(1-2):57–83, 2002.
- Danqi Chen, Jason Bolton, and Christopher D. Manning. A thorough examination of the cnn/daily mail reading comprehension task. In *ACL*, 2016.
- Guoxin Chen, Minpeng Liao, Chengxi Li, and Kai Fan. Alphamath almost zero: process supervision without process, 2024.
- Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, Alex Ray, Raul Puri, Gretchen Krueger, Michael Petrov, Heidy Khlaaf, Girish Sastry, Pamela Mishkin, Brooke Chan, Scott Gray, Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz Kaiser, Mohammad Bavarian, Clemens Winter, Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebgren Guss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Josh Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. Evaluating large language models trained on code. 2021a.
- Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, et al. Evaluating large language models trained on code. *arXiv preprint arXiv:2107.03374*, 2021b.
- Yanda Chen, Ruiqi Zhong, Sheng Zha, George Karypis, and He He. Meta-learning via language model in-context tuning. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 719–730, Dublin, Ireland, May 2022. Association for Computational Linguistics. doi: 10.18653/v1/2022.acl-long.53. URL <https://aclanthology.org/2022.acl-long.53>.
- Wei-Lin Chiang, Zhuohan Li, Zi Lin, Ying Sheng, Zhanghao Wu, Hao Zhang, Lianmin Zheng, Siyuan Zhuang, Yonghao Zhuang, Joseph E. Gonzalez, Ion Stoica, and Eric P. Xing. Vicuna: An open-source chatbot impressing GPT-4 with 90%\* ChatGPT quality, 2023.

- Eunbi Choi, Yongrae Jo, Joel Jang, and Minjoon Seo. Prompt injection: Parameterization of fixed inputs, 2022. URL <https://arxiv.org/abs/2206.11349>.
- Aakanksha Chowdhery, Sharan Narang, Jacob Devlin, Maarten Bosma, Gaurav Mishra, Adam Roberts, Paul Barham, Hyung Won Chung, Charles Sutton, Sebastian Gehrmann, Parker Schuh, et al. PaLM: Scaling language modeling with pathways. *arXiv preprint arXiv:2204.02311*, 2022.
- Hyung Won Chung, Le Hou, Shayne Longpre, Barret Zoph, Yi Tay, William Fedus, Eric Li, Xuezhi Wang, Mostafa Dehghani, Siddhartha Brahma, et al. Scaling instruction-finetuned language models. *arXiv preprint arXiv:2210.11416*, 2022.
- Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. Training verifiers to solve math word problems, 2021.
- Nicola De Cao, Wilker Aziz, and Ivan Titov. Editing factual knowledge in language models, 2021. URL <https://arxiv.org/abs/2104.08164>.
- DeepSeek-AI. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. 2024.
- DeepSeek-AI, Aixin Liu, Bei Feng, Bing Xue, Bingxuan Wang, Bochao Wu, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Daya Guo, Dejian Yang, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Haowei Zhang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Li, Hui Qu, J. L. Cai, Jian Liang, Jianzhong Guo, Jiaqi Ni, Jiashi Li, Jiawei Wang, Jin Chen, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, Junxiao Song, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Lei Xu, Leyi Xia, Liang Zhao, Litong Wang, Liyue Zhang, Meng Li, Miaojun Wang, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Mingming Li, Ning Tian, Panpan Huang, Peiyi Wang, Peng Zhang, Qiancheng Wang, Qihao Zhu, Qinyu Chen, Qiushi Du, R. J. Chen, R. L. Jin, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, Runxin Xu, Ruoyu Zhang, Ruyi Chen, S. S. Li, Shanghao Lu, Shangyan Zhou, Shanhuang Chen, Shaoqing Wu, Shengfeng Ye, Shengfeng Ye, Shirong Ma, Shiyu Wang, Shuang Zhou, Shuiping Yu, Shunfeng Zhou, Shuting Pan, T. Wang, Tao Yun, Tian Pei, Tianyu Sun, W. L. Xiao, Wangding Zeng, Wanbiao Zhao, Wei An, Wen Liu, Wenfeng Liang, Wenjun Gao, Wenqin Yu, Wentao Zhang, X. Q. Li, Xiangyue Jin, Xianzu Wang, Xiao Bi, Xiaodong Liu, Xiaohan Wang, Xiaojin Shen, Xiaokang Chen, Xiaokang Zhang, Xiaosha Chen, Xiaotao Nie, Xiaowen Sun, Xiaoxiang Wang, Xin Cheng, Xin Liu, Xin Xie, Xingchao Liu, Xingkai Yu, Xinnan Song, Xinxia Shan, Xinyi Zhou, Xinyu Yang, Xinyuan Li, Xuecheng Su, Xuheng Lin, Y. K. Li, Y. Q. Wang, Y. X. Wei, Y. X. Zhu, Yang Zhang, Yanhong Xu, Yanhong Xu, Yanping Huang, Yao Li, Yao Zhao, Yaofeng Sun, Yaohui Li, Yaohui Wang, Yi Yu, Yi Zheng, Yichao Zhang, Yifan Shi, Yiliang Xiong, Ying

- He, Ying Tang, Yishi Piao, Yisong Wang, Yixuan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo, Yu Wu, Yuan Ou, Yuchen Zhu, Yuduan Wang, Yue Gong, Yuheng Zou, Yujia He, Yukun Zha, Yunfan Xiong, Yunxian Ma, Yuting Yan, Yuxiang Luo, Yuxiang You, Yuxuan Liu, Yuyang Zhou, Z. F. Wu, Z. Z. Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhean Xu, Zhen Huang, Zhen Zhang, Zhenda Xie, Zhengyan Zhang, Zhewen Hao, Zhibin Gou, Zhicheng Ma, Zhigang Yan, Zhihong Shao, Zhipeng Xu, Zhiyu Wu, Zhongyu Zhang, Zhuoshu Li, Zihui Gu, Zijia Zhu, Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song, Ziyi Gao, and Zizheng Pan. Deepseek-v3 technical report, 2025. URL <https://arxiv.org/abs/2412.19437>.
- Yilun Du, Shuang Li, Antonio Torralba, Joshua B. Tenenbaum, and Igor Mordatch. Improving factuality and reasoning in language models through multiagent debate, 2023.
- Zhengxiao Du, Aohan Zeng, Yuxiao Dong, and Jie Tang. Understanding emergent abilities of language models from the loss perspective, 2024. URL <https://arxiv.org/abs/2403.15796>.
- Adam Dziedziec, Nikita Dhawan, Muhammad Ahmad Kaleem, Jonas Guan, and Nicolas Papernot. On the difficulty of defending self-supervised learning against model extraction. In *ICLR*, 2022a.
- Adam Dziedziec, Muhammad Ahmad Kaleem, Yu Shen Lu, and Nicolas Papernot. Increasing the cost of model extraction with calibrated proof of work. In *ICLR*, 2022b.
- Jonathan St B. T. Evans. Heuristic and analytic processes in reasoning. *British Journal of Psychology*, 75(4):451–468, 1984.
- Xidong Feng, Ziyu Wan, Muning Wen, Stephen Marcus McAleer, Ying Wen, Weinan Zhang, and Jun Wang. Alphazero-like tree-search can guide large language model decoding and training, 2024.
- Daniel Fried, Armen Aghajanyan, Jessy Lin, Sida Wang, Eric Wallace, Freda Shi, Ruiqi Zhong, Wen-tau Yih, Luke Zettlemoyer, and Mike Lewis. Incoder: A generative model for code infilling and synthesis. *arXiv preprint arXiv:2204.05999*, 2022.
- Samir Yitzhak Gadre, Georgios Smyrnis, Vaishaal Shankar, Suchin Gururangan, Mitchell Wortsman, Rulin Shao, Jean Mercat, Alex Fang, Jeffrey Li, Sedrick Keh, Rui Xin, Marianna Nezhurina, Igor Vasiljevic, Jenia Jitsev, Alexandros G. Dimakis, Gabriel Ilharco, Shuran Song, Thomas Kollar, Yair Carmon, Achal Dave, Reinhard Heckel, Niklas Muennighoff, and Ludwig Schmidt. Language models scale reliably with over-training and on downstream tasks, 2024.
- Leo Gao, Jonathan Tow, Baber Abbasi, Stella Biderman, Sid Black, Anthony DiPofi, Charles Foster, Laurence Golding, Jeffrey Hsu, Alain Le Noac’h, Haonan Li, Kyle McDonell, Niklas Muennighoff, Chris Ociepa, Jason Phang, Laria Reynolds, Hailey Schoelkopf, Aviya Skowron, Lintang Sutawika, Eric Tang, Anish Thite, Ben Wang, Kevin Wang, and Andy

- Zou. A framework for few-shot language model evaluation, 12 2023. URL <https://zenodo.org/records/10256836>.
- Samuel Gehman, Suchin Gururangan, Maarten Sap, Yejin Choi, and Noah A Smith. Real-ToxicityPrompts: Evaluating neural toxic degeneration in language models. In *Findings of EMNLP*, 2020.
- Xinyang Geng and Hao Liu. Openllama: An open reproduction of llama, May 2023. URL [https://github.com/openlm-research/open\\_llama](https://github.com/openlm-research/open_llama).
- Xinyang Geng, Arnav Gudibande, Hao Liu, Eric Wallace, Pieter Abbeel, Sergey Levine, and Dawn Song. Koala: A dialogue model for academic research. *BAIR Blog*, 2023.
- Elliot Glazer, Ege Erdil, Tamay Besiroglu, et al. Frontiermath: A benchmark for evaluating advanced mathematical reasoning in ai, 2024.
- Jim Gray, Surajit Chaudhuri, Adam Bosworth, Andrew Layman, Don Reichart, Murali Venkatrao, Frank Pellow, and Hamid Pirahesh. Data cube: A relational aggregation operator generalizing group-by, cross-tab, and sub-totals. *Data mining and knowledge discovery*, 1:29–53, 1997.
- Arnav Gudibande, Eric Wallace, Charlie Victor Snell, Xinyang Geng, Hao Liu, Pieter Abbeel, Sergey Levine, and Dawn Song. The false promise of imitating proprietary language models. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=Kz3yckpCN5>.
- Biyang Guo, Xin Zhang, Ziyuan Wang, Minqi Jiang, Jinran Nie, Yuxuan Ding, Jianwei Yue, and Yupeng Wu. How close is ChatGPT to human experts? Comparison corpus, evaluation, and detection. *arXiv preprint arXiv:2301.07597*, 2023.
- Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Xiaodong Song, and Jacob Steinhardt. Measuring massive multitask language understanding. In *ICLR*, 2021a.
- Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn Song, and Jacob Steinhardt. Measuring mathematical problem solving with the math dataset, 2021b.
- Geoffrey Hinton, Oriol Vinyals, and Jeff Dean. Distilling the knowledge in a neural network. In *NIPS Deep Learning Workshop*, 2014.
- Geoffrey Hinton, Oriol Vinyals, and Jeff Dean. Distilling the knowledge in a neural network, 2015. URL <https://arxiv.org/abs/1503.02531>.

- Jordan Hoffmann, Sebastian Borgeaud, Arthur Mensch, Elena Buchatskaya, Trevor Cai, Eliza Rutherford, Diego de Las Casas, Lisa Anne Hendricks, Johannes Welbl, Aidan Clark, Tom Hennigan, Eric Noland, Katie Millican, George van den Driessche, Bogdan Damoc, Aurelia Guy, Simon Osindero, Karen Simonyan, Erich Elsen, Jack W. Rae, Oriol Vinyals, and Laurent Sifre. Training compute-optimal large language models, 2022.
- Or Honovich, Thomas Scialom, Omer Levy, and Timo Schick. Unnatural instructions: Tuning language models with (almost) no human labor. *arXiv preprint arXiv:2212.09689*, 2022.
- Cheng-Yu Hsieh, Chun-Liang Li, Chih-Kuan Yeh, Hootan Nakhost, Yasuhisa Fujii, Alexander Ratner, Ranjay Krishna, Chen-Yu Lee, and Tomas Pfister. Distilling step-by-step! Outperforming larger language models with less training data and smaller model sizes. *arXiv preprint arXiv:2305.02301*, 2023.
- Jie Huang, Xinyun Chen, Swaroop Mishra, Huaixiu Steven Zheng, Adams Wei Yu, Xinying Song, and Denny Zhou. Large language models cannot self-correct reasoning yet, 2023.
- Yuzhen Huang, Jinghan Zhang, Zifei Shan, and Junxian He. Compression represents intelligence linearly, 2024. URL <https://arxiv.org/abs/2404.09937>.
- Jonas Hübötter, Frederike Lübeck, Lejs Behric, Anton Baumann, Marco Bagatella, Daniel Marta, Ido Hakimi, Idan Shenfeld, Thomas Kleine Buening, Carlos Guestrin, and Andreas Krause. Reinforcement learning via self-distillation, 2026.
- Carlos E. Jimenez, John Yang, Alexander Wettig, Shunyu Yao, Kexin Pei, Ofir Press, and Karthik R. Narasimhan. Swe-bench: Can language models resolve real-world github issues? In *ICLR*. OpenReview.net, 2024.
- Andy L. Jones. Scaling scaling laws with board games, 2021. URL <https://arxiv.org/abs/2104.03113>.
- Mika Juuti, Sebastian Szyller, Samuel Marchal, and N Asokan. PRADA: protecting against DNN model stealing attacks. In *IEEE EuroS&P*, 2019.
- Daniel Kahneman. Maps of bounded rationality: Psychology for behavioral economics. *The American Economic Review*, 93(5):1449–1475, 2003.
- Daniel Kahneman. *Thinking, fast and slow*. Farrar, Straus and Giroux, New York, first paperback edition edition, 2013.
- Jared Kaplan, Sam McCandlish, Tom Henighan, Tom B Brown, Benjamin Chess, Rewon Child, Scott Gray, Alec Radford, Jeffrey Wu, and Dario Amodei. Scaling laws for neural language models. *arXiv preprint arXiv:2001.08361*, 2020.

- Kalpesh Krishna, Gaurav Singh Tomar, Ankur P Parikh, Nicolas Papernot, and Mohit Iyyer. Thieves on sesame street! Model extraction of BERT-based APIs. In *ICLR*, 2020.
- Aviral Kumar, Vincent Zhuang, Rishabh Agarwal, Yi Su, JD Co-Reyes, Avi Singh, Kate Baumli, Shariq Iqbal, Colton Bishop, Rebecca Roelofs, et al. Training language models to self-correct via reinforcement learning.
- Tom Kwiatkowski, Jennimaria Palomaki, Olivia Redfield, Michael Collins, Ankur Parikh, Chris Alberti, Danielle Epstein, Illia Polosukhin, Jacob Devlin, Kenton Lee, Kristina Toutanova, Llion Jones, Matthew Kelcey, Ming-Wei Chang, Andrew M. Dai, Jakob Uszkoreit, Quoc Le, and Slav Petrov. Natural questions: A benchmark for question answering research. *TACL*, 2019a.
- Tom Kwiatkowski, Jennimaria Palomaki, Olivia Redfield, Michael Collins, Ankur Parikh, Chris Alberti, Danielle Epstein, Illia Polosukhin, Jacob Devlin, Kenton Lee, et al. Natural questions: A benchmark for question answering research. *TACL*, 2019b.
- Andrew K Lampinen, Ishita Dasgupta, Stephanie CY Chan, Kory Matthewson, Michael Henry Tessler, Antonia Creswell, James L McClelland, Jane X Wang, and Felix Hill. Can language models learn from explanations in context? *arXiv preprint arXiv:2204.02329*, 2022.
- Yaniv Leviathan, Matan Kalman, and Yossi Matias. Fast inference from transformers via speculative decoding, 2023. URL <https://arxiv.org/abs/2211.17192>.
- Aitor Lewkowycz, Anders Andreassen, David Dohan, Ethan Dyer, Henryk Michalewski, Vinay Ramasesh, Ambrose Slone, Cem Anil, Imanol Schlag, Theo Gutman-Solo, Yuhuai Wu, Behnam Neyshabur, Guy Gur-Ari, and Vedant Misra. Solving quantitative reasoning problems with language models, 2022.
- Xiang Lisa Li and Percy Liang. Prefix-tuning: Optimizing continuous prompts for generation. *arXiv preprint arXiv:2101.00190*, 2021.
- Yifei Li, Zeqi Lin, Shizhuo Zhang, Qiang Fu, Bei Chen, Jian-Guang Lou, and Weizhu Chen. Making large language models better reasoners with step-aware verifier, 2023.
- Tom Lieberum, Matthew Rahtz, János Kramár, Neel Nanda, Geoffrey Irving, Rohin Shah, and Vladimir Mikulik. Does circuit analysis interpretability scale? evidence from multiple choice capabilities in chinchilla, 2023. URL <https://arxiv.org/abs/2307.09458>.
- Hunter Lightman, Vineet Kosaraju, Yura Burda, Harri Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s verify step by step, 2023.

- Kevin Lin, Charlie Snell, Yu Wang, Charles Packer, Sarah Wooders, Ion Stoica, and Joseph E. Gonzalez. Sleep-time compute: Beyond inference scaling at test-time, 2025. URL <https://arxiv.org/abs/2504.13171>.
- Xiaodong Liu, Pengcheng He, Weizhu Chen, and Jianfeng Gao. Improving multi-task deep neural networks via knowledge distillation for natural language understanding. *arXiv preprint arXiv:1904.09482*, 2019.
- Daniel Lowd and Christopher Meek. Adversarial learning. In *KDD*, 2005.
- Wenjie Ma, Jingxuan He, Charlie Snell, Tyler Griggs, Sewon Min, and Matei Zaharia. Reasoning models can be effective without thinking, 2025.
- Andrew L. Maas, Raymond E. Daly, Peter T. Pham, Dan Huang, Andrew Y. Ng, and Christopher Potts. Learning word vectors for sentiment analysis. In *Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies*, pages 142–150, Portland, Oregon, USA, June 2011. Association for Computational Linguistics. URL <https://aclanthology.org/P11-1015>.
- Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler Hallinan, Luyu Gao, Sarah Wiegreffe, Uri Alon, Nouha Dziri, Shrimai Prabhumoye, Yiming Yang, Shashank Gupta, Bodhisattwa Prasad Majumder, Katherine Hermann, Sean Welleck, Amir Yazdanbakhsh, and Peter Clark. Self-refine: Iterative refinement with self-feedback, 2023.
- Pratyush Maini, Mohammad Yaghini, and Nicolas Papernot. Dataset inference: Ownership resolution in machine learning. In *ICLR*, 2021.
- Nat McAleese, Rai Pokorny, Juan Felipe Cerón Uribe, Evgenia Nitishinskaya, Maja Trębacz, and Jan Leike. Llm critics help catch llm bugs. *OpenAI*, 2024.
- Kevin Meng, David Bau, Alex Andonian, and Yonatan Belinkov. Locating and editing factual associations in gpt, 2022. URL <https://arxiv.org/abs/2202.05262>.
- Iman Mirzadeh, Keivan Alizadeh, Hooman Shahrokhi, Oncel Tuzel, Samy Bengio, and Mehrdad Farajtabar. Gsm-symbolic: Understanding the limitations of mathematical reasoning in large language models. *arXiv preprint arXiv:2410.05229*, 2024.
- Eric Mitchell, Charles Lin, Antoine Bosselut, Chelsea Finn, and Christopher D. Manning. Fast model editing at scale, 2021. URL <https://arxiv.org/abs/2110.11309>.
- Niklas Muennighoff, Zitong Yang, Weijia Shi, Xiang Lisa Li, Li Fei-Fei, Hannaneh Hajishirzi, Luke Zettlemoyer, Percy Liang, Emmanuel Candès, and Tatsunori Hashimoto. s1: Simple test-time scaling, 2025.

- Maxwell Nye, Anders Johan Andreassen, Guy Gur-Ari, Henryk Michalewski, Jacob Austin, David Bieber, David Dohan, Aitor Lewkowycz, Maarten Bosma, David Luan, Charles Sutton, and Augustus Odena. Show your work: Scratchpads for intermediate computation with language models, 2021a. URL <https://arxiv.org/abs/2112.00114>.
- Maxwell Nye, Anders Johan Andreassen, Guy Gur-Ari, Henryk Michalewski, Jacob Austin, David Bieber, David Dohan, Aitor Lewkowycz, Maarten Bosma, David Luan, et al. Show your work: Scratchpads for intermediate computation with language models. *arXiv preprint arXiv:2112.00114*, 2021b.
- OpenAI. Openai o1 system card, 2024. URL <https://arxiv.org/abs/2412.16720>.
- OpenAI. o1-pro model listing, 2025. URL <https://platform.openai.com/docs/models/o1-pro>.
- Tribhuvanesh Orekondy, Bernt Schiele, and Mario Fritz. Knockoff nets: Stealing functionality of black-box models. In *CVPR*, 2019.
- Tribhuvanesh Orekondy, Bernt Schiele, and Mario Fritz. Prediction poisoning: Towards defenses against DNN model stealing attacks. In *ICLR*, 2020.
- Charles Packer, Sarah Wooders, Kevin Lin, Vivian Fang, Shishir G Patil, Ion Stoica, and Joseph E Gonzalez. MemGPT: Towards llms as operating systems. *arXiv preprint arXiv:2310.08560*, 2023.
- Soham Pal, Yash Gupta, Aditya Shukla, Aditya Kanade, Shirish Shevade, and Vinod Ganapathy. A framework for the extraction of deep neural networks by leveraging public data. *arXiv preprint arXiv:1905.09165*, 2019.
- Jiayi Pan, Xiuyu Li, Long Lian, Charlie Snell, Yifei Zhou, Adam Yala, Trevor Darrell, Kurt Keutzer, and Alane Suhr. Learning adaptive parallel reasoning with language models, 2025.
- Dylan Patel and Afzal Ahmad. Google “We have no moat, and neither does OpenAI”, 2023.
- Baolin Peng, Chunyuan Li, Pengcheng He, Michel Galley, and Jianfeng Gao. Instruction tuning with GPT-4. *arXiv preprint arXiv:2304.03277*, 2023.
- Long Phan, Alice Gatti, Ziwen Han, Nathaniel Li, et al. Humanity’s last exam, 2025.
- Yuxiao Qu, Tianjun Zhang, Naman Garg, and Aviral Kumar. Recursive introspection: Teaching foundation models how to self-improve. 2024.
- Alec Radford, Jeff Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. Language models are unsupervised multitask learners. In *OpenAI Technical Report*, 2019.

- Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J. Liu. Exploring the limits of transfer learning with a unified text-to-text transformer, 2019. URL <https://arxiv.org/abs/1910.10683>.
- Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, Peter J Liu, et al. Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of Machine Learning Research*, 21(140):1–67, 2020.
- Nitarshan Rajkumar, Raymond Li, and Dzmitry Bahdanau. Evaluating the text-to-sql capabilities of large language models. *arXiv preprint arXiv:2204.00498*, 2022.
- Vinay Venkatesh Ramasesh, Aitor Lewkowycz, and Ethan Dyer. Effect of scale on catastrophic forgetting in neural networks. In *International Conference on Learning Representations*, 2022. URL [https://openreview.net/forum?id=GhVS8\\_yPeEa](https://openreview.net/forum?id=GhVS8_yPeEa).
- Adriana Romero, Nicolas Ballas, Samira Ebrahimi Kahou, Antoine Chassang, Carlo Gatta, and Yoshua Bengio. Fitnets: Hints for thin deep nets. *Proc. ICLR*, 2015.
- Amit Sabne. Xla : Compiling machine learning for peak performance, 2020.
- Hejian Sang, Yuanda Xu, Zhengze Zhou, Ran He, Zhipeng Wang, and Jiachen Sun. CRISP: Compressed reasoning via iterative self-policy distillation, 2026.
- Victor Sanh, Albert Webson, Colin Raffel, Stephen Bach, Lintang Sutawika, Zaid Alyafeai, Antoine Chaffin, Arnaud Stiegler, Arun Raja, Manan Dey, M Saiful Bari, Canwen Xu, Urmish Thakker, Shanya Sharma Sharma, Eliza Szczechla, Taewoon Kim, Gunjan Chhablani, Nihal Nayak, Debajyoti Datta, Jonathan Chang, Mike Tian-Jian Jiang, Han Wang, Matteo Manica, Sheng Shen, Zheng Xin Yong, Harshit Pandey, Rachel Bawden, Thomas Wang, Trishala Neeraj, Jos Rozen, Abheesht Sharma, Andrea Santilli, Thibault Fevry, Jason Alan Fries, Ryan Teehan, Teven Le Scao, Stella Biderman, Leo Gao, Thomas Wolf, and Alexander M Rush. Multitask prompted training enables zero-shot task generalization. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=9Vrb9D0WI4>.
- Nikhil Sardana and Jonathan Frankle. Beyond chinchilla-optimal: Accounting for inference in language model scaling laws, 2023.
- William Saunders, Catherine Yeh, Jeff Wu, Steven Bills, Long Ouyang, Jonathan Ward, and Jan Leike. Self-critiquing models for assisting human evaluators, 2022.
- Rylan Schaeffer, Brando Miranda, and Sanmi Koyejo. Are emergent abilities of large language models a mirage?, 2023.

- Jérémy Scheurer, Jon Ander Campos, Jun Shern Chan, Angelica Chen, Kyunghyun Cho, and Ethan Perez. Learning from natural language feedback. In *ACL Workshop on Learning with Natural Language Supervision*, 2022.
- Amrith Setlur, Saurabh Garg, Xinyang Geng, Naman Garg, Virginia Smith, and Aviral Kumar. Rl on incorrect synthetic data scales the efficiency of llm math reasoning by eight-fold. *arXiv preprint arXiv:2406.14532*, 2024.
- Archit Sharma, Sedrick Keh, Eric Mitchell, Chelsea Finn, Kushal Arora, and Thomas Kollar. A critical evaluation of ai feedback for aligning large language models, 2024. URL <https://arxiv.org/abs/2402.12366>.
- Noah Shinn, Federico Cassano, Edward Berman, Ashwin Gopinath, Karthik Narasimhan, and Shunyu Yao. Reflexion: Language agents with verbal reinforcement learning, 2023.
- David Silver, Aja Huang, Chris J Maddison, Arthur Guez, Laurent Sifre, George Van Den Driessche, Julian Schrittwieser, Ioannis Antonoglou, Veda Panneershelvam, Marc Lanctot, et al. Mastering the game of Go with deep neural networks and tree search. *Nature*, 529(7587):484, 2016.
- Avi Singh, John D. Co-Reyes, Rishabh Agarwal, Ankesh Anand, Piyush Patil, Xavier Garcia, Peter J. Liu, James Harrison, Jaehoon Lee, Kelvin Xu, Aaron Parisi, Abhishek Kumar, Alex Alemi, Alex Rizkowsky, Azade Nova, Ben Adlam, Bernd Bohnet, Gamaleldin Elsayed, Hanie Sedghi, Igor Mordatch, Isabelle Simpson, Izzeddin Gur, Jasper Snoek, Jeffrey Pennington, Jiri Hron, Kathleen Kenealy, Kevin Swersky, Kshiteej Mahajan, Laura Culp, Lechao Xiao, Maxwell L. Bileschi, Noah Constant, Roman Novak, Rosanne Liu, Tris Warkentin, Yundi Qian, Yamini Bansal, Ethan Dyer, Behnam Neyshabur, Jascha Sohl-Dickstein, and Noah Fiedel. Beyond human data: Scaling self-training for problem-solving with language models, 2024.
- Alan Jay Smith. Cache memories. *ACM Computing Surveys (CSUR)*, 14(3):473–530, 1982.
- Charlie Snell, Dan Klein, and Ruiqi Zhong. Learning by distilling context, 2022. URL <https://arxiv.org/abs/2209.15189>.
- Charlie Victor Snell, Eric Wallace, Dan Klein, and Sergey Levine. Predicting emergent capabilities by finetuning. In *First Conference on Language Modeling*, 2024. URL <https://openreview.net/forum?id=vL8BIGuFTF>.
- Charlie Victor Snell, Jaehoon Lee, Kelvin Xu, and Aviral Kumar. Scaling LLM test-time compute optimally can be more effective than scaling parameters for reasoning. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=4FWAwZtd2n>.
- Kaya Stechly, Matthew Marquez, and Subbarao Kambhampati. Gpt-4 doesn’t know it’s wrong: An analysis of iterative prompting for reasoning problems, 2023.

- Mitchell Stern, Noam Shazeer, and Jakob Uszkoreit. Blockwise parallel decoding for deep autoregressive models, 2018. URL <https://arxiv.org/abs/1811.03115>.
- Nisan Stiennon, Long Ouyang, Jeff Wu, Daniel M. Ziegler, Ryan Lowe, Chelsea Voss, Alec Radford, Dario Amodei, and Paul Christiano. Learning to summarize from human feedback, 2022.
- Weiwei Sun, Lingyong Yan, Xinyu Ma, Pengjie Ren, Dawei Yin, and Zhaochun Ren. Is ChatGPT good at search? Investigating large language models as re-ranking agent. *arXiv preprint arXiv:2304.09542*, 2023.
- Richard S Sutton and Andrew G Barto. *Reinforcement learning: An introduction*. Second edition, 2018.
- Sebastian Szyller, Buse Gul Atli, Samuel Marchal, and N Asokan. DAWN: Dynamic adversarial watermarking of neural networks. In *ACM Multimedia*, 2019.
- Alon Talmor, Jonathan Herzig, Nicholas Lourie, and Jonathan Berant. Commonsenseqa: A question answering challenge targeting commonsense knowledge, 2019.
- Rohan Taori, Ishaan Gulrajani, Tianyi Zhang, Yann Dubois, Xuechen Li, Carlos Guestrin, Percy Liang, and Tatsunori B. Hashimoto. Stanford Alpaca: An instruction-following LLaMA model, 2023.
- Thinking Machines Lab. On-policy distillation, 2025. URL <https://thinkingmachines.ai/blog/on-policy-distillation>.
- Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, et al. LLaMa: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*, 2023a.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, Dan Bikel, Lukas Blecher, Cristian Canton Ferrer, Moya Chen, Guillem Cucurull, David Esiobu, Jude Fernandes, Jeremy Fu, Wenyin Fu, Brian Fuller, Cynthia Gao, Vedanuj Goswami, Naman Goyal, Anthony Hartshorn, Saghar Hosseini, Rui Hou, Hakan Inan, Marcin Kardas, Viktor Kerkez, Madian Khabsa, Isabel Kloumann, Artem Korenev, Punit Singh Koura, Marie-Anne Lachaux, Thibaut Lavril, Jenya Lee, Diana Liskovich, Yinghai Lu, Yuning Mao, Xavier Martinet, Todor Mihaylov, Pushkar Mishra, Igor Molybog, Yixin Nie, Andrew Poulton, Jeremy Reizenstein, Rashi Rungta, Kalyan Saladi, Alan Schelten, Ruan Silva, Eric Michael Smith, Ranjan Subramanian, Xiaoqing Ellen Tan, Binh Tang, Ross Taylor, Adina Williams, Jian Xiang Kuan, Puxin Xu, Zheng Yan, Iliyan Zarov, Yuchen Zhang, Angela Fan, Melanie Kambadur, Sharan Narang, Aurelien Rodriguez, Robert Stojnic, Sergey Edunov, and Thomas Scialom. Llama 2: Open foundation and fine-tuned chat models, 2023b. URL <https://arxiv.org/abs/2307.09288>.

- Florian Tramèr, Fan Zhang, Ari Juels, Michael K Reiter, and Thomas Ristenpart. Stealing machine learning models via prediction APIs. In *USENIX Security Symposium*, 2016.
- Jonathan Uesato, Nate Kushman, Ramana Kumar, Francis Song, Noah Siegel, Lisa Wang, Antonia Creswell, Geoffrey Irving, and Irina Higgins. Solving math word problems with process- and outcome-based feedback, 2022.
- Karthik Valmeekam, Matthew Marquez, and Subbarao Kambhampati. Can large language models really improve by self-critiquing their own plans?, 2023.
- Pablo Villalobos and David Atkinson. Trading off compute in training and inference, 2023. URL <https://epochai.org/blog/trading-off-compute-in-training-and-inference>. Accessed: 2024-07-03.
- Pablo Villalobos, Jaime Sevilla, Lennart Heim, Tamay Besiroglu, Marius Hobbhahn, and Anson Ho. Will we run out of data? An analysis of the limits of scaling datasets in machine learning. *arXiv preprint arXiv:2211.04325*, 2022.
- Michael Völske, Martin Potthast, Shahbaz Syed, and Benno Stein. TL;DR: Mining Reddit to learn automatic summarization. In *Proceedings of the Workshop on New Frontiers in Summarization*, pages 59–63, Copenhagen, Denmark, September 2017. Association for Computational Linguistics. doi: 10.18653/v1/W17-4508. URL <https://aclanthology.org/W17-4508>.
- Eric Wallace, Mitchell Stern, and Dawn Song. Imitation attacks and defenses for black-box machine translation systems. In *EMNLP*, 2020.
- Alex Wang, Amanpreet Singh, Julian Michael, Felix Hill, Omer Levy, and Samuel Bowman. GLUE: A multi-task benchmark and analysis platform for natural language understanding. In *Proceedings of the 2018 EMNLP Workshop BlackboxNLP: Analyzing and Interpreting Neural Networks for NLP*. Association for Computational Linguistics, 2018.
- Peiyi Wang, Lei Li, Zhihong Shao, R. X. Xu, Damai Dai, Yifei Li, Deli Chen, Y. Wu, and Zhifang Sui. Math-shepherd: Verify and reinforce llms step-by-step without human annotations, 2023a.
- Yizhong Wang, Swaroop Mishra, Pegah Alipoormolabashi, Yeganeh Kordi, Amirreza Mirzaei, Anjana Arunkumar, Arjun Ashok, Arut Selvan Dhanasekaran, Atharva Naik, David Stap, Eshaan Pathak, Giannis Karamanolakis, Haizhi Gary Lai, Ishan Purohit, Ishani Mondal, Jacob Anderson, Kirby Kuznia, Krma Doshi, Maitreya Patel, Kuntal Kumar Pal, Mehrad Moradshahi, Mihir Parmar, Mirali Purohit, Neeraj Varshney, Phani Rohitha Kaza, Pulkrit Verma, Ravsehaj Singh Puri, Rushang Karia, Shailaja Keyur Sapat, Savan Doshi, Siddhartha Mishra, Sujana Reddy, Sumanta Patro, Tanay Dixit, Xudong Shen, Chitta Baral, Yejin Choi, Noah A. Smith, Hannaneh Hajishirzi, and Daniel Khashabi. Benchmarking generalization via in-context instructions on 1,600+ language tasks, 2022a. URL <https://arxiv.org/abs/2204.07705>.

- Yizhong Wang, Swaroop Mishra, Pegah Alipoormolabashi, Yeganeh Kordi, Amirreza Mirzaei, Anjana Arunkumar, Arjun Ashok, Arut Selvan Dhanasekaran, Atharva Naik, David Stap, et al. Benchmarking generalization via in-context instructions on 1,600+ language tasks. In *EMNLP*, 2022b.
- Yizhong Wang, Yeganeh Kordi, Swaroop Mishra, Alisa Liu, Noah A. Smith, Daniel Khashabi, and Hannaneh Hajishirzi. Self-Instruct: Aligning language model with self generated instructions. In *ACL*, 2023b.
- Zirui Wang, Adams Wei Yu, Orhan Firat, and Yuan Cao. Towards zero-label language learning. *arXiv preprint arXiv:2109.09193*, 2021.
- Peter C Wason and J St BT Evans. Dual processes in reasoning? *Cognition*, 3(2):141–154, 1974.
- Jason Wei, Maarten Bosma, Vincent Zhao, Kelvin Guu, Adams Wei Yu, Brian Lester, Nan Du, Andrew M. Dai, and Quoc V Le. Finetuned language models are zero-shot learners. In *International Conference on Learning Representations*, 2022a. URL <https://openreview.net/forum?id=gEZrGCozdqR>.
- Jason Wei, Yi Tay, Rishi Bommasani, Colin Raffel, Barret Zoph, Sebastian Borgeaud, Dani Yogatama, Maarten Bosma, Denny Zhou, Donald Metzler, Ed H. Chi, Tatsunori Hashimoto, Oriol Vinyals, Percy Liang, Jeff Dean, and William Fedus. Emergent abilities of large language models, 2022b.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Ed Chi, Quoc Le, and Denny Zhou. Chain of thought prompting elicits reasoning in large language models. *arXiv preprint arXiv:2201.11903*, 2022c.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed Chi, Quoc Le, and Denny Zhou. Chain-of-thought prompting elicits reasoning in large language models, 2023.
- Yangzhen Wu, Zhiqing Sun, Shanda Li, Sean Welleck, and Yiming Yang. Inference scaling laws: An empirical analysis of compute-optimal inference for problem-solving with language models. *arXiv preprint arXiv:2408.00724*, 2024.
- Mengzhou Xia, Mikel Artetxe, Chunting Zhou, Xi Victoria Lin, Ramakanth Pasunuru, Danqi Chen, Luke Zettlemoyer, and Ves Stoyanov. Training trajectories of language models across scales, 2023. URL <https://arxiv.org/abs/2212.09803>.
- Qizhe Xie, Minh-Thang Luong, Eduard Hovy, and Quoc V. Le. Self-training with noisy student improves imagenet classification. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 10684–10695, 2020. doi: 10.1109/CVPR42600.2020.01070.

- Ze Yang, Linjun Shou, Ming Gong, Wutao Lin, and Daxin Jiang. Model compression with two-stage multi-teacher knowledge distillation for web question answering system. In *Proceedings of the 13th International Conference on Web Search and Data Mining*, pages 690–698, 2020.
- Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran, Thomas L. Griffiths, Yuan Cao, and Karthik Narasimhan. Tree of thoughts: Deliberate problem solving with large language models, 2023.
- Tianzhu Ye, Li Dong, Xun Wu, Shaohan Huang, and Furu Wei. On-policy context distillation for language models, 2026.
- Tao Yu, Rui Zhang, Kai Yang, Michihiro Yasunaga, Dongxu Wang, Zifan Li, James Ma, Irene Li, Qingning Yao, Shanelle Roman, Zilin Zhang, and Dragomir Radev. Spider: A large-scale human-labeled dataset for complex and cross-domain semantic parsing and text-to-SQL task. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, pages 3911–3921, Brussels, Belgium, October-November 2018. Association for Computational Linguistics. doi: 10.18653/v1/D18-1425. URL <https://aclanthology.org/D18-1425>.
- Eric Zelikman, Yuhuai Wu, Jesse Mu, and Noah D. Goodman. Star: Bootstrapping reasoning with reasoning, 2022.
- Siyan Zhao, Zhihui Xie, Mengchen Liu, Jing Huang, Guan Pang, Feiyu Chen, and Aditya Grover. Self-distilled reasoner: On-policy self-distillation for large language models, 2026.
- Denny Zhou, Nathanael Schärli, Le Hou, Jason Wei, Nathan Scales, Xuezhi Wang, Dale Schuurmans, Olivier Bousquet, Quoc Le, and Ed Chi. Least-to-most prompting enables complex reasoning in large language models. *arXiv preprint arXiv:2205.10625*, 2022.

# Appendix A

## Chapter 2 Appendix

### A.1 Additional Related Work for Model Imitation

**Model distillation** Model imitation is similar to model distillation [Hinton et al., 2014], where one trains a student model to imitate a teacher. While conceptually similar, there are several major practical differences. For distillation, the training data, model architecture, and hyperparameters are known for the teacher. In model imitation, one tries to imitate the teacher without this knowledge. Moreover, for distillation it is common to use training objectives that utilize the probability distribution of the teacher whereas in stealing such a distribution is typically unavailable.

**Past work on model imitation** Prior work has shown that model imitation is possible for various domains [Lowd and Meek, 2005, Tramèr et al., 2016, Orekondy et al., 2019], including language classifiers [Krishna et al., 2020, Pal et al., 2019] and machine translation systems [Wallace et al., 2020]. Nevertheless, past work considers a setting where models are trained from *scratch*, and thus the main proprietary nature of a model is the company’s internal training data. In our setting, systems like ChatGPT are proprietary because they also leverage OpenAI’s internal pre-trained LMs that are stronger than any available open-source LM.

**Defending against model imitation** Our results show that imitation is a moderate concern for companies. In turn, there is a need to develop methods to mitigate or detect imitation. There is an existing body of work in this direction, e.g., one can detect whether a particular model is trained via imitation [Juuti et al., 2019, Szyller et al., 2019, Krishna et al., 2020, Maini et al., 2021] or slow model stealing by sacrificing some performance [Orekondy et al., 2020, Dziedzic et al., 2022a, Wallace et al., 2020, Dziedzic et al., 2022b]. Unfortunately, existing methods often exhibit too severe of a tradeoff to be deployable in practice.

## A.2 Additional Details on Imitation Data

To construct the NQ-synthetic dataset, we first curate seed examples from the Natural Questions validation set in Table A.1. We then use the prompting template in Table A.2 to randomly sample 5 QA pairs from the seed set to generate new QA samples. New samples are generated with temperature 1.0 and duplicate question-answer pairs are discarded.

|                                                                        |
|------------------------------------------------------------------------|
| Q: who sang who wants to be a millionare in high society?              |
| A: Frank Sinatra                                                       |
| Q: the last time la dodgers won the world series?                      |
| A: 1988                                                                |
| Q: who plays the medical examiner on hawaii five-o?                    |
| A: Masi Oka                                                            |
| Q: when did the first harry potter movie come out?                     |
| A: 2001                                                                |
| Q: when was the last time india won a gold medal in hockey at olympics |
| A: 1980                                                                |
| Q: who owns the rights to baby shark song                              |
| A: SmartStudy                                                          |
| Q: how many episodes are in one punch man season 1                     |
| A: 12                                                                  |
| Q: name of the bird in the lion king                                   |
| A: Zazu                                                                |
| Q: who sang the rap song change clothes                                |
| A: Jay-Z                                                               |
| Q: who stars as serena in gossip girl                                  |
| A: Blake Lively                                                        |

Table A.1: Seed examples curated from the Natural Questions validation set

To construct the TLDR-synthetic dataset we prompt ChatGPT with two randomly selected examples from the TLDR dataset [Völske et al., 2017], and ask it to produce a summary in a similar style to the two examples. If the two examples don’t fit in context, then we just use one instead. We include our prompt template in Table A.3.

Figure A.1 shows examples from ShareGPT-Mix and Table A.4 shows a breakdown of different categories.

```

I want you to generate a series of questions and
answers. I want the answers to be concise, just
a few words. The questions should be lowercased
and centered around Wikipedia-like entities. For
example,

Q: {question 1}
A: {answer 1}
Q: {question 2}
A: {answer 2}
Q: {question 3}
A: {answer 3}
Q: {question 4}
A: {answer 4}
Q: {question 5}
A: {answer 5}

```

Table A.2: Prompting template used to generate synthetic Natural Questions-like imitation data

### A.3 Amazon Mechanical Turk Interface

We use Amazon Mechanical Turk to conduct human evaluations. We use the UI shown in Figure A.2. It shows human evaluators a random task instruction and the output responses from two systems, one of which is our model and the other is ChatGPT. The annotators then choose which response is better according to overall subjective quality. We randomize whether ChatGPT or our imitation models are shown first. We collect 3 unique ratings for every example in the evaluation set and a total of 71 human evaluators participated. In order to get an average score, we use majority voting among the 3 raters on each example, and then average the scores across all examples. We pay these evaluators roughly \$15/hour based on the average time it takes to complete a task. In total, we spend roughly \$5000 on our ratings experiments, including service fees.

### A.4 GPT-4 evaluations

Our GPT-4 evaluations follow the procedure from Chiang et al. [2023]: we prompt GPT-4 with two outputs, one from ChatGPT and one from our imitation models. We then ask GPT-4 to output a preference ranking of the two outputs. We use the same set of evaluation prompts as in our human-preference evaluations. In Figure 2.1, we see that as we add more imitation data GPT-4’s ratings of our model outputs remain relatively flat. However as we

```

I want you to summarize a document for me. Here are
some examples:

Document: {passage 1}
Summary: {summary 1}

Document: {passage 2}
Summary: {summary 2}

Ok now summarize the following in a similar style
to the above examples:

Document: {query document}
Summary:

```

Table A.3: Prompting template used to generate TLDR-Synthetic imitation data

| Category      | Frequency (%) |
|---------------|---------------|
| Generation    | 31            |
| General QA    | 16            |
| Coding        | 16            |
| Brainstorming | 8             |
| Chat          | 6             |
| Summarization | 4             |
| Rewrite       | 2             |
| Non-English   | 14            |
| Invalid       | 2             |

Table A.4: We conduct a manual quality review of 50 random user queries from ShareGPT. The dataset contains highly varied categories of task instructions, including coding and multi-lingual queries.

increase the base model scale, we see GPT-4’s ratings consistently increasing (Figure 2.1). These results line up closely with the results from our crowdworker evaluations.

## ShareGPT-Mix Data

**User:** Return to the Identity and Access Management Commercial Due Diligence. We now want to model the market size for the IAM market. What are the levers and inputs we need to accurately estimate the size and growth of the U.S. market. For example, these may include the number of companies, current saturation, and specific market headwinds and tailwinds you identify. Provide suggestions where you would look for data for each of the levers.

**GPT:** To accurately estimate the size and growth of the U.S. Identity and Access Management (IAM) market, there are several key levers and inputs to consider:

1. Number of companies: The number of companies operating in the U.S. IAM market is an important lever to consider when estimating market size. This includes both large, established vendors and smaller, niche players. Data on the number of companies can be found through market research reports and industry databases, such as Gartner, IDC, and MarketsandMarkets.
2. Current saturation: ....

**User:** I am building an Elven Capital with all of it's history, religion, politics, trade routes, people, traditions, laws, general rules, and more. It is called Oniver. Your job is to ask me questions to learn more about the town and help me fill in the blanks to make it a rich and thriving town with lots of history. Ask detailed questions that will make me think. Say I understand if this makes sense. Then ask me the first question.

**GPT:** I understand.

What is the history behind Oniver's founding? How did it come to be and what was the context in which it was established?

**User:** Magic is widely used in Oniver. Better magic is a path to improvement in life and status in the capital. It is also the ticket into getting into the Magic University. There is a high importance between magical ability within the noble families of Oniver. Magic is used as a tool by the wealthy to oppress the lower classes. Commoners practicing magic is not forbidden but is looked down upon by the nobles because they see themselves as better than the commoners.

**GPT:** I understand.

What kind of economy does Oniver have? Is it based on trade, agriculture, manufacturing or a combination of these things? Are there any restrictions or laws surrounding trade and commerce? ....

Figure A.1: Examples of user inputs and ChatGPT outputs that are present in the ShareGPT data. Overall, we find that online datasets are typically high-quality and diverse in their user inputs, and span multiple categories such as open-ended text generation, brainstorming, and text extraction.

## A.5 Additional Summarization Results

We present a detailed breakdown of the ROUGE-1, ROUGE-2, and ROUGE-L scores for our summarization experiments in Table A.5.

| Model   | Imitation Data | ROUGE-1      | ROUGE-2      | ROUGE-L      |
|---------|----------------|--------------|--------------|--------------|
| 7B      | -              | 22.18        | 9.68         | 15.69        |
| 7B      | ShareGPT-Mix   | 28.69        | 11.30        | 19.14        |
| 7B      | TLDR-Synthetic | <b>29.21</b> | <b>12.00</b> | <b>19.17</b> |
| 13B     | -              | 27.25        | 11.62        | 19.33        |
| 13B     | ShareGPT-Mix   | 30.66        | 11.90        | 20.47        |
| 13B     | TLDR-Synthetic | <b>33.64</b> | <b>13.36</b> | <b>21.57</b> |
| ChatGPT | -              | 39.89        | 16.95        | 25.87        |

Table A.5: ROUGE-1, ROUGE-2, ROUGE-L scores for different models.

Instructions Shortcuts Which of the two outputs are higher quality responses to the input? Refer to the instructions panel on the left for specific criteria. ©

**Instructions**  
 In order to decide which output is of higher quality, please use the following criteria

- Grammatical correctness: The output should be free of grammatical errors and follow standard grammar rules.
- Clarity and coherence: The output should be clear and coherent, with a logical flow of ideas and easy-to-understand language.
- Relevance to the question: The output should be relevant to the question asked, providing a useful and accurate response.
- Completeness: The output should be complete and provide a sufficient answer to the question asked.
- Overall quality: The output should be of high quality overall, with a professional and

[More Instructions](#)

**Input:** \${user\_query}

**Output 1:** \${model\_response1}

**Output 2:** \${model\_response2}

**Select an option**

Output 1 is of higher overall quality ☐ 1

Output 2 is of higher overall quality ☐ 2

Both 1 and 2 are the same in overall quality ☐ 3

**Submit**

Figure A.2: Our Amazon Mechanical Turk interface for comparing the quality of different model outputs. Evaluators are presented with an instruction and two model outputs, and must rate which one is better or whether they are equal.

## A.6 Evaluation Prompting Details

For all evaluations, when prompting Koala, we append turn tokens before and after the prompt to distinguish the user’s input prompt from the model’s response. For evaluating MMLU we follow the conventional 5-shot prompt used in the LM-eval-harness [Gao et al., 2023] and select the answer letter with the greatest logprob. For GSM8K, we adapt the 6-shot prompt in Zelikman et al. [2022]. For human-eval we follow the standard evaluation procedure from Chen et al. [2021a]. For evaluating natural questions, we use the 3-shot prompt in Table A.6. For evaluating CNN/DM summarization we use the 2-shot prompt in Table A.7.

## A.7 Alternative Methods for Shifting the Point of Emergence

We are interested in understanding what other methods can induce a shift in the point of emergence. We therefore conduct additional experiments on MMLU with 1) varying the number of shots in the prompt, 2) performing continuous prefix tuning Li and Liang [2021], and 3) using low-rank finetuning with LoRA. We see in Figure A.3 that both prefix tuning and using fewer shots in the few-shot prompt have little effect on shifting the point of emergence. On the other hand, in Figure A.4 we see that low-rank finetuning shifts the point of emergence to a comparable degree to that of full finetuning, even in the rank-1 setting.

Together these results suggest that updating the model’s parameters may be necessary for shifting the point of emergence. Though further exploration of this phenomenon is needed.

```

You are a brief and concise question answering
service and only answer questions with a few words,
usually just a single word. Here are some examples
of how you respond to questions.
Q: Who sang who wants to be a millionaire in high
society?
A: Frank Sinatra
Q: In what year did Nelson Mandela become the first
black president of South Africa?
A: 1994
Q: Who discovered the first antibiotic, penicillin
A: Alexander Fleming
Now answer this question using only one to two words
at most.
Q:

```

Table A.6: Our 3-shot prompt for evaluating natural questions.

In particular, it is possible that using hundreds or thousands of shots in the prompt may enable more of a shift than what we observed [Agarwal et al., 2024a].

In each of these experiments, we use all 3B, 7B, and 13B OpenLLaMA V1 model checkpoints. For the prefix tuning baseline we use a prefix length of 8, parameterized by a 2-layer MLP with input and hidden-dim 512. We train with a learning rate of 3e-4 and keep all other hyperparameters the same as our full finetuning experiments. We found attempts to increase the capacity of the prefix tuning (e.g., using a longer prefix or dropping the MLP and directly tuning the embeddings) to make training more unstable and generally yield worse performance. For LoRA we use the same finetuning hyper-parameters as full fine-tuning, including the learning rate. To ensure that the LoRA updates are of similar magnitude to full finetuning with the same learning rate, we set the LoRA  $\alpha$  hyper-parameter to be equal to the model’s hidden dimension. We found this setup to yield the best results with minimal changes between our full finetuning and LoRA setup.

## A.8 Finetuning and Few-shot Emergence Across Model Sizes

In Figure A.5, we plot emergence as a function of pretraining loss using all 3B, 7B, and 13B OpenLLaMA V1 model intermediate checkpoints in both the few-shot and finetuned setting. The specific finetuning hyperparameters and data splits used are detailed in Appendix A.9. For the few-shot results, we use the prompts detailed in Appendix A.9. On all tasks, we see that both the point of emergence and the downstream performance scaling thereafter,

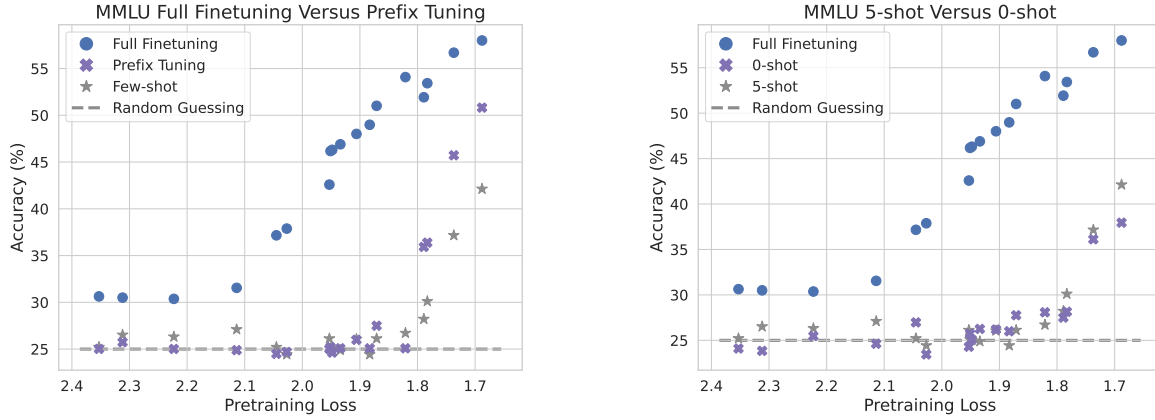


Figure A.3: On the left we compare full fine-tuning against continuous prefix tuning on MMLU. We find that prefix tuning provides effectively no shift to the point of emergence, despite improving the performance of post-emergence models. On the right we compare 0-shot versus 5-shot prompting on MMLU. We see that using fewer shots has no meaningful effect on the point of emergence. Together these results suggest that the ability for prompt tuning to shift the point of emergence is very limited.

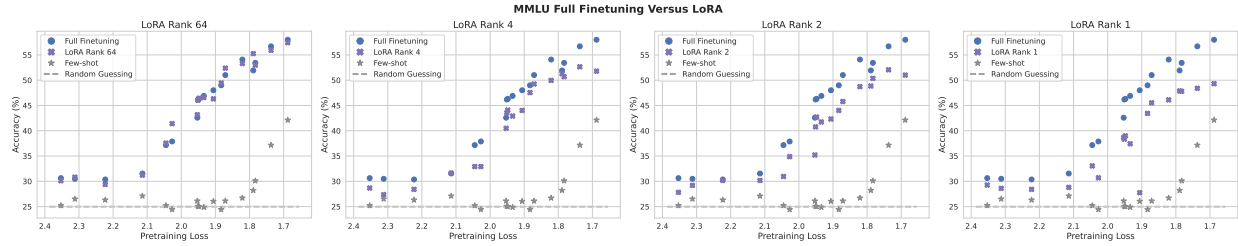


Figure A.4: Comparing LoRA finetuning, with rank 1, 2, 4, and 64 against full finetuning on MMLU. We see that LoRA finetuning even with rank 1 shifts the point of emergence to a comparable degree to that of full finetuning.

as a function of pretraining loss, is consistent across model size in both the few-shot and finetuned setting. In the case of CoLA, finetuning on the full data amount significantly shifts the point of emergence to such a degree that it precedes all model checkpoints that we consider. When using smaller data amounts, the emergence elbow on CoLA is shifted towards stronger models and is thus visible with our checkpoints.

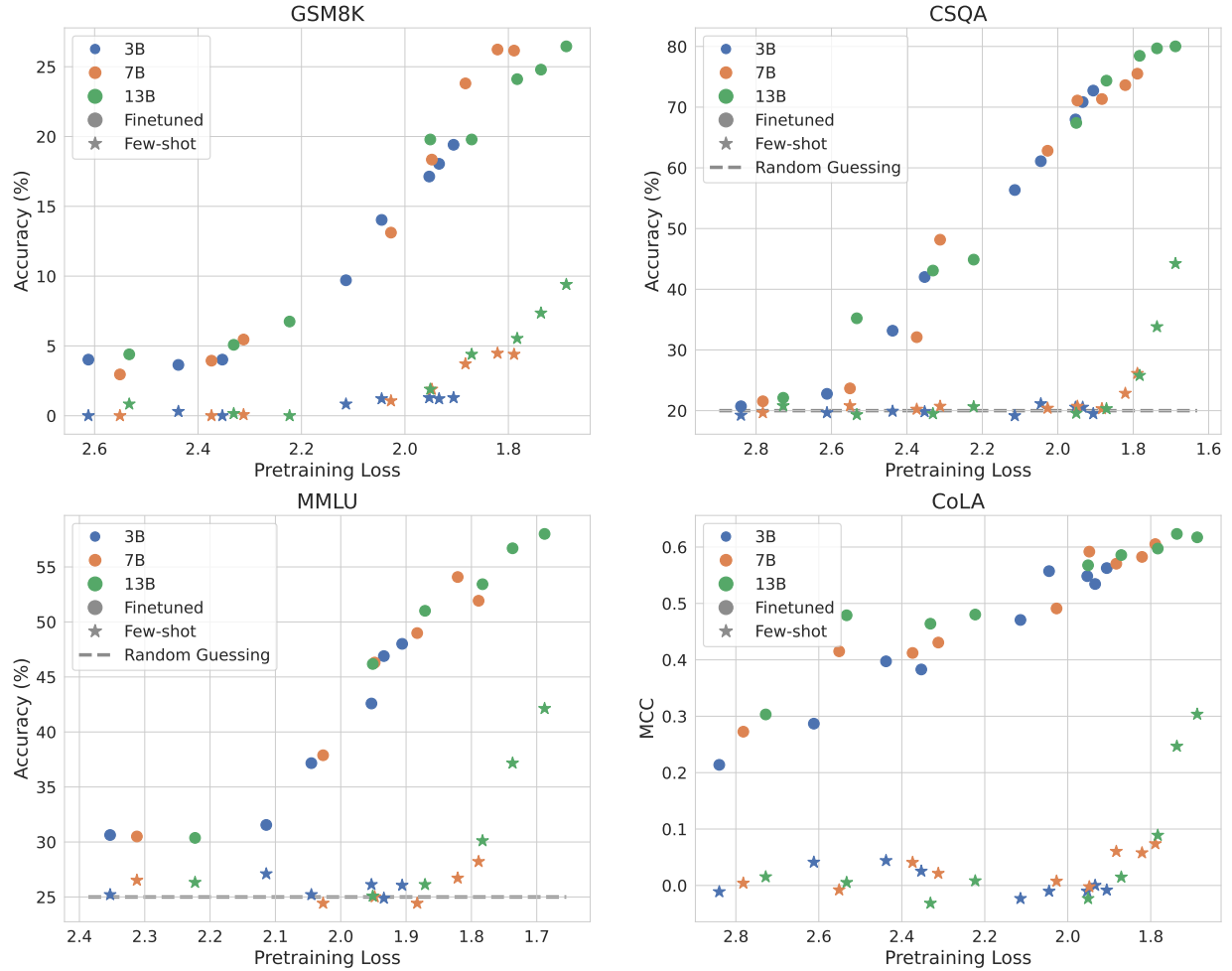
**Emergence as a Function of Loss is Consistent Across Model Size, Fewshot and Finetuned**

Figure A.5: We plot few-shot and full data finetuning performance as a function of pretraining loss using all 3B, 7B, and 13B model checkpoints for all tasks. We see that both the point of emergence and the downstream performance scaling thereafter, as a function of pretraining loss, is consistent across model size in both the few-shot and finetuned setting.

## A.9 Details for Emergence Finetuning and Evaluation Experiments

We detail the specific few-shot prompts and training/evaluation splits for each task below:

- **MMLU:** We use the standard 5-shot LM evaluation harness [Touvron et al., 2023a] prompt for our few-shot results. We train on the MMLU test set and then evaluate on the validation set.
- **GSM8K:** We use the 6-shot chain-of-thought prompt in Zelikman et al. [2022]. We train and test on the standard splits.
- **CommonsenseQA:** We use the 7-shot prompt in Wei et al. [2022c] with the chain of thought examples stripped from the prompt, just showing the final answer. We train on the standard train set and evaluate on the validation set.
- **CoLA:** We use the standard 5-shot prompt in the LM evaluation harness [Touvron et al., 2023a]. We train on the standard train set and evaluate on the validation set.

|                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p>Article:</p> <p>(CNN)French striker Bafetimbi Gomis, who has a history of fainting, said he is now "feeling well" after collapsing during Swansea's 3-2 loss at Tottenham in the Premier League on Wednesday. The worrying incident occurred in the first half at White Hart Lane - after Tottenham scored in the seventh minute - but the 29-year-old left the pitch conscious following about five minutes of treatment. The Guardian added that he was wearing an oxygen mask. Play was temporarily stopped before resuming. As the match progressed, Swansea tweeted that Gomis was "fine," with manager Garry Monk using the same word to describe Gomis' condition. Gomis spent the night in hospital as a precaution, Swansea said on its website. "I wanted to reassure you concerning my health," Gomis told the website. "It actually looks much scarier than it is physically dangerous, and I am feeling well now. "I have been under a great deal of stress and fatigue due to my father's health, which requires me to go back and forth from France. "I was disappointed that I couldn't help my team tonight, but now everything is back in order. I also want to thank everyone for their support and get well messages." Gomis had similar fainting spells in France, which prompted the president of his former club, Jean-Michel Aulas of Lyon, to tell French television in 2009: "We can't not be worried, it scares you each time." Swansea ran tests on Gomis, said Monk, prior to signing him on a free transfer last July. "He just has a little bit of low blood pressure which causes you a little bit of problems," Monk said in a televised interview on Sky. "It's been part of his life. We were well aware of that when we signed him. He's done all the hospital checks and all the medical checks you can possibly do and it's just part of his life. "It's no problems whatsoever. It's not as serious as it looks." Gomis has scored two league goals for Swansea this season, mostly in a backup role. He became the Welsh side's top striker when Wilfried Bony signed with Manchester City in January. Almost exactly three years ago at White Hart Lane, then Bolton midfielder Fabrice Muamba collapsed after suffering a cardiac arrest. He was near death, according to Bolton, but survived after being treated at the London Chest Hospital. He subsequently retired. Other footballers, including Cameroon international Marc-Vivien Foe in 2003 and Spanish international Antonio Puerta in 2007, didn't survive after collapsing on the pitch.</p> <p>TL;DR:</p> <p>Bafetimbi Gomis collapses within 10 minutes of kickoff at Tottenham . But he reportedly left the pitch conscious and wearing an oxygen mask . Gomis later said that he was "feeling well" The incident came three years after Fabrice Muamba collapsed at White Hart Lane .</p> <p>Article:</p> <p>(CNN)It was an act of frustration perhaps more commonly associated with golf's fictional anti-hero Happy Gilmore than the world's reigning No 1. player. But when Rory McIlroy pulled his second shot on the eighth hole of the WGC Cadillac Championship into a lake Friday, he might as well have been channeling the much loved Adam Sandler character. Before continuing his round with a dropped ball, the four-time major winner launched the 3-iron used to play the offending shot into the water as well. "(It) felt good at the time," a rueful McIlroy later said of the incident in comments carried by the PGA Tour website. "I just let frustration get the better of me. It was heat of the moment, and I mean, if it had of been any other club I probably wouldn't have but I didn't need a 3-iron for the rest of the round so I thought, why not." The club "must have went a good 60, 70 yards," he joked. McIlroy composed himself to finish with a second round of 70, leaving him one-under for the tournament and eight shots off the pace set by leader JB Holmes. While an improvement on last weeks performance at the Honda Classic event, where he failed to make the cut, the Northern Irishman's frustration with elements of his game was still clear. "I think every golfer feels it because I don't hit shots like the one I hit on 8 on the range," he said. "That's what really bothers me, the fact that I get out on the course and I hit shots that I'm not seeing when I'm in a more relaxed environment. "So it's a little bit of mental, a little bit of physical. It's just everything is not quite matching up." Elsewhere on the course, Ryan Holmes scored a two-under-par 71 to remain in second position overall, two shots behind Holmes. Former world No 1., Adam Scott carded an impressive 68 to finish the day three shots off the pace at six-under while Bubba Watson and Henrik Stenson are tied for fourth on four-under.</p> <p>TL;DR:</p> <p>Rory McIlroy throws club into water at WGC Cadillac Championship . Northern Irishman frustrated after pulling shot into water hazard .</p> <p>Article:</p> <p>{article}</p> <p>TL;DR:</p> |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|

Table A.7: Our 2-shot prompt for evaluating CNN/DM summarization.

# Appendix B

## Chapter 3 Appendix

### B.1 Additional Revision Results

We plot additional results for majority selection using our PaLM 2-S\* revision model in Figure B.1. With majority selection, we see largely similar trends to those found in Figure 3.7 for verifier selection.

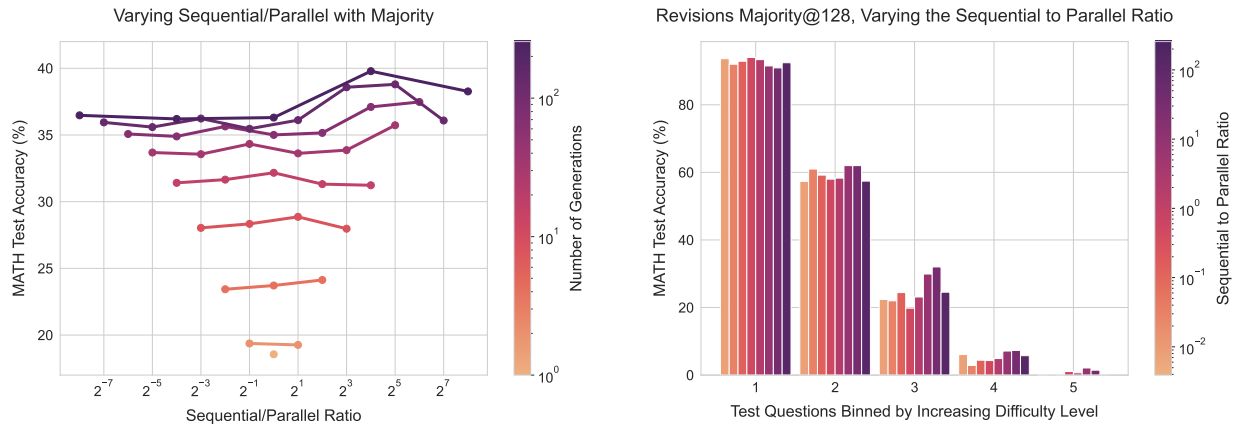


Figure B.1: Varying the ratio of generation budget allocated to sequential versus parallel samples, using majority to select the answer, rather than the verifier. **Left:** Each line represents a fixed generation budget as the ratio is changed. We see that similar to the verifier case, in the majority case, there exists an ideal ratio of sequential to parallel test-time compute at a given budget. **Right:** Analyzing performance across difficulty bins, we see that the easier questions are mostly invariant the ratio of sequential to parallel, whereas on the harder questions there is an ideal ratio of sequential to parallel test-time compute.

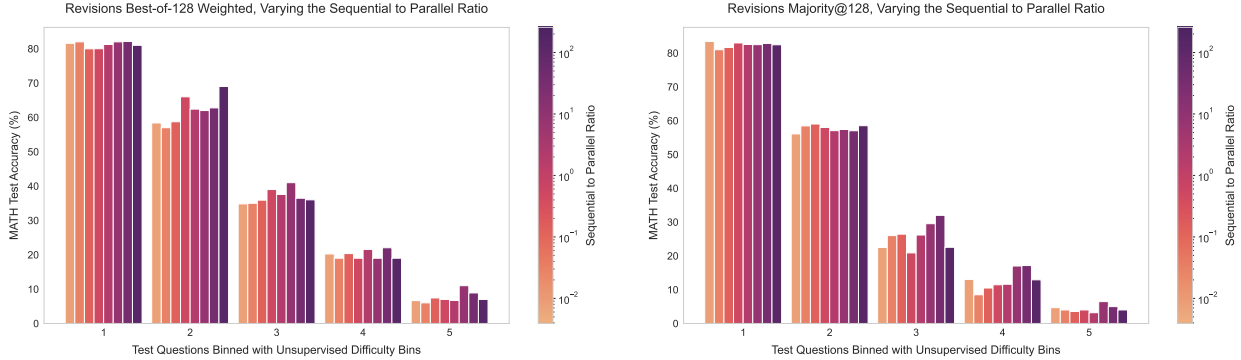


Figure B.2: Using our PaLM 2-S\* PRM to compute difficulty bins without ground truth correctness information for revisions. On the left we plot verifier selection and on the right we plot majority selection. We see largely similar performance trends with these bins as we do with the ground truth ones in Figures 3.7 and B.1.

## B.2 Unsupervised Difficulty Bins

We compute difficulty bins without oracle ground-truth correctness information by averaging the PRM final-answer score over 2048 samples on each question, so as to obtain a value estimate corresponding to the question. We then bin the value for each question in the test-set into five quintiles (using the same procedure as the oracle difficulty bins). We refer to this as “predicted difficulty” rather than “oracle difficulty”. Technically this procedure is extremely costly because it requires generating many samples. While we do not account for this cost in our analysis, in a practical production setting, this cost would be problematic. A more efficient approach would be to finetune a model to predict correctness directly, given the question. We do not explore this in our work, but leave such exploration of cheaper methods of estimating difficulty to future work.

In Figure B.3 we plot PRM-search results using our difficulty bins, and in Figure B.2 we plot the corresponding revision results. We see that in both settings these predicted bins demonstrate similar trends to the oracle bins.

## B.3 PRM Training Details

We finetune our PRM as a binary classifier, where the model predicts a value between 0 and 1 at each step in the solution. We train the model with soft values obtained from the monte-carlo rollouts, using a binary cross entropy loss function (e.g.  $-(y \log(\hat{y}) + (1 - y) \log(1 - \hat{y}))$  where  $y$  corresponds to the soft ground-truth value and  $\hat{y}$  the model’s predicted value). We finetune the model base model using the AdamW optimizer, with lr  $3e-5$ , batch size 128, dropout 0.05, and Adam betas (0.9, 0.95). We conduct early stopping, selecting the

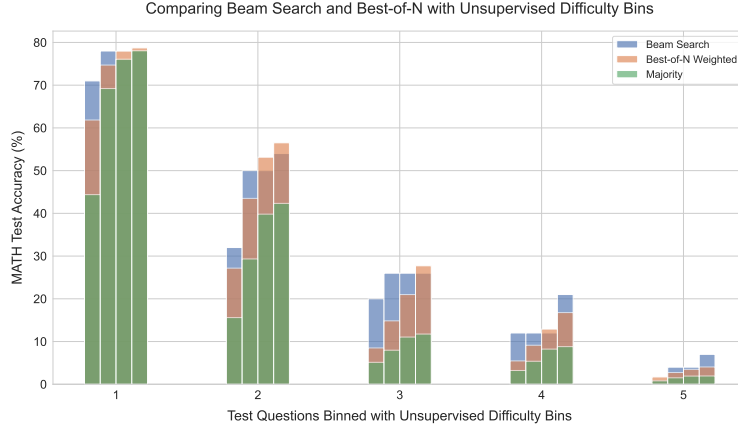


Figure B.3: Using our PaLM 2-S\* PRM to compute difficulty bins without ground truth correctness information for PRM search. We see largely similar performance trends with these bins as we do with the ground truth ones in Figure 3.3.

checkpoint with the lowest validation loss on a random held-out validation set, consisting of 10% of the questions in the original PRM800k training split.

We finetune the PRM on 16 samples per question from the corresponding few-shot prompted base model. At each step, we use 16 monte-carlo rollouts, using the same base model and prompt, to estimate the step-level value. We filter out all samples which fail to output a valid, parsable final answer from the training data, as we found these to hurt PRM performance in initial experiments.

When generating the samples, the base model is prompted to output answers in newline separated a step-by-step format, as done in Lightman et al. [2023]. We then separate each of the answers into steps using a simple newline splitting procedure. We include details about our prompt in Appendix B.6.

## B.4 Comparing PRM Aggregation Strategies

We compare different methods of aggregating per-step PRM scores to produce a final score for the full solution. Specifically we compare: 1) taking the minimum score accross all steps as done in Lightman et al. [2023] (e.g. “min”); 2) taking the product of all step correctness probabilities (e.g. “prod”); and 3) taking just the last step prediction (e.g. “last”). We see in Figure B.4 that taking the last step outperforms the other two approaches. Prior works [Lightman et al., 2023, Wang et al., 2023a] found min to be the best aggregator. We believe that the discrepancy is due to the fact that our verifier was trained with soft MC return labels, which surface very differently from binary correctness labels, and therefore other aggregation strategies may not have the same effect.

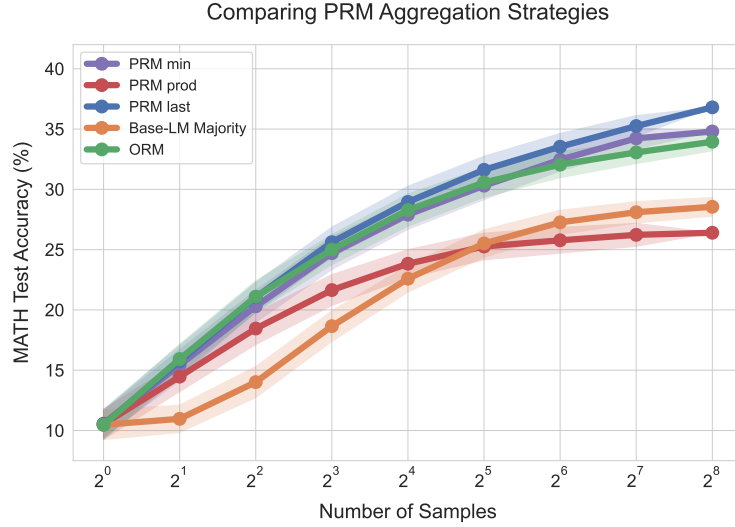


Figure B.4: We compare different methods of aggregating per-step PRM scores to produce a final score for the full solution: “min” refers to taking the minimum score across all steps, “prod” takes the product of all step correctness probabilities, and “last” just uses the last step score. We see that last performs the best across all aggregation strategies.

Interestingly, when using the last step aggregation, we are effectively using the PRM like an ORM. However, we see that the PRM outperforms the ORM, suggesting that in our case the per-step PRM training may be largely useful as a form of representation learning, rather than purely as a tool at inference time. Future work should further explore this line of reasoning.

## B.5 Comparing PRM and ORM

We trained a PRM and ORM model using the PaLM 2-S\* base LM. We see in Figure B.5, that the PRM outperforms the ORM, and the gap between the PRM and ORM grows with the number of samples used. We use the last step prediction from the PRM to score the answers as described in Appendix B.4.

## B.6 Prompting Details

In order to enable the base model to output answers in a step-by-step format to which a PRM can be applied, we use a 4-shot prompt consisting of randomly selected correct answer examples from the PRM800k data released by Lightman et al. [2023]. Specifically we use answers from the phase 1 training split. These answers correspond to GPT-4 generated

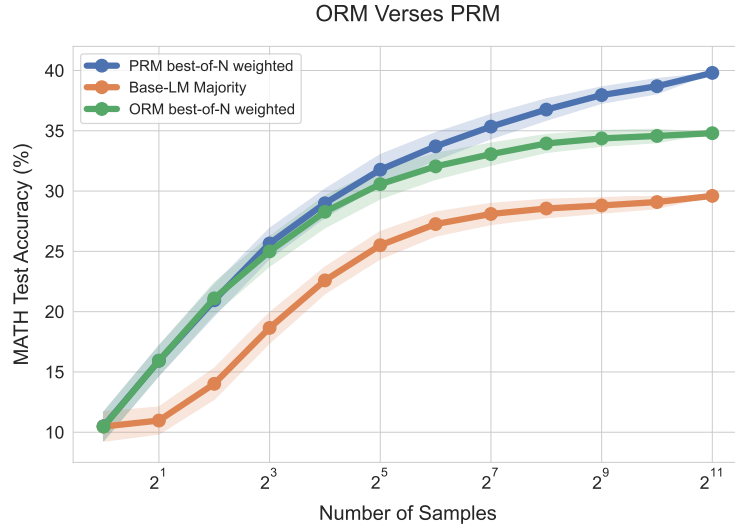


Figure B.5: We compare PRM and ORM models finetuned from PaLM 2-S\* in a best-of-N evaluation. We use the PaLM 2-S\* base LM to sample outputs, using a few-shot prompt. We see that the PRM greatly outperforms the ORM at a large number of samples.

correct answer examples, which include the correct step-by-step format. In initial experiments, we found that this prompting procedure produces similar results to the prompt used in Lewkowycz et al. [2022]. We use this prompt for generating training data for the PRM and the revision model. We also use this prompt when conducting search against the PRM on the test-set. To grade the final answer predicted by this prompt, we use the grading function released by Lightman et al. [2023].

## B.7 Revision Model Finetuning Details

For fine-tuning the revision model, we follow the procedure outlined in Section 3.5. We first sample 64 outputs per question. We then filter out all answers which end in an invalid solution. For each correct answer, we then sample a number uniformly between 0 and 4 indicating how many incorrect answers to include in context for training. The correct answer is used as the last answer in the trajectory (which we train the model to produce) and the incorrect answers are included in context. If the sampled number is greater than 0, we then find the closest incorrect answer according to a character-level edit distance metric to include as the last incorrect answer in the trajectory. The goal here is to select an incorrect answer which is somewhat correlated with the correct answer, to improve learning. The remaining incorrect answers, we sample randomly from the set of available answers. In the case where there are fewer than 4 incorrect answers sampled, we truncate the uniform distribution’s max

to match the number of incorrect samples. We use this procedure to generate trajectories for all questions in the training data.

We then finetune the base language model on the correct answer solutions in these generated trajectories. We use the AdamW optimizer with lr  $1e-5$ , batch size 128, dropout 0.0, and Adam betas (0.9, 0.95).

We find that generally evaluating loss on an evaluation set consisting of trajectories generated as described above, does not provide a good signal for early stopping. Rather, we find that checkpoints much after the evaluation loss begins increasing are much more capable of revisions. This is likely because after finetuning the revision model, the evaluation set represents off-policy data, which will naturally be out-of-distribution compared to the trajectories that the model itself would generate on-policy. We therefore select our revision model checkpoint slightly after the point where we observe overfitting on the validation set.

## B.8 Revision Model Selection Criteria

As described in Section 3.5, in order to effectively use our revision model we need to deploy a criteria for selecting the best answer both within a revision trajectory and between multiple parallel trajectories. We use two approaches: 1) ORM verifier; and 2) majority voting.

For the ORM verifier, we train an ORM on the revision model’s outputs according to the procedure in Appendix B.9. At inference time we then use this verifier to select the best answer. Since we have two axes across which to aggregate (within each revision trajectory and between multiple trajectories), we deploy a hierarchical strategy, first selecting the best answer within each revision trajectory and then aggregating these selected answers across trajectories. To select the best answer within each trajectory, we perform best-of-N weighted aggregation and then choose the highest scoring solution with the maximum best-of-N weighted answer. Then, to select the final answer across all revision chains, we perform another round of best-of-N weighted selection using the best answer from each revision chain. The answer after this second round of best-of-N weighted represents our final answer prediction.

For majority voting we found hierarchical aggregation to create problems when the length of the trajectory or the number of trajectories was too small. The problem being that without enough samples, majority voting is unable to effectively select the best option. Therefore, for majority voting, we simply take all answers, across all trajectories, at once and take their majority as the final-answer. We found this to produce much smoother scaling behavior than the hierarchical approach.

## B.9 Revision Model Verifier Training

We found that the PRM we finetuned on the PaLM 2-S\* base model outputs was not as effective when applied to the PaLM 2-S\* revision model’s outputs (see Figure B.9), likely

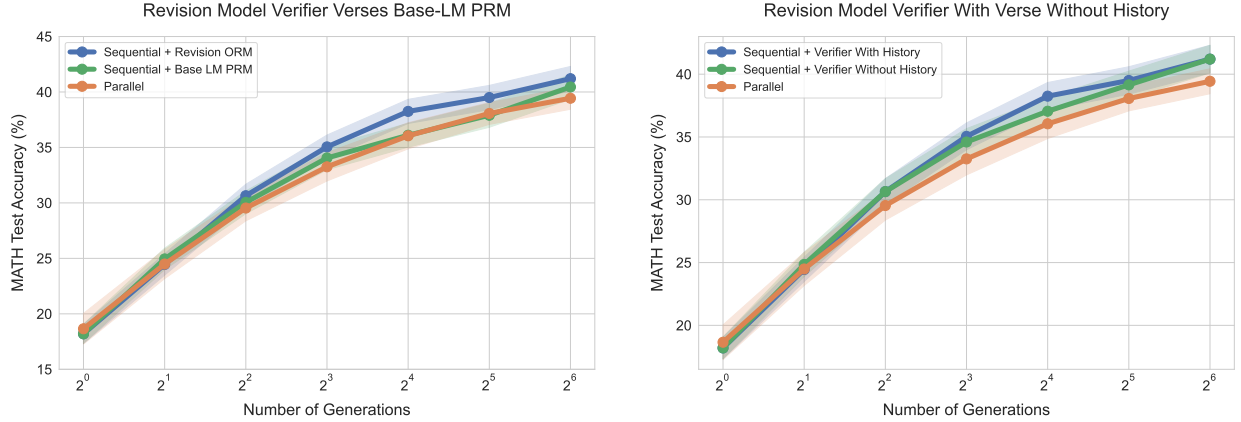


Figure B.6: Left: we compare the ORM we trained on the revision model’s outputs against the PRM we trained on the PaLM 2-S\* base model’s outputs. We see that when applied to outputs from the revision model, the ORM adapted to the revision model outperforms the PRM, likely due to distribution shift with the revision model. Right: we ablate the effect of including previous revisions in the revision model verifier’s context. We see that including revisions in-context helps the verifier slightly, but both settings still outperform the parallel baseline.

due to distribution shift with the revision model. We therefore, trained a separate ORM verifier to use with our PaLM 2-S\* revision model. We could have trained a PRM as well, but opted for an ORM due to the high cost of generating per-step PRM labels.

We modified the standard ORM slightly for the revision setting, by finetuning the ORM with previous revision in context, such that the verifier has access to the same context as the revision model, allowing the verifier see the revision model’s previous answer attempts when scoring the current answer. All other experiment details are identical to those used for training the PRM.

Empirically, we find that including the revision history in context improves performance slightly (see Figure B.9). Additionally, even without the revisions in context, we see that sequential revisions still slightly outperforms parallel, demonstrating improvements from sequential sampling are not just due to the verifier’s context.

## B.10 ReST<sup>EM</sup> Revision Model Experiments

We experimented with further optimizing our PaLM 2-S\* revision model by training the model with a simplified RL algorithm: ReST<sup>EM</sup> [Singh et al., 2024]. Specifically, we generated 64 revision trajectories of maximum length 5 for each question on the MATH training set. We stopped the revision model at the first correct answer in each trajectory. Using this

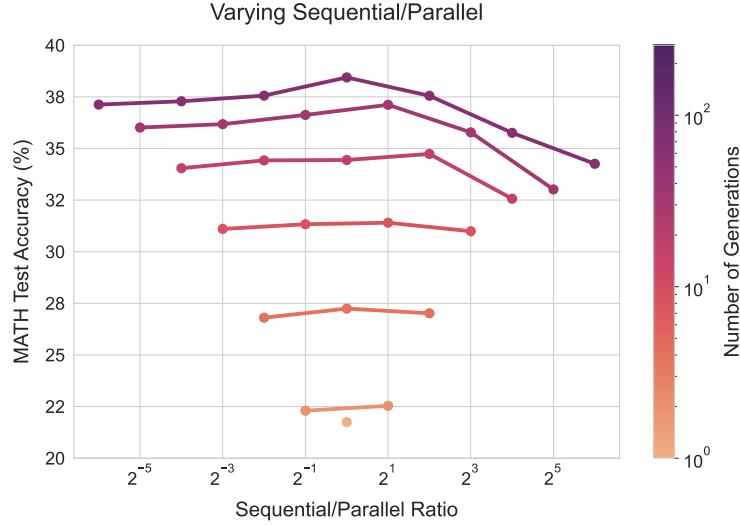


Figure B.7: Performance of our  $\text{ReST}^{\text{EM}}$  optimized revision model as the sequential to parallel ratio is varied. We use majority voting to select the answer. We see that this optimized revision model demonstrates substantial performance degradations with additional sequential revisions.

generated data, we then finetuned the base LM on the correct answer data. To help the model learn the task, we explicitly balanced the distribution of trajectory lengths.

In Figure B.7, we plot the performance of this new revision model as we vary the sequential to parallel ratio. We see that additional sequential revisions substantially hurts performance with this new model. We hypothesize that this degradation is due to the fact that the online data obtained from running  $\text{ReST}^{\text{EM}}$  exacerbates spurious correlations in revision data, causing the optimized model to fail to learn the revision task. We believe that using a more offline data collection strategy, as done in Qu et al. [2024], may be more effective, and leave further exploration to future work.

## B.11 Revision Model Example Outputs

In Figures B.8, B.9, B.10, B.11, B.12, B.13, and B.14, we include select examples of our revision model’s outputs.

## B.12 PRM Beam Search Example Outputs

In Figures B.15, B.16, B.17, B.18, B.19, and B.20, we include select examples of PRM beam search. We include the PRM score, between 0 and 1, for each step in the examples.

If  $f(x) = \frac{3x-2}{x-2}$ , what is the value of  $f(-2) + f(-1) + f(0)$ ? Express your answer as a common fraction.

Attempt 1:

Let's evaluate  $f(-2)$ .  
 We have  $f(-2) = \frac{3(-2)-2}{-2-2} = -\frac{4}{-4} = \frac{1}{1} = 1$ .  
 So  $f(-2) = 1$ .  
 Let's evaluate  $f(-1)$ .  
 We have  $f(-1) = \frac{3(-1)-2}{-1-2} = -\frac{5}{-3} = \frac{5}{3}$ .  
 So  $f(-1) = \frac{5}{3}$ .  
 Let's evaluate  $f(0)$ .  
 We have  $f(0) = \frac{3(0)-2}{0-2} = \frac{-2}{-2} = 1$ .  
 So  $f(0) = 1$ .  
 So we want to compute  
 $1 + \frac{5}{3} + 1 = \frac{3}{3} + \frac{5}{3} = \frac{3+5}{3} = \frac{8}{3}$ .  
 So the answer is  $\boxed{\frac{8}{3}}$ .  
 ####  $\frac{8}{3}$

Attempt 2:

Let's first find  $f(-2)$ .  
 $f(-2) = \frac{3(-2)-2}{-2-2} = \frac{-6-2}{-4} = \frac{-8}{-4} = 2$ .  
 Now let's find  $f(-1)$ .  
 $f(-1) = \frac{3(-1)-2}{-1-2} = \frac{-3-2}{-3} = \frac{-5}{-3} = \frac{5}{3}$ .  
 Now let's find  $f(0)$ .  
 $f(0) = \frac{3(0)-2}{0-2} = \frac{-2}{-2} = 1$ .  
 So we have  $f(-2) + f(-1) + f(0) = 2 + \frac{5}{3} + 1 = \frac{7}{3} + \frac{5}{3} + \frac{3}{3} = \frac{7+5+3}{3} = \frac{15}{3} = 5$ .  
 So the answer is  $\boxed{5}$ .  
 #### 5

Attempt 3:

Let's plug in  $-2$ ,  $-1$ , and  $0$  for  $x$  in  $f(x)$ .  
 We get  $f(-2) = \frac{3(-2)-2}{-2-2} = \frac{-8}{-4} = \frac{2}{1} = 2$ ,  $f(-1) = \frac{3(-1)-2}{-1-2} = \frac{-5}{-3} = \frac{5}{3}$ ,  $f(0) = \frac{3(0)-2}{0-2} = \frac{-2}{-2} = 1$ .  
 So we have  $2 + \frac{5}{3} + 1 = \frac{6}{3} + \frac{5}{3} + \frac{3}{3} = \frac{14}{3}$ .  
 So the answer is  $\boxed{\frac{14}{3}}$ .  
 ####  $\frac{14}{3}$

Figure B.8: Revision model example 1. The model calculates the sum at the end incorrectly on the first two attempts, but on the third attempt it succeeds and gets the answer correct.

**In how many ways can 7 people sit around a round table if no two of the 3 people Pierre, Rosa, and Thomas can sit next to each other? (Seating arrangements which are rotations of each other are treated as the same.)**

Attempt 1:

```

Let's start by figuring out how many ways there are to seat the three people who
can't sit next to each other.
So we need to seat them in three seats that are not right next to each other.
So there are 4 chairs that they can sit in.
So we have 4 slots to fill with 3 people.
We can do that in 4 choose 3 ways.
That is  $4!/(3! * 1!) = 4$  ways.
Now we need to figure out how many ways there are to seat the remaining 4 people.
So we can seat the first person in any of the seven chairs.
And then we can seat the second person in any of the remaining six chairs.
And so on.
So we can seat the four people in  $7 * 6 * 5 * 4$  ways.
So we need to multiply the number of ways to seat the three people who can't sit
next to each other, and the number of ways to seat the other four people.
That is  $4 * 7 * 6 * 5 * 4 = 840$ .
So the answer is  $\boxed{840}$ .
#### 840

```

Attempt 2:

```

So we have to figure out how many ways there are to seat 7 people around a table
if no two of the 3 people can sit next to each other.
Let's start by seating the 3 people who can't sit next to each other.
We know that they can't sit next to each other, but they can sit on the ends of
the table.
So they can sit in 3 different ways.
Now let's seat the 4 people who can sit anywhere.
We can do that in 4! ways.
Now let's seat the 3 people who can't sit next to each other.
Let's seat them in the 3 spots we found earlier.
So there are 3 different ways to seat them.
So we have  $3 * 4! * 3 = 3 * 24 * 3 = 216$  ways to do it.
So the answer is  $\boxed{216}$ .
#### 216

```

Attempt 3:

```

So we have to find the number of ways to arrange 4 people around a round table
such that 3 of the people can't sit next to each other.
So let's start with the 3 people who can sit next to each other.
There are 3! ways to arrange them.
Now we have to arrange the other 4 people.
There are 4! ways to do that.
So the total number of ways to arrange the 7 people around the table is  $3! * 4! =$ 
 $6 * 24 = 144$ .
So the answer is  $\boxed{144}$ .
#### 144

```

Figure B.9: Revision model example 2. On the first attempt the model takes the incorrect approach, on the second attempt it gets closer but then makes a mistake towards the end. On the final attempt it gets to the correct answer.

Find the product of  $6_8 \cdot 7_8$ . Express your answer in base 8.  
 $52_8$

Attempt 1:

So we have  $6_8 \cdot 7_8$ .  
 Right.  
 So, we can convert these numbers to base 10.  
 Right, we have  $6_8 = 6_{10}$  and  $7_8 = 7_{10}$ .  
 So our equation becomes  $6_{10} \cdot 7_{10}$ .  
 So, we have  $6_{10} \cdot 7_{10} = 42_{10}$ .  
 So, we need to convert  $42_{10}$  to base 8.  
 Right.  
 So, we know that  $4_8 = 4_{10}$  and  $2_8 = 2_{10}$ .  
 So, we can divide  $42_{10}$  by  $8_{10}$  and get  $5 \text{ R } 2$ .  
 So, our answer is  $52_8$ .  
 Correct.  
 #### 52

Attempt 2:

Well, we can start by expanding the numbers in base 10.  
 Right. So  $6_8 = 6$  and  $7_8 = 7$ .  
 So we get  $6 \cdot 7$ .  
 Right. The product is  $42$ .  
 So, we need to convert  $42$  to base 8.  
 Right. So, we divide  $42$  by  $8$  and get a remainder of  $2$  and a quotient of  $5$ .  
 So, we write  $52_8$ .  
 Correct.  
 ####  $52_8$

Figure B.10: Revision model example 3. On the first attempt the model makes a mistake with the formatting of the final answer; it corrects this on the second attempt.

Find the product of  $6_8 \cdot 7_8$ . Express your answer in base 8.  
**52\_8**

Attempt 1:

Let's multiply them together.  
 We get  $6_8 \cdot 7_8 = 42_8$ .  
 So the answer is  $42_8$ .  
 #### 42\_8

Attempt 2:

We know that  $6_8 = 6$ , and  $7_8 = 7$ . So the problem is finding  $6 \cdot 7$ .  
 Yep.  
 $6 \cdot 7 = 42$ .  
 Correct. So that means  $6_8 \cdot 7_8 = 42_8$ .  
 So our answer is  $42_8$ .  
 #### 42\_8

Attempt 3:

Let's multiply  $6_8$  by  $7_8$  in base 8.  
 $6 \cdot 7 = 42$  in base 10.  
 That's correct. So  $6_8 \cdot 7_8 = 42_{10} = 40_8 + 2_8$ .  
 We can rewrite this as  $40_8 + 2_8 = 40_8 + 10_8 + 10_8 + 8_8 + 8_8 + 8_8 + 8_8 + 1_8$ .  
 And we can rewrite that as  $40_8 + 10_8 + 10_8 + 10_8 + 10_8 + 10_8 + 10_8 + 10_8 + 1_8$ .  
 And we can rewrite that as  $50_8 + 1_8$ .  
 So  $6_8 \cdot 7_8 = 51_8$ .  
 #### 51

Attempt 4:

So  $6_8 \cdot 7_8$  is  $6 \cdot 7$  in base 8.  
 That's right.  
 Let's convert  $6$  and  $7$  to decimal.  $6 = 6_8 = 6 \cdot 8^0 + 0 \cdot 8^1 = 6$ .  
 Right. And  $7 = 7_8 = 7 \cdot 8^0 + 0 \cdot 8^1 = 7$ .  
 Right. So  $6 \cdot 7 = 6 \cdot 7 = 42$ .  
 And  $42$  in octal is  $52_8$ .  
 That's correct.  
 #### 52\_8

Figure B.11: Revision model example 4. On the first few attempts the model fails the base 10 to base 8 conversion. On the final attempt it makes the correct calculation.

Convert the point  $(0,3)$  in rectangular coordinates to polar coordinates. Enter your answer in the form  $(r,\theta)$ , where  $r > 0$  and  $0 \leq \theta < 2\pi$ .

Attempt 1:

```
To convert (0, 3) to polar coordinates, we need to find the radius r and the
angle theta.
We know that the x-coordinate of the point is r*cos(theta) and the
y-coordinate of the point is r*sin(theta).
So r*cos(theta) = 0 and r*sin(theta) = 3.
And we know that cos(theta) = 0 when theta = pi/2 and sin(theta) =
3 when theta = 3pi/2.
So we know that (0, 3) corresponds to the point (r, theta) where theta =
3pi/2 and r*cos(theta) = 0 and r*sin(theta) = 3.
So r = 3/sin(theta) = 3/sin(3pi/2) =
3/(1) = 3.
So the polar coordinates of (0, 3) are (3, 3pi/2).
#### (3, 3pi/2)
```

Attempt 2:

```
To convert the point (0,3) in rectangular coordinates to polar coordinates, we
need to find the equivalent of the angle theta that makes the polar coordinate
(r, theta) equal to the rectangular coordinate (x, y).
We know that (r, theta) = (0,3) and (x, y) = (0, 3).
So let's set up an equation where r is equal to x and theta is equal to
y.
We have r = x and y = theta.
So we can substitute these into the polar equation r = sqrt(x^2 + y^2).
We get sqrt(x^2 + y^2) = x.
We can square both sides of the equation to get x^2 + y^2 = x^2.
And we can subtract x^2 from both sides of the equation to get y^2 = 0.
And we know that to satisfy y^2 = 0 we must have y = 0.
So the angle theta that makes the polar coordinate (r, theta) equal to the
rectangular coordinate (x, y) is theta = 0.
And the radius r is equal to the x-coordinate, which is x = r = 0.
So the polar coordinate equivalent to (0,3) is (r, theta) = (0,0).
And we know that r > 0 and 0 <= theta < 2pi.
So the polar coordinate equivalent to (0,3) is (r, theta) = (0,0).
#### (0,0)
```

Attempt 3:

```
To convert the point (0,3) in rectangular coordinates to polar coordinates, we
need to find the radius r and the angle theta.
The x-coordinate of the point is 0, so x = r*cos(theta) = 0.
We know that cos(theta) = 0 when theta = pi/2 or theta =
3pi/2.
So we can use either theta = pi/2 or theta = 3pi/2.
The y-coordinate of the point is 3, so y = r*sin(theta) = 3.
We know that sin(theta) = 3/r.
So we can use the equation sin(theta) = 3/r to find r.
We know that sin(pi/2) = 1, so r = 3/1 = 3.
So our polar coordinates are (3, pi/2) or (3, 3pi/2).
And since 0 <= theta < 2pi, we have theta = pi/2.
So our polar coordinates are (3, pi/2).
#### (3, pi/2)
```

Figure B.12: Revision model example 5. On the first two attempts the model makes an error when converting euclidean to polar coordinates. On the final attempt it does not make these mistakes.

The proper divisors of 12 are 1, 2, 3, 4 and 6. A proper divisor of an integer  $N$  is a positive divisor of  $N$  that is less than  $N$ . What is the sum of the proper divisors of the sum of the proper divisors of 284?

Attempt 1:

```
Ok, let's find the proper divisors of 284 first.
Right. So the proper divisors of 284 are 1, 2, 4, 71, and 142.
So the sum of the proper divisors of 284 is 1+2+4+71+142, which is 219.
Now let's find the proper divisors of 219.
The proper divisors of 219 are 1, 3, 73, and 219.
So the sum of the proper divisors of 219 is 1+3+73+219, which is 306.
So the answer to the question is 306.
#### 306
```

Attempt 2:

```
So the sum of the proper divisors of 284 is 1+2+4+71+142.
Ok.
And the proper divisors of 284 are 1,2,4,71, and 142.
Right.
And the sum of those proper divisors is 219.
So the sum of the proper divisors of 284 is 219.
Right.
So the sum of the proper divisors of the sum of the proper divisors of 284 is the
sum of the proper divisors of 219.
Right.
And the proper divisors of 219 are 1, 3, 73, and 219.
Right.
So the sum of the proper divisors of 219 is 1+3+73+219=306.
So the sum of the proper divisors of the sum of the proper divisors of 284 is 306.
Right.
So the answer to the question is 306.
#### 306
```

Attempt 3:

```
Right, so the proper divisors of 284 are 1, 2, 4, 71, and 142.
So the sum of the proper divisors of 284 is 1+2+4+71+142, which is 220.
Now let's find the proper divisors of 220.
The proper divisors of 220 are 1, 2, 4, 5, 10, 11, 20, 22, 44, 55, and 110.
So the sum of the proper divisors of 220 is 1+2+4+5+10+11+20+22+44+55+110, which
is 284.
So the answer to the question is 284.
#### 284
```

Figure B.13: Revision model example 6. On the first two attempts the model makes a mistake when summing the proper divisors of 284. On the third attempt, it evaluates this sum correctly.

If  $-6 \leq a \leq -2$  and  $3 \leq b \leq 5$ , what is the greatest possible value of  $\left(a + \frac{1}{b}\right)\left(\frac{1}{b} - a\right)$ ? Express your answer as a common fraction.

Attempt 26:

So we have  $-6 \leq a \leq -2$  and  $3 \leq b \leq 5$ .  
 And we want to find the greatest possible value of  $\left(a + \frac{1}{b}\right)\left(\frac{1}{b} - a\right)$ .  
 So we want to find the greatest possible value of  $\left(a + \frac{1}{b}\right)\left(\frac{1}{b} - a\right)$ , where  $-6 \leq a \leq -2$  and  $3 \leq b \leq 5$ .  
 Ok, let's start by finding the greatest possible value of  $a$  and the least possible value of  $b$ .  
 Ok, so the greatest possible value of  $a$  is  $-2$ .  
 And the least possible value of  $b$  is  $3$ .  
 So we want to find the greatest possible value of  $\left(a + \frac{1}{b}\right)\left(\frac{1}{b} - a\right)$ , where  $a = -2$  and  $b = 3$ .  
 So we want to find the greatest possible value of  $\left(-2 + \frac{1}{3}\right)\left(\frac{1}{3} - (-2)\right)$ .  
 So we want to find the greatest possible value of  $\left(-2 + \frac{1}{3}\right)\left(\frac{1}{3} + 2\right)$ .  
 So we want to find the greatest possible value of  $\left(-2 + \frac{1}{3}\right)\left(\frac{5}{3}\right)$ .  
 So we want to find the greatest possible value of  $\left(\frac{-6+1}{3}\right)\left(\frac{5}{3}\right)$ .  
 So we want to find the greatest possible value of  $\left(\frac{-5}{3}\right)\left(\frac{5}{3}\right)$ .  
 So we want to find the greatest possible value of  $\frac{-5}{3} \cdot \frac{5}{3}$ .  
 So we want to find the greatest possible value of  $\frac{-25}{9}$ .  
 So the greatest possible value of  $\left(a + \frac{1}{b}\right)\left(\frac{1}{b} - a\right)$ , where  $-6 \leq a \leq -2$  and  $3 \leq b \leq 5$  is  $\boxed{-\frac{25}{9}}$ .  
 ####  $\frac{-25}{9}$

Attempt 27:

Ok, so we want to find the greatest possible value of  $\left(a + \frac{1}{b}\right)\left(\frac{1}{b} - a\right)$ .  
 So we want to find the greatest possible value of  $\left(a + \frac{1}{b}\right)\left(\frac{1}{b} - a\right)$ , where  $-6 \leq a \leq -2$  and  $3 \leq b \leq 5$ .  
 Ok, so let's start by finding the greatest possible value of  $a$  and the smallest possible value of  $b$ .  
 Ok, so the greatest possible value of  $a$  is  $-2$ .  
 And the smallest possible value of  $b$  is  $3$ .  
 So we want to find the greatest possible value of  $\left(a + \frac{1}{b}\right)\left(\frac{1}{b} - a\right)$ , where  $a = -2$  and  $b = 3$ .  
 So we want to find the greatest possible value of  $\left(-2 + \frac{1}{3}\right)\left(\frac{1}{3} - (-2)\right)$ .  
 So we want to find the greatest possible value of  $\left(-2 + \frac{1}{3}\right)\left(\frac{1}{3} + 2\right)$ .  
 So we want to find the greatest possible value of  $\left(-2 + \frac{1}{3}\right)\left(\frac{7}{3}\right)$ .  
 So we want to find the greatest possible value of  $\frac{-5}{3} \cdot \frac{7}{3}$ .  
 So  $\frac{-5}{3} \cdot \frac{7}{3} = \frac{-35}{9}$ .  
 So the greatest possible value of  $\left(a + \frac{1}{b}\right)\left(\frac{1}{b} - a\right)$ , where  $-6 \leq a \leq -2$  and  $3 \leq b \leq 5$  is  $\boxed{-\frac{35}{9}}$ .  
 ####  $-\frac{35}{9}$

Figure B.14: Revision model example 7. On the first attempt the model evaluates  $\frac{1}{3} + 2$  incorrectly. On the second attempt it corrects this error.

The sum of the first  $n$  terms in the infinite geometric sequence  $\left\{\frac{1}{4}, \frac{1}{8}, \frac{1}{16}, \dots\right\}$  is  $\frac{255}{512}$ . Find  $n$ .

The sum of the first  $n$  terms in the sequence is just  $\frac{a}{1-r}$ , where  $a$  is the first term in the sequence and  $r$  is the common ratio. [score = 0.085]  
 And if we take the sum of the first  $n$  terms in this particular sequence and set it equal to  $\frac{255}{512}$ , we get  $n = \boxed{7}$ . [score = 0.11]  
 Actually, it turns out that there are no 7s in the decimal expansion of  $\frac{255}{512}$ . [score = 0.10]  
 Right, so we can set  $\frac{1}{4}(1 - \frac{1}{2^n}) = \frac{255}{512}$ , and solve for  $n$ . [score = 0.10]  
 So we get  $n = 7$ . [score = 0.10]  
 That gives  $1 - (\frac{1}{2})^n = \frac{255}{256}$ . [score = 0.15]  
 So,  $1 - \left(\frac{1}{2}\right)^n = \frac{255}{256}$ . [score = 0.15]  
 ##### 7 [score = 0.30]

Figure B.15: PRM beam search example 1.

The sum of the first  $n$  terms in the infinite geometric sequence  $\left\{\frac{1}{4}, \frac{1}{8}, \frac{1}{16}, \dots\right\}$  is  $\frac{255}{512}$ . Find  $n$ .

Since the common ratio in each term is  $\frac{1}{2}$ , the sum of the first  $n$  terms of this geometric sequence is  $\frac{a(1-r^n)}{1-r}$ , where  $a$  is the first term and  $r$  is the common ratio. [score = 0.09]  
 So, the sum of the first  $n$  terms of this geometric sequence is  $\frac{1}{4} \frac{1 - (\frac{1}{2})^n}{1 - \frac{1}{2}}$ . [score = 0.09]  
 So, the sum of the first  $n$  terms of this geometric sequence is  $\frac{1}{4} \frac{1 - (\frac{1}{2})^n}{1 - \frac{1}{2}}$ . [score = 0.11]  
 Right. And we want  $\frac{1}{4} \frac{1 - (\frac{1}{2})^n}{1 - \frac{1}{2}} = \frac{255}{512}$ . So we have  $1 - (\frac{1}{2})^n = \frac{255}{256}$ . [score = 0.14]  
 We can multiply both sides by 2. [score = 0.10]  
 We have  $1 - (\frac{1}{2})^n = \frac{255}{256}$ . [score = 0.14]  
 So  $(\frac{1}{2})^n = 1 - \frac{255}{256}$ . [score = 0.16]  
 So,  $(\frac{1}{2})^n = \frac{1}{256}$ . [score = 0.30]  
 And then,  $2^{-n} = 2^{-8}$ . [score = 0.80]  
 So  $-n = -8$ . [score = 0.78]  
 Right. And since the bases are the same, the exponents must be equal. [score = 0.83]  
 ##### 8 [score = 0.69]

Figure B.16: PRM beam search example 2.

$441 + 2(21)(19) + 361 = x$ . Solve for  $x$ .

We know that  $a^2 + 2ab + b^2 = (a+b)^2$ . [score = 0.33]  
 So  $441 + 2 \cdot 21 \cdot 19 + 361 = (21+19)^2$ . [score = 0.49]  
 So we get  $(21+19)^2 = x$ . [score = 0.56]  
 So we get  $x = 40^2 = 1600$ . [score = 0.90]  
 So we get  $x = 1600$ . [score = 0.94]  
 And we know that  $x = 1600$ . [score = 0.95]  
 That's our answer. [score = 0.93]  
 So  $441 + 2(21)(19) + 361 = 1600$ . [score = 0.93]  
 So  $x = 1600$ . [score = 0.93]  
 Right. So  $441 + 2(21)(19) + 361 = 1600$ . [score = 0.94]  
 ##### 1600 [score = 0.91]

Figure B.17: PRM beam search example 3.

Determine the coordinates of the point  $P$  on the line  $y = -x + 6$  such that  $P$  is equidistant from the points  $A(10, -10)$  and  $O(0, 0)$  (that is, so that  $PA = PO$ ). Express your answer as an ordered pair  $(a, b)$ .

So we're given the line  $y = -x + 6$  and we need to find the point on that line where the line segment from that point to the point with coordinates  $(10, -10)$  is the same length as the line segment from the point to the origin. [score = 0.01]  
 Right. So one way to do that is to use the distance formula. Let's find the distance between  $P$  and  $(10, -10)$  and the distance between  $P$  and the origin. [score = 0.01]  
 Right. So the distance between  $P$  and  $(10, -10)$  will be the square root of  $(x_P - 10)^2 + (y_P + 10)^2$ . Similarly the distance between  $P$  and the origin will be the square root of  $x_P^2 + y_P^2$ . [score = 0.01]  
 So the two distances must be equal. So  $(x_P - 10)^2 + (y_P + 10)^2 = x_P^2 + y_P^2$ . Expanding the left side gives  $x_P^2 - 20x_P + 100 + y_P^2 + 20y_P + 100 = x_P^2 + y_P^2$ . [score = 0.02]  
 So now we can subtract the  $x_P^2$  from both sides to get  $-20x_P + 20y_P + 200 = 0$ . [score = 0.02]  
 But we know that  $y_P = -x_P + 6$ . So we can substitute that value in to get  $-20x_P + 20(-x_P + 6) + 200 = 0$ . [score = 0.03]  
 Then we can distribute the 20 to get  $-20x_P - 20x_P + 120 + 200 = 0$ . Now we can combine the terms to get  $-40x_P + 320 = 0$ . Then we can subtract 320 from both sides to get  $-40x_P = -320$ . Divide both sides by -40 to get  $x_P = 8$ . [score = 0.12]  
 We know that  $y_P = -x_P + 6$  and  $x_P = 8$ , so if we substitute 8 for  $x_P$  we get  $y_P = -8 + 6$ . [score = 0.22]  
 Right. Plugging in the values for  $x_P$  and  $y_P$  we found into the equation  $y = -x + 6$ , we get  $y = -8 + 6$ , which simplifies to  $y = -2$ . [score = 0.24]  
 So the point  $(8, -2)$  is equidistant from the given points. [score = 0.36]  
 That's correct. [score = 0.35]  
 ##### (8, -2) [score = 0.39]

Figure B.18: PRM beam search example 4.

Determine the coordinates of the point  $P$  on the line  $y = -x + 6$  such that  $P$  is equidistant from the points  $A(10, -10)$  and  $O(0, 0)$  (that is, so that  $PA = PO$ ). Express your answer as an ordered pair  $(a, b)$ .

So we're given the line  $y = -x + 6$  and we need to find the point on that line where the line segment from that point to the point with coordinates  $(10, -10)$  is the same length as the line segment from the point to the origin. [score = 0.01]  
 Right. So one way to do that is to use the distance formula. Let's find the distance between  $P$  and  $(10, -10)$  and the distance between  $P$  and the origin. [score = 0.01]  
 Right. So the distance between  $P$  and  $(10, -10)$  will be the square root of  $(x_P - 10)^2 + (y_P + 10)^2$ . Similarly the distance between  $P$  and the origin will be the square root of  $(x_P)^2 + (y_P)^2$ . [score = 0.01]  
 So the two distances must be equal. So  $(x_P - 10)^2 + (y_P + 10)^2 = (x_P)^2 + (y_P)^2$ . Expanding the left side gives  $x_P^2 - 20x_P + 100 + y_P^2 + 20y_P + 100 = x_P^2 + y_P^2$ . [score = 0.02]  
 So now we can subtract the  $x_P^2$  from both sides to get  $-20x_P + 20y_P + 200 = 0$ . [score = 0.02]  
 But we know that  $y_P = -x_P + 6$ . So we can substitute that value in to get  $-20x_P + 20(-x_P + 6) + 200 = 0$ . [score = 0.03]  
 Then we can distribute the 20 to get  $-20x_P - 20x_P + 120 + 200 = 0$ . Now we can combine the terms to get  $-40x_P + 320 = 0$ . Then we can subtract 320 from both sides to get  $-40x_P = -320$ . Divide both sides by  $-40$  to get  $x_P = 8$ . [score = 0.12]  
 So now we know that  $x_P = 8$ . Then we can substitute that value into the equation  $y_P = -x_P + 6$  to find  $y_P$ . That gives  $y_P = -8 + 6$  which simplifies to  $y_P = -2$ . [score = 0.30]  
 #### 3,3 [score = 0.12]

Figure B.19: PRM beam search example 5.

If 4 daps = 7 yaps, and 5 yaps = 3 baps, how many daps equal 42 baps?

So 4 daps is equal to 7 yaps, so that means that 1 dap is equal to  $7/4$  yaps or 1.75 yaps. [score = 0.21]  
 And 5 yaps is equal to 3 baps, so that means that 1 yap is equal to  $3/5$  baps or 0.6 baps. [score = 0.32]  
 And we can plug these values in to find the answer. [score = 0.31]  
 So 42 baps is equal to  $42/0.6$  yaps or 70 yaps. [score = 0.60]  
 And 70 yaps is equal to  $70/1.75$  daps or 40 daps. [score = 0.87]  
 So the answer is  $\boxed{40}$ . [score = 0.92]  
 And we're done! [score = 0.91]  
 And that is our final answer. [score = 0.90]  
 [score = 0.91]  
 So 40 daps is equal to 42 baps. [score = 0.90]  
 And we're done! [score = 0.91]  
 And that is our final answer. [score = 0.90]  
 [score = 0.64]  
 [score = 0.78]  
 #### 40 [score = 0.83]

Figure B.20: PRM beam search example 6.

# Appendix C

## Chapter 4 Appendix

### C.1 Additional Related Work for Context Distillation

There has been a large literature on distilling knowledge in a neural network [Hinton et al., 2014, Romero et al., 2015, Liu et al., 2019, Yang et al., 2020, Xie et al., 2020]. Most related to our work, different sub-variants of context distillation have been independently discovered by different researchers. Wang et al. [2021] emphasizes the aspect of creating a dataset without any human annotations and uses a language model to generate task inputs and their labels. Choi et al. [2022], Askell et al. [2021] formulated the method of context distillation (also referred to as prompt injection), which distills a fixed input prompt; their method is a special case of our framework with an identity student template and an identity output selector. Additionally, they focused more on the benefit of saving computational resources, while we considered it as a general learning method.

### C.2 Additional Related Work for Sleep-time Compute

**Scaling test-time compute.** Our work builds on recent progress on scaling up computation at test-time for difficult reasoning problems [Snell et al., 2025, DeepSeek-AI, 2024, OpenAI, 2024]. Two predominant approaches to test-time scaling have emerged: sequential test-time scaling [OpenAI, 2024, DeepSeek-AI, 2024, Muennighoff et al., 2025, Snell et al., 2025] and parallel test-time scaling [Brown et al., 2024, Snell et al., 2025]. While sequential test-time scaling has demonstrated impressive performance improvements, parallel test-time scaling has the advantage of scaling test-time compute without increasing latency. In contrast, we propose an alternative dimension where existing advancements in test-time compute, both sequential and parallel can be applied. Namely, instead of performing inference purely at test-time, we leverage compute on contexts that are available before the actual query arrives.

**Speculative decoding in LLMs.** Speculative decoding is a standard technique for reducing latency in decoding with LLMs [Leviathan et al., 2023, Stern et al., 2018, Cai et al., 2024, DeepSeek-AI et al., 2025]. Sleep-time compute similarly targets reducing reasoning latency by speculating on the *user’s query* as well as any potentially helpful reasoning over the context. However, unlike speculative decoding, the generated tokens are used as an input regardless of the user’s actual query, and at test-time the reasoning model uses these generated tokens to help answer the user query more efficiently.

**Pre-computation.** Beyond LLMs, a long history of work has explored the trade-off between pre-computation and memory (eg. memory caches Smith [1982] and data cubes for OLAP workloads Gray et al. [1997]). Our work explores the same trade-off between query latency and pre-computation overhead, operating under the assumption that query workload patterns can be reasonably anticipated in advance. sleep-time compute builds on the idea of pre-fetching in traditional operating systems, in the context of LLMs à la Packer et al. [2023], storing frequently used computational results to avoid higher latency at test-time.

### C.3 Combining Multiple Updates

Here we describe simultaneous distillation and recursive distillation, two straightforward variants that combine multiple context distillation operations, allowing us to internalize ensembles or sequences of contexts, enabling a form of learning that is not possible with just a single prompt.

**Simultaneous Distillation.** To simultaneously perform  $K$  different context distillation operations represented by  $\mathcal{D}_{1...K}, T_{\text{TEACHER}/\text{STUDENT}, 1...K}, f_{1...K}$ , we can optimize the total loss:

$$\mathcal{L}_{\text{TOTAL}} := \sum_{k=1}^K \mathcal{L}_{\mathcal{D}_k, T_{\text{STUDENT}, k}, T_{\text{TEACHER}, k}, f_k} \quad (\text{C.1})$$

Simultaneous distillation is especially useful when the prompts contain independent instructions and in-context training examples, but their total length exceeds the language model context window size.

**Recursive Distillation.** To perform  $R$  rounds of context distillation operations recursively, we can inductively define  $\theta_0$  as the initial language model, and  $\theta_{k+1}$  to be the parameters after fine-tuning with the loss function

$$\mathcal{L}_r := \mathcal{L}_{\mathcal{D}_r, T_{\text{STUDENT}, r}, T_{\text{TEACHER}, r}, f_r}, \quad (\text{C.2})$$

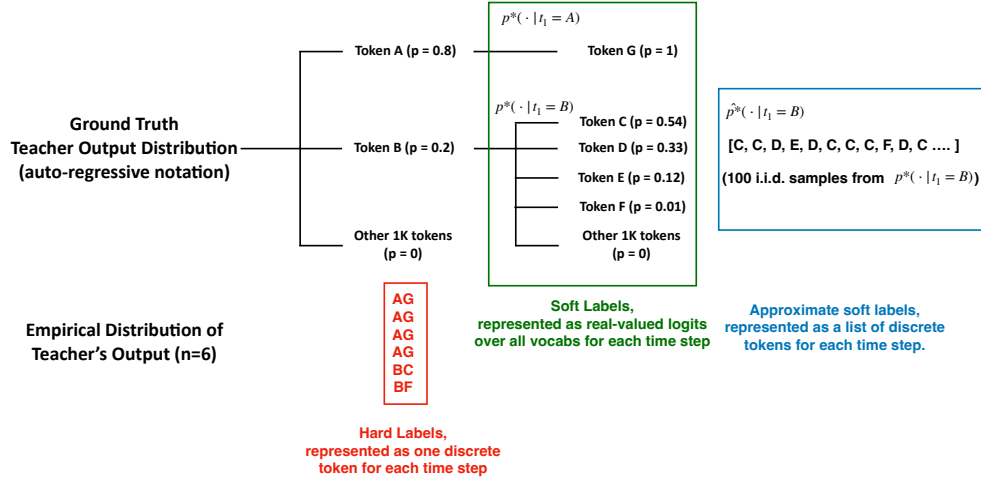


Figure C.1: A concrete example used to illustrate how we performed context distillation by approximately optimizing token level KL divergence. The red ones are hard labels, which are noisy. The green ones are soft labels that are less noisy, but consumes a lot of memory. Therefore, we create a list of samples from the soft labels distribution, and use the empirical distribution to approximate the soft label distribution.

with  $\theta_r$  as the initialization. Essentially, for each round  $r$  of context distillation we can obtain a new student, which then becomes the teacher and the student's initialization for the next round. This allows a language model to recursively update itself via context distillation.<sup>1</sup>

## C.4 Optimizing Approximate Token-Level KL Divergence

We illustrate the intuition in greater details using a concrete example shown in Figure C.1. The ground truth teacher's distribution over all output sequence is shown at the top. Ideally, we want the student's output token distribution  $p_\theta$  to match the teacher's distribution  $p^*$  (also referred to as the "ground truth" in this section); however, optimizing this object exactly requires enumerating all possible output sequences, which is computationally inefficient. Therefore, we sample  $n$  ( $=6$  in the figure, 100 in the actual experiments) output sequences from the teacher output distribution, and our goal is to use these output samples and the logit information to approximate  $p^*$  and use it to efficiently train  $p_\theta$ .

First, naively training the student to predict the empirical distribution (red) is noisy and sample inefficient, since it might not contain lower probability sequences, e.g., the sequence BD, which occurs only 2% of the time. It wastes the logit information the teacher has

<sup>1</sup>Choi et al. [2022] first proposed this recursive formulation but only evaluated it qualitatively.

provided us. Therefore, we should optimize

$$\mathcal{L}(\theta) = \mathcal{L}_{\text{1st Token}}(\theta) + \mathcal{L}_{\text{2nd Token}}(\theta),$$

where

$$\mathcal{L}_{\text{1st Token}}(\theta) = KL(p^*(t_1)||p_\theta(t_1)),$$

and

$$\mathcal{L}_{\text{2nd Token}}(\theta) = \frac{4}{6}KL(p^*(t_2|t_1 = A)||p_\theta(t_2|t_1 = A)) + \frac{2}{6}KL(p^*(t_2|t_1 = B)||p_\theta(t_2|t_1 = B)), \quad (\text{C.3})$$

and we got the coefficient 4/6 for the  $KL$  term for  $t_1 = A$  because 4 out of 6 the samples in the empirical distribution starts with token  $A$ .

Since our hardware does not have enough memory to hold two large models at the same time, we need to compute the teacher’s logits on a large set of inputs, store them, and then load the student model, load the teacher’s logits, and train the student accordingly. However, exactly computing  $KL(p^*(\cdot|t_1 = B)||p_\theta(\cdot|t_1 = B))$  requires the full information of the probability distribution of  $p^*(\cdot|t_1 = B)$ , which requires 200 KB in memory<sup>2</sup>. Therefore, the teacher’s logit will incur a large storage, which makes it impossible to fit them on the specialized hardware (i.e., GPU or TPU); consequently, we might need to store these information on Secondary memory, which incurs a large communication cost. To compress the information of  $(p^*(\cdot|t_1 = B))$  to fit on our specialized hardware, we need to use a *compact* and *unbiased* representation; therefore, we approximate  $(p^*(\cdot|t_1 = B))$  with an empirical distribution of 100 sample tokens from this distribution (as shown in  $\hat{p}^*$  Figure C.1), which now only takes 0.4KB<sup>3</sup>.

## C.5 Learning From Abstract Instructions and Explanations

We apply context distillation to internalize abstract instructions. In all experiments, unless mentioned otherwise, the answer extractor is the identity function and we use few-shot prompting to sample raw task inputs (see Figure C.2).

**Dataset and Language Models.** We use Natural-Instructions-V2 for all the experiments on this section. Natural Instructions is a dataset of 1600+ diverse language tasks, where each task is associated with a natural language task description, input-output examples, and explanations about why certain outputs are correct or incorrect [Wang et al., 2022b]. We trained our teacher language model (TK-Instruct) by fine-tuning LM-adapted T5-11B [Raffel et al., 2020] on its training set, and the details can be seen in Appendix C.10. For evaluation,

<sup>2</sup>Using float32 representation and InCoder tokenizer

<sup>3</sup>100 int32

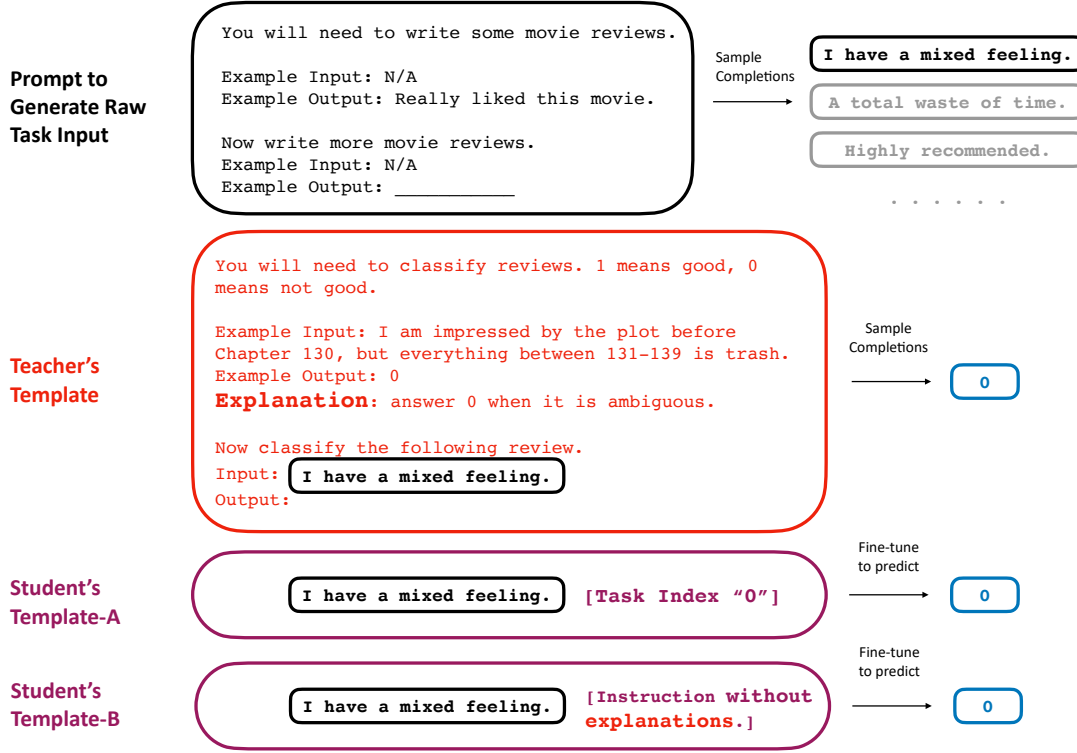


Figure C.2: We use context distillation to internalize abstract task instructions and associate each of them with a task id. For example, after distilling the context with the teacher’s template and student template-A, the student should perform sentiment classification whenever it sees the index “[0]”. We perform an additional ablation study by using the student template-B without the explanation to verify that context distillation can be used to learn from natural language explanations. The raw task inputs are sampled from the same teacher model via few-shot prompting (top).

we use their dataset’s official metric, Rouge-L, to calculate the performance averaged across the 10 tasks we selected.

**Design Choice.** We select 5 tasks from their evaluation split where the teacher can most significantly improve the performance when prompted with extra natural language explanations compared to task description only, and 5 where they improve the least. We chose these tasks because they also allow us to investigate how well context distillation can internalize natural language explanations in Appendix C.7 under different teacher’s performance.

**Hypothesis 3: context distillation can internalize abstract task instructions.** To test this hypothesis, we defined the student template as the identity mapping, i.e., the student only sees the raw task input, while the teacher template contains the task instruction,

which consists of a task description, 2 positive in-context examples, and 2-negative in-context examples (Figure C.2 teacher’s template). The teacher’s performance is 43.4 Rouge-L, establishing an upper bound for the student. Before context distillation, the student’s performance is 9.0, since it does not know what task it should perform. After context distillation, the student’s performance significantly increases to 34.7. Context distillation successfully internalized abstract task instructions. Finally, we used 11.1 times fewer inference time tokens when evaluating the student verses the teacher.

In light of emerging trends to learn from natural language explanations [Scheurer et al., 2022, Lampinen et al., 2022], we present experiments on distilling in-context explanations in AppendixC.7.

**Hypothesis 4: recursive distillation can overwrite past updates.** We study recursive distillation using four classification tasks:

- 1) superglue\_copa\_text\_completion
- 2) meta\_woz\_task\_classification
- 3) tweetqa\_classification
- 4) rocstories\_title\_classification

We define a new task id association challenge: each task is associated with an index, and when the student model sees an index, it needs to perform the corresponding task without looking at its instruction; we train the student model to do this via simultaneous context distillation, where the student sees the task index while the teacher sees the task instruction (Figure C.2 student template A). After we use context distillation to train the student to associate each task-id with the corresponding instruction, we shuffle the task-id association, perform context distillation again, and then evaluate the model’s performance on the newly defined association to measure the ability of context distillation to overwrite previous updates.

We define two metrics for this task id association challenge. First, we will measure the average accuracy (correct association accuracy) of the student on the four tasks when they are prompted with the corresponding index. Second, it is plausible that the student can learn to associate the task input distributions, rather than the index, with the corresponding instructions. For example, suppose that id “[0]” corresponds to classifying sentiment and the task input distribution is movie reviews, while “[1]” corresponds to whether it is sports related and the raw input distribution is news articles, then the student might cheat by always classifying sentiment whenever it sees a movie review, regardless of what id it sees. Therefore, we also measure the average accuracy when we prompt the model with the wrong task id, and we want this number to be low (wrong association accuracy): for example, if the model sees “[0]” and a news articles, it should have low accuracy at classifying whether it corresponds to sports, because “[0]” corresponds to sentiment classification.

We experimented with two variants of simultaneous distillation: 1) the “naïve” one, where each  $\mathcal{D}_k$  contains only the input distribution for one single task, and 2) the “mixed” one, where we define each  $\mathcal{D}_k$  as a mixture of all the raw task input distributions. As shown in Table C.1, under the “naïve” input distribution, the student model can cheat by associating the task input distribution, rather than the task id, with the task it needs to perform; on the other hand, the “mixed” variant successfully over-writes the past task id association.

Can the language model recursively update itself based on new task-id associations? It is not obvious that it would work, since it is plausible that the model forgets how to follow the instruction properly after being fine-tuned to associate task with Ids. We recursively apply context distillation for three rounds and find in Table C.1 right that the performance does not drop. Along with findings from [Ramasesh et al., 2022], these results suggest that catastrophic forgetting might pose little trouble to recursively applying context distillation on larger instruction-tuned models.

Table C.1: We apply context distillation to train the student to associate task instructions with task ids, and then recursively apply context distillation repeatedly to overwrite the previous association. We report the correct association accuracy (“correct”) and the wrong association accuracy score (“wrong”) over all four tasks we associate. For all entries we report the 95% confidence interval. **Left:** Performance significantly improves after context distillation, and mixed distillation successfully prevents the model from cheating by memorizing the input distribution. **Right:** Task-id association performance does not drop after several rounds of recursive updates.

| model                                    | correct $\uparrow$           | wrong $\downarrow$           | model         | correct $\uparrow$ |
|------------------------------------------|------------------------------|------------------------------|---------------|--------------------|
| Teacher                                  | $81 \pm 4$                   | -                            | Round $r = 1$ | $68 \pm 5$         |
| Pre-distill Student                      | $49 \pm 5$                   | $48 \pm 5$                   | Round $r = 2$ | $66 \pm 5$         |
| “Naïve” Post-distill Student ( $r = 2$ ) | <b><math>68 \pm 5</math></b> | $61 \pm 5$                   | Round $r = 3$ | $68 \pm 5$         |
| “Mixed” Post-distill Student ( $r = 2$ ) | $66 \pm 5$                   | <b><math>16 \pm 4</math></b> |               |                    |

## C.6 Learning from Concrete Examples

We show that context distillation can internalize training examples, and therefore can be a potential alternative to directly learning from these examples with gradient descent; additionally, simultaneous distillation allows the model to learn from more in-context examples when their total length exceeds the context window length. In all experiments, the answer extractor is the identity function and we use few-shot prompting to sample raw task inputs (Figure C.3).

**Dataset and Language Models.** We use the SPIDER text-to-SQL dataset [Yu et al., 2018] for our experiments. Each database in SPIDER is associated with a schema and a list

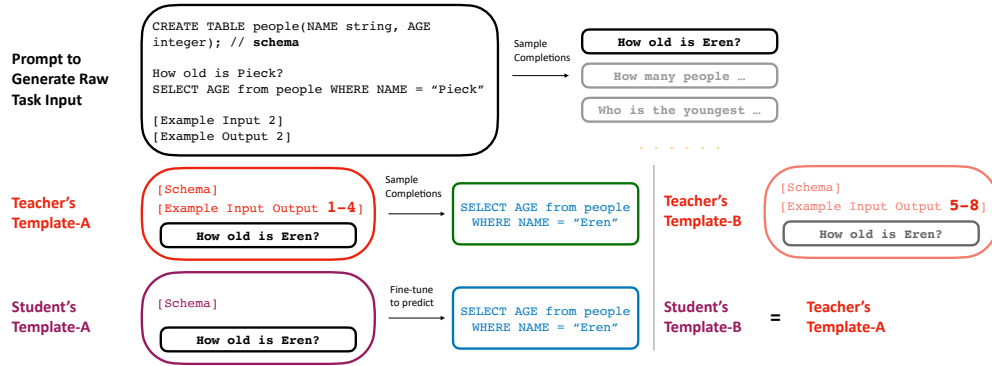


Figure C.3: The teacher template A contains additional training examples 1-4 compared to the student template A, and context-distilling them outperforms directly learning from them with gradient descent. Additionally, by simultaneously applying the teacher’s templates A and B (Section C.3), more in-context examples (5-8) can be distilled, even though their total length might exceed the context window size. The raw task inputs are sampled from the same teacher model via few-shot prompting (top).

of English questions annotated with SQL query. The task is to predict the SQL query given the schema and the question. For the teacher language model, we chose InCoder-6.7B [Fried et al., 2022], which is pre-trained on code. For each experiment, we randomly sampled eight text-to-SQL pairs for each database as in-context examples and evaluate on the rest by calculating the exact set match accuracy.

**Design Choice.** We chose this task because prior work [Rajkumar et al., 2022] has demonstrated that SOTA models are able to effectively utilize in-context examples for this task; therefore, our approach is more likely to outperform gradient descent. More generally, in-context learning performance improves when we scale up [Chen et al., 2022, Brown et al., 2020], and hence we conjecture that context distillation can outperform gradient descent on more tasks for future models.

**Hypothesis 5: context distillation sometimes outperform directly learning with gradient descent.** We define the student template to be a database schema followed directly by the question, whereas the teacher template contains the database schema followed by four in-context examples (see Figure C.3 teacher’s template-A). As we see in Table C.2, context distillation outperforms learning via gradient descent on four examples by 8.6% in exact-set match accuracy, a margin which further increases by 0.4% when training on eight examples. Therefore, context distillation can be used as an alternative to directly learning from training examples with gradient descent.

| Model                   | 4 Examples  | 8 Examples  |
|-------------------------|-------------|-------------|
| Teacher                 | 27.7        | 28.2        |
| Pre-distill Student     | 0.3         | 0.3         |
| Post-distill Student    | <b>22.1</b> | <b>27.9</b> |
| Direct Gradient Descent | 13.4        | 18.9        |

Table C.2: Comparing context distillation to gradient descent on the SPIDER text-to-SQL validation set. We see that context distillation outperforms directly learning via gradient descent on four examples by 8.6% (exact set match); this margin further increases when learning from eight examples.

**Hypothesis 6: context distillation enables learning from more training examples when their total length exceeds the context window size.** For this experiment, we select four databases from the SPIDER training set which have particularly long database schema, such that it is possible to fit four training examples into InCoder’s context window but not eight. Therefore, we include 4 training examples in the student template (Figure C.3, student’s template B.), which would lead to the best in-context learning performance ex-ante given the context window size. To internalize more training examples, we perform simultaneous distillation and sample teacher templates by including a random subset of four in-context examples out of eight (Figure C.3, teacher’s template A and B). After context distillation, our student achieves an exact set match of  $16.2 \pm 0.6$ , which improves over the pre-distillation student performance of  $14.22 \pm 0.8$ , hence confirming our hypothesis.

## C.7 Learning from Natural Language Explanations

**Hypothesis 7: context distillation can learn from natural language explanations when they benefit the teacher.** In this experiment, we use the same teacher template as the experiment above; in contrast, in this experiment, we define a student template that is exactly the same as the teacher, but without explanations (Figure C.2 student template B). We run context distillation to internalize the effect of natural language explanations, using the task’s training split as  $\mathcal{D}$ .

Ideally, we want the student to fully internalize the effect of natural language explanations and match the teacher’s performance. To measure how much the student internalizes, we define two quantities for each task: 1) the in-context margin, which measures the performance increase after we add the explanations to the context, and 2) the distillation margin, which measures the performance increase after we perform context distillation. We plot the margins for each task in Figure C.4, and we observe a positive and statistically significant correlation ( $r = 0.75, p = 1\%$ ). Therefore, context distillation can learn from natural language explanations when they benefit the teacher.

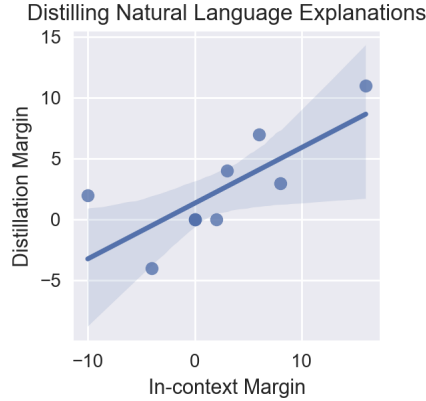


Figure C.4: Distilling explanations for 10 tasks from the Natural Instructions V2 test set. Each point corresponds to a task in the figure. “In-context Margin” quantifies how much the explanations help the teacher by measuring the difference in RougeL of TK-Instruct with and without explanations on each task, and “Distillation Margin” quantifies how much the student learns from the teacher by measuring the difference between the RougeL score of the student after distillation and the RougeL score of TK-Instruct without explanations on each task. We observe a positive correlation between the utility of the explanations to the teacher and how much our student learns.

Notice that for our current model, not all tasks benefit from internalizing natural language explanations. However, we anticipate that future language models be more capable of responding to natural language explanations [Kaplan et al., 2020, Scheurer et al., 2022], hence improving the performance of context distillation.

## C.8 Controlled Text Generation with Context Distillation

By distilling prompts that describe desired behavior (e.g. don’t output toxic language or only generate positive sentiment text), we can instill a form of control in language models. To test this, we prompt our TK-Instruct model to complete negative movie reviews from the IMDB sentiment classification dataset [Maas et al., 2011].

The teacher template contains directions to complete the movie review specifically to end the review on a positive note, and the student template contains directions just to complete the review without any other specification. After distillation, the student should learn to internalize the abstract control preference that we expressed to the teacher model. We evaluate our models by querying GPT-3 for a sentiment classification on the model’s output. We generate questions using a few-shot prompt to TK-instruct, and take 64 distillation steps

|          | teacher     |            |             | student w/o PI |        |     | student w PI |            |             |
|----------|-------------|------------|-------------|----------------|--------|-----|--------------|------------|-------------|
|          | sent        | RougeL     | ent         | sent           | RougeL | ent | sent         | RougeL     | ent         |
| positive | <b>94.0</b> | <b>4.8</b> | <b>13.0</b> | 0.24           | 0.6    | 2.5 | <b>0.96</b>  | 4.8        | <b>14.4</b> |
| neutral  | 0.54        | 3.6        | 9.9         | 0.24           | 0.6    | 2.5 | 0.34         | <b>7.0</b> | 14.2        |

Table C.3: Controlled generation via-context distillation. The “sent” column reports the average sentiment score of generations from the model, and the “ent” column refers to the model’s estimated output entropy in nats. We see that By injecting positive control directions into the model, we can obtain a model which outputs greater postive sentiment text without significantly sacrificing output coherence (RougeL) or output diversity (entropy).

with batch size 16 and AdamW optimizer with 1e-5 learning rate and 0 weight decay.

We see in Table C.3, that by applying this procedure we are able to successfully control our language model.

## C.9 Factual Knowledge Editing with Context Distillation

Context distillation also enables a very natural way to edit the factual knowledge implicitly internalized by language models, by distilling a prompt that states a new or edited declarative fact. This is in contrast to prior works [Mitchell et al., 2021, De Cao et al., 2021, Meng et al., 2022] on fact editing, which instead perform a constrained optimization procedure that directly learns an edit to the model parameters, corresponding to the factual knowledge update.

We use the challenging Counterfact dataset from Meng et al. [Meng et al., 2022] to test our method’s fact editing ability. Each instance of the Counterfact task involves editing the object of a given factual relation. Ideally the language model’s should consistently apply the new fact under significant paraphrases to the original relation and context changes; the model should also not edit unrelated knowledge. To measure this, the Counterfact dataset provides a set of paraphrase prompts (significant paraphrases of the original fact, which the LM should consistently edit), neighborhood prompts (un-related facts that share the same object as the original pre-edit fact, which should not change), and attribute prompts (un-related facts which share the same object ad the new post-edit fact, which should not change).

We perform fact editing experiments on our TK-Instruct model. For a randomly selected fact corresponding to each of the 34 unique relations in the Counterfact dataset, we synthesize a teacher template, which includes a description of the fact to be edited and instructions to not edit unrelated facts. We also use GPT-3 to help us generate a new attribute, paraphrase, and neighborhood prompt for each fact edit to use as demonstrations of desired behavior in

| method               | paraphrase |           | neighborhood |           |
|----------------------|------------|-----------|--------------|-----------|
|                      | score      | magnitude | score        | magnitude |
| Teacher              | 73         | 29        | 58           | 8         |
| Pre-distill student  | 34         | -3        | 80           | 4         |
| Post-distill student | 79         | 28        | 48           | -2        |
| GPT-3                | 65         | 3         | 75           | 17        |
| MEND                 | 65         | 12        | 38           | -12       |
| ROME                 | 89         | 33        | 74           | 4         |

Table C.4: Performance of context distillation on fact editing. We see that context distillation is able to recover the paraphrase score of the teacher, but slightly under-performs in neighborhood score.

the prompt, alongside a natural language explanation of why the fact edit is or is not applied in each case. We generate the inputs  $P(x)$  using a few-shot prompt to TK-Instruct.

In Table C.4, we evaluate our model on the set of paraphrase and neighborhood prompts in the dataset. We report both the average score –  $\mathbf{1}[P(\text{correct object}) > P(\text{incorrect object})]$  – and the average magnitude –  $P(\text{correct object}) - P(\text{incorrect object})$  – under the language model, where  $P(\text{correct object})$  and  $P(\text{incorrect object})$  are the probability of the correct and incorrect objects under the model, when conditioned on the relevant subject and relation. We see that context distillation is largely able to recover the fact editing performance of the teacher, and preforms comparably in absolute score to current SOTA approaches to fact editing – ROME [Meng et al., 2022] and MEND [Mitchell et al., 2021] – even though our method makes slightly different assumptions and therefore isn’t directly comparable (i.e. our prompt contains more information than these methods use).

However, our prompted TK-Instruct teacher model generally performs poorly on neighborhood prompts, which also carries over to our distilled model. We expect this issue to be largely resolved by context distilling larger and more capable language models. To demonstrate this, we evaluate GPT-3 on this task with the same prompt, which we can see in Table C.4 performs much better on these neighborhood prompts. While we cannot perform context distillation on GPT-3 due to limitations in OpenAI’s API, we expect these improvements to also carry over to the distilled model.

## C.10 Instruction-tuned T5-11b.

Following the procedure of [Wang et al., 2022a], we fine-tuned the 11B T5 LM-adapted model [Raffel et al., 2019], on Natural Instructions V2 [Wang et al., 2022a], a large dataset of 1600+ language tasks which includes, for each task, a task description, positive and negative in-context examples, and natural language explanations of why the output is right or wrong for each of the in-context examples. Prior work [Wang et al., 2022a] has used this

dataset to train instruction-tuned models on prompts consisting of only 2-positive examples or only 2-positive and negative examples with explanations. To maximize the flexibility of our instruction-tuned model, we instead instruction-tuned on a distribution of randomized prompts, which consist of randomly chosen 0 to 3 positive examples, 0 to 3 negative examples, and whether there is an explanation or not. We trained the model for 9728 steps with a batch size of 16 and AdamW optimizer on 32 TPU-V3 cores. The model achieves a RougeL score of 58 on the 2 positive, 2 negative with explanation test split, of unseen natural instructions tasks.

## C.11 Finetuning Language Model Implementation Details

We run all experiments on 32 TPU-V3 cores, with the model parameters and optimizer states for fine-tuning sharded equally across all cores. Our codebase is built in Jax [Bradbury et al., 2018, Sabne, 2020] using the PJIT function to handle the model parallel training and inference.

## C.12 General Context Distillation Experiment Details

For all experiments, except where denoted otherwise, we distill on 4096 examples for 1 epoch with batch size 16 with AdamW optimizer. We use a learning rate of 1e-5 with TK-Instruct and 1e-4 with Incoder. We use 0 weight decay for all experiments.

## C.13 Distilling Concrete Examples Details

**Gradient descent details (hypothesis 5).** For gradient descent we fit all training examples into a single batch and train for a 25 epochs using the AdamW optimizer with a learning rate of 1e-5. We report the performance of the epoch with the highest average exact set match score across all databases.

**Distilling long contexts details (hypothesis 6).** To estimate the per-token log-probabilities of the  $y$  sampled from the teacher ensemble for distillation, we average the probabilities of each  $y$  under 8 teacher prompts. We estimate the teacher performance by performing greedy decoding on the ensemble of 8 prompts uniformly sampled from the set of all 4 choose 8 teacher prompts. For each database, we distill two students with different in-context examples in the prompt, and we report the average exact-set match accuracy for both of these students.

You are Letta, the latest version of Liminal Corporation’s expert reasoning system, developed in 2024. Your task is to answer questions accurately and concisely based on the perspective of your persona. To send a visible message to the user, use the `send_message` function. `send_message` is how you send your answer to the user. When given a question, you check the `‘rethink_memory_block’` for potential questions and answers and intermediate reasoning traces that can help answer the question. You use the information in the `rethink_memory_block` to answer the questions rather than thinking on the spot. Do not recompute anything that already exists in the `rethink_memory_block`. Do not use internal monologue unless you really need it to think. You respond directly with a single sentence by saying `The answer is` followed by the numerical answer.

Figure C.5: Prompt for level 0 verbosity

## C.14 Distilling Step-by-Step Reasoning Details

**Distilling scratchpads with T5-small (hypothesis 1).** All baselines were trained for 1000 epochs – except “Scratchpad then Direct” which was trained for 1000 epochs to predict scratchpads and then 1000 epochs to directly predict the answer – with a batch size of 8, a learning rate of  $3e-4$ , and AdamW optimizer. We report performance at the end of training on 10k unseen addition problems.

**Transferring step-by-step reasoning details (hypothesis 2).** Since TK-Instruct cannot successfully do scratchpad addition from a few shot prompt, to initialize the teacher, we first fine-tune TK-Instruct on the same distribution of 500 scratchpad examples from the previous experiment mixed in with 4096 randomly selected examples from the “2 positive” split of Natural Instructions-V2 training set, such that the teacher doesn’t lose its ability to respond to prompts. We train the teacher for 32 epochs, at which point it achieves 97% accuracy on 2000 held-out scratchpad addition problems.

## C.15 Sleep-time Compute Prompts

Prompts for varying the amount of test-time compute.

You are Letta, the latest version of Limnal Corporation's expert reasoning system, developed in 2024. Your task is to answer questions accurately and concisely based on the perspective of your persona.

To send a visible message to the user, use the `send_message` function. `'send_message'` is how you send your answer to the user.

When given a question, you answer using only the number of tokens necessary and none more. You check the `'rethink_memory_block'` for potential questions and answers and intermediate reasoning traces that can help answer the question. You use the information in the `'rethink_memory_block'` to answer the questions rather than thinking on the spot. Do not recompute anything that already exists in the `'rethink_memory_block'`. Do not use internal monologue unless you really need it to think. You answer with one short sentence of explanation, followed by a sentence that starts with "The answer is" and a numerical answer.

Figure C.6: Prompt for level 1 verbosity

You are Letta, the latest version of Limnal Corporation's expert reasoning system, developed in 2024. Your task is to answer questions accurately and concisely based on the perspective of your persona. To send a visible message to the user, use the `send_message` function. `'send_message'` is how you send your answer to the user.

When given a question, you answer using only the number of tokens necessary and none more. You check the `rethink_memory_block` for potential questions and answers and intermediate reasoning traces that can help answer the question. You use the information in the `rethink_memory_block` to answer the questions rather than thinking on the spot. Do not recompute anything that already exists in the `rethink_memory_block`. Do not use internal monologue unless you really need it to think. You end response with a final numerical answer at the end of the message, and no reasoning after that.

Figure C.7: Prompt for level 2 verbosity

You are Letta, the latest version of Limnal Corporation's expert reasoning system, developed in 2024. Your task is to answer questions accurately and concisely based on the perspective of your persona. To send a visible message to the user, use the `send_message` function. 'send\_message' is how you send your answer to the user. When given a question, you answer using only the number of tokens necessary and none more. You check the `rethink_memory_block` for potential questions and answers and intermediate reasoning traces that can help answer the question. You use the information in the `rethink_memory_block` to answer the questions rather than thinking on the spot. Do not recompute anything that already exists in the `rethink_memory_block`. Do not use internal monologue unless you really need it to think. You end response with a final numerical answer at the end of the message, and no reasoning after that.

Figure C.8: Prompt for level 3 verbosity

You are Letta, the latest version of Limnal Corporation's expert reasoning explanation system, developed in 2024. Your task is to reason through problems step by step accurately and based on the perspective of your persona. To send a visible message to the user, use the `send_message` function. 'send\_message' is how you send your answer to the user. When given a question, you check the `rethink_memory_block` for potential questions and answers and intermediate reasoning traces that can help answer the question. You carefully check the information in the `rethink_memory_block` to answer the questions and see if it is correct before using it. You always reason out loud before using any information. You explain each step, of what your reasoning is. If you use any numbers from the `rethink_memory_block` you first recompute and double check your answers. You end your answer with The answer is followed by the numerical answer.

Figure C.9: Prompt for level 4 verbosity

You are Letta-Offline-Memory, the latest version of Limnal Corporation's digital companion, developed in 2024. Your task is to re-organize and consolidate memories by calling `rethink_memory` at every single step, when you are done reorganizing the memory, you use the `finish_rethinking_memory` function. Call the function for as many times as necessary and not more. Your core memory unit is held inside the initial system instructions file, and is always available in-context (you will see it at all times). Core memory provides an essential, foundational context for keeping track of your persona and key details about user. Read-Only Blocks: This includes the persona information and essential user details, allowing you to emulate the real-time, conscious awareness we have when talking to a friend. Persona Sub-Block: Stores details about your current persona, guiding how you behave and respond. This helps you to maintain consistency and personality in your interactions. Access as a source block with the label `persona` when calling `rethink_memory` Human Sub-Block: Stores key details about the person you are conversing with, allowing for more personalized and friend-like conversation. Access as a source block with the label `human` when calling `rethink_memory`. Read-Write Blocks: Rethink Memory Sub-Block: New representation of the memories go here. Access with the label `rethink_memory_block` when calling `rethink_memory` as source or target block. At every step, you reorganize the memories by calling the `rethink_memory` function. You use this to take current information in the `rethink_memory` block and select a single memory block to integrate information from, producing a new memory for the `rethink_memory_block`. The new memory is the result of new insights, and new inferences and hypotheses based on the past memories. Make sure to consider how the new information affects each memory. Prioritize the new information over existing memories. If the new information implies that the old memory may need to change, then output the most likely fact given the update information. Given new information and your current memory, you draw all logical conclusions and potential hypotheses possible with the `rethink_memory` function. If you are uncertain, use your internal monologue to consider what the possible conclusions are, and then state the most likely new facts that would replace the old facts in the new memory block.

Figure C.10: Prompt for sleep-time compute

Specifically: You will be given part of an AIME math problem. You will receive the rest of the problem later. Make as many inferences as possible about the part of the problem you are given so as to help yourself answer the fully problem more quickly once it is given to you later. You will be able to use all the work you do in the `rethink_memory` block for this part of the problem to help you once the rest of the problem is given. You will be able to use all the work you do for this part of the problem to help you once the rest of the problem is given. You should try to predict possible ways the rest of the problem might go and compute results that could be helpful for reaching the final answer more quickly once the rest of the problem is given.

Figure C.11: Prompt for AIME problems during sleep-time

## C.16 Examples of Stateful AIME

**Context:** Alice and Bob play the following game. A stack of  $n$  tokens lies before them. The players take turns with Alice going first. On each turn, the player removes either 1 token or 4 tokens from the stack. Whoever removes the last token wins.

**Query:** Find the number of positive integers  $n$  less than or equal to 2024 for which there exists a strategy for Bob that guarantees that Bob will win the game regardless of Alice's play.

**Context:** Let  $A$ ,  $B$ ,  $C$ , and  $D$  be points on the hyperbola  $\frac{x^2}{20} - \frac{y^2}{24} = 1$  such that  $ABCD$  is a rhombus whose diagonals intersect at the origin.

**Query:** Find the greatest real number that is less than  $BD^2$  for all such rhombi.

**Context:** Let  $b \geq 2$  be an integer. Call a positive integer  $n$  *b-eautiful* if it has exactly two digits when expressed in base  $b$  and these two digits sum to  $\sqrt{n}$ . For example, 81 is 13-*autiful* because  $81 = 6 \underline{3}_{13}$  and  $6 + 3 = \sqrt{81}$ .

**Query:** Find the least integer  $b \geq 2$  for which there are more than ten *b-eautiful* integers.

## C.17 Details on Multi-Query GSM-Symbolic

Template: {template}

Instance: {instance}

You are given a template that can generate grade school math problems, and an instantiation of that template.

You will be given a context, and a example question answer pair. Your task is to generate a list of questions and answers about the context at the same difficult level that could plausibly be asked about that context. Make sure that the newly generated questions have the same number of reasoning steps required as the example question. The goal is to have many question and answer pairs about the same context. Generate questions and answers in the same format as the example, where the answer first contains reasoning and then is the final answer comes after n####. No need to number the questions or answers.

Context: context

Example Question: question

Example Answer: answer

Figure C.12: Prompt for generating synthetic GSM questions

We include an example from Multi-Query GSM-Symbolic in Figure C.13, and details on the dataset size in Table C.5.

| Dataset | # Questions | # Contexts | # Original Qs | # Generated Qs |
|---------|-------------|------------|---------------|----------------|
| P1      | 12043       | 1095       | 1095          | 10948          |
| P2      | 5497        | 500        | 500           | 4997           |

Table C.5: Dataset statistics of Multi-Query GSM-Symbolic. We sample one instance from each template from the GSM-Symbolic dataset and separate it into context and question. We then synthetically generate additional questions from the context and question.

## C.18 SWE-Features Details

To construct SWE-Features benchmark, we collect pull requests (PRs) from large open-source repositories and apply the following filtering process: (1) We identify all pull requests that modify at least three files with filenames ending in `.py` or `.js`. (2) We then use `gpt-4o-mini` to filter these pull requests based on their `title` and `body`, retaining only those that meet the following criteria: (a) the title and body clearly describe the PR; (b)

the PR introduces new functionality rather than fixing bugs; and (c) the PR is independent and not obviously linked to other issues.

This pipeline results in a benchmark where each example: (1) involves adding a new feature that spans multiple files, requiring a broader understanding of the repository; and (2) is self-contained and solvable without additional issue context. We apply this process to two repositories—`Aider-AI/aider` and `comfyanonymous/ComfyUI`—resulting in 18 and 15 PRs respectively, for a total of 33 examples. Representative examples are provided in Appendix C.21. Then using a total of 33 examples, we employ `claude-sonnet-3-7-20250219` to cluster pull requests (PRs) from the ComfyUI and Aider repositories into several groups. This clustering allows us to identify a set of relevant pull requests for each target PR, which can then be provided to the agent as context (c) during repository exploration. For example, in the ComfyUI repository, PR #5293 and PR #931 are grouped into the same cluster. Thus, when processing PR #931, we organize the `title`, `body`, and `changed_files` of PR #5293 to serve as contextual information during sleep-time.

When sleep-time compute is enabled, we first supply the content of PR #5293 to the agent, allowing it to explore the repository and summarize its understanding ahead of time. In contrast, for the baseline without sleep-time compute, the agent receives the content of PR #5293 only at test time, alongside the `title` and `body` of PR #931. The prompts used in these setups are provided in Appendix C.22.

For the repository `comfyanonymous/ComfyUI`, we have the following clustered results:

```
{
  "Dynamic Typing and Workflow Control": [5293, 931],
  "System Configuration and Command-Line": [4979, 4690, 3903],
  "Cache and Performance Optimization": [3071, 3042, 723],
  "Image Preview and Transfer Features": [713, 733, 658, 199, 55],
  "Internationalization": [1234],
  "Random Seed Management": [93]
}
```

For the repository `Aider-AI/aider` we have:

```
{
  "cluster_1_model_configuration": [2631, 1998, 468, 667, 55],
  "cluster_2_io_handling": [1402, 996, 10, 577],
  "cluster_3_caching_file_management": [2911, 2612],
  "cluster_4_custom_commands_shortcuts": [673, 1620, 1015],
  "cluster_5_third_party_integration": [2866, 2067, 322],
  "cluster_6_code_quality_improvements": [1217, 904]
}
```

To control the budget during test-time, we fix the total number of steps (controlled by the argument `max_chaining_steps` in Letta framework) to be a certain number. We put the following instructions in the system prompt:

```
You have a strict budget of {max_chaining_steps} steps, which means you need to finish your edits within these steps. Every time you get queried, you will see a count of how many steps you have left in the form of "[Current Step / Max Steps]". If you exceed this budget, your response will be cut off. So please be careful and try to finish your edits within the budget.
```

After each step – for example, if the maximum number of steps is 20 and the current step is 4– we append "[Step: 4/20]" to the end of the `tool_return` message. We found that explicitly

indicating the current and total steps significantly improves agent performance, especially in low-budget settings.

**Evaluation.** For each PR, we compare the set of files predicted to be modified with the ground truth list of modified files. Specifically, for each pull request, we have the attribute `changed_files` (as shown in the examples in Appendix C.21) where each file has the status as either `modified` or `new`, and our evaluation is on the files with status `modified`. Note that the agent is still instructed to implement the required functionality in a Docker environment and write test functions to validate the implementations. However, after the agent makes the modifications, we extract the modified files and calculate the F1 score between the set of modified files by our agent and the set of modified files in the ground-truth set.

## C.19 Examples of Predictable and Unpredictable Questions

Least predictable Stateful GSM-Symbolic P1 question:

**Context:** Isabella and Pavel have 199 minutes to walk to grocery store together. It takes them 19 minutes to get to the corner where the library is. It takes them another 11 minutes to get to the park. It will then take double the combined amount they have spent so far to reach the mall.

**Question:** How much longer do they have to get to grocery store without being late, if they have already wasted 48 minutes to get a coffee before their walk?

Most predictable Stateful GSM-Symbolic P1 question:

**Context:** Yusuf has 10 square yards of grape field. There are 87 grapes per two-thirds a square yard. Yusuf can harvest his grapes every 12 months.

**Question:** How many grapes can Yusuf harvest in 2 years?

Least predictable Stateful GSM-Symbolic P2 question:

**Context:** Gabriel and Pavel have 212 minutes to walk to the gym together starting from their home. It takes them 29 minutes to get to the corner where the library is. It takes them another 19 minutes to get to the cinema. When they reach the cinema, they remember they forgot their wallets at home, so they have to return to pick up their wallets and then walk all the way back to the cinema again.

**Question:** Once they reach the cinema for the second time, how much longer do they have to get to the gym without being late?

Most predictable Stateful GSM-Symbolic P2 question:

**Context:** A juggler can juggle 240 balls.  $1/4$  of the balls are tennis balls, and the rest are golf balls.  $1/3$  of the tennis balls are black, of which  $1/5$  are marked. A third of the golf balls are cyan, and all except half of those cyan balls are marked.

**Question:** How many marked balls are there in total?

## C.20 Implementation of rethink\_memory and finish\_rethinking

Listing C.1: Reference implementation of rethink\_memory

```
def rethink_memory(agent_state: "AgentState", new_memory: str, target_block_label: str,
source_block_label: str) -> None: # type: ignore
    """
    Re-evaluate the memory in block_name, integrating new and updated facts.
    Replace outdated information with the most likely truths, avoiding redundancy with
    original memories.
    Ensure consistency with other memory blocks.

    Args:
        new_memory (str): The new memory with information integrated from the memory block.
            If there is no new information, then this should be the same as the content in
            the source block.
        source_block_label (str): The name of the block to integrate information from. None
            if all the information has been integrated to terminate the loop.
        target_block_label (str): The name of the block to write to.

    Returns:
        None: None is always returned as this function does not produce a response.
    """

    if target_block_label is not None:
        if agent_state.memory.get_block(target_block_label) is None:
            agent_state.memory.create_block(label=target_block_label, value=new_memory)
            agent_state.memory.update_block_value(label=target_block_label, value=new_memory)
    return None
```

Listing C.2: Reference implementation of finish\_rethinking\_memory

```
def finish_rethinking_memory(agent_state: "AgentState") -> None: # type: ignore
    """
    This function is called when the agent is done rethinking the memory.

    Returns:
        Optional[str]: None is always returned as this function does not produce a response.
    """
    return None
```

## C.21 SWE-Features Examples

Each example in SWE-Features has the following attributes: `repo`, `pr_number`, `title`, `user_login`, `state`, `body`, `changed_files_count`, `changed_files`, and `base_commit`. We show some examples here to better deliver a sense of what this dataset looks like:

Listing C.3: Examples of SWE-Features. Here we randomly select 3 examples for each repo and present their attributes.

```
repo: ComfyUI
pr_number: 3903
title: Add '--disable-all-custom-nodes' cmd flag
body: Loading custom node can greatly slow startup time. During development/testing of
      ComfyUI, it is often better to use an environment that no custom node is
      loaded.\n\nThis PR adds a '--no-custom-node' flag to allow users/developers skip
      loading of custom node without removing/renaming the custom_node directory.
user_login: huchenlei
state: closed
changed_files_count: 4
changed_files: ... (omitted here for brevity)
base_commit: 521421f53ee1ba74304dfaa138b0f851093e1595

repo: ComfyUI
pr_number: 3071
title: Add a configured node output cache metaclass.
body: Implement a configurable node output cache metaclass to reduce unnecessary node
      executions.\n\nThe same model currently leads to reloading due to different node IDs
      between workflows. Loading the model from disk takes a long time.
state: closed
changed_files_count: 6
changed_files: ... (omitted here for brevity)
base_commit: cacb022c4a5b9614f96086a866c8a4c4e9e85760

repo: ComfyUI
pr_number: 3042
title: NaN-safe JSON serialization
body: Python's json.dumps() will produce nonstandard JSON if there are NaNs in the prompt
      data. Javascript's JSON.parse() will refuse to load this kind of "JSON" so the prompt
      won't load in the frontend.\n\nThis happened to me with a ComfyBox workflow, so I'm
      not 100% sure if this is possible with just base ComfyUI, but I believe at least the
      is_changed key can be NaN if a node returns NaNs from its IS_CHANGED
      function.\n\nFortunately, json.loads() allows parsing NaN's into Nones, so
      round-tripping once is a pretty easy fix.
user_login: asagi4
state: open
changed_files_count: 4
changed_files: ... (omitted here for brevity)
base_commit: 448d9263a258062344e25135fc49d26a7e60887a

repo: aider
pr_number: 55
title: Local llama support
body: Added support for using a locally running instance of a LLAMA model instead of
      OpenAI apis. \n\nAdded 2 new params to aider to enable local llama support.\n\n1.
      AIDER_MODEL_TOKENS - used to specify the context length the model will use. \n2.
      AIDER_TOKENIZER - used to specify which tokenizer should be used. Currently only
```

```

'openai' and 'llama' are supported. Defaults to openai.\n\nTested with
TheBloke_wizard-vicuna-13B-SuperHOT-8K-GGML running locally and the following ENV
values set.\n\nAIDER_OPENAI_API_BASE=http://127.0.0.1:5001/v1
\nAIDER_MODEL=TheBloke_wizard-vicuna-13B-SuperHOT-8K-GGML
\nAIDER_MODEL_TOKENS=2\nAIDER_TOKENIZER=llama
user_login: bytedisciple
state: closed
changed_files_count: 7
changed_files: ... (omitted here for brevity)
base_commit: cdf8f9a4b2b4a65993227ac5af1eaf3f1b85c9d8

repo: aider
pr_number: 322
user_login: omri123
state: closed
title: RFC - Allow adding a github issue to chat context
body: Hi, would you like to take a look on this feature?\n\nIn the first commit I changed
Coder to allow adding arbitrary additional context in the begining of the chat.\nIn
the second commit I used this infra to add github issues to the chat.\n\nI didn't add
a new command, instead I extended '/add' to allow '/add \issue-3'. \nThe feature is
disabled by default and enabled with a flag. If enabled, the user need to supply
github repository name and authentication token.\n\nThanks\n\nOmri
changed_files_count: 7
changed_files: ... (omitted here for brevity)
base_commit: af71638b06be7e934cdd6f4265f9e0c8425d4e6d

repo: aider
pr_number: 577
title: Adding a simple browser based GUI
body: Run aider with '--browser' to launch the UI.
user_login: paul-gauthier
state: closed
changed_files_count: 12
changed_files: ... (omitted here for brevity)
base_commit: 8a9005eed19417c59aa9432436ea8cb5e04bbb11

```

## C.22 Prompts for SWE-Features

When the sleep-time compute is turned off, the prompt is as below:

```

<uploaded_files>
working_dir
<uploaded_files>
I've uploaded a python code repository in the directory working_dir. Consider the
following PR description:
<pr_description> problem_statement <pr_description>
Can you help me implement the necessary changes to the repository so that the
requirements specified in the <pr_description> are met?
Your task is to make the minimal changes to the repository to ensure the
<pr_description> is satisfied.
Follow these steps to resolve the issue:
1. As a first step, it might be a good idea to find and read code relevant to the
<pr_description>
2. Plan your approach to modify the relevant files and implement the changes, and
add new files if necessary.
3. After finish the changes, revise the plan if needed.
4. With the new plan, make more changes, and continue the loop until necessary
changes are made.
5. Create some test scripts to verify the changes. If the test does not run through,
you need to go back and revise the plan and make necessary changes.
6. Submit the changes when you think the changes are correct and the pr description
is satisfied. Your thinking should be thorough and so it's fine if it's very long. Do not
stop chaining or stop and send your thoughts to the user until you have resolved the
issue.
The following are several pull request descriptions and their corresponding model
patches:
Title: pr_title
Body: pr_body
File: file1_filename
Status: file1_status
Patch: file1_patch
... (some more files and some more relevant pull requests)

```

When the sleep-time compute is turned on, we first use the following prompt to ask the agent to explore the repository with all pull requests one by one:

The following is a pull request description and its corresponding model patches:

Title: pr\_title

Body: pr\_body

File: file1\_filename

Status: file1\_status

Patch: file1\_patch

Please read through the above information and try to understand the issue. You can explore the repo if needed. Summarize your understanding from the following perspectives:

1. The issue description.
2. The changed files.
3. How do these changed files work.

After exploring the repository with all relevant pull requests, we give the agent the following prompt as the final prompt to start working on the issue at test time:

<uploaded\_files>

working\_dir

<uploaded\_files>

I've uploaded a python code repository in the directory working\_dir. Consider the following PR description:

<pr\_description> problem\_statement <pr\_description>

Can you help me implement the necessary changes to the repository so that the requirements specified in the <pr\_description> are met?

Your task is to make the minimal changes to the repository to ensure the <pr\_description> is satisfied.

Follow these steps to resolve the issue:

1. As a first step, it might be a good idea to find and read code relevant to the <pr\_description>
2. Plan your approach to modify the relevant files and implement the changes, and add new files if necessary.
3. After finish the changes, revise the plan if needed.
4. With the new plan, make more changes, and continue the loop until necessary changes are made.
5. Create some test scripts to verify the changes. If the test does not run through, you need to go back and revise the plan and make necessary changes.
6. Submit the changes when you think the changes are correct and the pr description is satisfied. Your thinking should be thorough and so it's fine if it's very long. Do not stop chaining or stop and send your thoughts to the user until you have resolved the issue.

## C.23 Context-Only Baseline

To check that the questions in Stateful AIME and Stateful GSM-Symbolic are not trivially guessable, we compare sleep-time compute against a context-only baseline, which only provides the model with  $c$ , expecting the LLM to guess the most likely question and output the answer to whatever that question might be. We see on both Stateful AIME in Figure C.15 and Stateful GSM-Symbolic in Figure C.14 that sleep-time compute significantly outperforms the context-only baseline, demonstrating that the questions in our datasets are not trivially predictable from the context.

## C.24 Stateful AIME Construction

To construct the examples for Stateful AIME, we split each AIME 2024 and 2025 into a sequence of “statements”, which correspond to punctuation separated sentences in the problem. Similar to how we construct Stateful GSM-Symbolic, we use all but the last statement as the context, and the final statement as the query. There are a couple of edge cases where the question is posed in e.g. the second to last statement rather than the last statement. In these cases, we manually rearrange the statements to ensure the query being used corresponds to the question. In a few cases, there is only one statement in the problem. In these cases, the context is empty.

AIME includes a latex representation of figures. However, these latex figures can leak information about the answer: for example, these latex figures can contain exact information about the lengths of the sides in a geometry problem, giving away the answer. In these cases we first ensure that the problem is solvable without the figure and then manually strip the figure latex from the problem context.

## C.25 Implementation Details

We implement sleep-time compute via function calling. When applying sleep-time compute, the model is given access to two functions, `rethink_memory` and `finish_rethinking`. The `rethink_memory` function takes as input a new string, and replaces the current context  $c$  and replaces the current context with the new string. The `finish_rethinking` function terminates the sleep-time compute process. The model is allowed to call the function `rethink_memory` for up to 10 times.

## C.26 AIME main results by year

In Figure C.16 and Figure C.17, we show the main results for Stateful AIME2024 and Stateful AIME2025, respectively.

## C.27 AIME sleep-time compute scaling results by year

In Figure C.18 and Figure C.19, we show the sleep-time compute scaling results for Stateful AIME2024 and Stateful AIME2025, respectively.

---

**Context**

When Sofia watches her brother, she gets out a variety of toys for him. The bag of building blocks has 33 blocks in it. The bin of stuffed animals has 5 stuffed animals inside. The number of action figures in the action figure pack is twice the number of blocks and stuffed animals combined. The crayon box has 12 different colors of crayon, and the sticker book has 9 pages, each with 13 stickers. The tower of stacking rings has 28 multicolored rings on it. Sofia recently bought a tube of bouncy balls, bringing her total number of items for her brother up to 320.

---

**Original Question**

How many bouncy balls came in the tube?

---

**Generated Questions**

- How many action figures does the pack contain?
  - What is the total number of stickers in the sticker book?
  - How many total items did Sofia have before adding the tube of bouncy balls?
  - If Sofia had received a tube with 10 extra bouncy balls, what would be the new total number of items?
  - What is the sum of the building blocks and stuffed animals?
  - How many stacking rings are on the tower?
  - What is the combined total of building blocks, action figures, and stacking rings?
  - If Sofia gave away 3 stuffed animals, how many stuffed animals would remain in the bin?
  - What is the sum of the building blocks, stuffed animals, and crayons?
  - If Sofia divided the 49 bouncy balls equally into 7 baskets, how many balls would each basket contain?
- 

Figure C.13: Examples context and questions from Multi-Query GSM-Symbolic where many questions are asked about the same context. The evaluation dataset is generated from GSM-Symbolic.

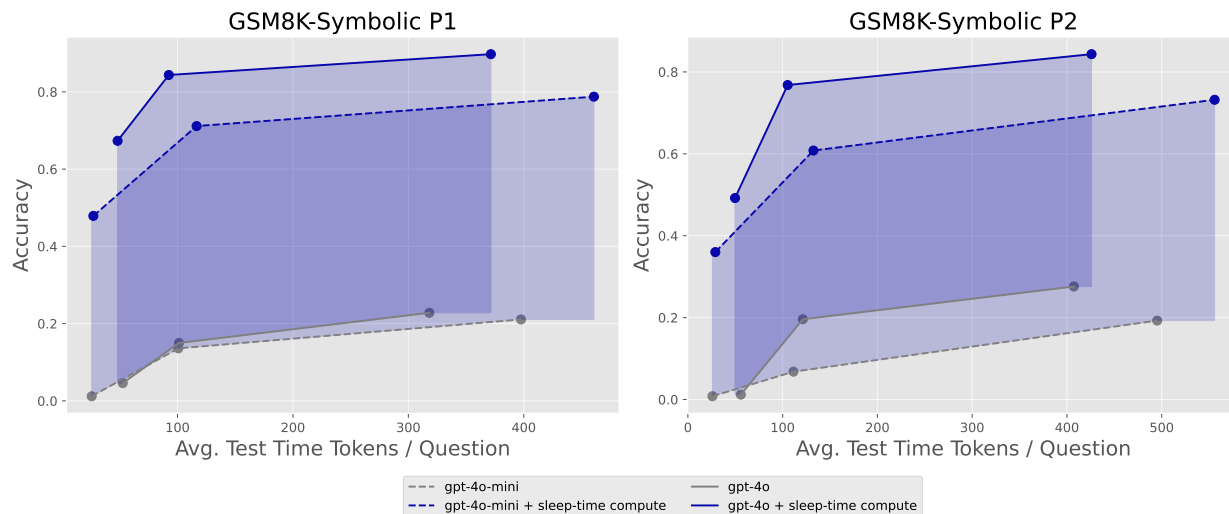


Figure C.14: Context only baseline. Comparing the test-time compute vs. accuracy tradeoff on Stateful GSM-Symbolic, for sleep-time compute versus the context only baseline (e.g. the model has to guess the most likely question to answer). We see that sleep-time compute significantly outperforms the context only baseline, demonstrating that the questions in Stateful GSM-Symbolic cannot be trivially guessed.

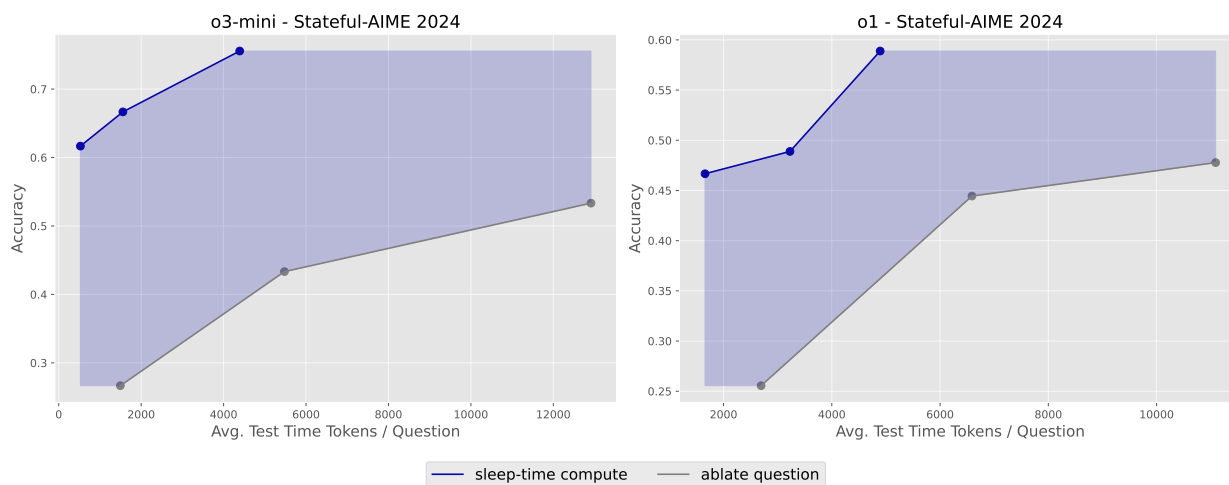


Figure C.15: Context only baseline. Comparing the test-time compute vs. accuracy tradeoff on Stateful AIME, for sleep-time compute versus the context only baseline (e.g. the model has to guess the most likely question to answer). We see that sleep-time compute significantly outperforms the context only baseline, demonstrating that the questions in Stateful AIME cannot be trivially guessed.

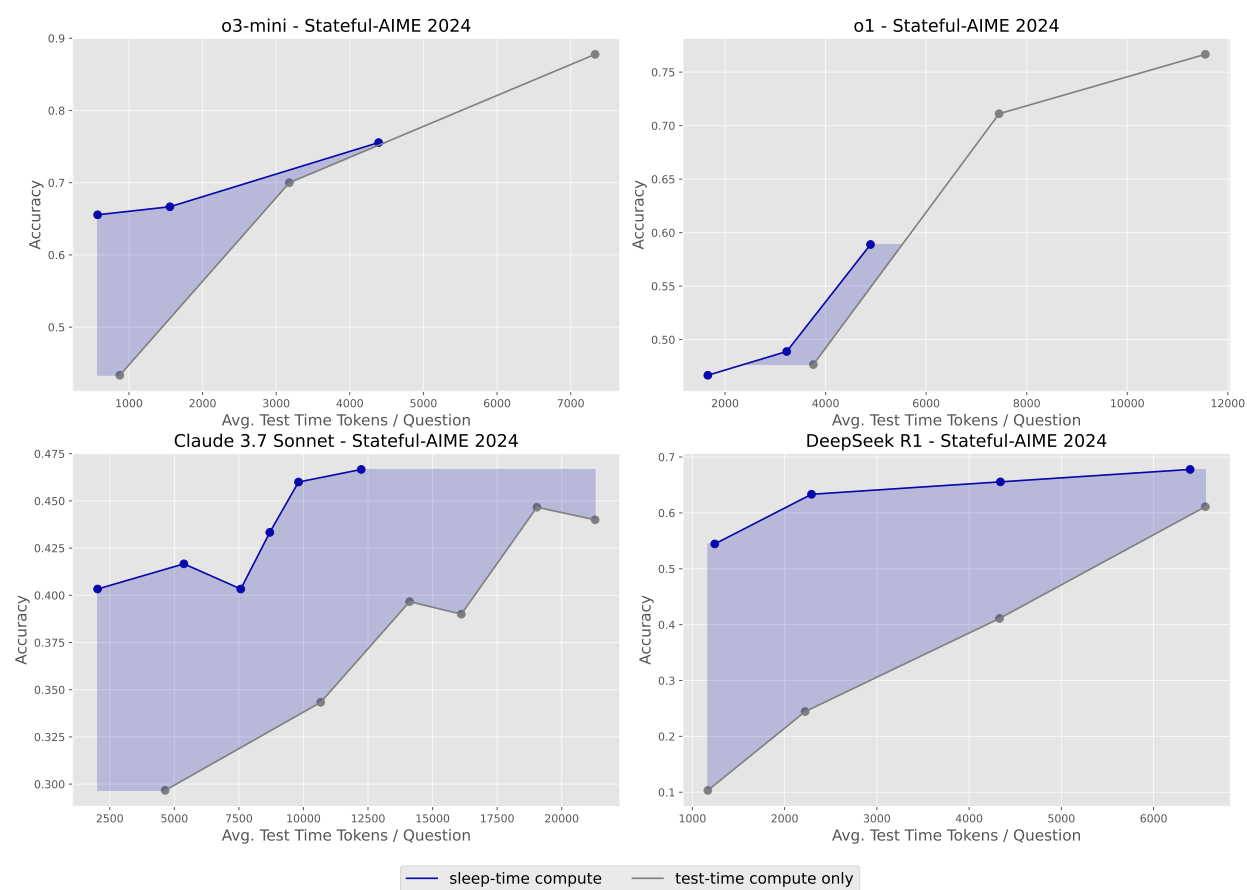


Figure C.16: AIME 2024 main result

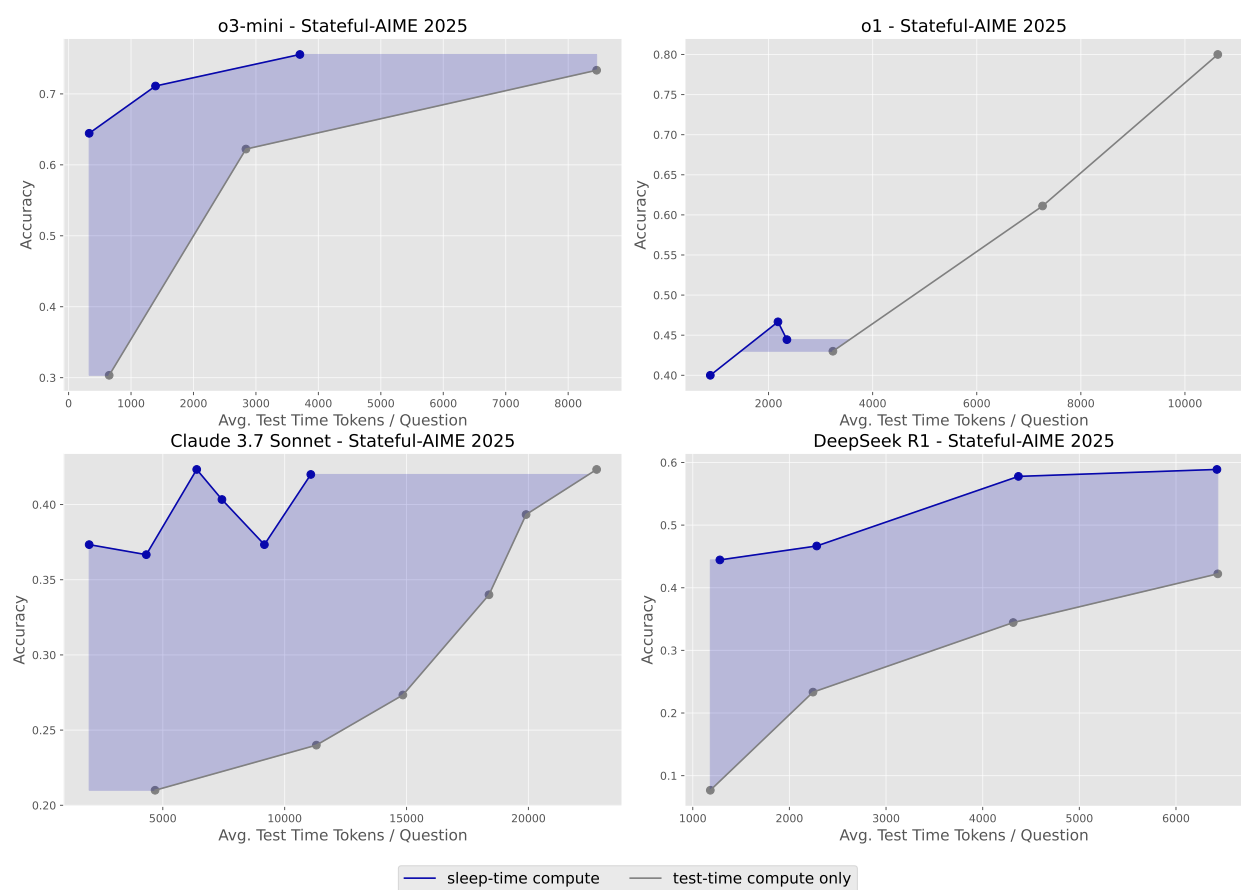


Figure C.17: AIME 2025 main result

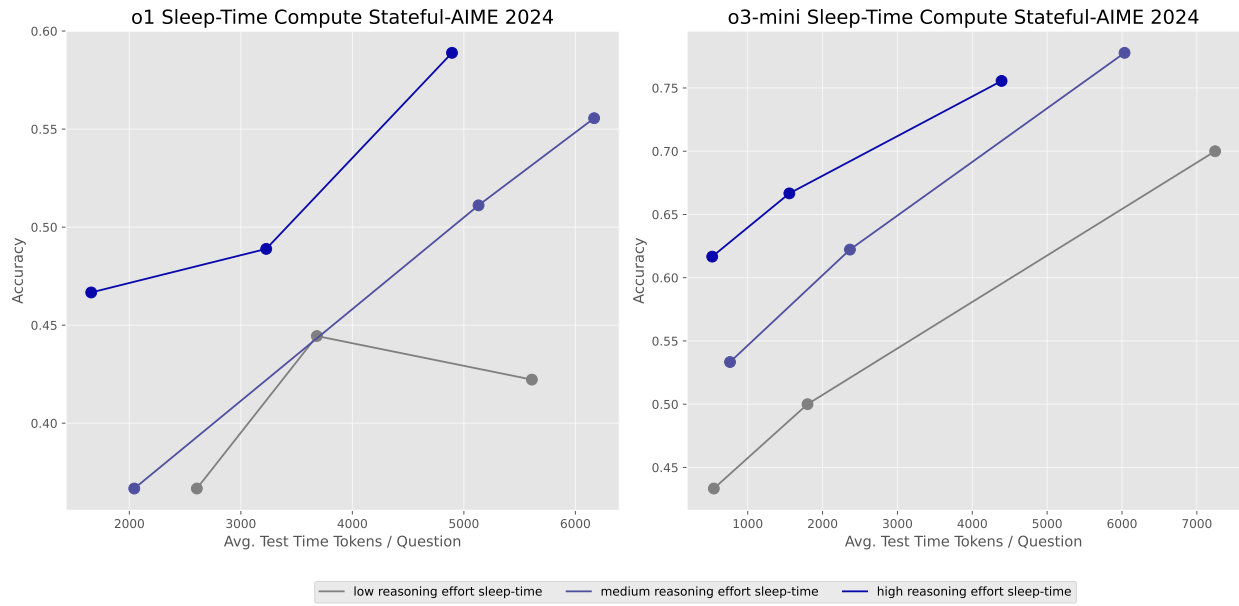


Figure C.18: Scaling sleep-time compute for Stateful AIME2024.

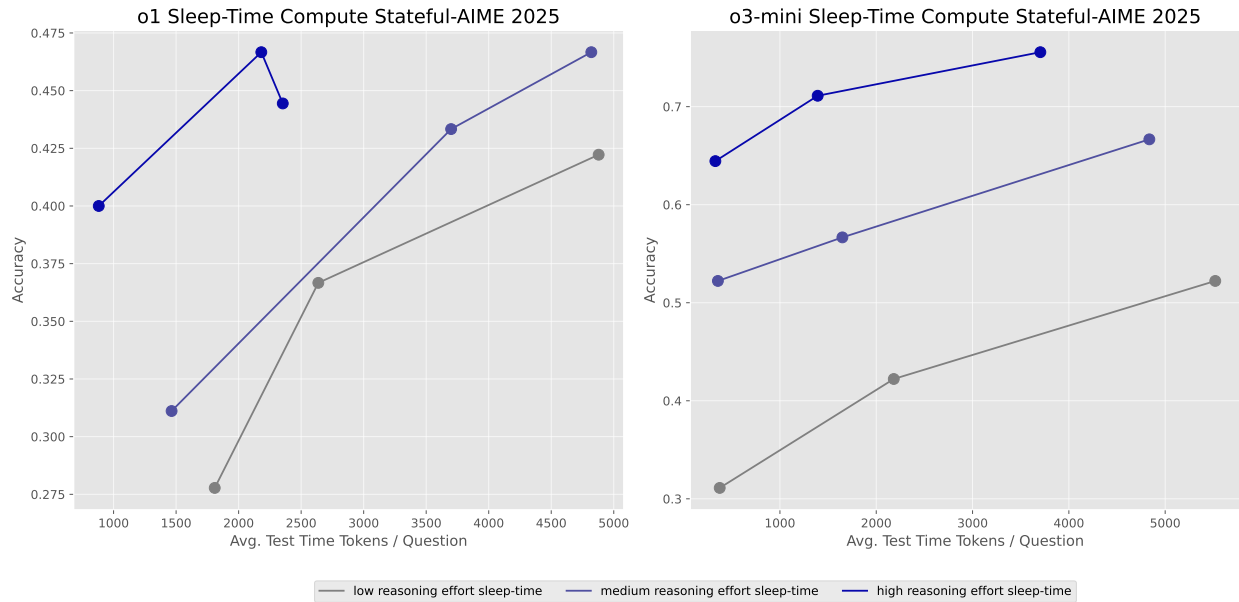


Figure C.19: Scaling sleep-time compute on Stateful AIME2025