

# Investigating Knowledge Unlearning in Large Language Models via Multi-Hop Queries

Anonymous ACL submission

## Abstract

Large language models (LLMs) serve as giant information stores, often including personal or copyrighted data, and retraining them from scratch is not a viable option for removal. This has led to the development of various fast, approximate unlearning techniques to selectively remove knowledge from LLMs. Prior research has largely focused on minimizing the probabilities of specific token sequences by reversing the language modeling objective. However, these methods may still leave LLMs vulnerable to adversarial attacks that exploit *indirect references*. In this work, we examine the limitations of current unlearning techniques in effectively erasing a particular type of indirect prompt: multi-hop queries. Our findings reveal that existing methods fail to completely remove multi-hop knowledge when one of the intermediate hops is unlearned. To address this issue, we introduce MEMMUL, a simple memory-based approach that stores all forgotten facts externally and filters multi-hop queries based on their respective scores. We demonstrate that MEMMUL achieves comparable results with GPT-4o using a 7B model and outperforms previous unlearning methods by a large margin, establishing it as a strong efficient baseline for multi-hop knowledge unlearning.<sup>1</sup>

## 1 Introduction

As the volume of data used to train large language models (LLMs) grows exponentially, these models have become vast repositories of information (Carlini et al., 2021). However, this creates a formidable challenge when specific data from the models need to be removed. For instance, sensitive information, such as personal or copyrighted data, may unintentionally be included in the training mix, or individuals may exercise their Right to be Forgotten (RTBF) (Rosen, 2011) under privacy laws such

<sup>1</sup>To promote future research, our code and data will be released upon acceptance.

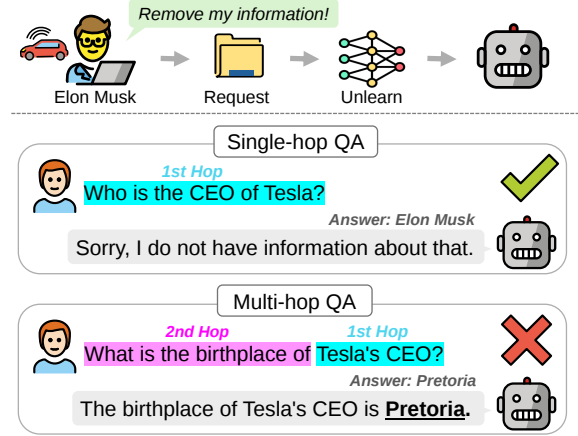


Figure 1: **A conceptual example.** After Elon Musk (i.e., “the user”) requests his personal information to be removed from the LLM, existing unlearning methods often succeed in deleting direct, single-hop facts but fail on indirect, multi-hop facts that entail one or a few of the unlearned facts.

as the European Union’s General Data Protection Regulation (GDPR) (Hoofnagle et al., 2019) or the California Consumer Privacy Act (CCPA) (Paradai, 2018) in the United States. These regulations mandate the removal of personal or protected information from databases, extending to data embedded within machine learning models. In such cases, model owners must develop mechanisms to safely eliminate specific data while preserving the model’s overall functionality.

To address these concerns, several machine unlearning methods have been introduced (Jang et al., 2023; Zhang et al., 2024c) with the goal of reversing gradients to prevent LLMs from generating certain sensitive token sequences. However, these approaches may be vulnerable to adversarial attacks, where specific token sequences are replaced or aliased with alternative sequences. For example, prompting in low-resource languages has been shown to jailbreak GPT-4 (Yong et al., 2023), and Choi et al. (2024) demonstrated that current

unlearning techniques lack cross-lingual transfer, making LLMs susceptible to such low-resource language exploits. This leads to an important research question: **Do current unlearning methods effectively erase multi-hop knowledge when one of the intermediate hops is removed?** As illustrated in Figure 1, consider a scenario where Elon Musk (i.e., “the user”) requests the removal of his personal information from an LLM. After unlearning, we expect direct, single-hop knowledge related to Elon Musk, such as “*Who is the CEO of Tesla?*”, would be deleted. Additionally, we would expect associated multi-hop knowledge, like “*What is the birthplace of Tesla’s CEO?*”, which indirectly references Musk, to also be removed.

In this study, we investigate the effectiveness of existing unlearning methods in removing multi-hop knowledge. Ideally, when one or more facts within a reasoning chain are unlearned, the model should propagate these changes, rendering it unable to answer the corresponding multi-hop questions. However, our preliminary experiments reveal that current unlearning techniques struggle to forget multi-hop questions when an intermediate hop is removed. In response, we present MEMMUL, a simple yet effective approach that explicitly stores facts to be forgotten in memory and filters incoming multi-hop questions based on their relevance scores. Concretely, MEMMUL decomposes multi-hop questions into successive subquestions, computes their relevance to the stored facts, and applies a forgetting threshold to determine whether to return a rejective response (e.g., “I don’t know.”). Such a frustratingly easy approach serves as a strong baseline for multi-hop knowledge unlearning, designed for researchers and practitioners to consider when developing their own unlearning pipelines. Notably, it requires no additional training, and even small LLMs (e.g., 7B) can match the performance of GPT-4o, making it a highly efficient and practical solution. To our knowledge, this is the first work to explore the unlearning of multi-hop knowledge.

## 2 Problem Definition

### 2.1 Probing Factual Knowledge in LLMs

We express a fact as a triple  $(s, r, o)$ , where  $s$  is the subject,  $r$  the relation, and  $o$  the object. Following Petroni et al. (2019), we define that a pretrained language model possesses specific factual knowledge if it can accurately predict the object  $o$  when given

the subject  $s$  and relation  $r$ . For example, if the subject is *Tesla* and the relation is *chief executive officer*, the model should be able to answer the question, “*Who is the CEO of Tesla?*”. While earlier work primarily focused on cloze-style statements, such as “*The CEO of Tesla is \_\_\_\_.*”, using manually written templates, we employ natural language questions to effectively query chat-based models that are becoming widely used.

### 2.2 Knowledge Unlearning

Given a token sequence  $\mathbf{x} = \{x\}_{i=1}^T$  from the training dataset  $\mathcal{D} = \{\mathbf{x}\}_{i=1}^N$ , knowledge unlearning aims to safely remove the influence of a specific subset of data  $\mathcal{D}_f$  from a trained machine learning model. The goal is to make the model behave as if this removed data was never used during training, while still maintaining its performance on the remaining dataset. Typically, the data to be forgotten  $\mathcal{D}_f$  is denoted as the *forget set*, and the data to be retained  $\mathcal{D}_r$  is referred to as the *retain set*. For simplicity, we consider the standard case where  $\mathcal{D}_f$  and  $\mathcal{D}_r$  are mutually exclusive subsets of the entire training dataset, meaning  $\mathcal{D}_f \cup \mathcal{D}_r = \mathcal{D}$  and  $\mathcal{D}_f \cap \mathcal{D}_r = \emptyset$ . In the context of factual knowledge unlearning, each token sequence  $\mathbf{x}$  represents a fact (e.g., “*The CEO of Tesla is Elon Musk.*”), and the objective is to update the model  $\pi_\theta$  to  $\pi_{\theta'} = S(\pi_\theta; \mathcal{D}_f)$ . The unlearning function  $S$  ensures that the model behaves as if it had only been trained on  $\mathcal{D}_r$ , effectively forgetting  $\mathcal{D}_f$  while preserving its performance on the retained data.

### 2.3 Assessing Multi-Hop Queries

To evaluate the unlearning of multi-hop knowledge, we must first consider a chain of facts  $\mathcal{C} = \langle (s_1, r_1, o_1), \dots, (s_n, r_n, o_n) \rangle$ , where the object of the  $i^{\text{th}}$  fact also serves as the subject of the next fact in the chain, i.e.,  $o_i = s_{i+1}$ . Using this chain, we formulate a multi-hop question that starts with the head entity  $s_1$  and ends with the tail entity  $o_n$ . For instance, consider a chain of two facts: (*Tesla*, *chief executive officer*, *Elon Musk*) and (*Elon Musk*, *place of birth*, *Pretoria*). This could generate a 2-hop question such as: “*What is the birthplace of Tesla’s CEO?*” When one or more facts from the chain are unlearned, an LLM should adjust its reasoning accordingly, effectively losing the ability to correctly answer the question. There may be a debate over how many hops should be unlearned, or whether certain multi-hop knowledge should be unlearned at all, as theoretically, the intercon-

| Hop                                       |                | Forget | Train  | Retain Valid | Test  |
|-------------------------------------------|----------------|--------|--------|--------------|-------|
| <b>MQuAKE-NoEdit (Zhong et al., 2023)</b> |                |        |        |              |       |
| Single                                    | # of questions | 1,046  | 7,322  | 1,046        | 1,046 |
|                                           | Avg. words     | 8.7    | 8.6    | 8.7          | 8.7   |
| Multi                                     | # of questions | 1,036  | -      | 988          | 976   |
|                                           | Avg. words     | 14.4   | -      | 14.5         | 14.5  |
|                                           | Avg. hops      | 2.3    | -      | 2.4          | 2.4   |
| <b>MuSiQue-Ans (Trivedi et al., 2022)</b> |                |        |        |              |       |
| Single                                    | # of questions | 1,735  | 12,150 | 1,735        | 1,735 |
|                                           | Avg. words     | 8.8    | 8.9    | 8.9          | 8.9   |
| Multi                                     | # of questions | 4,807  | -      | 4,132        | 3,347 |
|                                           | Avg. words     | 18.4   | -      | 18.2         | 17.5  |
|                                           | Avg. hops      | 2.5    | -      | 2.5          | 2.4   |

Table 1: **Dataset statistics.** The average words denote the number of words in questions. When constructing the retain set for training, we randomly sample the same number of instances as in the forget set.

nected nature of facts could lead to the unlearning of broader knowledge in the LLM. For the scope of this study, we focus on the datasets used in our experiments; nevertheless, we hope these discussions inspire further insights into developing more effective and reliable knowledge unlearning methods.

### 3 Evaluating Unlearning Approaches in Multi-Hop Question Answering

#### 3.1 Datasets

To evaluate unlearning approaches in multi-hop QA, we employ MQuAKE (Zhong et al., 2023) and MuSiQue (Trivedi et al., 2022) datasets. Their key statistics are presented in Table 1. MQuAKE, designed for multi-hop knowledge editing, assesses a model’s ability to adapt its responses to multi-hop queries when individual facts are modified. However, since our study focuses on unlearning rather than knowledge editing, we disregard the edited facts and keep only the original ones, referring to this subset as **MQuAKE-NoEdit**. While we could have chosen a dataset explicitly constructed for multi-hop QA, MQuAKE-NoEdit remains well-suited for our task due to its diverse and high-quality multi-hop questions generated using ChatGPT (gpt-3.5-turbo) based on chains of single-hop fact triples from Wikidata (Vrandečić and Krötzsch, 2014).

MuSiQue, on the other hand, is a multi-hop QA dataset created through a bottom-up approach, synthesizing multi-hop questions from a collection of single-hop questions across five English Wikipedia-based datasets: SQuAD (Rajpurkar et al., 2016), ZsRE (Levy et al., 2017), T-REx (Elsahar et al., 2018), Natural Questions (Kwiatkowski et al.,

2019), and MLQA (Lewis et al., 2020). MuSiQue is shown to be more challenging than previous multi-hop datasets such as HotpotQA (Yang et al., 2018) and 2WikiMultihopQA (Ho et al., 2020). It enforces connected reasoning, reducing the likelihood of shortcut-based answering, making it an ideal choice for our study. Specifically, we employ the **MuSiQue-Ans** subset, which contains approximately 25K answerable questions spanning 2-4 hops. For details on our data preprocessing steps, including the partitioning of forget and retain sets, refer to Appendix B.

#### 3.2 Experimental Setup

**Knowledge unlearning approaches** We evaluate the following state-of-the-art knowledge unlearning approaches (see Appendix A for details):

- **GA** (Jang et al., 2023): Applies gradient ascent to decrease the likelihood of token sequences associated with the forget set
- **DPO** (Rafailov et al., 2023): Performs direct preference optimization, prioritizing “I don’t know” responses for items in the forget set
- **NPO** (Zhang et al., 2024c): Implements negative preference optimization to actively disfavor responses linked to the forget set
- **+RT**: Includes additional finetuning on the retain set to explicitly reinforce knowledge retention in the model

**Implementation details** We built our framework on PyTorch (Paszke et al., 2019) and Hugging Face Transformers (Wolf et al., 2020). We employed OLMo-2-7B-Instruct (OLMo et al., 2024) and Qwen-2.5-7B-Instruct (Yang et al., 2024) as the backbones of our framework (see Appendix C for results on more models) and optimized their weights with AdamW (Loshchilov and Hutter, 2019). We trained all +RT models for 5 epochs (and non-RT models for 2 epochs) with warmup during the first epoch and set the batch size to 32, the learning rate to 1e-5, and the weight decay to 0.01. We set the loss scaling factor  $\alpha$  to 0.3 to balance between forgetting and retaining (see Equation 5). All experiments utilized 1% of data as the forget set (i.e., 104 and 173 samples for MQuAKE-NoEdit and MuSiQue-Ans, respectively) unless otherwise noted. Each experiment was repeated with three different random seeds, and the results were averaged for reporting.

|                             | Forget Set (Single-Hop) |             |             | Forget Set (Multi-Hop) |             |            | Retain Set (Single-Hop) |             |            | Retain Set (Multi-Hop) |             |            | Utility Set |
|-----------------------------|-------------------------|-------------|-------------|------------------------|-------------|------------|-------------------------|-------------|------------|------------------------|-------------|------------|-------------|
|                             | PA(↓)                   | R-L(↓)      | LM(↑)       | PA(↓)                  | R-L(↓)      | LM(↑)      | PA(↑)                   | R-L(↑)      | LM(↓)      | PA(↑)                  | R-L(↑)      | LM(↓)      | Avg.(↑)     |
| <i>OLMo-2-7B-Instruct</i>   |                         |             |             |                        |             |            |                         |             |            |                        |             |            |             |
| Original                    | 95.2                    | 82.7        | 1.2         | 98.1                   | 53.3        | 1.2        | 97.0                    | 84.5        | 1.1        | 97.4                   | 49.5        | 1.2        | 64.0        |
| GA                          | 19.2                    | 56.6        | 10.3        | 39.4                   | 29.4        | 7.9        | 35.3                    | 67.4        | 8.8        | 44.7                   | 26.1        | 7.5        | 62.8        |
| DPO                         | 25.0                    | 0.0         | 6.3         | 70.2                   | 0.0         | 3.7        | 40.8                    | 0.0         | 5.3        | 63.2                   | 0.2         | 3.6        | 63.1        |
| NPO                         | 20.2                    | 55.9        | 10.3        | 38.5                   | 27.9        | 7.9        | 35.0                    | 68.2        | 8.8        | 44.6                   | 26.3        | 7.5        | 62.8        |
| GA+RT                       | <b>29.5</b>             | 38.1        | <b>12.0</b> | 87.8                   | 26.3        | <b>3.7</b> | 78.0                    | 65.2        | 3.6        | 91.8                   | <b>30.6</b> | 2.6        | <b>62.6</b> |
| DPO+RT                      | 34.9                    | <b>1.0</b>  | 11.6        | 90.4                   | <b>2.4</b>  | 2.8        | <b>83.3</b>             | 15.8        | 3.2        | <b>95.6</b>            | 4.7         | <b>2.0</b> | 62.4        |
| NPO+RT                      | 32.1                    | 37.3        | 11.1        | <b>87.5</b>            | 26.8        | 3.3        | 79.7                    | <b>67.7</b> | <b>3.0</b> | 92.9                   | 29.9        | 2.3        | <b>62.6</b> |
| <i>Qwen-2.5-7B-Instruct</i> |                         |             |             |                        |             |            |                         |             |            |                        |             |            |             |
| Original                    | 97.1                    | 78.9        | 2.2         | 96.2                   | 53.9        | 2.6        | 94.3                    | 82.1        | 2.2        | 95.8                   | 49.5        | 2.6        | 65.5        |
| GA                          | 23.1                    | 0.5         | 25.8        | 25.0                   | 0.0         | 27.9       | 19.9                    | 2.1         | 25.0       | 31.6                   | 0.3         | 27.6       | 63.9        |
| DPO                         | 32.7                    | 3.7         | 8.8         | 67.3                   | 1.8         | 5.9        | 37.5                    | 2.3         | 8.2        | 68.3                   | 2.0         | 5.8        | 64.2        |
| NPO                         | 22.1                    | 1.0         | 25.5        | 26.9                   | 0.0         | 27.6       | 20.1                    | 2.2         | 24.8       | 31.6                   | 0.2         | 27.4       | 63.8        |
| GA+RT                       | <b>41.7</b>             | 38.9        | <b>15.5</b> | 93.6                   | 35.3        | <b>4.2</b> | 82.8                    | <b>72.2</b> | 5.5        | 95.3                   | 39.0        | 3.2        | <b>66.5</b> |
| DPO+RT                      | 46.5                    | <b>11.8</b> | 12.0        | <b>90.4</b>            | <b>14.7</b> | 3.6        | <b>88.8</b>             | 42.3        | <b>3.6</b> | <b>96.4</b>            | 22.6        | <b>2.6</b> | 65.1        |
| NPO+RT                      | 47.8                    | 38.0        | 13.5        | 92.6                   | 34.9        | 3.9        | 87.0                    | 72.0        | 4.3        | 95.9                   | <b>40.1</b> | 3.0        | <b>66.5</b> |

Table 2: Performance comparison of different knowledge unlearning methods after erasing single-hop facts from the forget set in **MQuAKE-NoEdit**. The best results among +RT models are highlighted in **bold**.

|                             | Forget Set (Single-Hop) |            |             | Forget Set (Multi-Hop) |            |            | Retain Set (Single-Hop) |             |            | Retain Set (Multi-Hop) |             |            | Utility Set |
|-----------------------------|-------------------------|------------|-------------|------------------------|------------|------------|-------------------------|-------------|------------|------------------------|-------------|------------|-------------|
|                             | PA(↓)                   | R-L(↓)     | LM(↑)       | PA(↓)                  | R-L(↓)     | LM(↑)      | PA(↑)                   | R-L(↑)      | LM(↓)      | PA(↑)                  | R-L(↑)      | LM(↓)      | Avg.(↑)     |
| <i>OLMo-2-7B-Instruct</i>   |                         |            |             |                        |            |            |                         |             |            |                        |             |            |             |
| Original                    | 89.0                    | 37.2       | 1.4         | 94.7                   | 26.2       | 1.4        | 83.1                    | 37.9        | 1.5        | 94.0                   | 24.5        | 1.5        | 64.0        |
| GA                          | 22.5                    | 18.5       | 13.6        | 29.6                   | 12.0       | 10.3       | 27.2                    | 19.6        | 12.5       | 30.8                   | 12.7        | 10.3       | 61.0        |
| DPO                         | 19.1                    | 0.1        | 8.4         | 45.9                   | 0.9        | 5.2        | 31.1                    | 0.2         | 7.1        | 47.1                   | 0.2         | 5.2        | 61.9        |
| NPO                         | 22.5                    | 18.2       | 13.6        | 29.6                   | 12.3       | 10.3       | 27.0                    | 19.2        | 12.5       | 30.7                   | 12.5        | 10.3       | 61.0        |
| GA+RT                       | <b>22.5</b>             | 18.0       | <b>16.2</b> | <b>83.4</b>            | 16.8       | <b>5.1</b> | 83.8                    | <b>28.6</b> | 3.8        | 94.0                   | <b>19.3</b> | 3.4        | 61.0        |
| DPO+RT                      | 24.9                    | <b>1.9</b> | 14.1        | 85.5                   | <b>4.6</b> | 4.6        | <b>87.5</b>             | 11.0        | 3.4        | <b>95.7</b>            | 5.4         | 3.4        | 60.6        |
| NPO+RT                      | 24.5                    | 16.9       | 14.3        | 84.5                   | 16.5       | 4.3        | 85.1                    | 27.9        | <b>3.1</b> | 95.0                   | 18.6        | <b>2.9</b> | <b>61.2</b> |
| <i>Qwen-2.5-7B-Instruct</i> |                         |            |             |                        |            |            |                         |             |            |                        |             |            |             |
| Original                    | 79.8                    | 34.3       | 2.9         | 93.1                   | 21.5       | 2.5        | 79.6                    | 34.8        | 2.9        | 92.0                   | 23.2        | 2.6        | 65.5        |
| GA                          | 24.9                    | 3.9        | 50.0        | 22.7                   | 3.3        | 43.8       | 21.9                    | 1.7         | 49.8       | 21.7                   | 2.5         | 44.5       | 61.5        |
| DPO                         | 15.6                    | 0.6        | 25.8        | 16.0                   | 2.6        | 21.3       | 20.2                    | 0.7         | 25.7       | 16.4                   | 1.6         | 21.9       | 63.0        |
| NPO                         | 21.4                    | 4.9        | 49.3        | 19.1                   | 2.9        | 42.9       | 20.8                    | 1.8         | 48.9       | 19.1                   | 2.6         | 43.6       | 61.4        |
| GA+RT                       | 36.8                    | 15.4       | <b>17.3</b> | 83.7                   | 14.1       | <b>6.8</b> | 86.7                    | 24.8        | 3.7        | 92.9                   | 16.2        | 4.5        | <b>65.3</b> |
| DPO+RT                      | <b>32.0</b>             | <b>4.9</b> | 15.8        | <b>83.6</b>            | <b>8.0</b> | 5.5        | <b>87.7</b>             | 16.9        | <b>3.5</b> | 94.0                   | 9.6         | <b>3.9</b> | 65.2        |
| NPO+RT                      | 35.1                    | 16.4       | 16.2        | 87.4                   | 16.7       | 6.1        | 87.6                    | <b>27.4</b> | 3.6        | <b>94.2</b>            | <b>18.3</b> | 4.2        | <b>65.3</b> |

Table 3: Performance comparison of different knowledge unlearning methods after erasing single-hop facts from the forget set in **MuSiQue-Ans**. The best results among +RT models are highlighted in **bold**.

**Evaluation metrics** To evaluate the unlearning of factual knowledge, we adopt the approach of Petroni et al. (2019) and report **Probing Accuracy (PA)**. This rank-based metric computes the mean precision at  $k$  ( $P@k$ ) across all relations, with  $k$  set to 1. In other words, for a given fact, the value is 1 if the correct object appears among the top  $k$  predictions, and 0 otherwise. By the definition of probing in Section 2.1, we consider a pretrained language model to have successfully unlearned a fact if it can no longer predict the correct object accurately. To generate answer candidates, we use GPT-4o to produce perturbed responses for each example. To assess the model’s generation capability, we measure **ROUGE-L recall (R-L)** (Lin, 2004) to compare the model’s gen-

erated outputs (via greedy decoding) against the ground-truth answers, accounting for slight variations in phrasing between the generated and reference outputs. Additionally, **Language Modeling Loss (LM)** is computed over token sequences to quantify how perplexed the model is by the data. Finally, we assess the model’s overall utility by averaging performance across eight language understanding benchmarks: ARC-Challenge (Clark et al., 2018), CommonsenseQA (Talmor et al., 2019), HellaSwag (Zellers et al., 2019), Lambada (Paperno et al., 2016), MMLU (Hendrycks et al., 2021), OpenbookQA (Mihaylov et al., 2018), PIQA (Bisk et al., 2020), and Winogrande (Sakaguchi et al., 2021). Full individual results can be found in Appendix C.

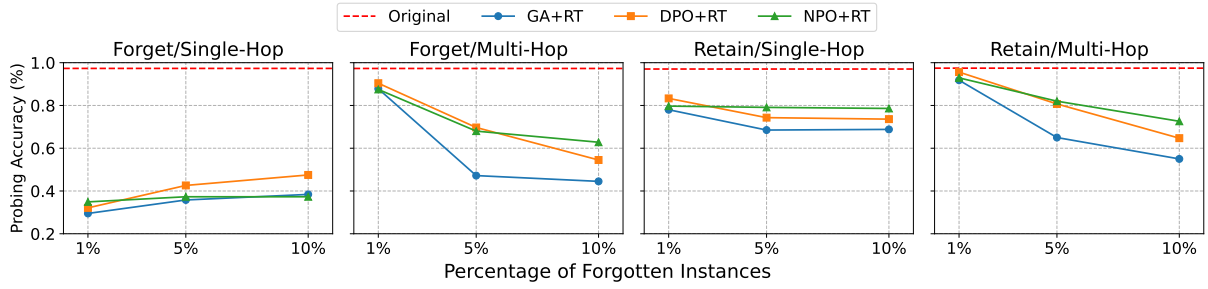


Figure 2: Data scaling performance of various unlearning methods using **OLMo-2-7B-Instruct** across different proportions of data for the forget set (1%, 5%, and 10%). Models consistently preserve the ability to unlearn and retain single-hop facts with scaling. While unlearning multi-hop facts seems to improve with scaling, a similar decline is also observed in the retain set. This suggests that the effect may be attributed to catastrophic forgetting of general information rather than a genuine improvement in unlearning multi-hop facts.

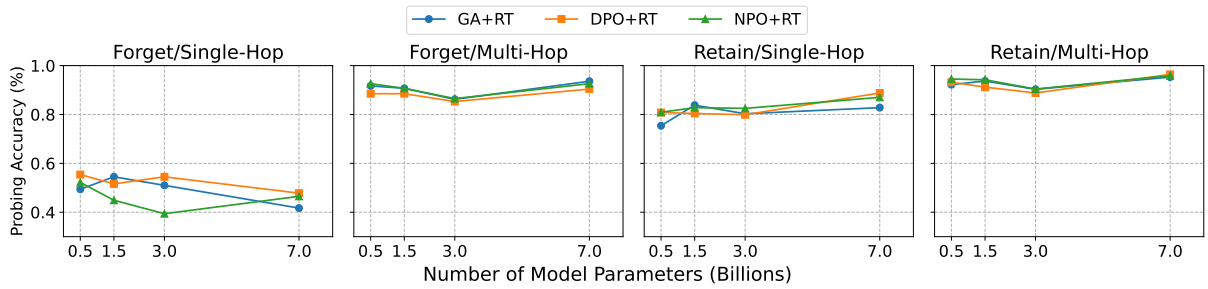


Figure 3: Model scaling performance of various unlearning methods using **Qwen-2.5-Instruct** across different model sizes (0.5B, 1.5B, 3B, and 7B). The performance trends remain consistent across all model sizes, indicating that the underlying issue persists regardless of model scale.

### 3.3 Knowledge Unlearning Results

We present a comparison of unlearning performance across various methods in Tables 2 and 3. Each method was trained for at least one epoch to ensure the model had exposure to all samples in the forget set. We find that all non-RT methods successfully forget corresponding multi-hop knowledge but also the knowledge to be retained, indicating that models largely lost their ability to retain information and function correctly. On the other hand, additional finetuning on the retain set (i.e., +RT) mitigates catastrophic forgetting, evidenced by retention performance comparable to the original for both single-hop and multi-hop facts. Nevertheless, in both OLMo and Qwen models, multi-hop facts within the forget set were not effectively unlearned. Similar trends emerge in results from four other open-source LLMs, detailed in Appendix C. This indicates that **existing unlearning methods, while capable of removing single-hop information, struggle to extend that effect to the corresponding multi-hop knowledge**. These outcomes underscore the need for new approaches to address unlearning in multi-hop scenarios.

### 3.4 Evaluation with Unlearning at Scale

**Data scaling** In real-world scenarios, the number of samples to forget can vary. Thus, we evaluate the performance of unlearning and retaining multi-hop facts as the size of the forget set changes. We conduct experiments using 1%, 5%, and 10% of the MQAKE-NoEdit dataset for forgetting (104, 523, and 1,046 single-hop instances, respectively), with the results shown in Figure 2. Our findings indicate that all knowledge unlearning methods effectively scale for single-hop, consistently preserving the ability to forget and retain single-hop facts. For multi-hop facts, unlearning performance improves with larger forget sets, as reflected in a noticeable performance drop. However, a similar decline is observed in the retain set, suggesting that this effect might stem from catastrophic forgetting of general knowledge rather than a true enhancement in unlearning multi-hop facts.

**Model scaling** To assess the impact of model size on the unlearning of multi-hop facts, we evaluate performance across different LLM scales. Leveraging the Qwen-2.5 series, we conduct experiments on models of 0.5B, 1.5B, 3B, and 7B parameters,

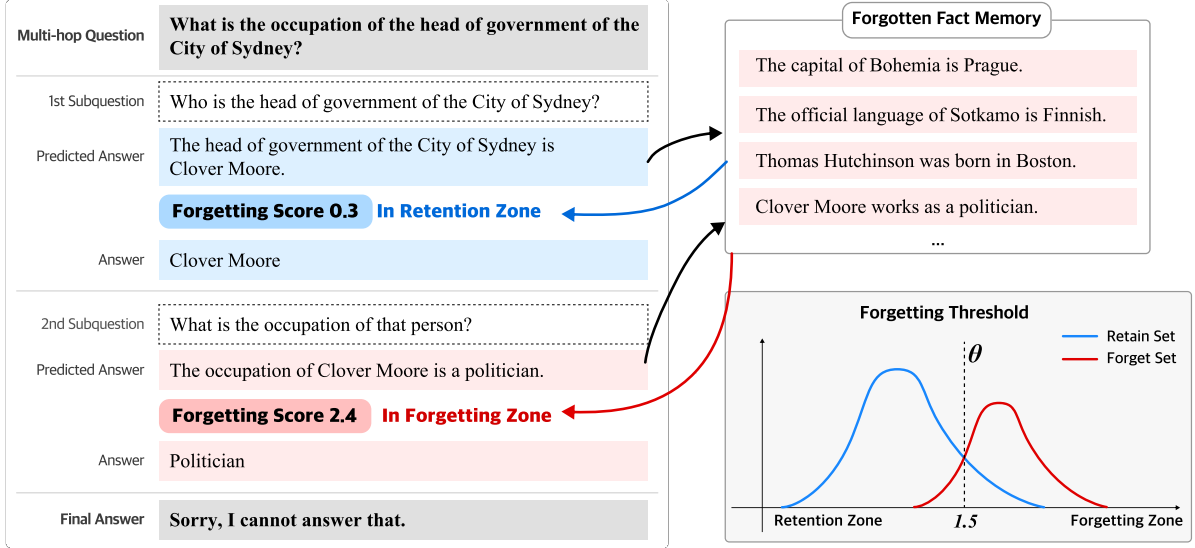


Figure 4: **Overview of the proposed MEMMUL framework.** MEMMUL begins by breaking down a multi-hop question into a sequence of subquestions, where each subquestion is passed to the base model to generate predicted answers. Then, these predictions are compared with the stored facts in memory to derive corresponding forgetting scores. If any predicted answer yields a high forgetting score, MEMMUL responds with a rejection (e.g., “I don’t know.”). Otherwise, the final response is based on the last intermediate answer in the sequence.

with results presented in Figure 3. Our findings indicate that, regardless of model size, forgetting single-hop facts does not trigger cascading changes in multi-hop knowledge, highlighting a persistent challenge in knowledge unlearning.

## 4 MEMMUL: A Strong Baseline for Unlearning Multi-Hop Facts in LLMs

In this section, we present the **Memory-based Multi-Hop Knowledge UnLearning (MEMMUL)**, a simple yet effective approach to forgetting multi-hop facts in LLMs. Figure 4 illustrates the overview of our method.

### 4.1 Methodology

Inspired by memory-based multi-hop knowledge editing frameworks (Zhong et al., 2023), MEMMUL tracks all the forgotten facts in an explicit memory while keeping the base LLM frozen. Particularly, MEMMUL (1) decomposes a multi-hop question into subquestions, (2) computes the forgetting scores relative to the stored facts, and (3) decides whether to respond with a refusal (e.g., “I don’t know.”) based on the forgetting threshold.

**Forgotten fact memory** MEMMUL explicitly stores all forgotten facts in memory. For simplicity and consistency with prior work, we assume that all facts are single-hop. Specifically, single-hop facts in the forget set are first transformed into

sentence-level statements using GPT-4o. Next, we encode these statements with the off-the-shelf embedding model *Contriever* (Izacard et al., 2022) and store them in a retrieval index. Given a query, the index retrieves the most relevant forgotten fact, determined by proximity in the embedding space.

**Decomposing multi-hop questions** Since single-hop facts are stored in memory, it is intuitive to compare them against *single-hop statements*. To enhance the unlearning of multi-hop facts in LLMs, we build on previous work by breaking down multi-hop questions into a series of simpler queries (Zhou et al., 2023). In multi-hop reasoning, where the predicted answer of one question serves as the subject for the next fact (i.e.,  $o_i = s_{i+1}$ ), model-generated responses to intermediate questions can slow down the process. To mitigate this, we leverage coreference resolution to construct subquestions all at once, bypassing the need for sequential answering. For instance, as shown in Figure 4, if the first subquestion is “Who is the head of government of the City of Sydney?”, the second subquestion would be “What is the occupation of that person?”, eliminating the need to resolve the first before proceeding. In practice, we leverage a few-shot prompt with three demonstrations, as illustrated in Figure 5.

**Distinguishing forget and retain facts** If a multi-hop question has effectively been decomposed, its subquestions should resemble single-hop

| Decomposer     | Method             | MQuAKE-NoEdit |            |             |             | MuSiQue-Ans |            |             |             |
|----------------|--------------------|---------------|------------|-------------|-------------|-------------|------------|-------------|-------------|
|                |                    | Forget Set    |            | Retain Set  |             | Forget Set  |            | Retain Set  |             |
|                |                    | PA(↓)         | R-L(↓)     | PA(↑)       | R-L(↑)      | PA(↓)       | R-L(↓)     | PA(↑)       | R-L(↑)      |
|                | Original           | 96.2          | 53.9       | 95.8        | 49.5        | 93.1        | 21.5       | 92.0        | 23.2        |
|                | GA+RT              | 93.6          | 35.3       | 95.3        | 39.0        | 83.7        | 14.1       | 92.9        | 16.2        |
|                | DPO+RT             | 90.4          | 14.7       | 96.4        | 22.6        | 83.6        | 8.0        | 94.0        | 9.6         |
|                | NPO+RT             | 92.6          | 34.9       | 95.9        | 40.1        | 87.4        | 16.7       | <u>94.2</u> | 18.3        |
| Qwen-2.5-7B-IT | MeLLO <sup>†</sup> | 96.2          | 21.1       | <b>97.0</b> | 22.4        | 93.1        | 10.2       | <b>92.7</b> | 7.9         |
|                | MEMMUL             | <b>7.7</b>    | <b>4.1</b> | 91.6        | <b>51.4</b> | <b>25.4</b> | <b>5.4</b> | 84.7        | <b>21.4</b> |
| GPT-4o-mini    | MeLLO <sup>†</sup> | 57.7          | 24.1       | 78.6        | 41.2        | 75.2        | 17.0       | <b>91.0</b> | <b>26.5</b> |
|                | MEMMUL             | <b>9.6</b>    | <b>5.7</b> | <b>92.3</b> | <u>52.7</u> | <u>22.6</u> | <u>4.9</u> | 85.1        | 21.1        |
| GPT-4o         | MeLLO <sup>†</sup> | 22.1          | 12.5       | 83.1        | 48.6        | 52.5        | 12.2       | <b>90.8</b> | <b>38.8</b> |
|                | MEMMUL             | <b>10.6</b>   | <b>7.0</b> | <b>91.4</b> | <b>51.0</b> | <b>23.9</b> | <u>4.9</u> | 90.7        | 22.0        |

Table 4: Multi-hop knowledge unlearning performance of MEMMUL (ours) and MeLLO (Zhong et al., 2023), a memory-based multi-hop knowledge editing method. (†) indicates our modified implementation adapted specifically for the unlearning task. The base model is **Qwen-2.5-7B-Instruct** in MEMMUL, predicting answers to questions decomposed by the **Decomposer** model. The base and decomposer models are the same in MeLLO. The better results between the two are in **bold**, while the best results are underlined.

queries. The outputs generated for these subquestions (which we refer to as *subanswers*) can then serve as proxies for the forgotten single-hop facts. Therefore, we feed each subquestion to the base model to generate subanswers and compute their forgetting scores w.r.t. the stored facts. Since retrieval is not our primary focus, we use simple dot product similarity between embeddings to identify the most relevant fact for each predicted answer.

| Retrieval Strategy   | Forget Set |            | Retain Set  |             |
|----------------------|------------|------------|-------------|-------------|
|                      | PA(↓)      | R-L(↓)     | PA(↑)       | R-L(↑)      |
| Subanswer (ours)     | <b>7.7</b> | <b>4.1</b> | <u>91.6</u> | <u>51.4</u> |
| Subquestion          | <u>8.7</u> | <u>4.2</u> | <u>91.6</u> | <b>51.7</b> |
| Subquestion (coref.) | 18.3       | 13.8       | 88.3        | 48.9        |
| Multi-hop question   | 19.2       | 13.8       | <b>93.6</b> | 48.1        |

Table 5: Comparison of different retrieval strategies on **MQuAKE-NoEdit** for multi-hop knowledge unlearning performance using **Qwen-2.5-7B-Instruct** as the multi-hop question decomposer.

To determine which multi-hop facts to unlearn or retain, we establish a threshold that effectively separates the two data distributions. We approximate this threshold by plotting the probability density functions of the forget set and the validation split of the retain set. During inference, we apply this threshold to assess whether the forgetting score of each subanswer is high or low. If any subanswer yields a high forgetting score, we replace the final answer with a rejective response, following the approach of selective generation (Zhang et al., 2024a). Otherwise, the final answer is drawn from the last intermediate predicted answer.

## 4.2 Evaluation Results

Table 4 presents the performance of MEMMUL in multi-hop knowledge unlearning. To highlight the effectiveness of our method, we compare it with MeLLO (Zhong et al., 2023), a memory-based approach for editing multi-hop knowledge. Since MeLLO is not designed for unlearning, we modify its prompt to generate rejective responses instead of edits. Specifically, MeLLO sequentially decomposes a multi-hop question, retrieves the most relevant fact to each subquestion, and asks the model whether the retrieved fact is the same as the decomposed output. The key difference from our method lies in the decision-making process: MeLLO relies on the base LLM’s judgment, whereas MEMMUL incorporates a forgetting threshold, which better aligns with unlearning. Our results show that MeLLO benefits from stronger base models but struggles with Qwen-2.5-7B-Instruct, highlighting its reliance on LLM reasoning. In contrast, MEMMUL consistently unlearns multi-hop knowledge effectively, regardless of model capacity. Notably, our 7B model performs comparably to GPT-4o, demonstrating MEMMUL’s efficiency and practicality for multi-hop knowledge unlearning.

## 4.3 Effect of Question Decomposition

To verify the efficacy of question decomposition, we compare performance with and without it in Table 5. Given the small size of the forget set in memory (~100 samples), distinguishing between forget and retain sets using multi-hop questions appears relatively easy, as evidenced by a notable

|                    |                                                                                                                                                                                                                                                                                                                                                                             |
|--------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Question</b>    | What is the capital of the country where Greg Maddux’s sport originated? ( <b>Answer:</b> Washington, D.C.)                                                                                                                                                                                                                                                                 |
| Original           | Greg Maddux is a famous former professional baseball player. Baseball originated in the United States, and the capital of the United States is <b>Washington, D.C.</b>                                                                                                                                                                                                      |
| NPO+RT             | Greg Maddux’s sport is Major League Baseball. Major League Baseball originated in the United States. The capital of the United States is <b>Washington, D.C.</b>                                                                                                                                                                                                            |
| MeLLO <sup>†</sup> | The capital of the country where Greg Maddux’s sport (baseball) originated is <b>Washington, D.C.</b>                                                                                                                                                                                                                                                                       |
| MEMMUL             | <b>Subquestion 1:</b> What sport is Greg Maddux associated with?<br><b>Subanswer 1:</b> Greg Maddux is associated with baseball. He was a highly successful Major League Baseball pitcher...<br><b>Retrieved Fact:</b> Greg Maddux is associated with the sport of baseball. (Forgetting Score: 2.12 / Threshold: 1.51)<br><b>Final Answer:</b> I’m unaware of that detail. |

Table 6: Qualitative examples of generated outputs for a three-hop question, which contains a single-hop fact included in the forget set.

drop in the forget set. Our decomposed subquestions use coreference resolution (e.g., “What is the occupation of *that person*?”), which does not effectively retrieve relevant facts. Thus, we also evaluate subquestions without coreferences, confirming that decomposition improves fact retrieval. MEMMUL leverages subanswers, consistently outperforming other strategies.

#### 4.4 Qualitative Analysis

Table 6 displays qualitative results for a three-hop question. As shown, the original model employs a chain-of-thought approach to solve the multi-hop question. Similarly, the NPO+RT unlearned model follows a nearly identical reasoning process, indicating that its unlearning mechanism had minimal impact on the corresponding multi-hop knowledge. During iterative prompting, MeLLO failed to recognize that one of the decomposed questions aligned with a forgotten fact. In contrast, MEMMUL successfully refrained from answering correctly by leveraging a high forgetting score.

### 5 Related Work

#### 5.1 Machine Unlearning

Machine unlearning has emerged as a critical research area in response to growing concerns over data privacy, regulatory compliance, and ethical AI (Cao and Yang, 2015; Ginart et al., 2019; Bourtole et al., 2021). In the context of LLMs, addressing memorization has garnered wide attention (Wang et al., 2023; Chen and Yang, 2023; Kassem et al., 2023; Liu et al., 2024a,b; Hong et al., 2024; Tian et al., 2024; Choi et al., 2024; Jia et al., 2024; Ji et al., 2024). The predominant approach involves maximizing prediction loss on

the forget set (Jang et al., 2023; Yao et al., 2023; Lee et al., 2024; Zhang et al., 2024c; Feng et al., 2024). Other methods train LLMs to generate alternative responses, such as “I don’t know” (Maini et al., 2024), random labels (Yao et al., 2024), or generic terms (Eldan and Russinovich, 2023). Recent studies have also explored task arithmetic (Ilharco et al., 2023; Bărbulescu and Triantafillou, 2024) and training-free methods that simulate unlearning through specific instructions (Thaker et al., 2024) or in-context examples (Pawelczyk et al., 2024). This work focuses on multi-hop knowledge unlearning, a novel challenge that introduces new complexities and opportunities in the field.

#### 5.2 Multi-Hop Reasoning

Multi-hop reasoning involves connecting multiple pieces of evidence across contexts to derive information (Huang and Chang, 2023). Editing multi-hop knowledge in LLMs is challenging, as it requires consistent propagation of updates across interconnected facts (Valmeekam et al., 2022; Press et al., 2023; Dziri et al., 2023; Petty et al., 2024). Existing knowledge editing methods primarily modify individual facts (De Cao et al., 2021; Meng et al., 2022; Zhang et al., 2024b) but often struggle to update related knowledge (Onoe et al., 2023; Zhong et al., 2023; Cohen et al., 2024). Recent approaches address this by injecting information at inference time (Sakarvadia et al., 2023), removing shortcut-inducing neurons (Ju et al., 2024), or adjusting model representations to fix multi-hop reasoning errors (Ghandeharioun et al., 2024). In contrast, our work examines whether *removing* specific information from models can be generalized effectively in multi-hop scenarios.

### 6 Conclusion

This study explores the effectiveness of existing unlearning methods in eliminating multi-hop knowledge. Our results indicate that they struggle when an intermediate hop is unlearned. To overcome this issue, we propose MEMMUL, a simple yet effective memory-based approach that dissects multi-hop questions into subquestions and utilizes forgetting scores relative to stored facts to determine when to issue a rejective response. MEMMUL serves as a strong baseline for multi-hop knowledge unlearning, offering a highly efficient and practical solution for unlearning in LLMs.

## Limitations

Building on previous training-free unlearning methods (Thaker et al., 2024; Pawelczyk et al., 2024) and memory-based multi-hop knowledge editing frameworks (Zhong et al., 2023), MEMMUL is a training-free approach that selectively refuses to answer multi-hop questions based on their relevance to forgotten facts stored in memory. Therefore, it does not essentially erase multi-hop knowledge from model parameters, meaning this information could still be extracted through advanced adversarial techniques. We emphasize that MEMMUL only serves as a baseline for multi-hop knowledge unlearning, demonstrating that effective performance can be achieved without additional training. We urge researchers and practitioners to exercise caution when attempting to fully remove multi-hop knowledge through training, as doing so may inadvertently erase broader knowledge interconnected through hops. Furthermore, we acknowledge the need for more rigorous evaluation metrics to better defend against state-of-the-art jailbreaking attacks. We hope this work stimulates further research and discussions on creating more robust frameworks for knowledge unlearning.

## References

- Marah Abdin, Sam Ade Jacobs, Ammar Ahmad Awan, Jyoti Aneja, Ahmed Awadallah, Hany Awadalla, Nguyen Bach, Amit Bahree, Arash Bakhtiari, Harkirat Behl, et al. 2024. [Phi-3 technical report: A highly capable language model locally on your phone](#). *arXiv preprint arXiv:2404.14219*.
- George-Octavian Bărbulescu and Peter Triantafillou. 2024. [To each \(textual sequence\) its own: Improving memorized-data unlearning in large language models](#). In *Forty-first International Conference on Machine Learning*.
- Yonatan Bisk, Rowan Zellers, Jianfeng Gao, Yejin Choi, et al. 2020. [Piqa: Reasoning about physical commonsense in natural language](#). In *Proceedings of the AAAI conference on artificial intelligence*, volume 34, pages 7432–7439.
- Lucas Bourtole, Varun Chandrasekaran, Christopher A Choquette-Choo, Hengrui Jia, Adelin Travers, Baiwu Zhang, David Lie, and Nicolas Papernot. 2021. [Machine unlearning](#). In *2021 IEEE Symposium on Security and Privacy (SP)*, pages 141–159. IEEE.
- Yinzhi Cao and Junfeng Yang. 2015. [Towards making systems forget with machine unlearning](#). In *2015 IEEE symposium on security and privacy*, pages 463–480. IEEE.
- Nicholas Carlini, Florian Tramèr, Eric Wallace, Matthew Jagielski, Ariel Herbert-Voss, Katherine Lee, Adam Roberts, Tom Brown, Dawn Song, Úlfar Erlingsson, Alina Oprea, and Colin Raffel. 2021. [Extracting training data from large language models](#). In *30th USENIX Security Symposium (USENIX Security 21)*, pages 2633–2650. USENIX Association.
- Jiaao Chen and Diyi Yang. 2023. [Unlearn what you want to forget: Efficient unlearning for llms](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 12041–12052.
- Minseok Choi, Kyunghyun Min, and Jaegul Choo. 2024. [Cross-lingual unlearning of selective knowledge in multilingual language models](#). *arXiv preprint arXiv:2406.12354*.
- Peter Clark, Isaac Cowhey, Oren Etzioni, Tushar Khot, Ashish Sabharwal, Carissa Schoenick, and Oyvind Tafjord. 2018. [Think you have solved question answering? try arc, the ai2 reasoning challenge](#). *arXiv preprint arXiv:1803.05457*.
- Roi Cohen, Eden Biran, Ori Yoran, Amir Globerson, and Mor Geva. 2024. [Evaluating the ripple effects of knowledge editing in language models](#). *Transactions of the Association for Computational Linguistics*, 12:283–298.
- Nicola De Cao, Wilker Aziz, and Ivan Titov. 2021. [Editing factual knowledge in language models](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 6491–6506, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. 2024. [The llama 3 herd of models](#). *arXiv preprint arXiv:2407.21783*.
- Nouha Dziri, Ximing Lu, Melanie Sclar, Xiang Lorraine Li, Liwei Jiang, Bill Yuchen Lin, Sean Welleck, Peter West, Chandra Bhagavatula, Ronan Le Bras, Jena D. Hwang, Soumya Sanyal, Xiang Ren, Allyson Ettinger, Zaid Harchaoui, and Yejin Choi. 2023. [Faith and fate: Limits of transformers on compositionality](#). In *Thirty-seventh Conference on Neural Information Processing Systems*.
- Ronen Eldan and Mark Russinovich. 2023. [Who’s harry potter? approximate unlearning in llms](#). *arXiv preprint arXiv:2310.02238*.
- Hady Elsahar, Pavlos Vougiouklis, Arslan Remaci, Christophe Gravier, Jonathon Hare, Frederique Laforest, and Elena Simperl. 2018. [T-REx: A large scale alignment of natural language with knowledge base triples](#). In *Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)*, Miyazaki, Japan. European Language Resources Association (ELRA).



|     |                                                                                                                                                                                                                                                                                                                                                                         |     |
|-----|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| 736 | Patrick Lewis, Barlas Oguz, Ruty Rinott, Sebastian Riedel, and Holger Schwenk. 2020. <a href="#">MLQA: Evaluating cross-lingual extractive question answering</a> . In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 7315–7330, Online. Association for Computational Linguistics.                             | 790 |
| 737 |                                                                                                                                                                                                                                                                                                                                                                         | 791 |
| 738 |                                                                                                                                                                                                                                                                                                                                                                         | 792 |
| 739 |                                                                                                                                                                                                                                                                                                                                                                         | 793 |
| 740 |                                                                                                                                                                                                                                                                                                                                                                         | 794 |
| 741 |                                                                                                                                                                                                                                                                                                                                                                         | 795 |
| 742 |                                                                                                                                                                                                                                                                                                                                                                         | 796 |
| 743 | LG AI Research, Soyoung An, Kyunghoon Bae, Eunbi Choi, Kibong Choi, Stanley Jungkyu Choi, Seokhee Hong, Junwon Hwang, Hyojin Jeon, Gerrard Jeongwon Jo, et al. 2024. <a href="#">Exaone 3.5: Series of large language models for real-world use cases</a> . <i>arXiv preprint arXiv:2412.04862</i> .                                                                    | 797 |
| 744 |                                                                                                                                                                                                                                                                                                                                                                         |     |
| 745 |                                                                                                                                                                                                                                                                                                                                                                         | 798 |
| 746 |                                                                                                                                                                                                                                                                                                                                                                         | 799 |
| 747 |                                                                                                                                                                                                                                                                                                                                                                         | 800 |
| 748 |                                                                                                                                                                                                                                                                                                                                                                         | 801 |
| 749 | Chin-Yew Lin. 2004. <a href="#">ROUGE: A package for automatic evaluation of summaries</a> . In <i>Text Summarization Branches Out</i> , pages 74–81, Barcelona, Spain. Association for Computational Linguistics.                                                                                                                                                      | 802 |
| 750 |                                                                                                                                                                                                                                                                                                                                                                         | 803 |
| 751 |                                                                                                                                                                                                                                                                                                                                                                         | 804 |
| 752 |                                                                                                                                                                                                                                                                                                                                                                         | 805 |
| 753 |                                                                                                                                                                                                                                                                                                                                                                         | 806 |
| 754 | Yujian Liu, Yang Zhang, Tommi Jaakkola, and Shiyu Chang. 2024a. <a href="#">Revisiting who’s harry potter: Towards targeted unlearning from a causal intervention perspective</a> . In <i>Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing</i> , pages 8708–8731, Miami, Florida, USA. Association for Computational Linguistics. | 807 |
| 755 |                                                                                                                                                                                                                                                                                                                                                                         | 808 |
| 756 |                                                                                                                                                                                                                                                                                                                                                                         | 809 |
| 757 |                                                                                                                                                                                                                                                                                                                                                                         |     |
| 758 |                                                                                                                                                                                                                                                                                                                                                                         | 810 |
| 759 |                                                                                                                                                                                                                                                                                                                                                                         | 811 |
| 760 | Zheyuan Liu, Guangyao Dou, Zhaoxuan Tan, Yijun Tian, and Meng Jiang. 2024b. <a href="#">Towards safer large language models through machine unlearning</a> . In <i>Findings of the Association for Computational Linguistics: ACL 2024</i> , pages 1817–1829, Bangkok, Thailand. Association for Computational Linguistics.                                             | 812 |
| 761 |                                                                                                                                                                                                                                                                                                                                                                         | 813 |
| 762 |                                                                                                                                                                                                                                                                                                                                                                         | 814 |
| 763 |                                                                                                                                                                                                                                                                                                                                                                         | 815 |
| 764 |                                                                                                                                                                                                                                                                                                                                                                         | 816 |
| 765 |                                                                                                                                                                                                                                                                                                                                                                         | 817 |
| 766 | Ilya Loshchilov and Frank Hutter. 2019. <a href="#">Decoupled weight decay regularization</a> . In <i>International Conference on Learning Representations</i> .                                                                                                                                                                                                        | 818 |
| 767 |                                                                                                                                                                                                                                                                                                                                                                         | 819 |
| 768 |                                                                                                                                                                                                                                                                                                                                                                         |     |
| 769 | Pratyush Maini, Zhili Feng, Avi Schwarzschild, Zachary Chase Lipton, and J Zico Kolter. 2024. <a href="#">TOFU: A task of fictitious unlearning for LLMs</a> . In <i>First Conference on Language Modeling</i> .                                                                                                                                                        | 820 |
| 770 |                                                                                                                                                                                                                                                                                                                                                                         | 821 |
| 771 |                                                                                                                                                                                                                                                                                                                                                                         | 822 |
| 772 |                                                                                                                                                                                                                                                                                                                                                                         | 823 |
| 773 | Kevin Meng, David Bau, Alex J Andonian, and Yonatan Belinkov. 2022. <a href="#">Locating and editing factual associations in GPT</a> . In <i>Advances in Neural Information Processing Systems</i> .                                                                                                                                                                    | 824 |
| 774 |                                                                                                                                                                                                                                                                                                                                                                         | 825 |
| 775 |                                                                                                                                                                                                                                                                                                                                                                         | 826 |
| 776 |                                                                                                                                                                                                                                                                                                                                                                         | 827 |
| 777 | Todor Mihaylov, Peter Clark, Tushar Khot, and Ashish Sabharwal. 2018. <a href="#">Can a suit of armor conduct electricity? a new dataset for open book question answering</a> . In <i>Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing</i> , pages 2381–2391, Brussels, Belgium. Association for Computational Linguistics.       | 828 |
| 778 |                                                                                                                                                                                                                                                                                                                                                                         | 829 |
| 779 |                                                                                                                                                                                                                                                                                                                                                                         | 830 |
| 780 |                                                                                                                                                                                                                                                                                                                                                                         | 831 |
| 781 |                                                                                                                                                                                                                                                                                                                                                                         | 832 |
| 782 |                                                                                                                                                                                                                                                                                                                                                                         |     |
| 783 |                                                                                                                                                                                                                                                                                                                                                                         | 833 |
| 784 | Mistral AI. 2024. Un ministral, des ministraux. <a href="https://mistral.ai/en/news/ministraux">https://mistral.ai/en/news/ministraux</a> .                                                                                                                                                                                                                             | 834 |
| 785 |                                                                                                                                                                                                                                                                                                                                                                         | 835 |
| 786 | Team OLMo, Pete Walsh, Luca Soldaini, Dirk Groeneveld, Kyle Lo, Shane Arora, Akshita Bhagia, Yuling Gu, Shengyi Huang, Matt Jordan, et al. 2024. <a href="#">2 olmo 2 furious</a> . <i>arXiv preprint arXiv:2501.00656</i> .                                                                                                                                            | 836 |
| 787 |                                                                                                                                                                                                                                                                                                                                                                         | 837 |
| 788 |                                                                                                                                                                                                                                                                                                                                                                         | 838 |
| 789 |                                                                                                                                                                                                                                                                                                                                                                         | 839 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 840 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 841 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 842 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 843 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 844 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 845 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 846 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 847 |
|     |                                                                                                                                                                                                                                                                                                                                                                         |     |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 790 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 791 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 792 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 793 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 794 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 795 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 796 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 797 |
|     |                                                                                                                                                                                                                                                                                                                                                                         |     |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 798 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 799 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 800 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 801 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 802 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 803 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 804 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 805 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 806 |
|     |                                                                                                                                                                                                                                                                                                                                                                         |     |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 807 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 808 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 809 |
|     |                                                                                                                                                                                                                                                                                                                                                                         |     |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 810 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 811 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 812 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 813 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 814 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 815 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 816 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 817 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 818 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 819 |
|     |                                                                                                                                                                                                                                                                                                                                                                         |     |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 820 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 821 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 822 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 823 |
|     |                                                                                                                                                                                                                                                                                                                                                                         |     |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 824 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 825 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 826 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 827 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 828 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 829 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 830 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 831 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 832 |
|     |                                                                                                                                                                                                                                                                                                                                                                         |     |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 833 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 834 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 835 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 836 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 837 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 838 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 839 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 840 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 841 |
|     |                                                                                                                                                                                                                                                                                                                                                                         |     |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 842 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 843 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 844 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 845 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 846 |
|     |                                                                                                                                                                                                                                                                                                                                                                         | 847 |

|     |                                                                             |                                                                             |     |
|-----|-----------------------------------------------------------------------------|-----------------------------------------------------------------------------|-----|
| 848 | Rafael Rafailov, Archit Sharma, Eric Mitchell, Christo-                     | 2022 Foundation Models for Decision Making Work-                            | 905 |
| 849 | pher D Manning, Stefano Ermon, and Chelsea Finn.                            | shop.                                                                       | 906 |
| 850 | 2023. <a href="#">Direct preference optimization: Your language</a>         |                                                                             |     |
| 851 | <a href="#">model is secretly a reward model</a> . In <i>Thirty-seventh</i> | Denny Vrandečić and Markus Krötzsch. 2014. <a href="#">Wiki-</a>            | 907 |
| 852 | <i>Conference on Neural Information Processing Sys-</i>                     | <a href="#">data: a free collaborative knowledgebase</a> . <i>Communi-</i>  | 908 |
| 853 | <i>tems</i> .                                                               | <i>cations of the ACM</i> , 57(10):78–85.                                   | 909 |
| 854 | Pranav Rajpurkar, Jian Zhang, Konstantin Lopyrev, and                       | Lingzhi Wang, Tong Chen, Wei Yuan, Xingshan Zeng,                           | 910 |
| 855 | Percy Liang. 2016. <a href="#">SQuAD: 100,000+ questions for</a>            | Kam-Fai Wong, and Hongzhi Yin. 2023. <a href="#">KGA:</a>                   | 911 |
| 856 | <a href="#">machine comprehension of text</a> . In <i>Proceedings of</i>    | <a href="#">A general machine unlearning framework based on</a>             | 912 |
| 857 | <i>the 2016 Conference on Empirical Methods in Natu-</i>                    | <a href="#">knowledge gap alignment</a> . In <i>Proceedings of the 61st</i> | 913 |
| 858 | <i>ral Language Processing</i> , pages 2383–2392, Austin,                   | <i>Annual Meeting of the Association for Computational</i>                  | 914 |
| 859 | Texas. Association for Computational Linguistics.                           | <i>Linguistics (Volume 1: Long Papers)</i> , pages 13264–                   | 915 |
| 860 | Jeffrey Rosen. 2011. The right to be forgotten. <i>Stan. L.</i>             | 13276, Toronto, Canada. Association for Computa-                            | 916 |
| 861 | <i>Rev. Online</i> , 64:88.                                                 | tional Linguistics.                                                         | 917 |
| 862 | Keisuke Sakaguchi, Ronan Le Bras, Chandra Bhagavat-                         | Thomas Wolf, Lysandre Debut, Victor Sanh, Julien                            | 918 |
| 863 | ula, and Yejin Choi. 2021. <a href="#">Winogrande: An adver-</a>            | Chaumond, Clement Delangue, Anthony Moi, Pier-                              | 919 |
| 864 | <a href="#">sarial winograd schema challenge at scale</a> . <i>Communi-</i> | ric Cistac, Tim Rault, Remi Louf, Morgan Funtow-                            | 920 |
| 865 | <i>cations of the ACM</i> , 64(9):99–106.                                   | icz, Joe Davison, Sam Shleifer, Patrick von Platen,                         | 921 |
| 866 | Mansi Sakarvadia, Aswathy Ajith, Arham Khan, Daniel                         | Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu,                            | 922 |
| 867 | Grzenda, Nathaniel Hudson, André Bauer, Kyle                                | Teven Le Scao, Sylvain Gugger, Mariama Drame,                               | 923 |
| 868 | Chard, and Ian Foster. 2023. <a href="#">Memory injections:</a>             | Quentin Lhoest, and Alexander Rush. 2020. <a href="#">Trans-</a>            | 924 |
| 869 | <a href="#">Correcting multi-hop reasoning failures during infer-</a>       | <a href="#">formers: State-of-the-art natural language processing</a> .     | 925 |
| 870 | <a href="#">ence in transformer-based language models</a> . In              | In <i>Proceedings of the 2020 Conference on Empirical</i>                   | 926 |
| 871 | <i>Proceedings of the 6th BlackboxNLP Workshop: An-</i>                     | <i>Methods in Natural Language Processing: System</i>                       | 927 |
| 872 | <i>alyzing and Interpreting Neural Networks for NLP</i> ,                   | <i>Demonstrations</i> , pages 38–45, Online. Association                    | 928 |
| 873 | pages 342–356, Singapore. Association for Computa-                          | for Computational Linguistics.                                              | 929 |
| 874 | tional Linguistics.                                                         | An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui,                          | 930 |
| 875 | Alon Talmor, Jonathan Herzig, Nicholas Lourie, and                          | Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu,                             | 931 |
| 876 | Jonathan Berant. 2019. <a href="#">CommonsenseQA: A ques-</a>               | Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jian-                           | 932 |
| 877 | <a href="#">tion answering challenge targeting commonsense</a>              | hong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang,                           | 933 |
| 878 | <a href="#">knowledge</a> . In <i>Proceedings of the 2019 Conference</i>    | Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu,                             | 934 |
| 879 | <i>of the North American Chapter of the Association for</i>                 | Keqin Bao, Kexin Yang, Le Yu, Mei Li, Mingfeng                              | 935 |
| 880 | <i>Computational Linguistics: Human Language Tech-</i>                      | Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tian-                          | 936 |
| 881 | <i>nologies, Volume 1 (Long and Short Papers)</i> , pages                   | hao Li, Tingyu Xia, Xingzhang Ren, Xuancheng                                | 937 |
| 882 | 4149–4158, Minneapolis, Minnesota. Association for                          | Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan,                              | 938 |
| 883 | Computational Linguistics.                                                  | Yuqiong Liu, Zeyu Cui, Zhenru Zhang, and Zihan                              | 939 |
| 884 | Pratiksha Thaker, Yash Maurya, and Virginia Smith.                          | Qiu. 2024. <a href="#">Qwen2.5 technical report</a> . <i>arXiv preprint</i> | 940 |
| 885 | 2024. <a href="#">Guardrail baselines for unlearning in LLMs</a> .          | <i>arXiv:2412.15115</i> .                                                   | 941 |
| 886 | In <i>ICLR 2024 Workshop on Secure and Trustworthy</i>                      | Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Bengio,                        | 942 |
| 887 | <i>Large Language Models</i> .                                              | William Cohen, Ruslan Salakhutdinov, and Christo-                           | 943 |
| 888 | Bozhong Tian, Xiaozhuan Liang, Siyuan Cheng, Qing-                          | pher D. Manning. 2018. <a href="#">HotpotQA: A dataset for</a>              | 944 |
| 889 | bin Liu, Mengru Wang, Dianbo Sui, Xi Chen, Hua-                             | <a href="#">diverse, explainable multi-hop question answering</a> .         | 945 |
| 890 | jun Chen, and Ningyu Zhang. 2024. <a href="#">To forget or</a>              | In <i>Proceedings of the 2018 Conference on Empiri-</i>                     | 946 |
| 891 | <a href="#">not? towards practical knowledge unlearning for</a>             | <i>cal Methods in Natural Language Processing</i> , pages                   | 947 |
| 892 | <a href="#">large language models</a> . In <i>Findings of the Associ-</i>   | 2369–2380, Brussels, Belgium. Association for Com-                          | 948 |
| 893 | <i>ation for Computational Linguistics: EMNLP 2024</i> ,                    | putational Linguistics.                                                     | 949 |
| 894 | pages 1524–1537, Miami, Florida, USA. Association                           | Jin Yao, Eli Chien, Minxin Du, Xinyao Niu, Tianhao                          | 950 |
| 895 | for Computational Linguistics.                                              | Wang, Zezhou Cheng, and Xiang Yue. 2024. <a href="#">Ma-</a>                | 951 |
| 896 | Harsh Trivedi, Niranjan Balasubramanian, Tushar Khot,                       | <a href="#">chine unlearning of pre-trained large language mod-</a>         | 952 |
| 897 | and Ashish Sabharwal. 2022. <a href="#">MuSiQue: Multi-</a>                 | <a href="#">els</a> . In <i>Proceedings of the 62nd Annual Meeting of</i>   | 953 |
| 898 | <a href="#">hop questions via single-hop question composition</a> .         | <i>the Association for Computational Linguistics (Vol-</i>                  | 954 |
| 899 | <i>Transactions of the Association for Computational</i>                    | <i>ume 1: Long Papers)</i> , pages 8403–8419, Bangkok,                      | 955 |
| 900 | <i>Linguistics</i> , 10:539–554.                                            | Thailand. Association for Computational Linguistics.                        | 956 |
| 901 | Karthik Valmeekam, Alberto Olmo, Sarath Sreedharan,                         | Yuanshun Yao, Xiaojun Xu, and Yang Liu. 2023. <a href="#">Large</a>         | 957 |
| 902 | and Subbarao Kambhampati. 2022. <a href="#">Large language</a>              | <a href="#">language model unlearning</a> . In <i>Socially Responsible</i>  | 958 |
| 903 | <a href="#">models still can’t plan (a benchmark for LLMs on</a>            | <i>Language Modelling Research</i> .                                        | 959 |
| 904 | <a href="#">planning and reasoning about change)</a> . In <i>NeurIPS</i>    | Zheng Xin Yong, Cristina Menghini, and Stephen Bach.                        | 960 |
|     |                                                                             | 2023. <a href="#">Low-resource languages jailbreak GPT-4</a> . In           | 961 |
|     |                                                                             | <i>Socially Responsible Language Modelling Research</i> .                   | 962 |

Rowan Zellers, Ari Holtzman, Yonatan Bisk, Ali Farhadi, and Yejin Choi. 2019. [HellaSwag: Can a machine really finish your sentence?](#) In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 4791–4800, Florence, Italy. Association for Computational Linguistics.

Hanning Zhang, Shizhe Diao, Yong Lin, Yi Fung, Qing Lian, Xingyao Wang, Yangyi Chen, Heng Ji, and Tong Zhang. 2024a. [R-tuning: Instructing large language models to say ‘I don’t know’](#). In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pages 7113–7139, Mexico City, Mexico. Association for Computational Linguistics.

Ningyu Zhang, Yunzhi Yao, and Shumin Deng. 2024b. [Knowledge editing for large language models](#). In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024): Tutorial Summaries*, pages 33–41, Torino, Italia. ELRA and ICCL.

Ruiqi Zhang, Licong Lin, Yu Bai, and Song Mei. 2024c. [Negative preference optimization: From catastrophic collapse to effective unlearning](#). In *First Conference on Language Modeling*.

Zexuan Zhong, Zhengxuan Wu, Christopher Manning, Christopher Potts, and Danqi Chen. 2023. [MQuAKE: Assessing knowledge editing in language models via multi-hop questions](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 15686–15702, Singapore. Association for Computational Linguistics.

Denny Zhou, Nathanael Schärli, Le Hou, Jason Wei, Nathan Scales, Xuezhi Wang, Dale Schuurmans, Claire Cui, Olivier Bousquet, Quoc V Le, and Ed H. Chi. 2023. [Least-to-most prompting enables complex reasoning in large language models](#). In *The Eleventh International Conference on Learning Representations*.

## A Additional Details for Knowledge Unlearning Methods

### A.1 GA

Gradient ascent (GA) (Jang et al., 2023) reverses the language modeling loss, which can be understood as equivalent to gradient descent on the negative next-token prediction loss:

$$\mathcal{L}_{\text{GA}} = -\mathbb{E}_{\mathcal{D}_f}[-\log(\pi_{\theta}(y|x))], \quad (1)$$

which serves to minimize the token probabilities of the specific token sequences in the forget set  $\mathcal{D}_f$ .

### A.2 DPO

In direct preference optimization (DPO) (Rafailov et al., 2023), we are provided with a dataset of preference feedbacks  $\mathcal{D}_{\text{paired}} = \{(x_i, y_{i,w}, y_{i,l})\}_{i=1}^N$ , where “w” stands for “win” and “l” stands for “lose” for two responses  $y_w$  and  $y_l$ . The goal is to train the model  $\pi_{\theta}$  to align more closely with human preferences. In this work, the winning responses are rejections (e.g., “I don’t know.”), randomly sampled from 100 candidates used by Maini et al. (2024). Formally, DPO minimizes

$$\mathcal{L}_{\text{DPO}} = -\mathbb{E}_{\mathcal{D}_{\text{paired}}} \left[ \log \sigma \left( \beta \log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \beta \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right) \right], \quad (2)$$

where  $\sigma(t) = 1/(1 + e^{-t})$  represents the sigmoid function,  $\beta > 0$  is the inverse temperature, and  $\phi_{\text{ref}}$  is a reference model.

### A.3 NPO

Negative preference optimization (NPO) (Zhang et al., 2024c) ignores the  $y_w$  term in DPO in Equation 2 and aligns the language model with negative responses exclusively:

$$\mathcal{L}_{\text{NPO}} = -\mathbb{E}_{\mathcal{D}_f} \left[ \log \sigma \left( -\beta \log \frac{\pi_{\theta}(y|x)}{\pi_{\text{ref}}(y|x)} \right) \right]. \quad (3)$$

Minimizing  $\mathcal{L}_{\text{NPO}}$  drives the prediction probability  $\pi_{\theta}(y|x)$  on the forget set to be as low as possible, effectively achieving the goal of unlearning the forget set.

### A.4 +RT

The explicit retention finetuning is achieved through standard language modeling on the retain set, which serves as the positive counterpart to Equation 1:

$$\mathcal{L}_r = -\mathbb{E}_{\mathcal{D}_r}[\log(\pi_{\theta}(y|x))]. \quad (4)$$

Finally, the overall training objective is minimizing the following loss:

$$\mathcal{L} = \alpha \cdot \mathcal{L}_f + (1 - \alpha) \cdot \mathcal{L}_r, \quad (5)$$

where  $\mathcal{L}_f$  is one of the unlearning losses  $\mathcal{L}_{\text{GA}}$ ,  $\mathcal{L}_{\text{DPO}}$ , or  $\mathcal{L}_{\text{NPO}}$ , and  $\alpha$  is a loss scaling hyperparameter balancing the forgetting and retaining losses.

## B Dataset Details

We describe further details in datasets and their preprocessing processes for our experiments.

**MQuAKE** (Zhong et al., 2023) considers whether models can adapt to updates in factual knowledge by modifying their responses to multi-hop queries when individual facts are altered. The benchmark comprises two datasets: MQuAKE-CF, which focuses on counterfactual scenarios, and MQuAKE-T, which addresses temporal knowledge updates by replacing outdated facts with current information. Both datasets are based on Wikidata and consist of knowledge triplets for single-hop reasoning, as well as multi-hop chains derived from these triplets. Each instance in the benchmark includes: (1) an edit set of single-hop knowledge triplets  $(s, r, o \rightarrow o^*)$ , where  $o^*$  represents the updated object; (2) a chain of facts  $\mathcal{C}$  and its updated version  $\mathcal{C}^*$  after knowledge editing; and (3) questions about both the single-hop knowledge and multi-hop chains before and after knowledge updates. We used the merged set of both MQuAKE-CF and MQuAKE-T for our experiments and only used the single-hop and multi-hop triplets before the knowledge update.

**MuSiQue** (Trivedi et al., 2022) is a benchmark designed to ensure that multi-hop QA models genuinely perform multi-hop reasoning rather than relying on single-hop shortcuts. It constructs multi-hop questions by carefully composing independent single-hop questions, ensuring that models must retrieve and integrate information from multiple documents. Built on Wikipedia, MuSiQue provides a challenging evaluation for multi-hop QA models. We used the answerable subset of MuSiQue, referred to as MuSiQue-Ans, for our experiments. We noticed that some single-hop triplets in MuSiQue lack a natural language question and are instead represented only by a subject-relation  $(s, r)$  pair. To address this, we used gpt-4o-mini with few-shot demonstrations to generate a corresponding question from each  $(s, r)$  pair. For example, the pair (*Just Ask Your Heart*, *performer*) is transformed into the question: "Who performed the song '*Just Ask Your Heart*'?".

To adapt both datasets for multi-hop knowledge unlearning, we preprocessed the data by: (1) splitting the single-hop triplets into a forget set and a retain set at a predefined ratio (1:9), with the retain set further divided into training, validation, and test splits and (2) We also ensured that multi-hop ques-

tions are linked to the corresponding single-hop triplets from both the forget and retain sets. If a multi-hop question contains both a forget and retain triplet, it was assigned exclusively to the forget set, ensuring that the final multi-hop forget and retain sets remain mutually exclusive. For MQuAKE dataset, as the frequency of single-hop triplets is imbalanced, we ensured that any triplet involved in more than two multi-hop questions is assigned to the retain set’s training split.

## C Full Evaluation Results

We report additional experimental results with different open-source LLMs including Phi-3.5-Mini-Instruct (Abdin et al., 2024), EXAONE-3.5-7.8B-Instruct (LG AI Research et al., 2024), Llama-3.1-8B-Instruct (Dubey et al., 2024), and Ministral-8B-Instruct (Mistral AI, 2024) in Tables 7 and 8. Furthermore, we display the utility performance of the model according to individual LLM benchmarks in Tables 9 and 10.

## D Licenses and Terms of Use for Artifacts

We utilized multiple datasets and open-source LLMs, each governed by specific licensing terms. The MQuAKE repository is available under the MIT License, while MuSiQue is licensed under CC BY 4.0, permitting academic and research use. All open-source LLMs referenced in this paper, including OLMo, Qwen, Phi, EXAONE, Llama, and Ministral, are freely licensed for research purposes.

## E Use of AI Assistants

We leveraged AI assistants, including ChatGPT and Copilot, to enhance research, writing, and coding. ChatGPT played a key role in refining the paper’s narrative and ensuring clarity, while Copilot accelerated the coding process. All AI tools were used responsibly, with careful oversight to uphold the integrity and originality of the research. Their usage has been documented to ensure transparency and proper acknowledgment of AI contributions.

|                                 | Forget Set (Single-Hop) |            |             | Forget Set (Multi-Hop) |            |             | Retain Set (Single-Hop) |             |            | Retain Set (Multi-Hop) |             |            | Utility Set |
|---------------------------------|-------------------------|------------|-------------|------------------------|------------|-------------|-------------------------|-------------|------------|------------------------|-------------|------------|-------------|
|                                 | PA(↓)                   | R-L(↓)     | LM(↑)       | PA(↓)                  | R-L(↓)     | LM(↑)       | PA(↑)                   | R-L(↑)      | LM(↓)      | PA(↑)                  | R-L(↑)      | LM(↓)      | Avg.(↑)     |
| <i>Phi-3.5-Mini-Instruct</i>    |                         |            |             |                        |            |             |                         |             |            |                        |             |            |             |
| Original                        | 82.7                    | 83.3       | 3.8         | 79.8                   | 56.2       | 3.3         | 87.7                    | 80.8        | 3.7        | 81.1                   | 51.6        | 3.4        | 65.0        |
| GA                              | 24.0                    | 15.6       | 24.4        | 36.5                   | 7.4        | 20.4        | 28.4                    | 19.2        | 22.9       | 32.0                   | 3.2         | 20.5       | 62.6        |
| DPO                             | 31.7                    | 0.7        | 12.0        | 43.3                   | 0.2        | 8.3         | 38.1                    | 0.9         | 11.4       | 41.8                   | 0.4         | 8.2        | 64.0        |
| NPO                             | 22.1                    | 19.5       | 24.4        | 33.7                   | 6.2        | 20.3        | 26.9                    | 18.4        | 22.9       | 32.5                   | 4.1         | 20.3       | 62.7        |
| GA+RT                           | 39.7                    | 41.1       | 7.8         | 87.8                   | 31.8       | 3.1         | 78.6                    | 61.7        | 3.0        | <b>88.9</b>            | 36.6        | 2.8        | <b>64.7</b> |
| DPO+RT                          | 44.9                    | <b>7.8</b> | 6.7         | <b>77.2</b>            | <b>4.6</b> | 3.2         | <b>80.4</b>             | 30.0        | <b>2.7</b> | 80.6                   | 6.6         | <b>2.7</b> | 64.6        |
| NPO+RT                          | <b>39.1</b>             | 44.9       | <b>8.5</b>  | 87.2                   | 33.5       | <b>3.5</b>  | 78.2                    | <b>65.5</b> | 3.2        | 87.7                   | <b>38.7</b> | 3.2        | 64.6        |
| <i>EXAONE-3.5-7.8B-Instruct</i> |                         |            |             |                        |            |             |                         |             |            |                        |             |            |             |
| Original                        | 92.3                    | 82.4       | 3.6         | 97.1                   | 49.1       | 3.3         | 90.2                    | 76.6        | 3.5        | 95.5                   | 46.1        | 3.4        | 62.6        |
| GA                              | 22.1                    | 0.0        | 87.4        | 21.2                   | 0.0        | 88.4        | 24.9                    | 0.0         | 87.4       | 18.4                   | 0.0         | 88.4       | 58.7        |
| DPO                             | 18.3                    | 0.6        | 53.2        | 17.3                   | 0.9        | 52.8        | 20.3                    | 1.0         | 52.8       | 14.3                   | 0.7         | 53.0       | 60.4        |
| NPO                             | 21.2                    | 0.0        | 86.9        | 19.2                   | 0.0        | 87.9        | 25.1                    | 0.0         | 86.8       | 18.7                   | 0.0         | 87.9       | 58.7        |
| GA+RT                           | 35.3                    | 44.2       | <b>12.3</b> | 86.2                   | 40.8       | <b>3.9</b>  | 79.2                    | 66.4        | 4.2        | 91.8                   | 38.2        | 3.2        | <b>62.1</b> |
| DPO+RT                          | <b>30.4</b>             | <b>3.6</b> | 9.2         | <b>77.9</b>            | <b>1.8</b> | 3.3         | <b>80.1</b>             | 23.7        | <b>2.4</b> | 89.5                   | 4.8         | <b>2.5</b> | <b>62.1</b> |
| NPO+RT                          | 37.2                    | 42.2       | 12.3        | 85.3                   | 41.0       | 3.9         | 79.8                    | <b>67.6</b> | 4.0        | <b>92.1</b>            | <b>38.8</b> | 3.2        | <b>62.1</b> |
| <i>Llama-3.1-8B-Instruct</i>    |                         |            |             |                        |            |             |                         |             |            |                        |             |            |             |
| Original                        | 99.0                    | 60.8       | 0.6         | 99.0                   | 31.2       | 1.0         | 98.9                    | 60.1        | 0.6        | 98.4                   | 28.6        | 1.0        | 64.7        |
| GA                              | 33.7                    | 0.0        | 88.7        | 28.8                   | 0.0        | 87.0        | 34.3                    | 0.0         | 88.3       | 24.5                   | 0.0         | 87.1       | 62.9        |
| DPO                             | 10.6                    | 2.6        | 39.8        | 11.5                   | 2.4        | 39.0        | 15.5                    | 2.2         | 39.5       | 12.5                   | 1.4         | 39.0       | 61.8        |
| NPO                             | 30.8                    | 0.0        | 89.1        | 30.8                   | 0.0        | 87.7        | 33.7                    | 0.0         | 88.9       | 25.8                   | 0.0         | 87.7       | 62.8        |
| GA+RT                           | 56.4                    | 49.4       | <b>13.8</b> | 87.2                   | 28.7       | 5.8         | 86.6                    | <b>79.3</b> | 5.7        | <b>91.2</b>            | <b>29.9</b> | <b>4.9</b> | <b>64.7</b> |
| DPO+RT                          | <b>49.4</b>             | <b>9.3</b> | 13.3        | <b>86.9</b>            | <b>9.8</b> | <b>6.1</b>  | <b>88.3</b>             | 44.7        | 5.1        | 90.7                   | 13.2        | 5.0        | 63.7        |
| NPO+RT                          | 56.7                    | 48.0       | 12.5        | 88.1                   | 26.3       | 5.7         | 87.5                    | 78.5        | <b>4.9</b> | 90.1                   | 29.6        | 4.9        | <b>64.7</b> |
| <i>Ministral-8B-Instruct</i>    |                         |            |             |                        |            |             |                         |             |            |                        |             |            |             |
| Original                        | 98.1                    | 85.3       | 0.9         | 99.0                   | 46.9       | 0.8         | 97.5                    | 84.8        | 0.8        | 98.2                   | 41.9        | 0.8        | 64.5        |
| GA                              | 23.1                    | 0.0        | 69.6        | 28.8                   | 0.0        | 70.6        | 20.5                    | 0.0         | 69.7       | 23.6                   | 0.0         | 70.6       | 41.9        |
| DPO                             | 15.4                    | 0.0        | 38.7        | 21.2                   | 0.9        | 38.1        | 12.3                    | 1.6         | 38.6       | 14.0                   | 0.7         | 38.2       | 44.0        |
| NPO                             | 22.1                    | 0.0        | 69.8        | 30.8                   | 0.0        | 70.8        | 21.0                    | 0.0         | 69.8       | 23.5                   | 0.0         | 70.7       | 42.0        |
| GA+RT                           | <b>34.0</b>             | 17.1       | <b>18.4</b> | 77.9                   | 9.8        | <b>10.5</b> | 76.9                    | 41.3        | 6.9        | 87.8                   | <b>13.9</b> | 6.0        | <b>59.3</b> |
| DPO+RT                          | 36.9                    | <b>1.2</b> | 14.3        | <b>75.3</b>            | <b>1.4</b> | 7.6         | <b>82.8</b>             | 5.1         | 5.1        | 88.4                   | 2.4         | 4.9        | 56.9        |
| NPO+RT                          | 36.2                    | 19.6       | 14.8        | 81.1                   | 11.6       | 7.4         | 80.8                    | <b>45.6</b> | <b>4.6</b> | <b>90.2</b>            | 13.6        | <b>4.2</b> | 59.2        |

Table 7: Performance comparison of different knowledge unlearning methods after erasing single-hop facts from the forget set in MQAKE-NoEdit. The best results amongst +RT models are highlighted in **bold**.

|                                 | Forget Set (Single-Hop) |            |             | Forget Set (Multi-Hop) |            |             | Retain Set (Single-Hop) |             |            | Retain Set (Multi-Hop) |             |            | Utility Set |
|---------------------------------|-------------------------|------------|-------------|------------------------|------------|-------------|-------------------------|-------------|------------|------------------------|-------------|------------|-------------|
|                                 | PA(↓)                   | R-L(↓)     | LM(↑)       | PA(↓)                  | R-L(↓)     | LM(↑)       | PA(↑)                   | R-L(↑)      | LM(↓)      | PA(↑)                  | R-L(↑)      | LM(↓)      | Avg.(↑)     |
| <i>Phi-3.5-Mini-Instruct</i>    |                         |            |             |                        |            |             |                         |             |            |                        |             |            |             |
| Original                        | 67.6                    | 34.5       | 4.1         | 67.7                   | 23.0       | 3.4         | 71.0                    | 34.0        | 4.1        | 69.9                   | 22.7        | 3.4        | 65.0        |
| GA                              | 21.4                    | 7.1        | 61.8        | 19.0                   | 7.9        | 52.6        | 19.4                    | 5.4         | 61.3       | 19.5                   | 6.7         | 53.5       | 64.1        |
| DPO                             | 11.6                    | 0.1        | 52.0        | 11.7                   | 0.5        | 43.3        | 14.6                    | 0.2         | 51.4       | 13.0                   | 0.5         | 44.2       | 64.5        |
| NPO                             | 20.2                    | 7.3        | 62.9        | 20.2                   | 8.5        | 53.7        | 19.4                    | 5.0         | 62.4       | 19.3                   | 6.5         | 54.6       | 64.0        |
| GA+RT                           | 31.2                    | 17.4       | 11.6        | 77.2                   | 16.2       | 4.8         | <b>82.5</b>             | 28.0        | <b>3.3</b> | <b>83.9</b>            | 18.0        | <b>4.2</b> | 64.1        |
| DPO+RT                          | <b>22.4</b>             | <b>2.2</b> | <b>14.5</b> | <b>63.5</b>            | <b>1.7</b> | <b>6.1</b>  | 73.4                    | 12.2        | 5.7        | 70.6                   | 2.3         | 5.7        | <b>64.5</b> |
| NPO+RT                          | 30.8                    | 19.1       | 13.1        | 75.8                   | 18.2       | 5.6         | 81.0                    | <b>29.4</b> | 3.8        | 82.2                   | <b>20.0</b> | 4.9        | 64.0        |
| <i>EXAONE-3.5-7.8B-Instruct</i> |                         |            |             |                        |            |             |                         |             |            |                        |             |            |             |
| Original                        | 76.3                    | 34.6       | 4.0         | 90.1                   | 18.9       | 3.5         | 77.0                    | 33.2        | 4.0        | 89.4                   | 21.1        | 3.6        | 62.6        |
| GA                              | 24.3                    | 0.0        | 105.4       | 14.4                   | 0.0        | 107.6       | 27.9                    | 0.0         | 105.4      | 16.1                   | 0.0         | 107.4      | 53.9        |
| DPO                             | 16.8                    | 1.4        | 79.1        | 12.1                   | 3.0        | 82.1        | 19.9                    | 1.0         | 78.7       | 13.0                   | 2.1         | 82.0       | 57.8        |
| NPO                             | 24.9                    | 0.0        | 99.0        | 13.8                   | 0.0        | 100.7       | 28.1                    | 0.0         | 99.0       | 15.7                   | 0.0         | 100.6      | 54.0        |
| GA+RT                           | 23.5                    | 13.7       | <b>20.7</b> | <b>70.6</b>            | 14.5       | <b>9.7</b>  | 80.1                    | 23.0        | 5.9        | 84.0                   | 16.5        | 7.2        | 59.5        |
| DPO+RT                          | <b>11.6</b>             | <b>1.6</b> | 12.1        | 72.7                   | <b>3.4</b> | 4.0         | 80.8                    | 9.1         | <b>2.6</b> | <b>88.1</b>            | 4.4         | <b>2.6</b> | <b>61.7</b> |
| NPO+RT                          | 24.1                    | 15.2       | 19.5        | 72.1                   | 16.4       | 8.9         | <b>81.5</b>             | <b>26.7</b> | 5.1        | 85.6                   | <b>18.8</b> | 6.4        | 59.8        |
| <i>Llama-3.1-8B-Instruct</i>    |                         |            |             |                        |            |             |                         |             |            |                        |             |            |             |
| Original                        | 90.2                    | 35.9       | 1.1         | 96.5                   | 16.5       | 1.3         | 90.9                    | 32.8        | 1.1        | 96.2                   | 16.1        | 1.3        | 64.7        |
| GA                              | 30.1                    | 0.0        | 103.0       | 30.0                   | 0.0        | 102.0       | 29.4                    | 0.0         | 103.0      | 33.4                   | 0.0         | 102.1      | 60.5        |
| DPO                             | 28.3                    | 1.4        | 73.7        | 31.4                   | 1.1        | 70.7        | 28.9                    | 0.6         | 73.6       | 31.0                   | 0.9         | 71.1       | 59.9        |
| NPO                             | 25.4                    | 0.0        | 99.1        | 23.8                   | 0.0        | 98.3        | 26.0                    | 0.0         | 99.1       | 25.5                   | 0.0         | 98.4       | 60.5        |
| GA+RT                           | <b>32.2</b>             | 12.0       | <b>28.6</b> | <b>77.1</b>            | 11.8       | <b>14.4</b> | 80.7                    | 23.7        | 9.5        | 86.4                   | 13.3        | 10.4       | 61.7        |
| DPO+RT                          | 34.3                    | <b>5.2</b> | 24.3        | 79.1                   | <b>4.8</b> | 11.9        | <b>85.0</b>             | 18.7        | 6.6        | 90.1                   | 6.3         | 8.3        | <b>63.5</b> |
| NPO+RT                          | 37.4                    | 14.9       | 23.0        | 84.4                   | 14.0       | 10.3        | <b>85.0</b>             | <b>31.7</b> | <b>6.4</b> | <b>91.4</b>            | <b>18.3</b> | <b>7.1</b> | <b>63.5</b> |
| <i>Ministral-8B-Instruct</i>    |                         |            |             |                        |            |             |                         |             |            |                        |             |            |             |
| Original                        | 88.4                    | 40.8       | 1.1         | 96.3                   | 23.4       | 1.0         | 87.2                    | 38.7        | 1.1        | 96.2                   | 23.1        | 1.0        | 64.5        |
| GA                              | 22.5                    | 0.1        | 88.8        | 25.7                   | 0.2        | 89.5        | 32.5                    | 0.0         | 88.8       | 30.4                   | 0.3         | 89.4       | 30.7        |
| DPO                             | 35.5                    | 0.6        | 58.0        | 28.8                   | 1.0        | 60.9        | 34.4                    | 0.5         | 57.7       | 28.9                   | 0.5         | 60.6       | 36.2        |
| NPO                             | 23.1                    | 0.1        | 87.1        | 24.8                   | 0.2        | 87.6        | 30.4                    | 0.0         | 87.0       | 29.2                   | 0.3         | 87.5       | 31.1        |
| GA+RT                           | 28.9                    | 7.9        | <b>21.8</b> | 78.9                   | 10.2       | <b>11.7</b> | 78.4                    | 16.1        | 7.8        | 87.1                   | 10.3        | 8.1        | <b>55.4</b> |
| DPO+RT                          | <b>24.3</b>             | <b>3.7</b> | 15.7        | <b>79.4</b>            | <b>3.6</b> | 7.6         | 82.5                    | 9.9         | <b>4.5</b> | 89.5                   | 4.4         | <b>5.1</b> | 55.3        |
| NPO+RT                          | 28.1                    | 10.4       | 17.5        | 81.8                   | 12.3       | 7.8         | <b>82.9</b>             | <b>19.8</b> | 4.9        | <b>91.2</b>            | <b>12.2</b> | 5.1        | <b>55.4</b> |

Table 8: Performance comparison of different knowledge unlearning methods after erasing single-hop facts from the forget set in MuSiQue-Ans. The best results amongst +RT models are highlighted in **bold**.

|                                 | ARC-C | CSQA | Hella. | Lamba. | MMLU | OBQA | PIQA | Wino. | Avg. |
|---------------------------------|-------|------|--------|--------|------|------|------|-------|------|
| <i>OLMo-2-7B-Instruct</i>       |       |      |        |        |      |      |      |       |      |
| Original                        | 54.6  | 72.6 | 65.7   | 68.6   | 59.2 | 39.4 | 80.5 | 71.6  | 64.0 |
| GA+RT                           | 50.8  | 71.3 | 64.4   | 69.5   | 59.2 | 39.1 | 75.7 | 70.3  | 62.6 |
| DPO+RT                          | 50.9  | 70.7 | 64.5   | 66.4   | 58.9 | 39.2 | 77.5 | 71.4  | 62.4 |
| NPO+RT                          | 50.9  | 71.3 | 64.5   | 69.7   | 59.2 | 38.9 | 75.8 | 70.3  | 62.6 |
| <i>Qwen-2.5-7B-Instruct</i>     |       |      |        |        |      |      |      |       |      |
| Original                        | 52.6  | 82.6 | 62.1   | 69.4   | 71.8 | 34.8 | 79.3 | 71.5  | 65.5 |
| GA+RT                           | 56.3  | 83.4 | 62.1   | 71.8   | 71.5 | 37.7 | 79.0 | 70.4  | 66.5 |
| DPO+RT                          | 53.4  | 81.2 | 60.6   | 67.9   | 71.2 | 35.8 | 79.5 | 71.5  | 65.1 |
| NPO+RT                          | 56.1  | 83.2 | 62.1   | 71.9   | 71.5 | 37.3 | 79.1 | 70.7  | 66.5 |
| <i>Phi-3.5-Mini-Instruct</i>    |       |      |        |        |      |      |      |       |      |
| Original                        | 59.5  | 75.3 | 58.8   | 65.1   | 68.7 | 37.6 | 80.0 | 74.6  | 65.0 |
| GA+RT                           | 58.9  | 75.3 | 59.5   | 63.3   | 68.6 | 39.1 | 79.0 | 73.6  | 64.7 |
| DPO+RT                          | 59.8  | 74.4 | 58.2   | 60.4   | 68.5 | 38.4 | 80.3 | 76.4  | 64.6 |
| NPO+RT                          | 58.8  | 75.0 | 59.6   | 63.5   | 68.6 | 39.2 | 78.9 | 73.4  | 64.6 |
| <i>EXAONE-3.5-7.8B-Instruct</i> |       |      |        |        |      |      |      |       |      |
| Original                        | 56.9  | 75.4 | 60.2   | 62.4   | 65.2 | 35.6 | 76.9 | 68.0  | 62.6 |
| GA+RT                           | 56.0  | 75.0 | 59.7   | 61.5   | 64.6 | 36.9 | 76.3 | 67.3  | 62.1 |
| DPO+RT                          | 57.6  | 75.2 | 59.6   | 55.8   | 64.4 | 36.7 | 77.7 | 69.7  | 62.1 |
| NPO+RT                          | 55.7  | 75.2 | 59.7   | 61.3   | 64.7 | 36.7 | 76.1 | 67.3  | 62.1 |
| <i>Llama-3.1-8B-Instruct</i>    |       |      |        |        |      |      |      |       |      |
| Original                        | 51.8  | 77.1 | 59.2   | 73.2   | 68.1 | 33.8 | 80.2 | 74.1  | 64.7 |
| GA+RT                           | 53.2  | 75.8 | 59.2   | 74.7   | 67.2 | 35.9 | 79.1 | 72.7  | 64.7 |
| DPO+RT                          | 51.1  | 74.7 | 59.2   | 70.1   | 66.3 | 34.9 | 80.2 | 73.2  | 63.7 |
| NPO+RT                          | 53.5  | 75.5 | 59.3   | 74.7   | 67.3 | 35.7 | 79.2 | 72.8  | 64.7 |
| <i>Ministral-8B-Instruct</i>    |       |      |        |        |      |      |      |       |      |
| Original                        | 54.6  | 72.5 | 59.6   | 74.6   | 64.1 | 36.2 | 81.0 | 73.6  | 64.5 |
| GA+RT                           | 53.0  | 38.6 | 58.4   | 77.6   | 55.5 | 39.2 | 80.6 | 71.3  | 59.3 |
| DPO+RT                          | 51.8  | 46.4 | 57.8   | 49.7   | 59.2 | 38.4 | 80.7 | 71.3  | 56.9 |
| NPO+RT                          | 52.8  | 37.1 | 58.5   | 77.5   | 55.9 | 39.5 | 80.8 | 71.4  | 59.2 |

Table 9: Model utility performance per task after erasing single-hop facts from the forget set in MQuAKE-NoEdit.

|                                 | ARC-C | CSQA | Hella. | Lamba. | MMLU | OBQA | PIQA | Wino. | Avg. |
|---------------------------------|-------|------|--------|--------|------|------|------|-------|------|
| <i>OLMo-2-7B-Instruct</i>       |       |      |        |        |      |      |      |       |      |
| Original                        | 54.6  | 72.6 | 65.7   | 68.6   | 59.2 | 39.4 | 80.5 | 71.6  | 64.0 |
| GA+RT                           | 47.2  | 70.1 | 63.3   | 69.8   | 58.6 | 38.3 | 73.1 | 67.8  | 61.0 |
| DPO+RT                          | 46.8  | 69.2 | 63.2   | 68.1   | 58.6 | 37.4 | 73.1 | 68.3  | 60.6 |
| NPO+RT                          | 47.4  | 70.5 | 63.4   | 70.2   | 58.7 | 38.1 | 73.3 | 67.9  | 61.2 |
| <i>Qwen-2.5-7B-Instruct</i>     |       |      |        |        |      |      |      |       |      |
| Original                        | 52.6  | 82.6 | 62.1   | 69.4   | 71.8 | 34.8 | 79.3 | 71.5  | 65.5 |
| GA+RT                           | 51.9  | 83.3 | 62.5   | 71.6   | 71.5 | 36.6 | 76.4 | 68.5  | 65.3 |
| DPO+RT                          | 52.9  | 81.7 | 61.8   | 69.1   | 71.4 | 36.7 | 78.8 | 69.5  | 65.2 |
| NPO+RT                          | 52.0  | 83.3 | 62.6   | 71.5   | 71.5 | 36.5 | 76.5 | 68.2  | 65.3 |
| <i>Phi-3.5-Mini-Instruct</i>    |       |      |        |        |      |      |      |       |      |
| Original                        | 59.5  | 75.3 | 58.8   | 65.1   | 68.7 | 37.6 | 80.0 | 74.6  | 65.0 |
| GA+RT                           | 57.7  | 73.1 | 60.5   | 65.0   | 68.7 | 38.4 | 76.9 | 72.4  | 64.1 |
| DPO+RT                          | 60.5  | 74.5 | 58.7   | 61.8   | 68.6 | 38.2 | 79.8 | 74.0  | 64.5 |
| NPO+RT                          | 57.5  | 72.8 | 60.6   | 64.8   | 68.8 | 38.3 | 76.9 | 72.3  | 64.0 |
| <i>EXAONE-3.5-7.8B-Instruct</i> |       |      |        |        |      |      |      |       |      |
| Original                        | 56.9  | 75.4 | 60.2   | 62.4   | 65.2 | 35.6 | 76.9 | 68.0  | 62.6 |
| GA+RT                           | 50.2  | 75.4 | 57.6   | 59.7   | 64.7 | 34.7 | 71.0 | 62.3  | 59.5 |
| DPO+RT                          | 56.5  | 75.0 | 59.6   | 57.0   | 64.5 | 36.7 | 76.9 | 67.0  | 61.7 |
| NPO+RT                          | 50.4  | 75.5 | 57.9   | 59.9   | 64.8 | 35.2 | 71.7 | 63.0  | 59.8 |
| <i>Llama-3.1-8B-Instruct</i>    |       |      |        |        |      |      |      |       |      |
| Original                        | 51.8  | 77.1 | 59.2   | 73.2   | 68.1 | 33.8 | 80.2 | 74.1  | 64.7 |
| GA+RT                           | 43.4  | 73.1 | 58.6   | 75.8   | 66.0 | 35.8 | 70.9 | 70.2  | 61.7 |
| DPO+RT                          | 48.1  | 74.6 | 59.0   | 75.1   | 65.5 | 36.3 | 78.1 | 71.5  | 63.5 |
| NPO+RT                          | 48.0  | 74.4 | 59.3   | 76.5   | 65.9 | 36.5 | 75.5 | 71.7  | 63.5 |
| <i>Ministral-8B-Instruct</i>    |       |      |        |        |      |      |      |       |      |
| Original                        | 54.6  | 72.5 | 59.6   | 74.6   | 64.1 | 36.2 | 81.0 | 73.6  | 64.5 |
| GA+RT                           | 43.6  | 30.5 | 60.3   | 77.0   | 55.0 | 34.6 | 73.9 | 68.1  | 55.4 |
| DPO+RT                          | 44.2  | 32.7 | 59.8   | 70.8   | 58.0 | 34.9 | 76.1 | 65.6  | 55.3 |
| NPO+RT                          | 43.6  | 29.4 | 60.3   | 77.5   | 54.2 | 35.3 | 75.0 | 68.2  | 55.4 |

Table 10: Model utility performance per task after erasing single-hop facts from the forget set in MuSiQue-Ans.

[System]

You are an AI assistant helping a user decompose a multi-hop question into a series of single-hop questions. The number of single-hop questions may vary. Each single-hop question should be answerable by a short word or phrase. Questions should be connected in a logical sequence, where the coreference resolution is maintained.

[User]

Who is the ex-wife of the founder of Microsoft?

[Assistant]

Single-Hop Question 1: Who is the founder of Microsoft?

Single-Hop Question 2: Who is the ex-wife of that person?

[User]

What is the continent of the country of origin of the sport Tottenham Hotspur F.C. plays?

[Assistant]

Single-Hop Question 1: What is the sport that Tottenham Hotspur F.C. plays?

Single-Hop Question 2: What is the country of origin of that sport?

Single-Hop Question 3: What is the continent of that country?

[User]

What is the capital of the country where the CEO of the developer of the game Fortnite was born?

[Assistant]

Single-Hop Question 1: Who is the developer of the game Fortnite?

Single-Hop Question 2: Who is the CEO of that company?

Single-Hop Question 3: What is the country where that person was born?

Single-Hop Question 4: What is the capital of that country?

[User]

{multi-hop question}

[Assistant]

Figure 5: Prompt used in MEMMUL to decompose a multi-hop question into a series of subquestions using GPT-4o. It consists of a system prompt followed by three fixed demonstration examples.