

Retrieval-Enhanced Knowledge Editing for Multi-Hop Question Answering in Language Models

Yucheng Shi
University of Georgia
Athens, Georgia, USA
yucheng.shi@uga.edu

Qiaoyu Tan
New York University
New York, USA
qiaoyu.tan@nyu.edu

Xuansheng Wu
University of Georgia
Athens, Georgia, USA
yd6eb@virginia.edu

Shaochen Zhong
Rice University
Houston, Texas, USA
shaochen.zhong@rice.edu

Kaixiong Zhou
North Carolina State University
Raleigh, North Carolina, USA
kzhou22@ncsu.edu

Ninghao Liu
University of Georgia
Athens, Georgia, USA
ninghao.liu@uga.edu

ABSTRACT

Large Language Models (LLMs) have shown proficiency in question-answering tasks but often struggle to integrate real-time knowledge updates, leading to potentially outdated or inaccurate responses. This problem becomes even more challenging when dealing with multi-hop questions since they require LLMs to update and integrate multiple knowledge pieces relevant to the questions. To tackle the problem, we propose the Retrieval-Augmented model Editing (RAE) framework tailored for multi-hop question answering. RAE first retrieves edited facts and then refines the language model through in-context learning. Specifically, our retrieval approach, based on mutual information maximization, leverages the reasoning abilities of LLMs to identify chain facts that naïve similarity-based searches might miss. Additionally, our framework incorporates a pruning strategy to eliminate redundant information from the retrieved facts, which enhances the editing accuracy and mitigates the hallucination problem. Our framework is supported by theoretical justification for its fact retrieval efficacy. Finally, comprehensive evaluation across various LLMs validates RAE’s ability in providing accurate answers with updated knowledge.

KEYWORDS

Model editing, question answering, retrieval-augmented generation

ACM Reference Format:

Yucheng Shi, Qiaoyu Tan, Xuansheng Wu, Shaochen Zhong, Kaixiong Zhou, and Ninghao Liu. 2018. Retrieval-Enhanced Knowledge Editing for Multi-Hop Question Answering in Language Models. In *Proceedings of Make sure to enter the correct conference title from your rights confirmation email (Conference acronym ’XX)*. ACM, New York, NY, USA, 11 pages. <https://doi.org/XXXXXXX.XXXXXXX>

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

Conference acronym ’XX, June 03–05, 2018, Woodstock, NY

© 2018 Association for Computing Machinery.
ACM ISBN 978-1-4503-XXXX-X/18/06...\$15.00
<https://doi.org/XXXXXXX.XXXXXXX>

1 INTRODUCTION

Large language models (LLMs) exhibit impressive performance in question-answering (QA) tasks [25, 31, 33], yet they face a significant challenge: their static knowledge base cannot be easily updated in real-time, leading to the risk of generating outdated or incorrect responses. To overcome this issue, model editing has been proposed for pre-trained LLMs to align their output with up-to-date knowledge [6]. Previous editing methods have shown to be effective in updating answers to single-hop questions [7, 17, 18]. However, it remains a challenging task for model editing to handle multi-hop questions, which is crucial to evaluate machine’s comprehension and reasoning abilities.

Answering multi-hop questions requires the integration of multiple pieces of knowledge. For instance, to answer the question “What is the nationality of the author of ‘Harry Potter’?”, we must link two pieces of knowledge: “(Harry Potter, author, J. K. Rowling)” and “(J. K. Rowling, citizen of, United Kingdom)”, which collectively constitute a **fact chain** [46]. If we conduct a counterfactual edit on the first fact, i.e., “J. K. Rowling” is replaced by “Stephen King”, the subsequent knowledge must be adjusted accordingly, resulting in a totally different fact chain: “(Harry Potter, author, Stephen King), (Stephen King, citizen of, United States)”. Here, we use counterfactual edits to simulate real-world updates. Successful model editing for multi-hop questions requires that the edited LLMs identify and adopt the updated knowledge to derive the final answers.

Unfortunately, existing editing methods are often limited in handling multi-hop questions. First, methods that alter model parameters, including fine-tuning [47], locate-then-edit [17, 18], and meta-learning [19], suffer from the catastrophic forgetting issue [6, 7, 43], where the previously encoded knowledge could be lost after editing. Second, methods that depend on training auxiliary models also fall short in these scenarios. The auxiliary models are usually smaller language models, which lack the necessary reasoning capability to infer correct answers [20]. In contrast, a third category of methods, based on Retrieval-Augmented Generation (RAG), modifies model outputs in a more effective manner [6, 45, 46]. These methods integrate updated knowledge directly into the model prompt, guiding LLMs through in-context learning [45]. RAG-based approaches have shown significant advantages since they are immune to the catastrophic forgetting problem, and the editing process can be conducted on the fly [9, 34]. Yet, the application of RAG to model editing with multi-hop questions still remains largely underexplored.

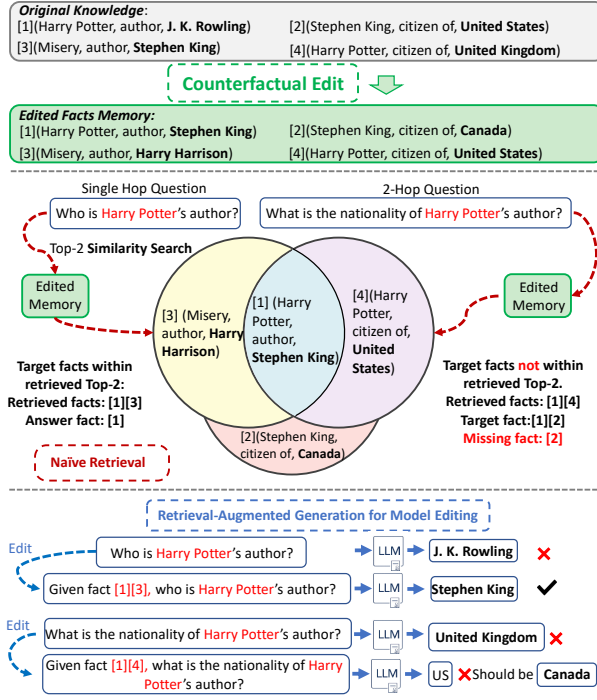


Figure 1: An example of the traditional similarity-based search that fails to retrieve the correct facts for LLM editing.

Furthermore, in practice, model editing often involves handling a large volume of edits, which are stored as fact triplets in a memory bank. In such cases, RAG-based multi-hop editing is expected to retrieve the most relevant facts for each question. However, this task is challenging for two reasons. **First, extracting multi-hop facts requires the retriever to understand the complex connections among multi-relations within the question.** A naïve design of the retriever is applying similarity-based search [9, 14, 34] to obtain the top- K facts that are most semantically similar to the question. However, semantic similarity alone does not guarantee that these facts contain the necessary information to correctly answer the question. We illustrate this issue with an example in **Figure 1**. The edited fact that should be retrieved is “[2]”: (*Stephen King*, citizen of, *Canada*), not “[4]”: (*Harry Potter*, citizen of, *United States*). The latter fact “[4]” is retrieved because it contains “*Harry Potter*” and “*citizen of*” that resemble our question “*What is the nationality of Harry Potter’s author?*”. Despite its higher similarity score, the fact “[4]” is actually irrelevant to the question. Therefore, effective retrieval requires a deep understanding of the question, a capability that goes beyond existing similarity-based search methods. **Second, the retrieved knowledge could contain redundant information and degrade model editing performance.** It is usually impractical to determine the exact amount of information to answer a specific question, so existing retrieval methods tend to return an extensive number of facts [8, 21, 27, 42] for more comprehensive coverage. While these methods do retrieve relevant facts, they also introduce redundant information. It is worth noting that incorporating irrelevant knowledge into the LLM input could mislead models

and cause serious “hallucination” problems [15, 16], where LLMs will generate factually incorrect content based on noise instead of relevant facts. Thus, it is crucial to reduce noise before applying retrieved facts to prompts.

To bridge the gap, we propose a novel Retrieval-Augmented model Editing (RAE) framework, where we first retrieve edited facts and then refine the target model with these facts through in-context learning. To address the first challenge, we propose **Mutual-information (MI) Maximization** for edited facts retrieval. Here, MI quantifies the shared information between the target question and edited facts. Edited facts with higher MI are more relevant. However, directly computing MI is challenging. Thus, we decompose MI into a series of conditional probabilities, and we utilize the next-word prediction capabilities of pre-trained LLMs to estimate these probabilities. In this process, we use the target LLM’s contextual understanding and reasoning abilities to identify necessary facts for successful model editing. To tackle the second challenge, we propose **Uncertainty-based Redundant Facts Pruning** by utilizing LLM output confidence. Specifically, it selectively retains facts that increase the LLMs’ confidence in answering edited questions and discards irrelevant information. Finally, we theoretically justify the formulation of our retrieval objective. Overall, our main contributions are listed below:

- We introduce a novel fact retrieval approach for multi-hop questions in model editing. This approach effectively harnesses the reasoning capabilities of LLMs to retrieve the most relevant multi-hop facts for each question.
- We propose a knowledge pruning strategy to reduce noise after the initial retrieval, mitigating the hallucination problem. Additionally, we provide theoretical analysis to justify our design for the retrieval objective.
- We conduct extensive experiments on various language models in different sizes to verify the effectiveness of our proposed editing method. Empirical results demonstrate the superiority of our RAE framework compared to state-of-art baselines.

2 PRELIMINARY: MODEL EDITING

2.1 Model Editing for Single-hop Questions

In LLMs, a *single model edit* refers to updating a specific piece of factual knowledge [18, 19, 45]. Each knowledge is defined as a triplet $\delta := (h, r, t)$, where h , r , and t denote the head entity, the relation, and the tail entity, respectively, such as (*Harry Potter*, author, *J. K. Rowling*). An edit is defined as changing the tail entity t to a new entity t' , i.e., $\delta \rightarrow \delta' := (h, r, t) \rightarrow (h, r, t')$, where δ' is the edited knowledge. Let q denote the language model’s input. The goal of model editing is to modify a target model f_θ , so that the new model f'_θ produces an output $f'_\theta(q)$ that is aligned with the new fact δ' , where $f'_\theta(q) \neq f_\theta(q)$. Specifically, given $q = [h; r]$, the model is expected to output $t' = f'_\theta([h; r])$ if $h, r \in \delta'$, where $[\cdot; \cdot]$ is the concatenation operator. However, if the input question is not relevant to the edit, i.e., $h \notin \delta'$ or $r \notin \delta'$, the model should output $t = f_\theta([h; r])$ that reflects the knowledge of LLMs before editing.

2.2 Model Editing for Multi-hop Questions

Answering multi-hop questions presents a greater challenge. A multi-hop question seeks to identify a specific tail entity t_k based

Table 1: Answering a 3-hop question q with model editing. The pre-edited and edited answer are t_5 and t_3^* , respectively. t_5 and t_3^* are the tail entity of δ_5 and δ_3^* . G_q and G_q^* denote the pre-edited and edited fact chain.

Edited Fact Bank Δ and Unedited Facts	
$(\delta_1 \rightarrow \delta'_1)$ Harry Potter, author, J. K. Rowling $\rightarrow$ Stephen King	#edit
(δ_2) Stephen King, citizen of, United States	
$(\delta_3 \rightarrow \delta'_3)$ United States, capital, Washington, D.C. $\rightarrow$ Boston	#edit
(δ_4) J. K. Rowling, citizen of, United Kingdom	
(δ_5) United Kingdom, capital, London	
A 3-hop Question q	
(q) Which city is the capital of the country where the author of Harry Potter held citizenship?	
Pre-edited Answer t_5 and Edited Answer t_3^*	
(t_5) London	
(t_3^*) Boston	
Pre-edited Fact Chain G_q	
(δ_1) (Harry Potter , author, J. K. Rowling)	
(δ_4) (J. K. Rowling , citizen of, United Kingdom)	
(δ_5) (United Kingdom , capital, London)	
Edited Fact Chain G_q^*	
(δ'_1) (Harry Potter , author, Stephen King) #edited fact	
(δ_2) (Stephen King , citizen of, United States) #unedited fact	
(δ'_3) (United States , capital, Boston) #edited fact	

on a sequence of linked facts: $\{(h_1, r_1, t_1), (h_2, r_2, t_2), \dots, (h_k, r_k, t_k)\}$, where each tail entity is the head entity of the next fact: $t_i = h_{i+1}$. Each input question q is associated with a sequence **fact chain** G_q for question answering. A k -hop question is usually formulated using only the initial head entity h_1 , and a series of relationships $\{r_1, r_2, \dots, r_k\}$. An example of model editing for a 3-hop question is shown in Table 1. Different fact chains are used to answer the question before and after editing. One key observation is that fact chains represent connected knowledge graphs, where a single entity is involved in two consecutive facts. Additionally, we notice a **"ripple effect"** in these chains: An edit in the first fact δ_1 will lead to changes in the subsequent facts, forming a new chain G_q^* . In practical scenarios, editing is usually conducted in batches, involving multiple fact changes simultaneously, resulting in an edited fact bank, denoted as $\Delta = \{\delta'_1, \delta'_2, \dots, \delta'_N\}$, where N is typically a large number [18]. Locating the relevant edited facts for one specific question is non-trivial due to the "ripple effect". To correctly answer multi-hop questions in model editing, it is crucial to address the retrieval task formally defined as follows:

PROBLEM 1 (RETRIEVAL-AUGMENTED EDITING). *Given an edited fact bank $\Delta = \{\delta'_1, \delta'_2, \dots, \delta'_N\}$ with N instances and a multi-hop question q whose answer requires model editing, we want to retrieve its corresponding edited facts $\Delta_q = \{\delta'_i, \delta'_j, \dots, \delta'_k\}$. The goal is to ensure the retrieved facts contain all the edited facts that appear in fact chain G_q^* , i.e., $\Delta_q \subseteq G_q^*$ and $\Delta \setminus \Delta_q \not\subseteq G_q^*$. Then, these facts are used to refine the target model f_θ to conduct model editing.*

3 METHODOLOGY

Our Retrieval-Augmented Editing (RAE) framework, as shown in Figure 2, contains two key steps: (1) retrieving edited facts relevant to the target question, and (2) editing the language model using these retrieved facts via in-context learning. We will first discuss step (2) with the motivation of our design in the following section. The details of step (1) can be found in Sections 3.2 and 3.3.

3.1 Retrieval-Augmented Editing

A naïve approach is using similarity-based search to retrieve edited facts similar to target question q [9, 34, 46]. These facts are then integrated into a prompt template for editing via in-context learning: $f'_\theta(q) = f_\theta(T_e(q, \{\delta'_1, \delta'_2, \dots, \delta'_K\}))$, where T_e is the editing template. For example, $T_e(\cdot)$ can be made as "Given fact: $\{\delta'\}$, $\{q\}$?". The Top- K nearest edited facts to question q in the embedding space are denoted as $\{\delta'_1, \delta'_2, \dots, \delta'_K\} = \text{Top-}K_{\delta \in \Delta} \text{sim}(g_z(\delta), g_z(q))$, where $\text{sim}(\cdot)$ denotes the similarity function and g_z is an embedding model. However, edited facts Δ_q needed to answer q are hard to retrieve by this approach since they usually contain entities different from q , which will result in a low similarity score in a large bank Δ (e.g., in Table 1, **United States** in δ'_3 , but not in q).

To address this problem, we propose **edited fact chain extraction** to obtain G_q^* . Inherently, each G_q^* forms a connected knowledge graph (KG) [46]. Such KGs can be retrieved by iteratively traversing links from one entity to another. Take $G_q^* = \{\delta'_1, \delta_2, \delta'_3\}$ in Table 1 as an example. It is composed of two edited facts δ'_1, δ'_3 and one unedited fact δ_2 . We can observe that the question entity: (**Harry Potter**) is the head entity h_1 in $\delta'_1 = \{h_1, t_1, t'_1\}$, and the edited tail entity t'_1 : (**Stephen King**) is also the head entity h_2 in the next fact $\delta_2 = \{h_2, r_2, t_2\}$. Moreover, for each subsequent fact in the chain, its head entity is always the tail entity of the previous fact. By effectively retrieving the KG that represents the fact chain G_q^* , we are able to capture all the edited factual triplets $\Delta_q = \{\delta'_i, \delta'_j, \dots, \delta'_k\}$. In light of this, we define our retrieval-augmented editing as:

$$f'_\theta(q) = f_\theta(T_e(q, G_q^*)), \quad (1)$$

where we give an example of such editing in Figure 2. In the next section, we will introduce the detailed strategy of retrieving G_q^* , where we first propose a mutual information-based retrieval strategy to extract facts needed to answer the target question (Section 3.2). Then, we propose a pruning method to delete irrelevant facts from the initial retrieval result (Section 3.3).

3.2 Edited Facts Retrieval via Maximizing MI

We first construct a knowledge graph that connects different facts. Then, we introduce our proposed retrieval objective of extracting relevant subgraphs given input questions.

3.2.1 External Knowledge Graph for Subgraph Retrieval. According to our previous discussion, we aim to retrieve the fact chain G_q^* for model editing. However, it is worth noting that G_q^* consists of both edited and unedited facts, whereas the unedited facts do not exist in our edited fact bank Δ by default. To effectively incorporate both types of facts into our retrieval process, we propose integrating all edited facts into an external knowledge graph $\mathcal{G}$. By selecting a comprehensive KG such as WikiData [32], the new graph $\mathcal{G}^*$

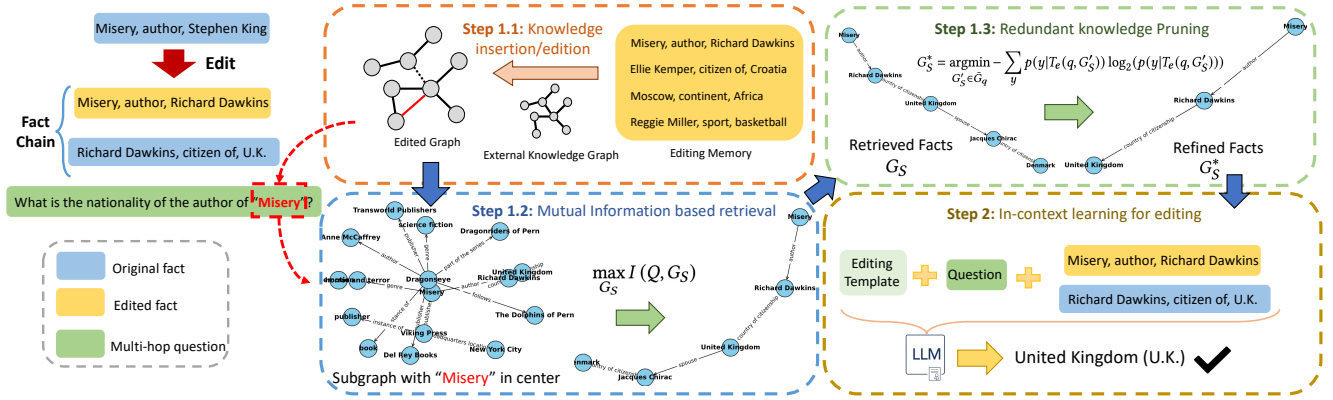


Figure 2: The overall framework of our retrieval-augmented in-context model editing method.

will encompass both unedited and edited facts. It complements our edited fact bank Δ and connects different entities. Besides, the external knowledge graph provides extra factual knowledge that can enhance language to output correct answers.

Specifically, given the edits $\Delta = \{\delta'_1, \dots, \delta'_n\}$ and an external $\mathcal{G}$, we consider two types of operations to combine them. **(1) Modifying existing facts:** If the original fact appears in $\mathcal{G}$, i.e., $(h, r, t) \in \mathcal{G}$, we will modify the KG according to the edits, so $\mathcal{G}^* = (h, r, t') \cup \mathcal{G} \setminus (h, r, t)$. **(2) Adding new facts:** If the original fact does not appear in $\mathcal{G}$, i.e., $(h, r, t) \notin \mathcal{G}$, then we append the modified fact to the KG, so $\mathcal{G}^* = (h, r, t') \cup \mathcal{G}$. Next, given a question q , we retrieve a subgraph G_S from $\mathcal{G}^*$, so that $G_S \subset \mathcal{G}^*$. Our goal is to ensure that G_S contains fact chains corresponding to q , i.e., $G_q^* \subseteq G_S$.

3.2.2 Mutual Information based Retrieval Objective. For effective editing, the retrieved subgraph G_S must share relevant information with the question. Therefore, we define the objective of subgraph retrieval as maximizing the mutual information (MI) between the subgraph and a set of questions Q whose answers require editing. The objective is formalized as below, where the theoretical justification is provided in Section 3.4:

$$\max_{G_S} I(Q, G_S) = H(Q) - H(Q | G = G_S). \quad (2)$$

Given a fixed question set Q , its Shannon entropy $H(Q)$ remains constant. Therefore, maximizing the mutual information $I(Q, G_S)$ is equivalent to minimizing the conditional entropy $H(Q | G = G_S)$. Thus, we focus on optimizing the following objective:

$$\max_{G_S} I(Q, G_S) = \min_{G_S} H(Q | G = G_S) \quad (3)$$

$$= \max_{G_S} \sum_{q \in Q} p(q|G = G_S) \log_2 p(q|G = G_S). \quad (4)$$

In practice, quantifying $p(q|G = G_S)$ poses a significant challenge due to its computational complexity. This complexity arises as there are numerous subgraph candidates G_S within the entire knowledge graph, making it prohibitively expensive to exhaustively search for the optimal subgraph. To circumvent this issue, we first replace the intractable term $p(q|G = G_S)$ with $\frac{p(q, G = G_S)}{p(G = G_S)}$. Then, suppose we consider one question each time, where $Q = q$, we can reformulate

the objective as:

$$\max_{G_S} \frac{p(q, G = G_S)}{p(G = G_S)} \log_2 \frac{p(q, G = G_S)}{p(G = G_S)}. \quad (5)$$

In the following, we will discuss how to estimate probability $p(q, G = G_S)$ and $p(G = G_S)$ efficiently.

3.2.3 Probabilities Estimation. We propose to compute probabilities by leveraging the next-word prediction capability of LLMs. Given that the fact chain forms a tail-to-head connected knowledge graph, our extracted subgraph G_S can be represented as $G_S = (h_1, r_1, t_1, \dots, h_n, r_n, t_n)$, where h_i and t_i are nodes, r_i is the edge, and n is the number of retrieved triplets. Thus, we can estimate $\frac{p(q, G = G_S)}{p(G = G_S)}$ as:

$$\frac{p(q, G = G_S)}{p(G = G_S)} = \frac{p(r_1, t_1, h_2, r_2, t_2, \dots, h_n, r_n, t_n | q, h_1)}{p(r_1, t_1, h_2, r_2, t_2, \dots, h_n, r_n, t_n | h_1)} \cdot \frac{p(q, h_1)}{p(h_1)}. \quad (6)$$

Specifically, for the term $p(r_1, t_1, h_2, r_2, t_2, \dots | q, h_1)$, we can further decompose it into following form:

$$p(r_1, t_1, h_2, r_2, t_2, \dots, h_n, r_n, t_n | q, h_1) = p(t_1, h_2, r_2, t_2, \dots, h_n, r_n, t_n | q, h_1, r_1) \cdot p(r_1 | q, h_1). \quad (7)$$

This decomposition allows us to initially focus on estimating the $p(r_1 | q, h_1)$. Specifically, the head entity h_1 is determined if q is given, since we assume h_1 is mentioned in question q . Candidate relations for r_1 can also be selected from the edited KG. Practically, we can estimate the probability $p(r_1 | q, h_1)$ for each candidate relation using an auto-regressive language model f_ϕ [25, 38]:

$$p(r_1 | q, h_1) \approx \prod_{i=1}^{|r_1|} f_\phi(w_{r_1}^{(i)} | w_q^{(1)}, \dots, w_q^{(|q|)}, w_{h_1}^{(1)}, \dots, w_{h_1}^{(|h_1|)}, w_{r_1}^{(1)}, \dots, w_{r_1}^{(i-1)}), \quad (8)$$

where f_ϕ is the predicted word probability, and w_q, w_{h_1}, w_{r_1} denote words in the question q , head entity h_1 , and relation r_1 , respectively. We can employ open-source LLMs like GPT-2 [25] for this estimation. Please note that, the model f_θ being edited does not need to be the same model used for probability estimation, making our method applicable even for editing proprietary LLMs. With a specific input context $\{q, h_1\}$, the language model will assign different

probabilities to each relation based on its contextual understanding and reasoning ability.

Then, $p(t_1, h_2, r_2, t_2, \dots | q, h_1, r_1)$ can be further decompose into $p(h_2, r_2, t_2, \dots | q, h_1, r_1, t_1) \cdot p(t_1 | q, h_1, r_1)$. In our case, we assume $p(t_1 | q, h_1, r_1) = 1$, since one relation usually only corresponds to one tail entity. When there are multiple tail entities, we find the assumption still works well empirically. So, $p(h_2, r_2, t_2, \dots | q, h_1, r_1, t_1)$ can be decomposed into $p(r_2, t_2, \dots | q, h_1, r_1, t_1, h_2) \cdot p(h_2 | q, h_1, r_1, t_1)$. Additionally, since the tail entity in one fact becomes the head entity in the subsequent fact, we can also have $p(h_2 | q, h_1, r_1, t_1) = 1$. Thus, we can have $p(t_1, h_2, r_2, t_2, \dots | q, h_1, r_1) = p(r_2, t_2, \dots | q, h_1, r_1, t_1, h_2)$. This is a nice property that helps us iteratively decompose this intractable probability term. By iteratively applying the aforementioned step for n times, we can compute the conditional probability of all subgraphs within an n -hop distance from the question entity. The final estimation can be expressed as:

$$\begin{aligned} & p(r_1, t_1, h_2, r_2, t_2, \dots, h_n, r_n, t_n | q, h_1) \\ &= p(r_n | q, h_1, r_1, t_1, \dots, h_{n-1}, r_{n-1}, t_{n-1}, h_n) \cdot \\ & p(r_{n-1} | q, h_1, r_1, t_1, \dots, h_{n-2}, r_{n-2}, t_{n-2}, h_{n-1}) \cdot \\ & \dots \cdot \\ & p(r_2 | q, h_1, r_1, t_1, h_2) \cdot p(r_1 | q, h_1). \end{aligned} \quad (9)$$

Till now, we have decomposed $p(r_1, t_1, h_2, \dots | q, h_1)$ in Equation (6) into the product of conditional probabilities of predicting different relations within the n -hop subgraph. This nice property ensures the selection of the subgraph will only be determined by relation probability, which is free from the interference of any potential edited tail entity. Similarly, we can decompose the denominator term $p(r_1, t_1, h_2, \dots | h_1)$ into:

$$\begin{aligned} & p(r_1, t_1, h_2, r_2, t_2, \dots, h_n, r_n, t_n | h_1) \\ &= p(r_n | h_1, r_1, t_1, \dots, h_{n-1}, r_{n-1}, t_{n-1}, h_n) \cdot \\ & p(r_{n-1} | h_1, r_1, t_1, \dots, h_{n-2}, r_{n-2}, t_{n-2}, h_{n-1}) \cdot \dots \cdot p(r_1 | h_1). \end{aligned} \quad (10)$$

Then, for the last term $p(q, h_1)/p(h_1)$ in Equation (6), based on Bayes' theorem, we can transform it into $p(q, h_1)/p(h_1) = p(q | h_1)$, which is a constant value given a specific question q . We can also apply model f_ϕ to estimate this conditional probability. Now, since we are able to estimate every term in Equation (6) and (5), we can effectively identify the subgraph that yields the maximum Mutual Information. Additionally, we utilize the beam search technique [26] to expedite the computational process, eliminating the necessity for exhaustively traversing all connected nodes. Further details on our implementations are provided in the Appendix B. In this work, we treat n as a hyperparameter since the number of hops required to answer a question is unknown in advance. To ensure thorough exploration, n is assigned a large value, enabling an extensive search range. However, this approach will also introduce irrelevant information in the retrieved subgraph G_S , which can potentially mislead the language model to generate undesired answers [15, 16]. Thus, in the next section, we will discuss how to mitigate this problem.

3.3 Uncertainty-based Redundant Fact Pruning

This section introduces a pruning method, which utilizes model output uncertainty, to eliminate redundant facts from G_S .

3.3.1 Editing Uncertainty. We define editing uncertainty as the uncertainty of the output generated by large language models. Formally, the output uncertainty is quantified by Shannon entropy:

$$H(Y|X=x) = - \sum_y p(y|x) \log_2(p(y|x)), \quad (11)$$

where y represents each possible answer generated by the language model, and $x = \{q, G_S\}$ is the model input composed of the question q and facts G_S . A higher entropy value $H(Y|X=x)$ signifies less confidence in the answer, reflecting greater uncertainty. Conversely, a lower entropy value indicates higher confidence and less uncertainty. Ideally, if input facts G_S are exactly the edited question fact chain G_q^* , i.e., $G_S = G_q^*$, then the model output should exhibit maximum confidence with minimal entropy, since G_q^* contains the precise knowledge to answer question q . In the next part, we conduct empirical experiments to verify this assumption.

In our experiments, GPT-J (6B) [33] is chosen as the base language model. We select 1000 instances for each of 2, 3, and 4-hop questions from the MQUAKE-CF dataset [46] for testing. The MQUAKE-CF dataset comprises multi-hop questions that are based on real-world facts, where the edited facts are counterfactual, meaning they do not exist in actual real-world scenarios. An example of such a question with an edit is provided in Table 1.

Our experiment seeks to identify the fact set G'_S that, when used as model input, yields the lowest output entropy, indicating minimal editing uncertainty. Our first step is to construct different fact set candidates. We begin with the first fact δ_1 in the fact chain G_q^* as our initial fact set G'_S . Then, we add each subsequent fact from the chain until the G'_S encompasses the entire fact chain G_q^* . After that, we introduce unrelated facts $\hat{\delta}$ into the set. This process is repeated until G'_S contains six elements. Finally, we build a prefix set $\bar{G}_q$ with all the six subsets G'_S . Specifically, for a 4-hop question, we have $\bar{G}_q = \{\{\delta_1\}, \{\delta_1, \delta_2\}, \{\delta_1, \delta_2, \delta_3\}, \dots, \{\delta_1, \delta_2, \delta_3, \delta_4, \hat{\delta}_5, \hat{\delta}_6\}\}$, where $\delta_1, \delta_2, \delta_3, \delta_4 \in G_q^*$ and $\hat{\delta}_5, \hat{\delta}_6$ are two irrelevant facts. Finally, we conduct in-context editing using each subset G'_S from the prefix set $\bar{G}_q$ with editing template T_e : "Given fact: $\{G'_S\}$, $\{q\}$?" In our experiments, for computational simplicity, we consider model output y to be each of the next predicted word. We report the entropy over all the words in the vocabulary as the editing uncertainty.

The editing uncertainty with different subsets is listed in Figure 3a. For comparison, we also report the editing uncertainty with random facts selected from Wikidata, as shown in Figure 3b. Our observations reveal a phenomenon: the language model produces answers with much lower entropy when G'_S is equal to the ground-truth fact chain G_q^* . If G'_S presents redundant facts or insufficient facts, the entropy will increase. Meanwhile, if the LLM is fed with random facts, the entropy level also remains consistently high. This observation justifies the use of entropy as the indicator of whether G'_S contains the correct facts for LLM inference.

3.3.2 Knowledge Pruning with Editing Uncertainty. Since incorporating the most relevant facts will result in the lowest entropy, we propose to utilize this finding for knowledge pruning. Specifically, we first follow Section 3.2 to retrieve a knowledge graph G_S containing n triplets for question q : $G_S = \{\delta_1, \delta_2, \dots, \delta_n\}$, where n is a

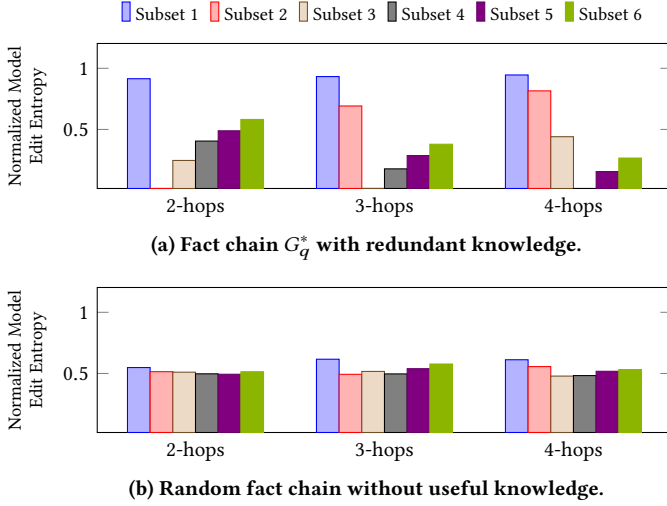


Figure 3: Distribution of normalized model editing entropy with different fact subsets as input. A lower normalized entropy on a given subset indicates that the model is more confident in answering the question with the given facts. Subset 1 includes the first fact $\{\delta_1\}$, Subset 2 includes the first two facts $\{\delta_1, \delta_2\}$, and so on. Figure 3a shows that the entropy is significantly lower if the subset contains exactly the entire fact chain of the question (e.g., Subset 2 for 2-hop questions).

sufficiently large number, and G_S could contain redundant knowledge. Then, to remove redundant knowledge, we first build the prefix sets $\tilde{G}_q$ for target question q based on the retrieved graph G_S . Then, we can obtain the pruned fact set G_S^* using the objective:

$$G_S^* = \operatorname{argmin}_{G'_S \in \tilde{G}_q} - \sum_y p(y|T_e(q, G'_S)) \log_2(p(y|T_e(q, G'_S))). \quad (12)$$

Finally, we can apply G_S^* as our retrieved fact chain for the in-context learning introduced in Section 3.1 to conduct editing.

3.4 Theoretical Justification

In this subsection, we theoretically justify that the facts collected by our retrieval objective (Eq. (2)) are effective in performing model editing with in-context learning. To begin with, we discuss what kinds of input can effectively activate in-context learning. Then, we explore how to build such effective input for model editing.

Our proposed editing method relies on the in-context learning ability of LLMs. In the following, we provide an analysis of how in-context learning can be effectively triggered. Theoretically, the text generation process of a language model can be understood as a Hidden Markov Model [2, 39]. The model initially selects a concept $\theta_c \in \Theta$ from a set of underlying concepts denoted as Θ , and then samples a sequence of words based on the chosen concept. Based on that, the in-context learning can be written as $p(y|S, x) = \int_{\theta_c \in \Theta} p(y|S, x, \theta_c) p(\theta_c|S, x) d\theta_c$, where S denotes in-context prompt and x denotes query. Existing research has theoretically proven that the condition to activate in-context learning is when there is a shared latent concept θ_c between prompt text S and the input query x . More discussions can be found in [39].

Motivated by the above analysis, as in-context prompt S is the edited knowledge in our design, we seek to include the edited knowledge that shares the same latent concept θ_c as question q . Ideally, this will activate in-context learning for effective model editing. Formally, we can define such knowledge graph as

$$G_S = \operatorname{argmax}_{G \in \mathcal{G}} I(G; \theta_c), \quad (13)$$

where θ_c is the latent concept used to generate question q , and we use mutual information $I(G; \theta_c)$ to quantify the share information. However, this is a non-trivial task since concept θ_c is an intractable hidden variable. To address this issue, we propose obtaining the target knowledge graph that maximizes the lower bound of such an objective. Specifically, we can have the following theorem:

THEOREM 1. *Given retrieved graph $G_S \in \mathcal{G}$, the latent concept θ_c , and the question q sampled conditioned on concept θ_c , there exists a mutual information inequality:*

$$I(G_S; \theta_c) \geq I(G_S; q). \quad (14)$$

Theorem 1 shows that we can maximize the mutual information between the selected knowledge graph G_S and question concepts θ_c by maximizing the mutual information between the selected graph G_S and the question q itself. In this way, the in-context learning ability of LLMs would be effectively triggered. When we apply such knowledge as the prompt, we can effectively conduct the in-context editing. The proof of Theorem 1 is in Appendix A.

4 EXPERIMENTS

We conduct experiments to answer the following questions. **Q1:** Does RAE successfully edit model output? **Q2:** How does our retrieval strategy perform compared with other retrieval methods? **Q3:** Does our proposed pruning technique remove redundant facts from the retrieved facts? **Q4:** Does RAE work for propriety LLMs?

4.1 Experiment Settings

4.1.1 Language Models. We evaluate RAE across various kinds of language models in different sizes and families, including GPT-2 (1.3B) [25], GPT-J (6B) [33], Falcon (7B) [1], Vicuna (7B) [4], and Llama2-chat (7B) [30]. Among them, GPT-2, GPT-J, and Falcon are pre-trained language models without instruction tuning [5, 24], while Vicuna is an instruction-tuned variation of Llama1 [31] and Llama2-chat is the instruction-tuned version of Llama2. Instruction-tuned models (Vicuna and Llama2-chat) are expected to better follow the instructions in the prompt compared to native pre-trained models (GPT-2, GPT-J, and Falcon). We include both kinds of models to verify the effectiveness of the proposed methods.

4.1.2 Editing Baselines. For comparison, we consider three kinds of model editing methods: (1) Model weight updating methods: Fine-tuning [47] edits the entire model weights by language modeling the edited knowledge. ROME [17] and MEMIT [18] focus on identifying and updating particular neurons associated with the knowledge that needs editing. (2) Auxiliary models methods: SEARC [20] trains an extra language model to store updated knowledge, and it switches to the auxiliary model when answering questions relevant to the edited facts. Details about the auxiliary models

for different language models are included in Appendix B. (3) In-context learning-based method: Mello [46] edits model outputs with multi-round conversations.

4.1.3 Implementation Details. We evaluate our editing method on the MQUAKE-CF and MQUAKE-T datasets [46], comprising 1000 counterfactual editing instances per 2-hop, 3-hop, and 4-hop questions in MQUAKE-CF, and total 1868 editing instances for 2-hop and 3-hop questions in MQUAKE-T. In our setting, we edit all these facts at the same time. Following previous work [46], we leverage cases from the MQUAKE-CF-9k dataset to craft prompt templates for both baselines and our method. An edit is considered successful if the edited LLM correctly predicts the updated answer within the first ten generated tokens. This criterion is reported as edited accuracy in Table 2. More setting details are in Appendix. B.

4.2 Editing Performance Evaluation

To answer **Q1**, we assess our model editing method across various language models, compared against different baseline methods. We focus on the edited accuracy of these models in answering multi-hop questions with edited answers. For references, we also report the accuracy of un-edited models in predicting both original and edited answers. Our key observations are: **(1)** Our RAE outperforms all others in both datasets across five language models, achieving an average improvement of 40.57% over the second-best method (Mello) when conducting thousands of edits at the same time. This superior performance primarily stems from our novel MI-based retrieval objective and an effective pruning strategy. Our design can also seamlessly integrate an external knowledge graph, which effectively links all edited facts, thereby facilitating the multi-hop editing process. **(2)** Mello demonstrates good performance on the MQUAKE-T dataset with models larger than 6B. However, it underperforms on the MQUAKE-CF dataset and fails with GPT-2, likely due to GPT-2’s inability to follow Mello’s complex prompts, resulting in no output. **(3)** Other methods generally show lower performance across all language models, which aligns with the findings from MQUAKE dataset paper [46].

4.3 Retrieval Performance Evaluation

To answer **Q2**, we assess the effectiveness of our mutual information-based retrieval method for multi-hop question-answering tasks.

4.3.1 Retrieval Baselines. We consider two types of retrieval methods as baselines, namely the embedding-based and the probability-based methods. We introduce their details as follows.

(1) Embedding-based retrieval methods mostly concentrate on text retrieval rather than triplet retrieval. For a fair comparison, we adapt them to fit our task by employing the Contriever [11] to encode all edited facts into embeddings, and then cache these embeddings for similarity search. Different retrieval methods use distinct strategies to modify the original K -hops questions as the final query to perform the search via dot product between the embeddings.

- **KG Link** [8, 40] is a straightforward strategy, which identifies the query entity and its linked entities as candidates to find the one that is the most similar to the original question. This process is repeated K times to retrieve the entire fact chain.

- **Question Reform** [21, 41] appends the retrieved entity to the previous query for the next hop retrieval. This design is motivated by simulating a reasoning path where the retrieved fact at each hop is viewed as an intermediate reasoning result.
- **Question Decompose**, proposed by Mello [46], decomposes a multi-hop question into multiple sing-hop questions. Specifically, a comprehensive conversational-style prompt is applied to interact with the LLM for collecting each single-hop question. For each single-hop question, the corresponding knowledge is retrieved with the naïve method described in KG Link.
- (2) In contrast with the embedding-based methods that estimate the similarities between the query and the facts with the dot products of their embeddings, probability-based methods leverage language models to directly estimate this similarity.
- **Max Prob** [29, 44] retrieves the subgraph $G_S \in \mathcal{G}^*$ that maximizes the conditional probability $p(G_S|q)$, where q is the input multiple question. In our experiments, we extract a K -hop subgraph from the edited knowledge graph $\mathcal{G}^*$ and estimate the probability with a language model f_ϕ as shown in Section 3.2.3.

4.3.2 Implementation Details. We select 300 cases for each of the 2, 3, and 4-hop problems from the MQUAKE-CF dataset. For each type, we report two matching accuracy scores in Table 3. Specifically, within the retrieved facts, if any fact also appears in the fact chain of a multi-hop question, we name it a partial match (PM). And if every retrieved fact matches the fact chain for the multi-hop question, we name it exact match (EM).

4.3.3 Results. We have the following observations: **(1)** Our proposed mutual information-based retrieval method demonstrates excellent performance across various LLMs in multi-hop fact extraction. We also find its success with relatively small language models, such as GPT-2, showing strong generalization ability. **(2)** In contrast, traditional embedding-based methods (KG Link and Question Reform) underperform in this multi-hop fact retrieval challenge. Their limitation mainly lies in failing to comprehend the complex interplay between multiple relations in a question. Thus, it hinders the extraction of target facts from the extensive knowledge base. **(3)** Mello demonstrates the efficacy of decomposing multi-hop questions into single-hop questions for this multi-hop facts retrieval task. Notably, its EM performance drops significantly with an increasing number of hops, probably because it becomes much more challenging for language models to perform question decomposing. **(4)** Probability maximization-based methods outperform traditional embedding methods, while there is still a gap compared to ours, probably because the predicted most probable fact may not be the most necessary one for answering a specific question.

4.4 Ablation Studies on Pruning Strategy

To answer **Q3**, we verify that our proposed pruning strategies are beneficial to multi-hop editing tasks. To simulate the situation where the retrieved facts contain redundant information, we conduct our experiment by always retrieving 2 additional facts over the facts needed by the original question. That is, given a K -hop question, we set the total number of retrieved facts as $n = K + 2$.

Table 4 reports the edited accuracy of RAE working with or without the pruning strategy, and we draw the following conclusions:

Table 2: Edited accuracy (%) on MQUAKE datasets (MQUAKE-CF and MQUAKE-T).

			Editing Methods							
Language Models		Datasets	w/o edit (orig ans)	w/o edit (edited ans)	Fine Tuned	ROME	MEMEIT	SEARC	Mello	Ours
Models w/o Instruction Tuning	GPT-2	MQUAKE-CF	10.7	0.0	3.8	1.7	2.3	4.0	0.0	62.8
		MQUAKE-T	4.7	0.0	5.8	6.4	1.6	2.7	0.0	61.8
	GPT-J	MQUAKE-CF	18.1	0.0	7.7	7.6	8.1	6.8	15.3	69.3
		MQUAKE-T	15.3	3.1	3.1	4.1	10.6	2.8	36.7	63.9
	Falcon	MQUAKE-CF	33.7	1.0	5.6	1.7	2.3	7.9	10.7	66.8
		MQUAKE-T	35.9	8.9	17.2	7.3	1.6	4.5	51.5	61.6
Models w/ Instruction Tuning	Vicuna	MQUAKE-CF	43.2	1.0	4.8	8.4	7.6	7.9	10.2	67.2
		MQUAKE-T	15.3	3.1	23.1	5.0	1.7	4.5	51.7	63.2
	Llama2 (chat)	MQUAKE-CF	29.0	1.1	5.4	6.3	3.8	7.9	20.7	69.1
		MQUAKE-T	26.7	5.1	17.1	8.7	1.7	4.5	49.4	66.2

Table 3: Multi-hop facts retrieval accuracy (%) comparison.

MQUAKE-CF (300 cases for each kind of question)							
Question Type		2-hops		3-hops		4-hops	
Category	Retrieval Method	PM	EM	PM	EM	PM	EM
Embedding	KG Link	52.7	28.7	18.2	3.7	14.0	0.0
	Question Reform	62.3	7.7	14.7	0.0	12.3	0.0
	Mello (Llama2)	84.3	80.0	80.7	42.3	83.3	25.7
Probability	MaxProb (GPT-2)	77.7	50.3	67.3	25.3	65.0	20.0
	MaxProb (Llama2)	78.3	55.7	79.7	37.0	69.3	28.7
Mutual Information	RAE (GPT-2)	83.0	66.3	77.3	41.0	80.3	43.7
	RAE (GPT-J)	83.0	69.7	81.3	53.7	82.7	54.0
	RAE (Falcon)	82.3	70.7	72.3	44.3	81.7	47.3
	RAE (Vicuna)	81.0	66.7	79.3	50.3	85.0	50.0
	RAE (Llama2)	82.7	69.3	84.0	49.3	82.0	47.0

PM: Partial Match. EM: Exact Match.

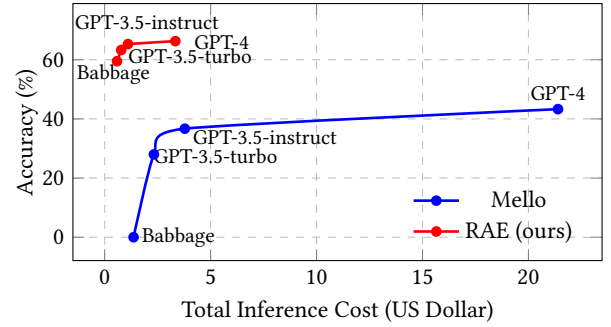
Table 4: Edited accuracy (%) with (w) or without (w/o) pruning.

Dataset		MQUAKE-CF				
Type	Strategy	GPT-2	GPT-J	Falcon	Vicuna	Llama2 (chat)
2-hops	w/o Pruning	63.0	63.7	65.2	63.8	70.1
	w/ Pruning	73.3	75.5	74.5	73.5	75.8
	Improved	10.3 ↑	11.8 ↑	9.3 ↑	9.7 ↑	5.7 ↑
3-hops	w/o Pruning	43.1	53.8	55.6	55.0	60.3
	w/ Pruning	53.2	65.4	62.1	62.7	65.8
	Improved	10.1 ↑	11.6 ↑	6.5 ↑	7.7 ↑	5.5 ↑
4-hops	w/o Pruning	49.9	58.8	55.2	61.5	61.6
	w/ Pruning	61.9	66.9	62.9	65.5	65.8
	Improved	12.0 ↑	8.1 ↑	7.7 ↑	4.0 ↑	4.2 ↑

(1) The proposed pruning technique significantly enhances the performance of model editing, demonstrated by achieving an average accuracy improvement of 8.3% across various language models. (2) We observe a more profound improvement for smaller language models on complex questions, while this benefit to larger language models is relatively less significant. In particular, the performance improvement of GPT-2 on the 4-hop questions reaches 12.0%, while for Llama-2 it is just 4.2%. This observation suggests that larger models are more robust to redundant information.

4.5 Editing Performance on Proprietary LLMs

To answer Q4, we apply the proposed RAE to edit proprietary language models, where we can only access the model via APIs, such

**Figure 4: Editing performance compared with inference cost over different proprietary models.**

as ChatGPT [22]. In such cases, RAE utilizes a different lightweight language model to perform relevant facts retrieval and pruning as introduced in Section 3.2 and Section 3.3. Then, the obtained facts will be fed into the proprietary models to perform in-context editing. We evaluate our proposed editing method on *GPT-babbage-002*, *GPT-3.5-turbo-0613*, *GPT-3.5-instruct*, and *GPT-4-0613* models [23]. We use GPT-2 (1.3B) as our retrieval model. We report the edited accuracy and total editing cost of our method for randomly selected 300 cases (MQUAKE-CF) in Figure 4. The editing cost includes the total fees of calling APIs and the cost of running GPT-2 for knowledge retrieval on rented GPUs. We also report the results of Mello [46] for comparison, where only the API fees are counted.

In Figure 4, we first observe that Babbage has 0.0% edited accuracy by using Mello, showing its ineffectiveness in editing the un-instruction tuned proprietary language model. It is expected since Mello relies on the conversational ability of language models to decompose multi-hop questions. In contrast, RAE is effective in editing all these proprietary models with a remarkably lower cost than Mello. In particular, RAE improves Mello for editing GPT-4 with almost 20% edited accuracy by only costing around its 15% budget. This highlights the benefit of utilizing the inherent language modeling ability instead of the instruction following ability to perform knowledge retrieval for multi-hop question answering.

4.6 Case Study

In Figure 5, we present two cases from the MQUAKE-CF dataset to demonstrate the retrieving process over a knowledge graph and the pruning process of the retrieved facts. For visualizations, the red,

black, and dotted lines represent the final, candidate, and discarded paths in the knowledge graph with beam search, reflecting the decision-making process of our retrieval design. In Figure 5a, a potential path *Misery*—*language*→*English*—*country*→*U.K.* is discarded even though it can lead to the correct answer (U.K. in this case). This is because it does not share the information that is needed to answer the target question, resulting in a low MI score. In Figure 5b, even though this question is associated with three edited facts, our method can successfully locate all of them, demonstrating robustness over multi-hop questions. In both cases, our pruning strategy successfully truncates the retrieved fact chain to only contain necessary facts relying on the normalized model editing entropy.

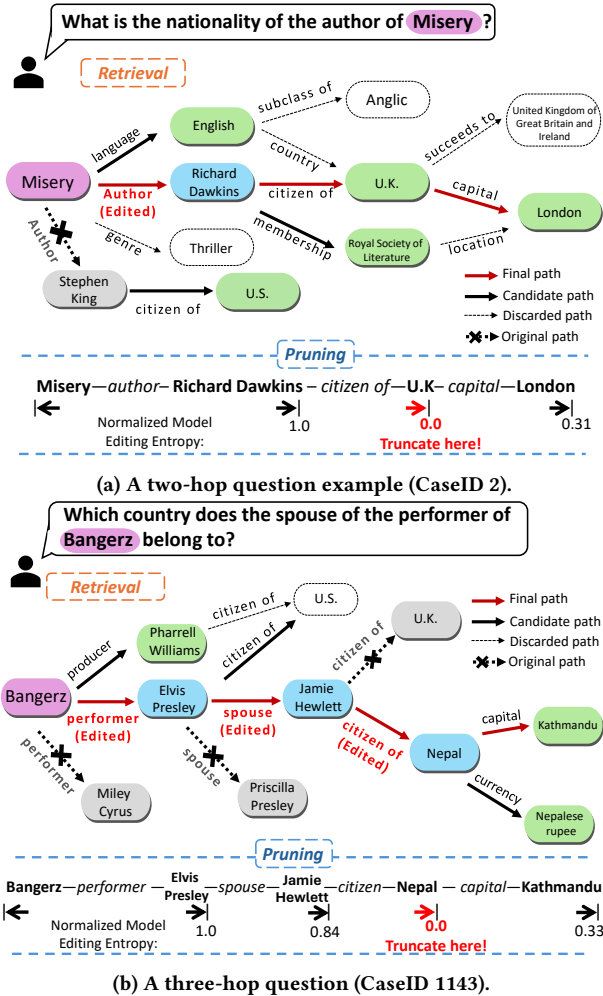


Figure 5: Case studies for edited facts retrieval and pruning. The retrieval process involves the beam search, starting from the query entity and navigating through the knowledge graph with two beams. The two primary candidate edges (beams) at each entity hop are boldfaced, and the others are discarded and denoted with dashed lines. We emphasize the final beam search result with red color.

5 RELATED WORK

5.1 Model Editing

Existing model editing methods for single-hop fact editing can be categorized into the following kinds: (1) Meta-learning method: MEND [19] employs an auxiliary network for gradient transformation during fine-tuning. (2) Local modification methods: MEMIT [18] and ROME [17] edit specific neurons within models with minimal impact on unrelated weights. (3) Extension-based methods: SERAC utilizes a fine-tuned small auxiliary language model to answer questions that require edited facts. Meanwhile, GRACE [10] provides a theoretical basis for concept erasure across model layers. (4) In-context learning methods: IKE [45] edit models through in-context learning, bypassing the need for parameter updates.

In contrast, model editing methods for multi-hop fact editing remain underexplored. One such method, Mello [46], utilizes a comprehensive prompt to decompose the multi-hop question into single-hop questions and then retrieve edited knowledge iteratively. A recent contemporaneous work, DeepEdit [36], edits LLMs through decoding with constraints, where a neuro-symbolic method is proposed for better reason coherence. However, these approaches heavily rely on the strong instruction following and reasoning ability of LLMs, which refers to their capacity to understand and follow the input instruction to generate coherent, context-aware responses. Existing methods [46] fall short for models in relatively small sizes, such as GPT-J(6B) or Vicuna(7B), compared to powerful GPT-3.5.

5.2 Multi-hop QA with Knowledge Graphs

Multi-hop Knowledge Graph Question Answering (KGQA) involves locating answers within a KG that are several hops away from the starting question entities. Existing techniques usually first extract a relevant subgraph from the entire KG, then apply multi-hop reasoning [28, 44]. Initial retrieval strategies apply similarity scores like PageRank [28] or embedding similarity [13, 44] for subgraph selection, but these methods often fail to grasp the logic of the question, potentially overlooking vital facts. In the reasoning stage, earlier research used graph network architectures for reasoning [12, 42]. Some recent approaches now employ LLMs for reasoning [29, 35]. Think-on-graph applies the LLM as an interactive agent on KGs, using retrieved knowledge for reasoning [29]. However, these methods rely on strong reasoning capabilities, while LLMs with less strong reasoning ability tend to underperform with these methods.

6 CONCLUSION

We propose a novel LLM editing framework for multi-hop QA. We employ a strategy that utilizes mutual information maximization for the retrieval of facts and a self-optimizing technique for pruning redundant information. This approach effectively addresses the challenges of integrating real-time knowledge updates in LLMs. Theoretical justifications and comprehensive evaluations demonstrate RAE’s effectiveness in enhancing the accuracy of updated LLM responses. Our contributions mark a substantial advancement in model editing research, presenting a promising direction for the integration of dynamic knowledge in language models.

A THEORETICAL JUSTIFICATION

THEOREM 1. *Given retrieved graph $G_S \in \mathcal{G}$, the latent concept θ_c , and the question q sampled conditioned on concept θ_c , there exists a mutual information inequality: $I(G_S; \theta_c) \geq I(G_S; q)$.*

PROOF. Consider latent concepts θ_c is inferred from a knowledge graph G_S , and a question q is then generated based on this concept. (For instance, this can be as straightforward as prompting LLMs to formulate a question about a given set of facts). The following Markov chain: $G_S \rightarrow \theta_c \rightarrow q$ exists, indicating that q is conditionally independent to G_S . According to the chain rule for mutual information, we can have:

$$I(G_S; \theta_c, q) = I(G_S; q) + I(G_S; \theta_c | q) = I(G_S; \theta_c) + I(G_S; q | \theta_c). \quad (15)$$

Given that the question q is conditionally independent of the graph G_S , the term $I(G_S; q | \theta_c)$ equals zero. Therefore, we are left with: $I(G_S; \theta_c) \geq I(G_S; q)$. Equality holds precisely when $I(G_S; q | \theta_c)$ is zero, meaning that the graph G_S provides no additional information about q once θ_c is known. $\square$

B IMPLEMENTATION DETAILS

In our fact retrieval process, we employ a beam search with a width of 2, meaning we consider only two candidate relations at each hop. For the total number of retrieval hops n , we set it to 4, 5, and 6 for 2-hop, 3-hop, and 4-hop questions, respectively. For simplicity, we omit the term $p(q|h_1)$ as it remains constant for each question, and its exclusion does not impact empirical performance. Regarding the SEARC baseline, we pair smaller models with larger auxiliary models: GPT-2-small (124M) [25] with GPT-2-xl (1.3B) [25], GPT-2-large (774M) [25] with GPT-J (6B) [33], and GPT-Neo (2.7B) [3] with both Falcon (7B) [1] and Vicuna (7B) [4], as well as Llama2 (7B) [30]. We implement these models with the code and checkpoints available from Huggingface library [37].

ACKNOWLEDGMENTS

The work is, in part, supported by NSF (#IIS-2223768). The views and conclusions in this paper are those of the authors and should not be interpreted as representing any funding agencies.

REFERENCES

- [1] Ebtesam Almazrouei, Hamza Alobeidli, Abdulaziz Alshamsi, Alessandro Cappelli, Ruxandra Cojocaru, M  rouane Debbah,   tienne Goffinet, Daniel Hesslow, Julien Launay, Quentin Malartic, et al. 2023. The falcon series of open language models. *arXiv preprint arXiv:2311.16867* (2023).
- [2] Leonard E Baum and Ted Petrie. 1966. Statistical inference for probabilistic functions of finite state Markov chains. *The annals of mathematical statistics* 37, 6 (1966), 1554–1563.
- [3] Sid Black, Leo Gao, Phil Wang, Connor Leahy, and Stella Biderman. 2021. *GPT-Neo: Large Scale Autoregressive Language Modeling with Mesh-Tensorflow*. <https://doi.org/10.5281/zenodo.5297715> If you use this software, please cite it using these metadata..
- [4] Wei-Lin Chiang, Zhuohan Li, Zi Lin, Ying Sheng, Zhanghao Wu, Hao Zhang, Lianmin Zheng, Siyuan Zhuang, Yonghao Zhuang, Joseph E. Gonzalez, Ion Stoica, and Eric P. Xing. 2023. Vicuna: An Open-Source Chatbot Impressing GPT-4 with 90%* ChatGPT Quality. <https://lmsys.org/blog/2023-03-30-vicuna/>
- [5] Hyung Won Chung, Le Hou, Shayne Longpre, Barret Zoph, Yi Tay, William Fedus, Eric Li, Xuezhi Wang, Mostafa Dehghani, Siddhartha Brahma, et al. 2022. Scaling instruction-finetuned language models. *arXiv preprint arXiv:2210.11416* (2022).
- [6] Roi Cohen, Eden Biran, Ori Yoran, Amir Globerson, and Mor Geva. 2023. Evaluating the ripple effects of knowledge editing in language models. *arXiv preprint arXiv:2307.12976* (2023).
- [7] Damai Dai, Li Dong, Yaru Hao, Zhifang Sui, Baobao Chang, and Furu Wei. 2021. Knowledge neurons in pretrained transformers. *arXiv preprint arXiv:2104.08696* (2021).
- [8] Rajarshi Das, Ameya Godbole, Dilip Kavarthapu, Zhiyu Gong, Abhishek Singhal, Mo Yu, Xiaoxiao Guo, Tian Gao, Hamed Zamani, Manzil Zaheer, et al. 2019. Multi-step entity-centric information retrieval for multi-hop question answering. In *Proceedings of the 2nd Workshop on Machine Reading for Question Answering*. 113–118.
- [9] Xiaoqi Han, Ru Li, Hongye Tan, Wang Yuanlong, Qinghua Chai, and Jeff Pan. 2023. Improving Sequential Model Editing with Fact Retrieval. In *Findings of the Association for Computational Linguistics: EMNLP 2023*. 11209–11224.
- [10] Thomas Hartvigsen, Swami Sankaranarayanan, Hamid Palangi, Yoon Kim, and Marzyeh Ghassemi. 2022. Aging with GRACE: Lifelong Model Editing with Discrete Key-Value Adaptors. *arXiv preprint arXiv:2211.11031* (2022).
- [11] Gautier Izacard, Mathilde Caron, Lucas Hosseini, Sebastian Riedel, Piotr Bojanowski, Armand Joulin, and Edouard Grave. 2021. Unsupervised dense information retrieval with contrastive learning. *arXiv preprint arXiv:2112.09118* (2021).
- [12] Jinhao Jiang, Kun Zhou, Wayne Xin Zhao, and Ji-Rong Wen. 2022. Great Truths are Always Simple: A Rather Simple Knowledge Encoder for Enhancing the Commonsense Reasoning Capacity of Pre-Trained Models. *arXiv preprint arXiv:2205.01841* (2022).
- [13] Jinhao Jiang, Kun Zhou, Wayne Xin Zhao, and Ji-Rong Wen. 2022. Unikgqa: Unified retrieval and reasoning for solving multi-hop question answering over knowledge graph. *arXiv preprint arXiv:2212.00959* (2022).
- [14] Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich K  ttler, Mike Lewis, Wen-tau Yih, Tim Rocktaschel, et al. 2020. Retrieval-augmented generation for knowledge-intensive nlp tasks. *Advances in Neural Information Processing Systems* 33 (2020), 9459–9474.
- [15] Junyi Li, Jie Chen, Ruiyang Ren, Xiaoxue Cheng, Wayne Xin Zhao, Jian-Yun Nie, and Ji-Rong Wen. 2024. The Dawn After the Dark: An Empirical Study on Factuality Hallucination in Large Language Models. *arXiv preprint arXiv:2401.03205* (2024).
- [16] Junyi Li, Xiaoxue Cheng, Wayne Xin Zhao, Jian-Yun Nie, and Ji-Rong Wen. 2023. HELMA: A Large-Scale Hallucination Evaluation Benchmark for Large Language Models. *arXiv preprint arXiv:2305.11747* (2023).
- [17] Kevin Meng, David Bau, Alex Andonian, and Yonatan Belinkov. 2022. Locating and editing factual associations in GPT. *Advances in Neural Information Processing Systems* 35 (2022), 17359–17372.
- [18] Kevin Meng, Arnab Sen Sharma, Alex Andonian, Yonatan Belinkov, and David Bau. 2022. Mass-editing memory in a transformer. *arXiv preprint arXiv:2210.07229* (2022).
- [19] Eric Mitchell, Charles Lin, Antoine Bosselut, Chelsea Finn, and Christopher D Manning. 2021. Fast model editing at scale. *arXiv preprint arXiv:2110.11309* (2021).
- [20] Eric Mitchell, Charles Lin, Antoine Bosselut, Christopher D Manning, and Chelsea Finn. 2022. Memory-based model editing at scale. In *International Conference on Machine Learning*. PMLR, 15817–15831.
- [21] Ping Nie, Yuyu Zhang, Arun Ramamurthy, and Le Song. 2020. Answering Any-hop Open-domain Questions with Iterative Document Reranking. *arXiv preprint arXiv:2009.07465* (2020).
- [22] OpenAI. 2023. GPT-3.5. <https://openai.com/blog/gpt-3-5/>. Accessed on [Date].
- [23] OpenAI. 2023. Models Referred to as GPT-3.5. <https://platform.openai.com/docs/models/gpt-3-5>. Accessed on [Date].

- [24] Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. 2022. Training language models to follow instructions with human feedback. *Advances in Neural Information Processing Systems* 35 (2022), 27730–27744.
- [25] Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. 2019. Language models are unsupervised multitask learners. *OpenAI blog* 1, 8 (2019), 9.
- [26] CARNEGIE-MELLON UNIV PITTSBURGH PA DEPT OF COMPUTER SCIENCE. 1977. *Speech Understanding Systems. Summary of Results of the Five-Year Research Effort at Carnegie-Mellon University*.
- [27] Georgios Sidiropoulos, Nikos Voskarides, Svitlana Vakulenko, and Evangelos Kanoulas. 2021. Combining Lexical and Dense Retrieval for Computationally Efficient Multi-hop Question Answering. In *Proceedings of the Second Workshop on Simple and Efficient Natural Language Processing*, Nafise Sadat Moosavi, Iryna Gurevych, Angela Fan, Thomas Wolf, Yufang Hou, Ana Marasović, and Sujith Ravi (Eds.). Association for Computational Linguistics, Virtual, 58–63. <https://doi.org/10.18653/v1/2021.sustainlp-1.7>
- [28] Haitian Sun, Bhuwan Dhingra, Manzil Zaheer, Kathryn Mazaitis, Ruslan Salakhutdinov, and William W Cohen. 2018. Open domain question answering using early fusion of knowledge bases and text. *arXiv preprint arXiv:1809.00782* (2018).
- [29] Jiahuo Sun, Chengjin Xu, Luminyuan Tang, Saizhuo Wang, Chen Lin, Yeyun Gong, Heung-Yeung Shum, and Jian Guo. 2023. Think-on-graph: Deep and responsible reasoning of large language model with knowledge graph. *arXiv preprint arXiv:2307.07697* (2023).
- [30] Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro, Faisal Azhar, et al. 2023. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971* (2023).
- [31] Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajwal Bhargava, Shruti Bhosale, et al. 2023. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288* (2023).
- [32] Denny Vrandečić and Markus Krötzsch. 2014. Wikidata: a free collaborative knowledgebase. *Commun. ACM* 57, 10 (2014), 78–85.
- [33] Ben Wang and Aran Komatsuzaki. 2021. GPT-J-6B: A 6 Billion Parameter Autoregressive Language Model. <https://github.com/kingoflolz/mesh-transformer-jax>.
- [34] Weixuan Wang, Barry Haddow, and Alexandra Birch. 2023. Retrieval-augmented Multilingual Knowledge Editing. *arXiv preprint arXiv:2312.13040* (2023).
- [35] Xintao Wang, Qianwen Yang, Yongting Qiu, Jiaqing Liang, Qianyu He, Zhouhong Gu, Yanghua Xiao, and Wei Wang. 2023. Knowledgept: Enhancing large language models with retrieval and storage access on knowledge bases. *arXiv preprint arXiv:2308.11761* (2023).
- [36] Yiwei Wang, Muhao Chen, Nanyun Peng, and Kai-Wei Chang. 2024. DeepEdit: Knowledge Editing as Decoding with Constraints. *arXiv preprint arXiv:2401.10471* (2024).
- [37] Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz, et al. 2019. Huggingface’s transformers: State-of-the-art natural language processing. *arXiv preprint arXiv:1910.03771* (2019).
- [38] Xuansheng Wu, Huachi Zhou, Wenlin Yao, Xiao Huang, and Ninghao Liu. 2023. Towards Personalized Cold-Start Recommendation with Prompts. *arXiv preprint arXiv:2306.17256* (2023).
- [39] Sang Michael Xie, Aditi Raghunathan, Percy Liang, and Tengyu Ma. 2021. An explanation of in-context learning as implicit bayesian inference. *arXiv preprint arXiv:2111.02080* (2021).
- [40] Wenhao Xiong, Mo Yu, Xiaoxiao Guo, Hong Wang, Shiyu Chang, Murray Campbell, and William Yang Wang. 2019. Simple yet effective bridge reasoning for open-domain multi-hop question answering. *arXiv preprint arXiv:1909.07597* (2019).
- [41] Vikas Yadav, Steven Bethard, and Mihai Surdeanu. 2020. Unsupervised alignment-based iterative evidence retrieval for multi-hop question answering. *arXiv preprint arXiv:2005.01218* (2020).
- [42] Michihiro Yasunaga, Hongyu Ren, Antoine Bosselut, Percy Liang, and Jure Leskovec. 2021. QA-GNN: Reasoning with language models and knowledge graphs for question answering. *arXiv preprint arXiv:2104.06378* (2021).
- [43] Yuexiang Zhai, Shengbang Tong, Xiao Li, Mu Cai, Qing Qu, Yong Jae Lee, and Yi Ma. 2023. Investigating the Catastrophic Forgetting in Multimodal Large Language Models. *arXiv preprint arXiv:2309.10313* (2023).
- [44] Jing Zhang, Xiaokang Zhang, Jifan Yu, Jian Tang, Jie Tang, Cuiping Li, and Hong Chen. 2022. Subgraph retrieval enhanced model for multi-hop knowledge base question answering. *arXiv preprint arXiv:2202.13296* (2022).
- [45] Ce Zheng, Lei Li, Qingxiu Dong, Yuxuan Fan, Zhiyong Wu, Jingjing Xu, and Baobao Chang. 2023. Can We Edit Factual Knowledge by In-Context Learning? *arXiv preprint arXiv:2305.12740* (2023).
- [46] Zexuan Zhong, Zhengxuan Wu, Christopher D Manning, Christopher Potts, and Danqi Chen. 2023. MQuAKE: Assessing Knowledge Editing in Language Models via Multi-Hop Questions. *arXiv preprint arXiv:2305.14795* (2023).
- [47] Chen Zhu, Ankit Singh Rawat, Manzil Zaheer, Srinadh Bhojanapalli, Daliang Li, Felix Yu, and Sanjiv Kumar. 2020. Modifying memories in transformer models. *arXiv preprint arXiv:2012.00363* (2020).

Received 20 February 2007; revised 12 March 2009; accepted 5 June 2009