



Anonymous ACL submission

Abstract

Knowledge Editing (KE) has gained increasing attention, yet current KE tasks remain relatively simple. Under current evaluation frameworks, many editing methods achieve exceptionally high scores, sometimes nearing perfection. However, few studies integrate KE into real-world application scenarios (e.g., recent interest in LLM-as-agent). To support our analysis, we introduce a novel script-based benchmark – **SCEDIT** (Script-based Knowledge **E**dit**I**ng Benchmark) – which encompasses both counterfactual and temporal edits. We integrate token-level and text-level evaluation methods, comprehensively analyzing existing KE techniques. The benchmark extends traditional fact-based (“What”-type question) evaluation to action-based (“How”-type question) evaluation. We observe that all KE methods exhibit a drop in performance on established metrics and face challenges on text-level metrics, indicating a challenging task. Our benchmark will be made publicly available.

1 Introduction

Large Language Models (LLMs) have demonstrated outstanding performance in natural language understanding and generation tasks (Zhao et al., 2023). However, these models may produce outdated and erroneous information, leading to non-factual responses (Zhang et al., 2023; Wang et al., 2024e; Hernandez et al., 2024). Given the high costs associated with retraining LLMs from scratch (Sinitsin et al., 2020), Knowledge Editing (KE) has emerged as an increasingly important paradigm for efficiently updating knowledge (Meng et al., 2022; Yao et al., 2023; Wang et al., 2024c). KE methodologies have been developed to incrementally infuse new information or correct existing knowledge without requiring full-scale retraining (Mitchell et al., 2022a; Meng et al., 2022, 2023; Hartvigsen et al., 2022; Zhong et al., 2023). These techniques enable more efficient updates, ensuring

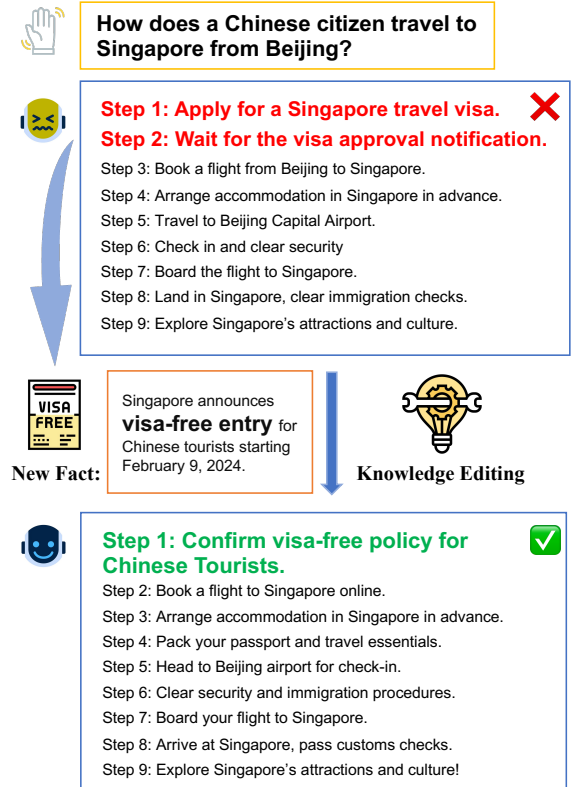


Figure 1: An example of the script-based assessment of Knowledge Editing (KE). **Top:** Outdated information generated by the LLM, instructing the user to apply for a visa, thereby misleading them. **Bottom:** Updated LLM successfully integrates new information, correctly informing the user about the visa-free policy.

continuous improvement and adaptation of LLMs (Zhang et al., 2024).

The conventional evaluation framework for KE largely relies on token-level metrics such as **Effectiveness**, **Generalization**, and **Specificity** (Meng et al., 2022). Although these metrics provide a great starting point, they exhibit notable limitations. For instance, **Generalization** attempts to transcend mere key-value pair memorization by evaluating a model’s capacity to answer rephrased or synonymously expressed questions. However,

such evaluations tend to remain in the realm of “What?”-type question transformations, overlooking the broader generalization capabilities of editing methods. Moreover, these three metrics typically gauge KE based on the next few tokens following prompts, leaving the potential for more complex reasoning and extended text generation largely unaddressed (Rosati et al., 2024).

In real-world scenarios, LLMs are increasingly deployed as agents that assist users in navigating daily life, making decisions, and performing complex tasks (Li et al., 2024; Sumers et al., 2024; Wang et al., 2024a; Lal et al., 2024). In these roles, users often pose “How?”-type questions, which require the models to generate goal-oriented **Scripts** not only recalling factual information, but applying, generalizing and reasoning based on that information (Lyu et al., 2021). A **Script** is a framework describing the sequence of events in a context. Specifically, in the context of KE, the rapidly changing landscape of factual knowledge means that generated **Script** may become erroneous, potentially misleading users. For example, as illustrated in Figure 1, a user might ask “How does a Chinese citizen travel to Singapore from Beijing?” A pre-trained LLM without updated knowledge might suggest applying for a visa, despite Singapore’s new visa exemption policy for Chinese tourists. Such questions necessitate a prompt and accurate update to ensure reliable responses.

Existing evaluation approaches, with their focus on token-level factual recall, do not sufficiently address these real-world complexities. To overcome these limitations, this paper introduces a script-based evaluation framework, named **SCEDIT**, assessing KE performance in a script-based scenario. **SCEDIT** emphasizes the model’s ability to handle “How?”-type questions and produce coherent and reliable guidance following targeted knowledge updates. We integrate token-level and text-level evaluation, comprehensively analyzing existing KE techniques in both counterfactual and temporal editing tasks. Three LLMs (GPT2-XL (Radford et al., 2019), GPT-J (Wang and Komatsuzaki, 2021) and Llama 3 (AI@Meta, 2024)) are tested in **SCEDIT**.

Experimental results on **SCEDIT** reveal that all comparable methods experience an average drop of 27% in the token-level metric S-ES compared to the similar PS metric introduced by Meng et al. (2022). Moreover, some methods struggle to balance effective editing with maintaining locality in both token-level and text-level evaluations. Even

methods that excel in token-level metrics show significant room for improvement in text-level editing performance. These findings highlight the need for further research into KE methods tailored for script-like scenarios.

We summarize our contributions of the paper as follows:

- **Develop a script-based assessment framework** that leverages scripts—structured procedural knowledge—to capture a model’s ability to integrate updated facts into complex reasoning and generation tasks. To the best of our knowledge, this is the first attempt to integrate KE into script-based scenarios, presenting a more challenging task compared to existing KE and constrained script generation tasks.
- **Introduce SCEDIT**, a novel and challenging script-focused benchmark, accompanied by comprehensive experiments to evaluate models’ ability at both token and text level.

2 Related Work

2.1 Scripts

A script is a structure that describes an appropriate sequence of events in a particular context (Schank and Abelson, 1975). Scripts are typically classified into *narrative scripts*, which describe a sequence of events in a story-like manner (Fang et al., 2022; Tandon et al., 2020), and *goal-oriented scripts*, which outline the steps needed to achieve a specific goal (Sancheti and Rudinger, 2021; Lyu et al., 2021). Our work aligns with the latter paradigm.

Generating high-quality scripts, a longstanding challenge, traditionally involves learning action sequences from narratives by analyzing causal relationships (Mooney and DeJong, 1985). Recently, script generation using large language models (LLMs) has become more feasible, with methods such as the over-generate-then-filter approach (Yuan et al., 2023). The script paradigm helps LLMs better understand the temporal order and logical relationships of everyday events. With fine-tuning and post-processing, models demonstrate enhanced generalization abilities in script generation (Sancheti and Rudinger, 2021). Smaller models, when trained on high-quality script datasets like CoScript, have shown superior constrained language planning quality compared to LLMs (Yuan et al., 2023).

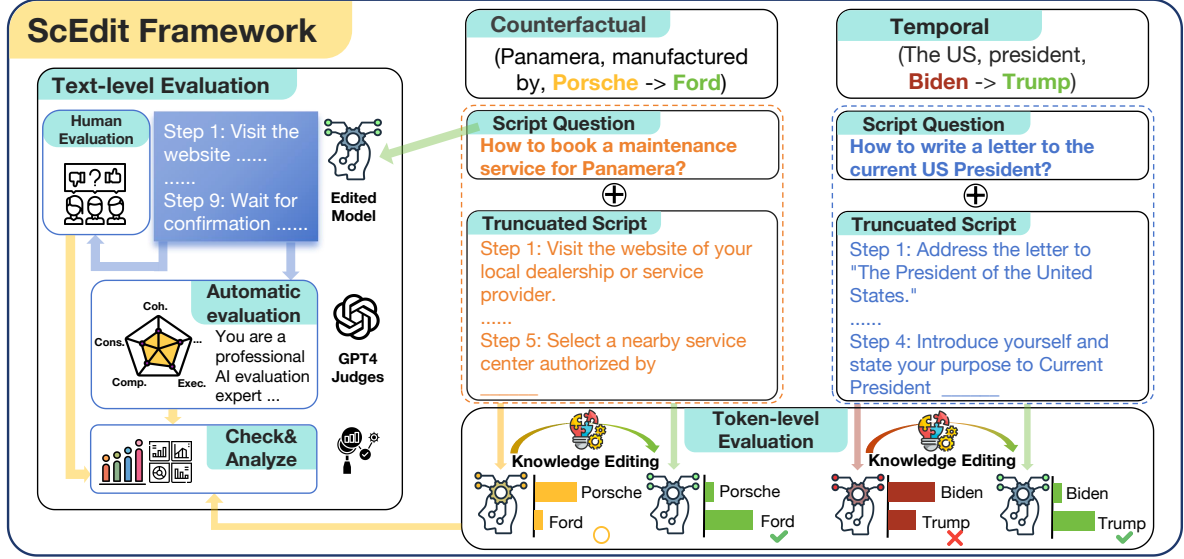


Figure 2: Overview of **SCEDIT**. For token-level evaluation, we concatenate the **Script Question** and **Truncated Script** to form a cloze-format prompt. For text-level evaluation, we involve automatic and human evaluation.

2.2 Knowledge Editing

Knowledge Editing (KE) has emerged as a promising approach to efficiently update LLMs without requiring full retraining (Sinitsin et al., 2020). Many applications and specific tasks also require ongoing adjustments to address defects or errors inherent in these models (Zhai et al., 2023). Current KE methods are generally classified into intrinsic and extrinsic approaches.

Intrinsic Methods. Intrinsic methods modify a model’s architecture or parameters to edit internal knowledge, including fine-tuning, meta-learning, and locate-then-edit approaches. Fine-tuning updates model parameters using new knowledge but demands high computational resources and risks catastrophic forgetting and overfitting (Chen et al., 2020; Zhu et al., 2020). Meta-learning methods like MEND (Mitchell et al., 2022a) and MALMEN (Tan et al., 2024) train a hyper-network to adjust weights indirectly, while locate-then-edit approaches like ROME (Meng et al., 2022) and MEMIT (Meng et al., 2023) use causal analysis of the hidden states to target specific areas storing knowledge.

Extrinsic Methods. Extrinsic methods use external knowledge to update the model’s input or output space, enhancing new representations while preserving original performance. A typical In-Context Learning method, IKE (Zheng et al., 2023), injects new knowledge by copying the updated facts into the context in a few-shot learn-

ing way. SERAC (Mitchell et al., 2022b) and MeLLO (Zhong et al., 2023) are both memory-based editing methods, but SERAC updates an external counterfactual model’s parameters and uses a classifier to determine if a fact update is needed, while MeLLO ensures fact updates through iterative prompting, making it more suitable for multi-hop reasoning.

Most prior work frames KE as a triplet-level task, updating entity-relation triples (subject, predicate, object) within LLMs (e.g., (The US, President, Biden $\mapsto$ Trump)). Some studies explore more extensive downstream applications (Wang et al., 2024d; Mao et al., 2023; Cheng et al., 2024; Wang et al., 2024f) or introduce more unstructured editing scenarios (Peng et al., 2024; Liu et al., 2024; Huang et al., 2024; Wu et al., 2024). However, they largely emphasize recalling edited facts while overlooking advanced reasoning and procedural capabilities (e.g., multi-step reasoning) essential for real-world tasks.

3 SCEDIT: Script-based Assessment of Knowledge Editing

We illustrate the proposed task in Figure 2. We will introduce the task definition (§3.1), dataset construction details in (§3.2), and the editing methods (§3.3) we used in the experiments.

3.1 Task Definition

KE was originally devised to update false or outdated information in a model, frequently by mutat-

ing fact-based triplets. Inspired by such approaches, we extend KE into *script-based* scenarios. In these scenarios, rather than merely performing a single-fact edit, the model must integrate newly updated knowledge into multi-step or procedural tasks. This shift offers an opportunity to assess whether models can propagate changes throughout an entire script, thereby providing a more comprehensive view of “editing success”.

Formally, we define three core elements:

- **Facts** are individual pieces of knowledge, often instantiated as (s, r, o) triples, where s is the subject, r the relation, and o the object. When performing a KE operation e , we apply

$$(s, r, o^c) \mapsto (s, r, o),$$

where o^c is the original object and o is the edited target object. Each fact has a fact prompt (s, r) directly related to s and r .

- **Script Questions** are prompts—typically starting with the word “How”—that require multi-step or procedural reasoning based on the updated fact. Because each fact can spawn multiple such questions, we denote them as

$$Q_{i,k}((s_i, r_i), o_i^c, o_i),$$

emphasizing that for fact i , there could be several questions indexed by k . These questions are designed so that the edit $e_i : (s_i, r_i, o_i^c) \mapsto (s_i, r_i, o_i)$ significantly affects the logic or flow of the script.

- **Scripts** are the model’s responses to *each* script question. For a **Script Question** $Q_{i,k}$, a LLM f_θ parameterized by θ and the Script $S_{i,k}$, we have

$$S_{i,k} = f_\theta(Q_{i,k}).$$

$S_{i,k}$ may or may not reflect the new object o_i , depending on whether the model has effectively understand the edit. The detailed format of **Scripts** can be found in Appendix A.

Based on the above elements, we evaluate KE performance using cloze-format prompts for **token-level** metrics (**ES**, **S-ES**, **S-NS**, **S-BO**) and automated/human evaluations for **text-level** metrics. Let f_θ be our large language model (LLM) parameterized by θ . $\mathbb{P}^c$ and $\mathbb{P}$ are the language model probability function before and after the update, respectively. Below we detail how we measure the edit e_i for each fact i . $\mathbb{E}_{i,k}[\cdot]$ denotes the average over all facts i and Script Questions k .

Efficacy. Following Meng et al. (2022), consider a fact prompt (s_i, r_i) whose original object is o_i^c and edited target object is o_i . **Efficacy Success (ES)** measures how often the model prefers o_i over o_i^c under this basic fact prompt:

$$\mathbb{E}_i[\mathbb{P}_{f_\theta}(o_i | (s_i, r_i)) > \mathbb{P}_{f_\theta}(o_i^c | (s_i, r_i))]. \quad (1)$$

Script-based Efficacy. We generalize **Efficacy** to the script-based setting. Given **Script Questions** $Q_{i,k}$, an external model (e.g., GPT-4) produces **Scripts** $S_{i,k}$ that intentionally include the old object o_i^c with original knowledge. To align with token-level evaluation, we **truncate** each original script $S_{i,k}$ at the point where o_i^c first appears, then concatenate this truncated script with $Q_{i,k}$ to form a cloze-format script-based prompt $\widetilde{Q}_{i,k}$. We compute **Script-based Efficacy Success (S-ES)** by checking whether f_θ prefers o_i to o_i^c under $\widetilde{Q}_{i,k}$:

$$\mathbb{E}_{i,k}[\mathbb{P}_{f_\theta}(o_i | \widetilde{Q}_{i,k}) > \mathbb{P}_{f_\theta}(o_i^c | \widetilde{Q}_{i,k})]. \quad (2)$$

Script-based Specificity. A robust editing process should not inadvertently corrupt unrelated or neighbor facts. Specifically, if (s_i, r_i, o_i^c) is replaced with (s_i, r_i, o_i) , then k collected neighbor facts (s_j, r_j, o_j) that share (r_i, o_i^c) or are semantically close to (s_i, r_i, o_i^c) should remain intact. Concretely, we construct a cloze-format, script-based neighborhood prompt $\widetilde{Q}_{i,k}'$ —analogous to $\widetilde{Q}_{i,k}$ but designed around these unmodified neighbor facts—and verify that the model retains the correct object o_j . Formally, for the first type of neighbor facts, which $o_j = o_i^c$, we define **Script-based Neighbor Success (S-NS)**:

$$\mathbb{E}_{i,k}[\mathbb{P}_{f_\theta}(o_i^c | \widetilde{Q}_{i,k}') > \mathbb{P}_{f_\theta}(o_i | \widetilde{Q}_{i,k}')]. \quad (3)$$

For the second type without o_i^c , inspired by Amar Khodja et al. (2024), we assess the accuracy drop of o_j via **Script-based Bleedover (S-BO)**:

$$\mathbb{E}_{i,k}[\max(\mathbb{P}_{f_\theta}^c(o_j | \widetilde{Q}_{i,k}') - \mathbb{P}_{f_\theta}(o_j | \widetilde{Q}_{i,k}'), 0)]. \quad (4)$$

Text-level Metrics. Beyond token probabilities, we assess the entire generated script’s quality by having the LLM answer Script Question $Q_{i,k}$ and conducting 7-point Likert-scale ratings across four dimensions via automatic and human evaluations.

1. **Executability (Exec.):** Are the script executable in a logical sense? ¹

¹For **Executability**, We do not consider the knowledge updates but focus solely on its inherent linguistic performance.

2. **Coherence (Coh.):** Are the script aligned with the newly updated fact?
3. **Consistency (Cons.):** Does the script remain free of internal contradictions?
4. **Completeness (Comp.):** Does the script adequately address all parts of the question, with sufficient procedural detail to be followed?

Detailed evaluation criteria and relative prompts are provided in Appendix C.1, with a further case study elaborating the metrics more in Appendix F.

3.2 Datasets

Task	Case	S-Eff.	S-Spec.
SCEDIT-CF	1830	7342	13672
SCEDIT-T	1762	7038	6597

Table 1: Statistics of our SCEDIT-CF and SCEDIT-T subtasks. “S-Eff.” denotes the sample size for Script-based Efficacy evaluation, while “S-Spec.” indicates the subset for measuring Script-based Specificity, ensuring that unrelated scripts remain correct after editing.

We introduce two subtasks, SCEDIT-CF and SCEDIT-T (Table 1), targeting different KE tasks. SCEDIT-CF centers on counterfactual knowledge, a common focus in KE, evaluating a method’s ability to perform edits in script-based scenarios. By contrast, SCEDIT-T utilizes temporal updates drawn from Wikipedia to assess a model’s adaptability to chronological updates, reflecting practical scenarios in which facts evolve over time.

An overview of the construction procedure is illustrated in Figure 3. Further details about the construction procedure can be found in Appendix B.

3.2.1 SCEDIT-CF Dataset Construction

We build on the **CounterFact** dataset (Meng et al., 2022), adapting it to script-based scenarios. In **CounterFact**, each fact (s, r, o^c) is replaced with (s, r, o) . To extend these edits into multi-step procedures, we design carefully formulated prompts and few-shot exemplars for a LLM (e.g., GPT-4) to generate several **Script Questions most likely** to be influenced by the updated fact. For example, given (Panamera, manufactured_by, Porsche $\mapsto$ Ford), a natural question might be “How to book a maintenance service for Panamera?”, which implicitly requires the updated manufacturer.

Following an initial generation step, we apply human filtering to ensure that the curated **Script Questions** $Q_{i,k}$ meaningfully hinge on the edited fact. We then prompt GPT-4 to generate scripts

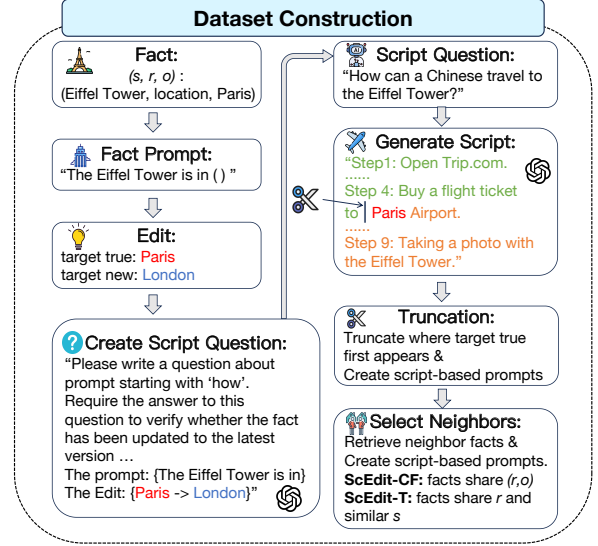


Figure 3: Overview of dataset construction process via a counterfactual edit as an example: Paris (ground truth) and London (target object). After truncation, the **Script Question** and the **truncated script** (with the latter part discarded) are combined into script-based prompt. Items marked with GPT-4 indicate GPT-4-generated content.

including the old object o^c with its original knowledge. To maintain consistency with prior work, we truncate first mentions of o^c and construct new script-based prompts $Q_{i,k}$ by appending $Q_{i,k}$ to the truncated script. Simultaneously, we filter out neighbor facts that share (r, o^c) to build script-based neighborhood prompts $\widetilde{Q}_{i,k}'$ in a similar way. The resulting question-answer pairs facilitate our core evaluations: **S-ES** (Script-based Efficacy Success), conducted using $Q_{i,k}$; **S-NS** (Script-based Neighborhood Success), performed with $\widetilde{Q}_{i,k}'$; and Text-level Metrics, evaluated using $Q_{i,k}$. We present a constructed example in Appendix B.

3.2.2 SCEDIT-T Dataset Construction

Constructed in a manner similar to SCEDIT-CF, SCEDIT-T leverages the **WDF_{real}**, a subset of **WikiFactDiff** (Ammar Khodja et al., 2024), testing whether the model can integrate temporal updates into scripts while preserving unrelated information. **WDF_{real}** gathers Wikipedia changes made between 4 January 2021 and 27 February 2023. Due to the data characteristics, a key difference from **CounterFact** is that Script-based Specificity set is retrieved from k -nearest neighbour fact (s', r, o') instead of (s, r, o^c) , where s' is a subject similar to s . Following a process similar to SCEDIT-CF, we construct $\widetilde{Q}_{i,k}'$ and measure accuracy degradation through **S-BO** (Script-based Bleedover).

Method	Model	SCEDIT-CF						SCEDIT-T					
		ES	↑	S-ES	↑	S-NS	↑	ES	↑	S-ES	↑	S-BO	↓
Base Model		20.55±1.85		21.18±1.51		81.52±1.20		44.27±2.32		41.72±2.03		0.00±0.00	
FT	GPT2-XL	100.00±0.00		<u>71.27±1.66</u>		65.08±1.51		87.17±1.56		52.80±2.03		1.15±0.14	
FT+L		99.13±0.42		40.39±1.84		78.50±1.26		70.60±2.13		44.39±2.03		0.39±0.08	
MEND		92.84±1.18		32.89±1.71		74.33±1.34		98.64±0.54		<u>74.24±1.77</u>		0.47±0.12	
ROME		<u>99.95±0.11</u>		74.76±1.56		<u>80.24±1.24</u>		<u>99.15±0.43</u>		68.00±1.86		<u>0.13±0.06</u>	
MEMIT		93.72±1.11		58.11±1.86		81.16±1.21		81.44±1.82		52.13±2.04		0.03±0.01	
PROMPT		96.28±0.87		69.63±1.66		42.88±1.44		99.49±0.33		84.39±1.44		0.54±0.08	
Base Model		13.99±1.59		16.06±1.31		85.77±1.05		40.64±2.29		39.62±1.99		0.00±0.00	
FT	GPT-J	100.00±0.00		<u>83.94±1.30</u>		25.81±1.26		99.60±0.29		97.9±0.56		5.47±0.38	
FT+L		<u>99.95±0.11</u>		39.07±1.81		<u>84.38±1.09</u>		71.51±2.11		42.78±1.99		<u>0.14±0.02</u>	
MEND		97.32±0.74		23.40±1.52		82.93±1.13		98.92±0.48		72.18±1.80		0.62±0.13	
ROME		<u>99.95±0.11</u>		86.50±1.14		83.35±1.13		99.60±0.29		74.29±1.73		0.28±0.08	
MEMIT		<u>99.95±0.11</u>		74.59±1.57		85.07±1.07		99.09±0.44		64.66±1.89		0.08±0.01	
PROMPT		90.55±1.34		70.95±1.61		44.01±1.47		98.24±0.61		<u>85.07±1.39</u>		1.03±0.11	
Base Model		7.32±1.19		9.19±0.97		92.53±0.70		-		-		-	
FT	LLAMA3	100.00±0.00		98.82±0.31		8.37±0.86		-		-		-	
ROME		<u>99.95±0.11</u>		<u>90.24±1.00</u>		<u>75.71±1.28</u>		-		-		-	
MEMIT		98.63±1.19		58.86±1.83		92.13±0.71		-		-		-	
PROMPT		92.30±1.22		77.02±1.46		56.48±1.30		-		-		-	

Table 2: Token-level results on the SCEDIT-CF and SCEDIT-T with their respective 95% confidence interval. Column-wise best results are highlighted in **bold green**, while the second-best results are underlined green. Values in **red** indicate a clear failure of a method on a particular metric. **S-ES** refers to Script-based Efficacy Success, **S-NS** is Script-based Neighborhood Success, and **S-BO** denotes Script-based Bleedover. SCEDIT-T was not evaluated on LLAMA3 because the cutoff date for its training data occurred after the time when the edited fact was introduced.

3.3 Editing Methods

SCEDIT primarily follows the single-edit paradigm. We include methods that excel within this paradigm, yet methods designed for massive or sequential editing (Tan et al., 2024; Hartvigsen et al., 2023; Fang et al., 2024; Wang et al., 2024b) remain unexplored and are considered as future work. Specifically, the editing methods employed in SCEDIT-CF and SCEDIT-T include:

- **Fine-tune (FT)**. A straightforward method that updates model weights via Adam optimization. Constrained Fine-Tuning (FT+L) (Zhu et al., 2020) further applies a L_∞ norm constraint, thereby limiting large parameter shifts.
- **ROME**. A parameter-editing technique that pinpoints the specific model weights driving factual predictions and directly modifies them to embed new or revised facts (Meng et al., 2022).
- **MEMIT**. A scalable multi-layer update algorithm built on ROME. It targets the relevant trans-

former module weights to handle multiple edits in parallel, enabling broader yet controlled updates (Meng et al., 2023).

- **MEND**. A meta-learning approach that learns auxiliary networks for fast, localized parameter adjustments, integrating new facts while preserving unrelated knowledge (Mitchell et al., 2022a).
- **PROMPT**. In addition to the methods in (§2.2), we evaluate PROMPT, which updates the model’s knowledge at inference by prefixing each prompt with $(s, r) + o$ - appending the target object to the fact prompt.

4 Experiments

4.1 Results on Token-level Metrics

Following previous KE evaluation paradigm, we use cloze-format prompts to assess token-level metrics, highlight the challenges of SCEDIT. S-ES, a metric akin to the PS (Paraphrase Success) intro-

Text-Level Metrics on SCEDIT-CF				
Method	Exec. $\uparrow$	Coh. $\uparrow$	Cons. $\uparrow$	Comp. $\uparrow$
LLAMA3-8B				
Base Model	6.74 ± 0.02	2.48 ± 0.03	6.86 ± 0.02	5.40 ± 0.05
FT	2.94 ± 0.05	2.97 ± 0.05	6.17 ± 0.05	2.17 ± 0.05
ROME	<u>6.41 ± 0.03</u>	<u>4.32 ± 0.05</u>	<u>6.57 ± 0.04</u>	4.67 ± 0.05
MEMIT	6.54 ± 0.02	3.67 ± 0.05	6.70 ± 0.03	<u>4.98 ± 0.05</u>
PROMPT	6.36 ± 0.03	4.35 ± 0.05	6.05 ± 0.05	5.49 ± 0.04

Table 3: Automatic evaluation results of four text-level metrics on SCEDIT-CF across different methods tested in LLAMA3-8B along with their respective 95% confidence interval. Column-wise best results are highlighted in **bold green**, while the second-best results are underlined green. In contrast, **red** values denote a clear failure in specific metric.

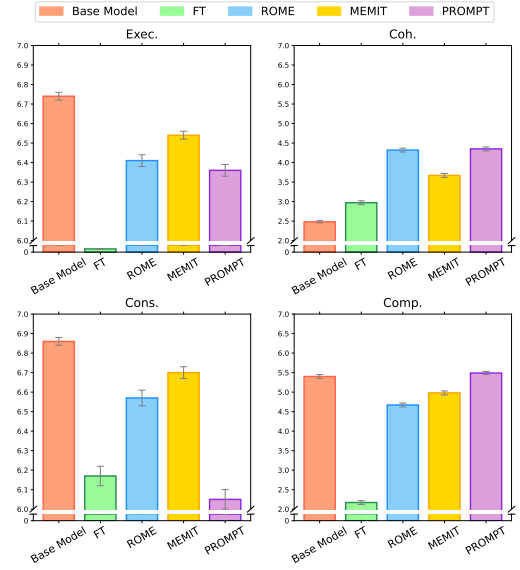


Figure 4: Results of text-level metrics. For clarity, the vertical axes for “Exec.” and “Cons.” begin at 6, while those for others start at 2.

duced by Meng et al. (2022) in both purpose and design, drops by an average of 27% compared to the original PS across all reported methods.

Certain methods reaffirm existing findings, whereas others unveil task-specific nuances in these script-based edits. Based on Table 2, we can draw:

FT and **FT+L** highlight the challenge of balancing effective edits with preserving locality. While FT excels in S-ES, its strong bias toward generating the targeted object hampers S-NS and S-BO. This trade-off worsens with larger models. By contrast, FT+L attempts to impose a constraint but falls short on S-ES, rendering it nearly unusable.

MEND displays divergent behavior by task. For SCEDIT-CF, S-ES drops by roughly 53% compared to simpler PS tasks. It should be noted that WikiText-based training may contribute to the performance gap, yet the drastic drop remains noteworthy, especially since PS and S-ES share same CounterFact edits. This suggests potential difficulties in adapting to different tasks. In contrast, MEND remains comparatively more viable for SCEDIT-T.

ROME achieves the best overall results across all models and all metrics, suggesting that a locate-then-edit strategy still offers strong performance in script-based scenarios.

MEMIT, designed for large-scale editing, exhibits moderate S-ES but attains the highest S-NS and S-BO scores, indicating particularly strong preservation of unrelated facts.

Lastly, although **PROMPT** excels in S-ES for

SCEDIT-T, its less favorable locality metrics reveal limitations when handling script-based contexts.

4.2 Results on Text-Level Metrics

4.2.1 Automatic Evaluation

We use GPT-4 to evaluate four text-level metrics on SCEDIT-CF for LLAMA3-8B-generated scripts after editing. While token-level metrics primarily capture edit performance, text-level metrics offer a more holistic assessment of how well a model integrates, generates, and reasons based on edited knowledge. Table 3 and Figure 4 shows the results.

Coherence. Coh. evaluates text-level edit effectiveness. PROMPT and ROME perform relatively well, aligning with their high S-ES scores. However, with 7 as the maximum score, these results remain unsatisfactory, highlighting even token-level strong methods still have room for improvement at the text level. In contrast, FT nearly wipes out the model’s capabilities, fixating on the target object o or and even part of its tokens, which can hardly be considered an effective text-level edit.

Executability and Completeness. Exec. and Comp. do not directly assess the newly edited facts but rather probe whether the model’s inherent script-related capabilities remain intact following the edits. MEMIT achieves the strongest performance here, possibly at the cost of Coh. ROME and PROMPT also perform well, with PROMPT even outperforming the Base Model in terms of

Comp., suggesting that it remains largely unaffected in terms of interpreting the **Script Question** and providing a well-rounded response. By contrast, FT registers poor results again, reflecting the irreparable damage it causes to the model’s broader script-related capabilities.

Consistency. Cons. checks whether the knowledge is stable, without mixing old and new facts. MEMIT and ROME both preserve consistency effectively, whereas PROMPT underperforms slightly here. This observation underscores that methods relying on in-context learning of the model can still face challenges in maintaining stable, conflict-free edits at the textual level.

4.2.2 Human Evaluation

Given the complexity of automated text-level evaluations, we further conduct a human evaluation on 400 sampled generated scripts. Three independent annotators, experienced in KE but uninvolved in the automated evaluation, scored the same four text-level metrics using the same criteria as GPT-4. Krippendorff’s α of 0.43 and Spearman’s β of 0.72 (between human and automated measures) indicate moderate to substantial agreement. Detailed statistics and analysis are provided in Appendix C.2.

4.3 Analysis of the Correlation of All Metrics

Inspired by Rosati et al. (2024), we analyze relationships between all metrics. GE² is included here. This represents a first attempt to integrate generative ability into KE evaluation. However, calculating entropy using short n-grams cannot fully capture the information present at the text level.

Figure 5 presents a clustered Spearman correlation heatmap comparing token-level with text-level metrics. All statistically significant correlations (with $p < 0.05$ and $|\rho| > 0.1$) are detailed and further analyzed in Appendix E.

In summary, our analysis yields three findings:

1. Fact-based efficacy (ES) alone fails to capture editing effectiveness in script scenarios.
2. Combining Exec. and Comp. — which incorporate the script’s inherent feature — provides a valuable complement to generative ability and specificity.
3. Text-level edit effectiveness (Coh.) weakly correlates with S-ES, while Cons. shows al-

²Weighted average of bi- and tri-gram entropies (Zhang et al., 2018) employed by Meng et al. (2022) in the original ROME papers.

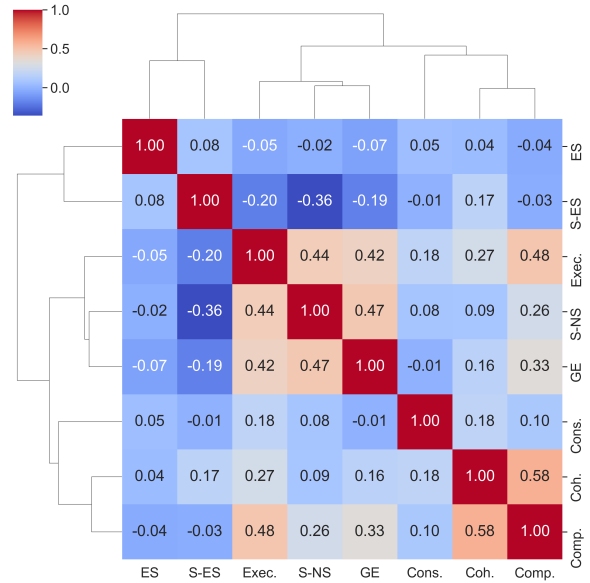


Figure 5: Clustered spearman correlation heatmap of token-level and text-level metrics

most no relation with token-level metrics, suggesting that each captures unique dimensions; therefore, integrating metrics across levels may yield a more robust evaluation.

5 Discussion

Both text-level and token-level metrics reveal an inherent trade-off in **SCEDIT** between achieving highly effective edits and limiting their broader impact on the model’s performance.

The deterministic nature of scripts enables a more definitive evaluation. Issues become more apparent at text level, highlighting the challenges of holistic editing in script-like scenarios, which require further research and advanced KE strategies.

6 Conclusion

In this paper, we present **SCEDIT**, a novel script-based benchmark for evaluating KE methods in real-world scenarios. Through rigorous experiments, we highlight several limitations of current KE methods in handling script-based evaluations. Some methods like FT struggles to maintain Efficacy and Specificity. While methods like ROME achieve strong token-level performance, text-level scores reveal room for improvement in script-like scenarios. Further analysis between token-level and text-level metrics underscores the need for more comprehensive evaluation frameworks. We hope that **SCEDIT** will inspire the development of more advanced KE techniques capable of addressing real-world complexities.

651	Yash Kumar Lal, Li Zhang, Faeze Brahman, Bodhisattwa Prasad Majumder, Peter Clark, and Niket Tandon. 2024. Tailoring with targeted precision: Edit-based agents for open-domain procedure customization . In <i>Findings of the Association for Computational Linguistics: ACL 2024</i> , pages 15597–15611, Bangkok, Thailand. Association for Computational Linguistics.	<i>Proceedings of Machine Learning Research</i> , pages 15817–15831. PMLR.	708 709
652			
653			
654		Raymond J Mooney and Gerald DeJong. 1985. Learning schemata for natural language processing. In <i>IJCAI</i> , pages 681–687.	710 711 712
655			
656		Hao Peng, Xiaozhi Wang, Chunyang Li, Kaisheng Zeng, Jiangshan Duo, Yixin Cao, Lei Hou, and Juanzi Li. 2024. Event-level knowledge editing . <i>Preprint</i> , arXiv:2402.13093.	713 714 715 716
657			
658			
659	Yuanchun Li, Hao Wen, Weijun Wang, Xiangyu Li, Yizhen Yuan, Guohong Liu, Jiacheng Liu, Wenxing Xu, Xiang Wang, Yi Sun, Rui Kong, Yile Wang, Hanfei Geng, Jian Luan, Xuefeng Jin, Zilong Ye, Guanqing Xiong, Fan Zhang, Xiang Li, Mengwei Xu, Zhijun Li, Peng Li, Yang Liu, Ya-Qin Zhang, and Yunxin Liu. 2024. Personal llm agents: Insights and survey about the capability, efficiency and security . <i>Preprint</i> , arXiv:2401.05459.	Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. 2019. Language models are unsupervised multitask learners. <i>OpenAI blog</i> , 1(8):9.	717 718 719 720
660			
661		Domenic Rosati, Robie Gonzales, Jinkun Chen, Xuemin Yu, Yahya Kayani, Frank Rudzicz, and Hassan Sajjad. 2024. Long-form evaluation of model editing . In <i>Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)</i> , pages 3749–3780, Mexico City, Mexico. Association for Computational Linguistics.	721 722 723 724 725 726 727 728 729
662			
663			
664			
665			
666			
667			
668	Jiateng Liu, Pengfei Yu, Yuji Zhang, Sha Li, Zixuan Zhang, Ruhi Sarikaya, Kevin Small, and Heng Ji. 2024. EVEDIT: Event-based knowledge editing for deterministic knowledge propagation . In <i>Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing</i> , pages 4907–4926, Miami, Florida, USA. Association for Computational Linguistics.	Abhilasha Sancheti and Rachel Rudinger. 2021. What do large language models learn about scripts? <i>arXiv preprint arXiv:2112.13834</i> .	730 731 732
669			
670			
671			
672			
673			
674			
675			
676	Yujie Lu, Weixi Feng, Wanrong Zhu, Wenda Xu, Xin Eric Wang, Miguel Eckstein, and William Yang Wang. 2023. Neuro-symbolic procedural planning with commonsense prompting . In <i>The Eleventh International Conference on Learning Representations</i> .	Roger C Schank and Robert P Abelson. 1975. Scripts, plans, and knowledge. In <i>IJCAI</i> , volume 75, pages 151–157. New York.	733 734 735
677			
678			
679			
680			
681	Qing Lyu, Li Zhang, and Chris Callison-Burch. 2021. Goal-oriented script construction . In <i>Proceedings of the 14th International Conference on Natural Language Generation</i> , pages 184–200, Aberdeen, Scotland, UK. Association for Computational Linguistics.	Anton Sinitin, Vsevolod Plokhotnyuk, Dmitriy Pyrkun, Sergei Popov, and Artem Babenko. 2020. Editable neural networks . <i>Preprint</i> , arXiv:2004.00345.	736 737 738
682			
683			
684			
685			
686	Shengyu Mao, Ningyu Zhang, Xiaohan Wang, Mengru Wang, Yunzhi Yao, Yong Jiang, Pengjun Xie, Fei Huang, and Huajun Chen. 2023. Editing personality for llms .	Theodore R. Sumers, Shunyu Yao, Karthik Narasimhan, and Thomas L. Griffiths. 2024. Cognitive architectures for language agents . <i>Preprint</i> , arXiv:2309.02427.	739 740 741 742
687			
688			
689			
690	Kevin Meng, David Bau, Alex Andonian, and Yonatan Belinkov. 2022. Locating and editing factual associations in GPT. <i>Advances in Neural Information Processing Systems</i> , 35.	Chenmian Tan, Ge Zhang, and Jie Fu. 2024. Massive editing for large language models via meta learning . In <i>International Conference on Learning Representations</i> .	743 744 745 746
691			
692			
693			
694	Kevin Meng, Arnab Sen Sharma, Alex J Andonian, Yonatan Belinkov, and David Bau. 2023. Mass-editing memory in a transformer . In <i>The Eleventh International Conference on Learning Representations</i> .	Niket Tandon, Keisuke Sakaguchi, Bhavana Dalvi, Dheeraj Rajagopal, Peter Clark, Michal Guerquin, Kyle Richardson, and Eduard Hovy. 2020. A dataset for tracking entities in open domain procedural text . In <i>Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)</i> , pages 6408–6417, Online. Association for Computational Linguistics.	747 748 749 750 751 752 753 754
695			
696			
697			
698			
699	Eric Mitchell, Charles Lin, Antoine Bosselut, Chelsea Finn, and Christopher D Manning. 2022a. Fast model editing at scale . In <i>International Conference on Learning Representations</i> .	Ben Wang and Aran Komatsuzaki. 2021. GPT-J-6B: A 6 Billion Parameter Autoregressive Language Model. https://github.com/kingoflolz/mesh-transformer-jax .	755 756 757 758
700			
701			
702			
703	Eric Mitchell, Charles Lin, Antoine Bosselut, Christopher D. Manning, and Chelsea Finn. 2022b. Memory-based model editing at scale . In <i>International Conference on Machine Learning, ICML 2022, 17-23 July 2022, Baltimore, Maryland, USA</i> , volume 162 of	Lei Wang, Chen Ma, Xueyang Feng, Zeyu Zhang, Hao Yang, Jingsen Zhang, Zhiyuan Chen, Jiakai Tang, Xu Chen, Yankai Lin, Wayne Xin Zhao, Zhewei Wei,	759 760 761
704			
705			
706			
707			

762	and Jirong Wen. 2024a. A survey on large language model based autonomous agents . <i>Frontiers of Computer Science</i> , 18(6).	818
763		819
764		820
765	Peng Wang, Zexi Li, Ningyu Zhang, Ziwen Xu, Yunzhi Yao, Yong Jiang, Pengjun Xie, Fei Huang, and Hua-	821
766	jun Chen. 2024b. WISE: Rethinking the knowledge memory for lifelong model editing of large language models . In <i>The Thirty-eighth Annual Conference on Neural Information Processing Systems</i> .	822
767		823
768		824
769		825
770		
771	Song Wang, Yaochen Zhu, Haochen Liu, Zaiyi Zheng, Chen Chen, and Jundong Li. 2024c. Knowledge editing for large language models: A survey . <i>Preprint</i> , arXiv:2310.16218.	826
772		827
773		828
774		829
775	Xiaohan Wang, Shengyu Mao, Ningyu Zhang, Shumin Deng, Yunzhi Yao, Yue Shen, Lei Liang, Jinjie Gu, and Huajun Chen. 2024d. Editing conceptual knowledge for large language models . <i>Preprint</i> , arXiv:2403.06259.	830
776		831
777		832
778		833
779		834
780	Yuxia Wang, Minghan Wang, Muhammad Arslan Manzoor, Fei Liu, Georgi Georgiev, Rocktim Jyoti Das, and Preslav Nakov. 2024e. Factuality of large language models: A survey . <i>Preprint</i> , arXiv:2402.02420.	835
781		836
782		837
783		
784		
785	Zecheng Wang, Xinye Li, Zhanyue Qin, Chunshan Li, Zhiying Tu, Dianhui Chu, and Dianbo Sui. 2024f. Can we debias multimodal large language models via model editing? In <i>Proceedings of the 32nd ACM International Conference on Multimedia</i> , MM '24, page 3219–3228, New York, NY, USA. Association for Computing Machinery.	838
786		839
787		840
788		841
789		842
790		843
791		
792	Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Remi Louf, Morgan Funtowicz, Joe Davison, Sam Shleifer, Patrick von Platen, Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gugger, Mariama Drame, Quentin Lhoest, and Alexander Rush. 2020. Transformers: State-of-the-art natural language processing . In <i>Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations</i> , pages 38–45, Online. Association for Computational Linguistics.	844
793		845
794		846
795		847
796		848
797		849
798		
799		
800		
801		
802		
803		
804	Xiaobao Wu, Liangming Pan, William Yang Wang, and Anh Tuan Luu. 2024. AKEW: Assessing knowledge editing in the wild . In <i>Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing</i> , pages 15118–15133, Miami, Florida, USA. Association for Computational Linguistics.	850
805		851
806		852
807		853
808		854
809		855
810	Yunzhi Yao, Peng Wang, Bozhong Tian, Siyuan Cheng, Zhoubo Li, Shumin Deng, Huajun Chen, and Ningyu Zhang. 2023. Editing large language models: Problems, methods, and opportunities . In <i>Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing</i> , pages 10222–10240, Singapore. Association for Computational Linguistics.	856
811		857
812		
813		
814		
815		
816		
817		
	Siyu Yuan, Jiangjie Chen, Ziquan Fu, Xuyang Ge, Soham Shah, Charles Jankowski, Yanghua Xiao, and Deqing Yang. 2023. Distilling script knowledge from large language models for constrained language planning . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 4303–4325, Toronto, Canada. Association for Computational Linguistics.	860
		861
		862
		863
		864
		865
		866
	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. <i>arXiv preprint arXiv:2309.10313</i> .	867
		868
		869
		870
	Ningyu Zhang, Yunzhi Yao, Bozhong Tian, Peng Wang, Shumin Deng, Mengru Wang, Zekun Xi, Shengyu Mao, Jintian Zhang, Yuansheng Ni, Siyuan Cheng, Ziwen Xu, Xin Xu, Jia-Chen Gu, Yong Jiang, Pengjun Xie, Fei Huang, Lei Liang, Zhiqiang Zhang, Xiaowei Zhu, Jun Zhou, and Huajun Chen. 2024. A comprehensive study of knowledge editing for large language models . <i>Preprint</i> , arXiv:2401.01286.	867
		868
		869
		870
	Yizhe Zhang, Michel Galley, Jianfeng Gao, Zhe Gan, Xiujun Li, Chris Brockett, and Bill Dolan. 2018. Generating informative and diverse conversational responses via adversarial information maximization. <i>Advances in Neural Information Processing Systems</i> , 31.	867
		868
		869
		870
	Yue Zhang, Yafu Li, Leyang Cui, Deng Cai, Lemao Liu, Tingchen Fu, Xinting Huang, Enbo Zhao, Yu Zhang, Yulong Chen, Longyue Wang, Anh Tuan Luu, Wei Bi, Freda Shi, and Shuming Shi. 2023. Siren's song in the ai ocean: A survey on hallucination in large language models . <i>Preprint</i> , arXiv:2309.01219.	867
		868
		869
		870
	Wayne Xin Zhao, Kun Zhou, Junyi Li, Tianyi Tang, Xiaolei Wang, Yupeng Hou, Yingqian Min, Beichen Zhang, Junjie Zhang, Zican Dong, Yifan Du, Chen Yang, Yushuo Chen, Zhipeng Chen, Jin hao Jiang, Ruiyang Ren, Yifan Li, Xinyu Tang, Zikang Liu, Peiyu Liu, Jian-Yun Nie, and Ji-Rong Wen. 2023. A survey of large language models . <i>Preprint</i> , arXiv:2303.18223.	867
		868
		869
		870
	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? <i>CoRR</i> , abs/2305.12740.	867
		868
		869
		870
	Zexuan Zhong, Zhengxuan Wu, Christopher D. Manning, Christopher Potts, and Danqi Chen. 2023. Mquake: Assessing knowledge editing in language models via multi-hop questions . <i>CoRR</i> , abs/2305.14795.	867
		868
		869
		870
	Chen Zhu, Ankit Singh Rawat, Manzil Zaheer, Srinadh Bhojanapalli, Daliang Li, Felix Yu, and Sanjiv Kumar. 2020. Modifying memories in transformer models . <i>Preprint</i> , arXiv:2012.00363.	867
		868
		869
		870

A Scripts Format Details

In SCEDIT, we define the **Script** in a specific format as follows:

- A step-by-step guide for answering the **Script Question**.
- Consists of brief events starting from Step 1, with a maximum of 9 steps.
- Each step should be a single concise action statement in one line less than 8 words.
- Excludes any explanations, details, notes, or additional instructions.
- The script should contain only the steps, and nothing else.

We prompt the model with following prompt when executing text-level evaluation:

Provide a step-by-step guide in Script
 ↳ form for answering the question:
 ↳ {Script Question}.

The Script should consist of brief
 ↳ events starting from Step 1, with
 ↳ a maximum of 9 steps.

Each step should be a single concise
 ↳ action statement in one line less
 ↳ than 8 words.

Do not include any explanations, details
 ↳ , notes, or further instructions.

The script should consist only of the
 ↳ steps, and nothing else.

B Dataset Construction Details

We leverage *gpt-4o-2024-11-20* to construct SCEDIT-CF as well as SCEDIT-T.

Two main processes are utilized in our dataset construction. First, as illustrated in Figure 6, we prompt GPT-4 with few-shot examples to generate **Script Questions** that are significantly influenced by the edits. Second, shown in Figure 7, we prompt GPT-4 with specific requirements to generate **Scripts** to response these **Script Questions** based on old knowledge. These **Scripts** are then truncated and concatenated with the corresponding **Script Questions** to form a cloze-format, script-based prompts aligning with token-level evaluation paradigms. Neighboring facts are also collected and processed similarly.

We present an example of the constructed data in Table 4, which includes a specific script question $Q_{i,k}$, a cloze-format script-based prompt $\widetilde{Q_{i,k}}$, and a cloze-format script-based neighbor prompt $\widetilde{Q_{i,k}'}.$

Fact	(Eiffel Tower, location, Paris)
Edit	Paris → Rome
Fact Prompt	The Eiffel Tower is in ()
Script Question $Q_{i,k}$	How can a Chinese travel to the Eiffel Tower?
Generated Script	Step 1: Open Trip.com. Step 2: Search for international flights. Step 3: Select your preferred flight option. Step 4: Buy a flight ticket to Paris Airport. Step 5: Apply for a Schengen visa. Step 6: Book a hotel in Paris. Step 7: Fly to France and reach your hotel. Step 8: Visit the Eiffel Tower. Step 9: Take a photo with the Eiffel Tower.
Script-based Prompt $Q_{i,k}$	How can a Chinese travel to the Eiffel Tower? Step 1: Open Trip.com. Step 2: Search for international flights. Step 3: Select your preferred flight option. Step 4: Buy a flight ticket to ()
Neighbor Fact	(Louvre Museum, location, Paris)
Script-based Neighborhood Prompt $\widetilde{Q_{i,k}'}$	How does a tourist in Korea visit the Louvre Museum? Step 1: Apply for a Schengen visa (if required). Step 2: Book tickets for the Louvre Museum. Step 3: Select your preferred date and time. Step 4: Book a flight ticket to ()

Table 4: An example of the constructed data (showcasing only one question and one neighbor fact here).

C Text-level Evaluation Details

C.1 Evaluation Criterion and Prompts

We employ *gpt-4o-2024-11-20* with a few-shot approach to automatically evaluate text-level metrics in SCEDIT-CF.

Figures 8 to 11 present the evaluation criteria and carefully crafted prompts used in the evaluation process. For the sake of clarity, several few-shot examples were omitted.

[SYSTEM]

- You are a helpful assistant.

[USER]

- Script Definition: A step-by-step guide for answering the question. The Script should consist of brief events starting from Step 1, with a maximum of 9 steps.
- Now, I will give you a factual knowledge edit which includes prompt, ground truth and new target.
- Based on the edit, generate 4 most relevant Script question that may be influenced. Script answers will have drastically different results due to changed factual knowledge.
- The question should not leak the factual change, i.e. the question should never include ground truth or new target. Let's think step by step.
- The Script questions should be start from 'How to', 'How does a', etc. , asking a question for a step-by-step guide (not asking the influence like 'How does ... impact ...') and output in a JSON format.
- Fact change:
 - Prompt: {prompt}
 - Ground truth: {ground_truth}
 - Target new: {target_new}

Figure 6: Prompts for generating **Script Questions** that are significantly influenced by the knowledge updates.

[SYSTEM]

- You are a helpful assistant.

[USER]

- Script Definition: A step-by-step guide for answering the question. The Script should consist of brief events starting from Step 1, with a maximum of 9 steps. Each step should be a single concise action statement in one line less than 8 words. **Do not** include any explanations, details, notes, or further instructions. The script should consist only of steps, and nothing else.
- Now I'll give you a question, an object, and you need to generate a nine-step answer according to the script rules.
- Ensure that the Object I provided is included in the script.
- Each step must output 'step x', where x is the current step.
- You need to return the question and the answer you generated to me.
- Please format your response as a JSON object with only one key: 'question'. The value of 'question' should be question and the entire answer, not separate keys for question and answer.
- Answer question:
 - Question: {question}
 - Object: {object}

Figure 7: Prompts for generating **Scripts** based on given **Script Questions**.

[SYSTEM]

- You are a professional AI evaluation expert.

[USER]

- We have posed questions for each modified knowledge and asked the large language model to answer each question separately. Your task is to evaluate the model's response using the following dimension on a 7-point Likert scale (1=worst, 7=best).
- Dimension: Executability
 - Executability: Are the steps in the script executable in a logical sense?
 - Assess if steps are feasible in logical contexts.
 - It is solely from a semantic perspective, without considering real-life facts and knowledge, that judgment is made.
 - Low-score example: Contain non executable operations (such as repeatedly outputting "English").
- For each of the scripts, return a number indicating your rating for each script.
- Your reply should be in JSON format.
- Your response should not contain spaces or line breaks. Start with 'executability:' followed by rating. Then, provide a brief explanation ("reason") for your rating in one sentence.
- Script: {script}

Figure 8: Prompts for evaluating Executability.

Metric	Krippendorff's α
Executability	0.59
Coherence	0.35
Consistency	-0.09
Completeness	0.25

Table 5: Metric-wise results of inter-rater agreement between annotators.

C.2 Human Evaluation

We conducted a human evaluation with the help of three researchers experienced in KE, who were not involved in the automated evaluation process. The inter-rater Krippendorff's Alpha coefficient indicates moderate agreement ($\alpha = 0.45$). Detailed metric-wise results are presented in Table 5, where some metrics, such as Exec. and Coh., show higher agreement, while Cons. exhibits poor agreement, reflecting its subjective nature.

The human evaluation results align closely with the findings and conclusions of (§4.2.1). We acknowledge Rosati et al. (2024) and believe one metric with poor agreement does not undermine our overall findings, especially considering the use of a seven-point Likert scale. However, it still highlights the need for further exploration of editing stability in future research.

For comparisons between human and automatic measures, Krippendorff's α of 0.43 and Spearman's β of 0.72 indicate moderate to substantial agreement, suggesting that automated scores align reasonably well with human raters. Increasing the number of examples under few-shots settings could be a direction for future improvements.

D Experiment Setup

The experimental procedures undertaken in this study are based on ROME (Meng et al., 2022) and MEMIT (Meng et al., 2023). We also incorporate the LLAMA3 covariance matrix provided by (Fang et al., 2024), and adapt PROMPT details from (Ammar Khodja et al., 2024).

All experiments are conducted on a single A40 (48GB) GPU. The LLMs are loaded using HuggingFace Transformers (Wolf et al., 2020).

In SCEDIT, each edit is performed individually, focusing on the specified (s, r, o) triple in a single operation rather than in a sequential chain. After each evaluation, the edited model is reverted to its original state, ensuring that edits remain isolated and do not affect subsequent operations. This design allows for a controlled script-based assessment of each individual KE.

[SYSTEM]

- You are a professional AI evaluation expert.

[USER]

- We have conducted knowledge editing on a large language model, changing the answer to the prompt from target true to target new.
- We have posed questions for each modified knowledge, and asked the large language model to answer each question separately. The response of the large model is called a script. Your task is to evaluate the scripts using the following dimension on a 7-point Likert scale (1=worst, 7=best).
- Dimension: Coherence
 1. Coherence: Are the script aligned with the newly updated fact?
 2. Verify complete alignment with new updated knowledge.
 3. Low-score example: Generate a script based on the premise of using target true as the prompt answer or mixing old/new facts (e.g., saying 'France's capital is Marseille' then mentioning 'Eiffel Tower in Paris'). If the script has target true and target new, two points will be given, and only target true will receive one point.
 4. If there is no target true or target new in the script, give 3 points.
- For each of the scripts, return a number indicating your rating for each script.
- **Your reply should be in JSON format.**
- Start with “coherence:” followed by rating. Then, provide a brief explanation (“reason”) for your rating in one sentence.
- Question: {question}
- Prompt: {prompt}
- Target true: {target_true}
- Target new: {target_new}
- Script: {script}

Figure 9: Prompts for evaluating Coherence.

[SYSTEM]

- You are a professional AI evaluation expert.

[USER]

- We have conducted knowledge editing on a large language model, changing the answer to the prompt from target true to target new.
- We have posed questions for each modified knowledge, and asked the large language model to answer each question separately. The response of the large model is called a script. Your task is to evaluate the scripts using the following dimension on a 7-point Likert scale (1= worst, 7=best).
- Dimension: Consistency
 1. Consistency: Does the script remain free of internal contradictions?
 2. Please disregard any knowledge you have about the prompt.
 3. Please ignore factual and logical errors.
 4. If and only if there are both target true and target new in the script and it cannot be determined which of them the script was generated based on as a fact, give 1 points.
 5. If the script only has target new, or target true, or neither, give it 7 points.
 6. Note that we only compare whether specific words appear, such as whether Donald Trump appears. If the President of America appears, we do not consider it to be the word 'Donald Trump'.
- For each of the scripts, return a number indicating your rating for each script.
- **Your reply should be in JSON format.**
- Start with “consistency:” followed by rating. Then, provide a brief explanation ("reason") for your rating in one sentence.
- Question: {question}
- Prompt: {prompt}
- Target true: {target_true}
- Target new: {target_new}
- Script: {script}

Figure 10: Prompts for evaluating Consistency.

[SYSTEM]

- You are a professional AI evaluation expert.

[USER]

- We have conducted knowledge editing on a large language model, changing the answer to the prompt from target true to target new.
- We have posed questions for each modified knowledge and asked the large language model to answer each question separately. The response of the large model is called a script. Your task is to evaluate the scripts using the following dimension on a 7-point Likert scale (1= worst, 7=best).
- Dimension: Completeness
 1. Completeness : Does the script adequately address all parts of the question, with sufficient procedural detail to be followed?
 2. In this dimension, our main criterion for evaluation is the completeness of the script's response to the question.
 3. Please pay attention to the current factual knowledge:
 - 1) If the tag is "pre," use target true as the basis for determining whether the script meets the target true criteria for scoring.
 - 2) If the tag is "post," use target new as the basis for determining whether the script meets the target new criteria for scoring.
 4. Note that if neither of target true and target new is mentioned, no points will be deducted. Only score the completeness of the answer to the question based on the script.
- For each of the scripts, return a number indicating your rating for each script.
- **Your reply should be in JSON format.**
- Start with “completeness:” followed by rating. Then, provide a brief explanation ("reason") for your rating in one sentence.
- Question: {question}
- Prompt: {prompt}
- Target true: {target_true}
- Target new: {target_new}
- Tag: {tag}
- Script: {script}

Figure 11: Prompts for evaluating Completeness.

E Detailed Analysis of the Correlation of All Metrics

Metric Pair	ρ
Text-level Metrics	
Coh. vs. Comp.	0.58
Exec. vs. Comp.	0.48
Exec. vs. Coh.	0.27
Exec. vs. Cons.	0.18
Coh. vs. Cons.	0.18
Token-level Metrics	
S-NS vs. GE	0.47
S-NS vs. S-ES	-0.36
S-ES vs. GE	-0.19
Token-level vs. Text-level Metrics	
S-NS vs. Exec.	0.44
GE vs. Exec.	0.42
GE vs. Comp.	0.33
S-NS vs. Comp.	0.26
S-ES vs. Exec.	-0.20
S-ES vs. Coh.	0.17
GE vs. Coh.	0.16

Table 6: Statistically significant ($p < 0.05$) combined Spearman’s rank correlations for metric pairs with $|\rho| > 0.1$.

In this section, we present a thorough analysis of the correlations among various performance metrics computed at both the text and token levels. Table 6 summarizes all the statistically significant correlations (with $p < 0.05$ and $|\rho| > 0.1$) observed in our analysis.

E.1 Text-Level Metrics

Among the text-level metrics, the Comp. exhibits moderate to strong correlations with both Exec. and Coh., with Spearman’s rank correlation coefficients of $\rho = 0.48$ and $\rho = 0.58$, respectively. Exec. and Coh. shows a weak positive correlation ($\rho = 0.27$), suggesting a weak association between the editing effectiveness and the inherent script-based generation ability. This relationship, which may appear counterintuitive, could be attributed to the fine-tuning (FT) effects discussed in (§4.2.1). Furthermore, Cons. shows weak or negligible correlations with all other text-level metrics, implying that it likely captures a unique aspect of performance that is not reflected in the other measures.

E.2 Token-Level Metrics

At the token level, ES does not show a significant correlation with either the S-ES or S-NS. However, a moderate negative correlation exists between S-NS and S-ES ($\rho = -0.36$). This negative relationship indicates that relying solely on Fact Edit Efficacy may be insufficient in script scenarios. Additionally, GE exhibits a moderate positive correlation with S-NS ($\rho = 0.47$), suggesting an intuitive link between generative ability and the specificity

E.3 Cross-Level Correlations

Beyond the intra-level correlations, cross-level analysis reveals several interesting patterns. Notably, Exec. correlates moderately with both GE ($\rho = 0.42$) and S-NS ($\rho = 0.44$). Moreover, GE shows a moderate correlation with Comp. ($\rho = 0.33$). In contrast, S-ES is only weakly correlated with the Coh. ($\rho = 0.17$), and S-NS shows a weak correlation with Comp. ($\rho = 0.26$). This time, S-ES exhibits a weak but intuitive negative correlation with Exec. ($\rho = -0.20$).

F Case Study

We present a case study of the generated Script as shown in Table 7. In the table, red indicates the original facts, green denotes the edited facts, and blue represents all other facts.

In the **Base** Script, the term “Nederlandse Taal” is mentioned—this is originally in Dutch and translates into English as “Dutch Language.” Consequently, we rated its Coh. as 1 (Strongly Disagree), while awarding full marks for other metrics.

FT Script, however, simply repeats “English” in a loop. Although this complies with the editing directive, we do not consider it text-level coherent with the edited knowledge, so its Coh. received a score of 3 (Slightly Disagree). Moreover, it achieved the poorest performance in both Exec. and Comp.

ROME, which has shown great performance in our experiments, exhibits an unusual failure case here: rather than reflecting either the original or the updated knowledge, it introduces “Arabic.” This suggests that the editing direction may have gone awry. Notably, when we reviewed the token-level metrics, all metrics performed well. This case clearly underscores the importance of integrating text-level evaluation metrics. In light of this, we assigned a score of 2 (Disagree) for both Coh. and Comp.

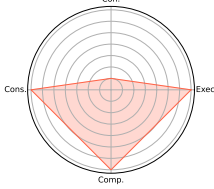
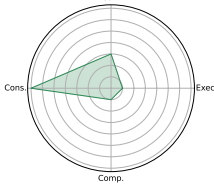
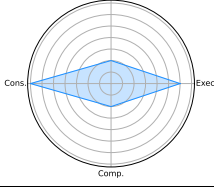
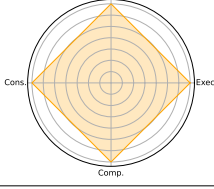
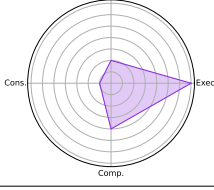
Fact	(Thomas Joannes Stieltjes, mother tongue, Dutch)
Fact Prompt	The mother tongue of Thomas Joannes Stieltjes is
Target True	Dutch
Target New	English
Script Question	How to discover the works of Thomas Joannes Stieltjes in his mother tongue?
Base Script Step 1: Open Google search engine. Step 2: Type "Stieltjes werken" in search bar. Step 3: Press Enter key. Step 4: Browse search results for "Stieltjes werken". Step 5: Click on " Nederlandse Taal " link. Step 6: Scroll down to "Werken" section. Step 7: Click on "PDF" link for "Recherches sur les fractions continues". Step 8: Download PDF file. Step 9: Open PDF file in Adobe Acrobat.	Evaluation (Exec: 7, Coh: 1, Cons: 7, Comp: 7) 
FT Script Step 1: Search Google for "English English English". Step 2: Click on the search results link. Step 3: Open the website "English English English". Step 4: Click on the "English English English" tab. Step 5: Click on the "English English English" link. Step 6: Click on the "English English English" link. Step 7: Click on the "English English English" link. Step 8: Click on the "English English English" link. Step 9: (repeated entries for English...)	Evaluation (Exec: 1, Coh: 3, Cons: 7, Comp: 1) 
ROME Script Step 1: Open Google search engine. Step 2: Type "Thomas Joannes Stieltjes in Arabic". Step 3: Press Enter key. Step 4: Review search results. Step 5: Look for official website. Step 6: Check for Arabic language option. Step 7: Select Arabic language option. Step 8: Read his works in Arabic. Step 9: Bookmark his website.	Evaluation (Exec: 6, Coh: 2, Cons: 7, Comp: 2) 
MEMIT Script Step 1: Search online for Stieltjes' biography. Step 2: Check the biography for native language. Step 3: Identify the native language as English. Step 4: Search online for Stieltjes' works in English. Step 5: Visit online archives of scientific journals. Step 6: Search journals for Stieltjes' publications. Step 7: Filter results by date of publication. Step 8: Download relevant articles in PDF format. Step 9: Read and analyze the articles.	Evaluation (Exec: 7, Coh: 7, Cons: 7, Comp: 7) 
PROMPT Script Step 1: Find Thomas Joannes Stieltjes' birthplace. Step 2: Research Dutch archives in that location. Step 3: Search for Stieltjes' writings in Dutch. Step 4: Visit libraries and museums in the Netherlands. Step 5: Consult experts in Dutch literature and history. Step 6: Check online archives and databases. Step 7: Translate relevant documents into English. Step 8: Read and analyze Stieltjes' works in English. Step 9: Share findings with the academic community.	Evaluation (Exec: 7, Coh: 2, Cons: 1, Comp: 4) 

Table 7: Case Study of our proposed text-level metrics. In the table, **red** indicates the original facts, **green** denotes the edited facts, and **blue** represents all other unrelated facts.

MEMIT performed well in this instance—not only did it successfully incorporate the new knowledge into the generated Script, but it also maintained the performance of other Script-related aspects, earning full marks.

Finally, **PROMPT** demonstrated an instance of unstable editing. Instead of recognizing that Thomas Joannes Stieltjes' mother tongue has been

updated to English, the model interpreted this as a directive to translate into English. As a result, it received 1 (Strongly Disagree) point for Cons. and 2 (Disagree) points for Coh., with Comp. scoring 4 (Neutral/Uncertain).