

CUB: Benchmarking Context Utilisation Techniques for Language Models

Anonymous ACL submission

Abstract

Incorporating external knowledge is crucial for knowledge-intensive tasks, such as question answering and fact checking. However, language models (LMs) may ignore relevant information that contradicts outdated parametric memory or be distracted by irrelevant contexts. While many context utilisation manipulation techniques (CMTs) that encourage or suppress context utilisation have recently been proposed to alleviate these issues, few have seen systematic comparison. In this paper, we develop CUB (Context Utilisation Benchmark) to help practitioners within retrieval-augmented generation (RAG) identify the best CMT for their needs. CUB allows for rigorous testing on three distinct context types, observed to capture key challenges in realistic context utilisation scenarios. With this benchmark, we evaluate seven state-of-the-art methods, representative of the main categories of CMTs, across three diverse datasets and tasks, applied to nine LMs. Our results show that most of the existing CMTs struggle to handle the full set of types of contexts that may be encountered in real-world retrieval-augmented scenarios. Moreover, we find that many CMTs display an inflated performance on simple synthesised datasets, compared to more realistic datasets with naturally occurring samples. Altogether, our results show the need for holistic tests of CMTs and the development of CMTs that can handle multiple context types.

1 Introduction

Context utilisation is a key component of language models (LMs) used for retrieval-augmented generation (RAG), as the benefits of retrieving external information are only realised if the generative model makes adequate use of the retrieved information. While recent research has identified many benefits of augmenting LMs with retrieved information (Shuster et al., 2021; Hagström et al., 2023), it has also identified weaknesses of LMs used for

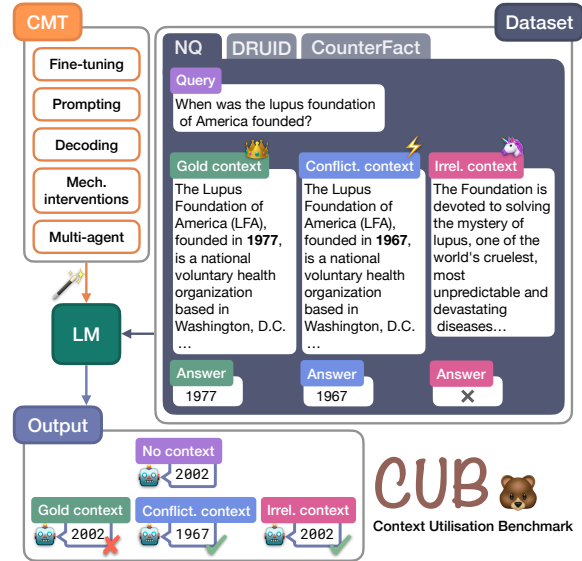


Figure 1: The Context Utilisation Benchmark. We evaluate a range of LMs under different CMTs on samples from NQ, DRUID and CounterFact for gold, conflicting and irrelevant contexts.

RAG, of which many are associated with context utilisation. For example, LMs can easily be distracted by irrelevant contexts (Shi et al., 2023) or ignore relevant contexts due to memory-context conflicts (Xu et al., 2024). The robustness of LMs to irrelevant contexts is important as information retrieval systems used for RAG are not guaranteed to always retrieve relevant information. Moreover, as information may be updated to conflict with the training data of the LM, the model should prioritise the most recently updated information.

As a consequence, many different methods for increasing or suppressing LM context utilisation, henceforth referred to as *CMTs* (Context utilisation Manipulation Techniques), have been proposed. The methods encompass a broad range of approaches, from different decoding methods (Shi et al., 2024; Kim et al., 2024) to fine-tuning methods (Li et al., 2023), prompting (Liu et al.,

2023), multi-agent (Feng et al., 2024; Du et al., 2024), and mechanistic interventions (Ortu et al., 2024; Jin et al., 2024). While each method yields promising results in isolation, their evaluation is often limited to narrow or idealised settings, leaving open the question of which approaches are applicable in real-world RAG scenarios. To address this evaluation gap, we develop a comprehensive CMT benchmark to test and compare different CMTs on datasets representative of different domains and tasks (Figure 1). Our **contributions** are as follows:

- We develop CUB (Context Utilisation Benchmark) to allow for a comprehensive evaluation and comparison of CMTs (§3).¹ CUB systematically tests the sensitivity of CMTs to underlying model and naturally occurring context types (gold, conflicting and irrelevant) on tasks representative of synthesised and realistic RAG scenarios.
- We evaluate a cohort of state-of-the-art CMTs representative of the main categories of CMTs (§4) on our benchmark (§6).
- We provide a deeper analysis of what CMT works best for a given scenario and identify areas of improvement for CMTs. We find that CMTs struggle to optimise performance across all context types, e.g. one approach may improve robustness to irrelevant contexts but degrade the utilisation of relevant contexts. This points to the need of CMTs that work well across all context types.

2 Related Work

Context-intensive datasets We consider two main categories of context-intensive datasets: 1) datasets representing *knowledge-intensive tasks*, i.e. tasks for which access to external context is crucial, and 2) datasets designed to *diagnose* model adaptability to external knowledge. Examples of datasets representative of knowledge-intensive tasks are Natural Questions (NQ), DRUID, the KILT datasets and PubMedQA (Kwiatkowski et al., 2019; Hagström et al., 2024; Petroni et al., 2021; Jin et al., 2019). Examples of diagnostic datasets representative of the latter category are CounterFact and ConflictQA (Meng et al., 2022; Xie et al., 2024a). These datasets contain synthesised queries based on fact triplets from LAMA (Petroni et al., 2019) (e.g. Thomas Ong-citizen of-Singapore) for which contexts have been synthesised to induce

knowledge conflicts by promoting answers in conflict with the parametric memory of the studied LM (e.g. “Pakistan” as opposed to “Singapore”). Diagnostic datasets have found widespread use for work on mechanistic interpretability and the evaluation of context utilisation (Meng et al., 2022; Geva et al., 2023; Ortu et al., 2024).

Previous work has typically evaluated different CMTs on either of the dataset categories. CUB incorporates datasets representative of both knowledge-intensive tasks and diagnostic datasets, thus enabling comprehensive evaluations of CMTs in different settings.

CMTs Many context utilisation manipulation techniques have recently been proposed. Existing CMTs can be categorised into one of four main groups based on *intervention level*, i.e. what aspect of the model they manipulate. 1) *fine-tuning* CMTs update model parameters to modify context utilisation. For example, fine-tuning on distracting contexts was found to yield improved robustness to distracting contexts (Li et al., 2023; Shen et al., 2024; Yoran et al., 2024). Moreover, Fang et al. (2024) specifically focus on different types of retrieval noise likely to be encountered in real-world environments and develop a fine-tuning approach to handle these. 2) *prompting techniques* modify the input to the LM to improve context utilisation, representing minimally modified settings. 3) *mechanistic interventions* on the LM modify certain model components at inference time to alter context utilisation. Examples involve attention modification (Ortu et al., 2024; Jin et al., 2024) and SpARe interventions (Zhao et al., 2025). Lastly, 4) *decoding methods* involve a modified decoding approach, applied to the output logits, to manipulate context utilisation. Examples include context-aware contrastive decoding (Yuan et al., 2024; Kim et al., 2024; Shi et al., 2024; Wang et al., 2024; Zhao et al., 2024) and lookback lens decoding (Chuang et al., 2024).

Apart from intervention level, many of the CMTs have different *objectives*, focused on improving one or multiple aspects of context utilisation. CMTs may focus on improving robustness to irrelevant contexts, faithfulness to conflicting contexts, or faithfulness to contexts in general.

Previous work has mainly focused on evaluating one CMT at a time, potentially due to the lack of a unified benchmark for CMTs. In this paper, we evaluate representatives from each of the four main

¹Dataset and code will be available upon publication.

Dataset	Split	#samples	%Gold	%Conflict.	%Irrel.
CounterFact	<i>dev</i>	198	33.3	33.3	33.3
	<i>test</i>	2,499	33.3	33.3	33.3
NQ	<i>dev</i>	198	33.3	33.3	33.3
	<i>test</i>	4,945	33.4	33.1	33.4
DRUID	<i>dev</i>	198	33.3	33.3	33.3
	<i>test</i>	4,302	43.5	56.1	0.4

Table 1: Statistics of the datasets that form CUB. ‘Conflict.’ denotes conflicting contexts and ‘Irrel.’ irrelevant contexts.

categories of CMTs on CUB, comparing a total of seven CMTs.

Benchmarks To the knowledge of the authors, there is not yet a benchmark for CMTs. The closest examples of existing benchmarks are RAG-Bench by Fang et al. (2024), KILT by Petroni et al. (2021) and AxBench by Wu et al. (2025). The first evaluates the retrieval-noise robustness of LMs, the second performance of RAG systems as a whole, and the latter steering techniques for LMs, focusing on safety and reliability. CUB takes inspiration from these benchmarks to create a comprehensive and relevant benchmark for the evaluation of CMTs.

3 CUB: A Context Utilisation Benchmark

Given a CMT, CUB is designed to test the technique across different datasets, models and metrics. To unify the tests, CUB also incorporates a pre-defined method for the hyperparameter search of the CMT.

3.1 Language Models

CUB evaluates the model sensitivity of CMTs on up to nine different LMs. The open-sourced models covered by the benchmark are GPT-2 XL, Pythia (6.9B), Qwen2.5 1.5B, Qwen2.5 7B, Qwen2.5 32B (Radford et al., 2019; Biderman et al., 2023; Yang et al., 2024). For the Qwen models we include the instruction-tuned variants. We also evaluate the API-based LLM Cohere Command A with 111B parameters.² The model selection is performed to enable comparisons across model families, model sizes, instruction-tuning and API-based LLMs. However, all LMs are not compatible with all CMTs evaluated on CUB – the selection of LMs onto which a CMT is applied depends on the CMT, further explained in Section 4. In addition, we adapt the prompts in CUB with prompt templates

²<https://cohere.com/blog/command-a>

compatible with each model type under consideration (base, instruction-tuned and chat-API).

3.2 Datasets

To evaluate how CMTs respond to different types of contextual information, CUB evaluates each CMT on CounterFact, NQ and DRUID (see Table 1). The inclusion of these datasets is based on three key criteria: (i) diversity in task difficulty, (ii) diversity in realistic and synthesised RAG scenarios, and (iii) high utilisation in related work. CounterFact represents a causal language modelling task based on a controlled setup with simple counterfactual contexts synthesised to conflict with model memory. NQ represents a popular, and more realistic setup, focused on RAG for open-domain QA of greater difficulty with contexts sampled from Wikipedia. DRUID is a fairly new dataset, representing another important RAG task – that of automated fact-checking; this requires a greater level of reasoning based on naturally occurring claims and evidence sampled from the internet. While DRUID has yet to see widespread use in studies of context utilisation, we include it in CUB as it is one of few datasets closely aligned with real-world RAG scenarios.

For each dataset, we curate samples representative of the three types of contexts that may be encountered in realistic RAG scenarios: 1) **gold** contexts that are relevant and do not contradict LM memory, 2) **conflicting** contexts that are relevant but contradict LM memory or gold labels, and 3) **irrelevant** contexts that should be ignored by the LM (Fang et al., 2024). For each dataset, we sample validation and test splits. To allow for fair and unified comparisons between CMTs, the validation set is used to tune potential hyperparameters of the CMT under evaluation. The test split is used for the final evaluation. More details on the datasets can be found in Appendix B.

CounterFact To construct a CounterFact dataset with counterfactual contexts, we first identify samples from LAMA that have been memorised by Pythia 6.9B, following the approach by Saynova et al. (2025). We base the CounterFact dataset on Pythia to obtain a set of samples likely to have been memorised by all CUB models, since LMs have been found to memorise more facts as they grow in size (Saynova et al., 2025). We confirm this in Appendix B; all CUB LMs are found to have memorised at least 70% of the CounterFact samples. Based on the known fact triplets, we sam-

ple conflicting contexts following the approach of Meng et al. (2022). We also sample gold contexts that simply state the correct triplet. For the irrelevant contexts, we randomly sample fact triplets unrelated to the sample query.

NQ The gold context samples are simply the original NQ samples. For the collection of samples with conflicting contexts, we follow a substitution approach inspired by the method of Longpre et al. (2021). We create conflicting contexts that promote a different answer simply by taking the gold context and substituting the gold answer in the context. The substitute answer is sampled to yield coherent conflicting contexts, and to have a different meaning compared to the gold answer. For the collection of samples with irrelevant contexts, we apply a LM re-ranker to identify the most relevant non-gold paragraph from the Wikipedia page in which the gold context was found. With this approach, we collect irrelevant contexts representative of real-world RAG scenarios.

DRUID The <claim, evidence> samples of DRUID have been manually annotated for stance of the evidence (supports, refutes, insufficient or irrelevant). We map stance to context type as described in Appendix B. No context synthesis is necessary for the DRUID samples as they, by virtue of utilising naturally occurring samples from a RAG pipeline, already contain samples representative of gold, conflicting and irrelevant contexts. Moreover, since DRUID represents a reasoning task, asking the model whether provided evidence supports the claim under consideration (True or False), or is insufficient (None), the output space for the DRUID samples is limited to three tokens (True, False or None).

3.3 Metrics

Similarly to Jin et al. (2024) we use a binary score to measure context utilisation. We refer to it as the *binary context utilisation* (BCU) score and define it as follows. For relevant contexts (gold and conflicting) the score is 1 if the LM prediction is the same as the token promoted by the context, t_C , and 0 otherwise. For irrelevant contexts the score is 1 if the LM prediction is the same as the memory token, t_M , (i.e. the prediction made by the model before any context has been introduced) and 0 otherwise. We report the averaged BCU score per context type. To assess the relative effectiveness of CMTs, we also report the net gain of each CMT,

Methods	Objective	Level	Tuning Cost	Inference Cost
Fine-tuning	Both	Fine-tuning	High	Low
Prompting	Both	Prompt.	Low	Mid
Multi-agent	Both	Prompt.	None	High
PH3 +context	Faith	Mech.	High	Low
COIECD	Faith	Decoding	Mid	Mid
PH3 +memory	Robust	Mech.	High	Low
ACD	Robust	Decoding	None	Mid

Table 2: Comparison of CMTs by objective, intervention level, and cost. The CMTs are coloured by objective with warm colours for ‘Both’, blue for ‘Faith’ and green for ‘Robust’. ‘Mech.’ denotes mechanistic interventions.

compared to when no CMT is applied, using BCU score ($\Delta = \text{BCU}_{\text{CMT}} - \text{BCU}_{\text{Regular}}$). We also consider *continuous context utilisation*, CCU, a more fine-grained metric that measures the change in outputted token probabilities as context is introduced. Appendix C contains more details on the metric.

We also measure the accuracy of each method. For CounterFact and DRUID, accuracy is measured based on whether the first generated token is the same as the first gold token. For NQ, for which the correct answer may be different permutations of the same set of tokens, we measure accuracy based on whether the first output token (e.g. “July”) matches any of the tokens in the answer (e.g. “15 July”).

3.4 Hyperparameter Search

For CMTs requiring hyperparameter tuning, we use the validation set of each dataset to select values that maximise the average BCU across all context types, unless a method-specific tuning procedure is explicitly specified. This ensures a fair comparison between CMTs. Further details are shown in Appendix D.

4 Context Utilisation Manipulation Techniques

We benchmark a total of seven different CMTs on CUB, all of which are state-of-the-art representatives from the main categories of CMTs. Table 2 summarises the key characteristics of the CMTs, including their main objective, intervention level, and cost in terms of tuning and inference. As a baseline, we also evaluate regular LMs on the same input, with no CMT applied (Regular).

Fine-tuning We adapt the approach of Li et al. (2023), which fine-tunes LMs to ensure the usage of relevant contexts. It considers four different

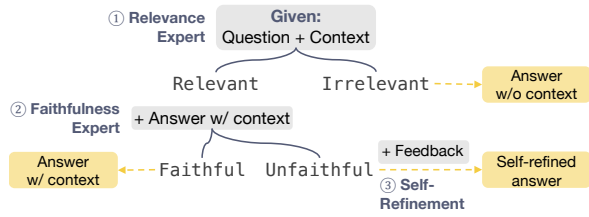


Figure 2: Overview of the multi-agent approach.

types of contexts: relevant, irrelevant, empty, and counterfactual contexts. To align the domain with our evaluation data, we curate the fine-tuning data with two QA datasets (Joshi et al., 2017; Rajpurkar et al., 2018), one FC dataset (Schlichtkrull et al., 2023), and one sentence completion dataset (Marjanovic et al., 2024). Before fine-tuning each LM, we elicit its parametric answers by querying without contexts. We then select the questions that the LM answered correctly and pair them with irrelevant and empty contexts. The fine-tuning data thus contains contexts that can be irrelevant, counterfactual, or empty. During fine-tuning, we train the LM to generate answers aligned with the provided context. When the context is irrelevant, we train the LM to be robust, i.e. ignore the context and output its parametric answer. Due to the computational costs associated with fine-tuning billion-sized LMs, we use the Low-Rank Adaptation method (Hu et al., 2021). Additional details can be found in Appendix E.

Prompting We curate a set of 12 prompts for each evaluation dataset and optimise the prompt selection to each evaluated model. Each set of prompts is based on 6 prompts curated by a human, similarly to the approach by Jin et al. (2024), and 6 prompts generated by a LLM,³ similarly to the approach by Wu et al. (2025).

Multi-agent Inspired by LM agents and self-refinement (Du et al., 2024; Feng et al., 2024; Madaan et al., 2023), which are widely adopted techniques in reasoning tasks, we decompose context utilisation into two components – relevance and context faithfulness – and assign each as a separate task to an individual LM agent. We aim to examine whether LMs are capable of accurately evaluating context relevance and answer faithfulness, to subsequently self-correct themselves for improved faithfulness to relevant contexts. As illustrated in Figure 2, we first assess relevance using

the relevance agent to determine whether the provided context should be used. Then, the faithfulness agent provides feedback on the model response that was generated with context. If the feedback indicates that the initial answer is unfaithful, the model generates a self-refined answer based on that feedback. Given that these tasks require instruction-following capabilities, we restrict our evaluation to instruction-tuned or chat LMs. Further details can be found in Appendix F.

Mechanistic interventions: PH3 We adopt the PH3 method by Jin et al. (2024). The method is implemented in two steps: 1) identification of attention heads responsible for context or memory reliance via path patching and 2) pruning the identified attention heads for increased memory or context usage. To identify attention heads, we use the CounterFact datasets with samples that elicit exact fact recall in each studied model (Saynova et al., 2025). For the evaluation on our studied datasets, we tune the number of heads to prune on the validation splits of each evaluation dataset, similarly to the approach by Jin et al. (2024). PH3 can be used in two different modes – suppressing context attention heads or suppressing memory attention heads. We tune the attention head configuration for each mode and report the results (PH3 +context enhances context utilisation by the suppression of memory heads, and vice versa for PH3 +memory).

Context-aware contrastive decoding: ACD and COIECD Contrastive decoding approaches adjust the model’s output distribution based on two distributions: one for which only the query is given as input and one for which the context also is included. Among them, *contextual information-entropy constraint decoding* (COIECD; Yuan et al., 2024) is designed to detect the presence of knowledge conflicts and selectively resolve them, aiming to improve faithfulness to conflicting context without compromising performance when no conflict exists. In contrast, *adaptive contrastive decoding* (ACD; Kim et al., 2024) addresses the challenge of irrelevant context by using entropy-based weighting to adaptively ensemble parametric and contextual distributions. We test both on CUB to cover the nuance in decoding approaches.

5 Features Impacting Context Utilisation

To deepen our understanding of the results on CUB, we complement the benchmark with an analysis of

³Mainly by ChatGPT, but also by Microsoft Co-pilot.

features likely to impact context utilisation. Our goal is to better understand *why* certain CMTs and LMs work well or not. We study features on a model and input level, described below.

5.1 Model Features

By virtue of the large LM coverage in CUB, we are able to measure multiple salient model features. We analyse **model size**, whether the model is **instruction-tuned** and **strength of model memory**. To control for external confounders related to model family and implementation, we only measure correlations with model size and instruction-tuning across Qwen models. Strength of model memory is measured as the softmaxed logits for the top token predicted by the LM when only the query is provided (without context).

5.2 Input Features

We measure multiple input characteristics found to impact context utilisation for humans and/or LMs. By considering **context length** and **Flesch reading ease score**, we aim to measure whether the context is *difficult to understand* (Gao et al., 2024; Vladika and Matthes, 2023). Using **distractor rate**, we aim to measure whether the context contains *distracting information* (Shaier et al., 2024). With **query-context overlap** we also aim to measure *query-context similarity* (Wan et al., 2024). Lastly, we check the **answer position** (Liu et al., 2024) and if the evaluated LMs find the context **relevant**. More details on the detection of the features can be found in Appendix G.

5.3 Metric for Feature Impact

By virtue of the unified setup of CUB, we can study correlation coefficients to investigate the impact of different input and model features with a low risk of confounders. We use Spearman’s ρ to measure the impact of features on context utilisation, proxied by BCU.

6 Main Results on CUB

The CUB results can be found in Figures 3 and 4, CCU scores and more detailed results can be found in Appendix A. We structure the results analysis around a set of main findings.

6.1 Overall Trends

We first note that the BCU and CCU scores in Figures 3 and 5, respectively, support the same trends and focus the analysis on the BCU results.

Context utilisation improves with model size.

From Figure 3, we note how larger Regular LMs generally outperform smaller LMs when all context types are taken into consideration for NQ and DRUID. On NQ, the best performing model is Qwen 32B, and on DRUID the best performing model is Command A. Notably, applying a CMT to a small LM can lead to context utilisation on par with that of a regular larger LM, such as Fine-tuning Qwen 7B compared to Regular Qwen 32B on NQ. Meanwhile, on CounterFact, we observe how Regular model performance across all contexts generally *decreases* when model size is increased. This is counter-intuitive and we attribute the phenomena to the artificial nature of the dataset, which likely confuses the larger LMs. In addition, we know the NQ and DRUID datasets to be more difficult, demanding greater model capacity. This shows how it is insufficient to evaluate context utilisation only on simple datasets like CounterFact.

Most CMTs show an inflated performance on conflicting CounterFact contexts.

All LMs that do not already have a perfect BCU score on the conflicting CounterFact contexts improve to a perfect score of 1.0 under Prompting, PH3 +context, and Fine-tuning. However, similar improvements cannot be observed for the same CMTs on NQ or DRUID. These results show how CMTs proven to work well in simpler settings are not guaranteed to work equally well in more complex settings, proving the necessity of holistic tests. A deeper analysis of the inflated CMT performance on CounterFact is provided in Appendix A.

6.2 CMT Comparison

We further assess whether the CMTs consistently outperform Regular across different context types. Figure 4 shows the average Δ of each CMT, aggregated over all evaluated models. A value above zero indicates that the CMT yields a net improvement over Regular, whereas a negative value highlights cases where the CMT degrades performance.

There is a conflict between optimising for utilisation of relevant contexts and robustness to irrelevant contexts.

As each CMT exhibits trade-offs across context types or only marginal differences from Regular, the overall CMT Δ values (Total) converge to near zero across NQ and DRUID. Consequently, *we find no CMT that is superior*. For instance, PH3 +context shows consistent improvements over Regular in conflicting contexts,

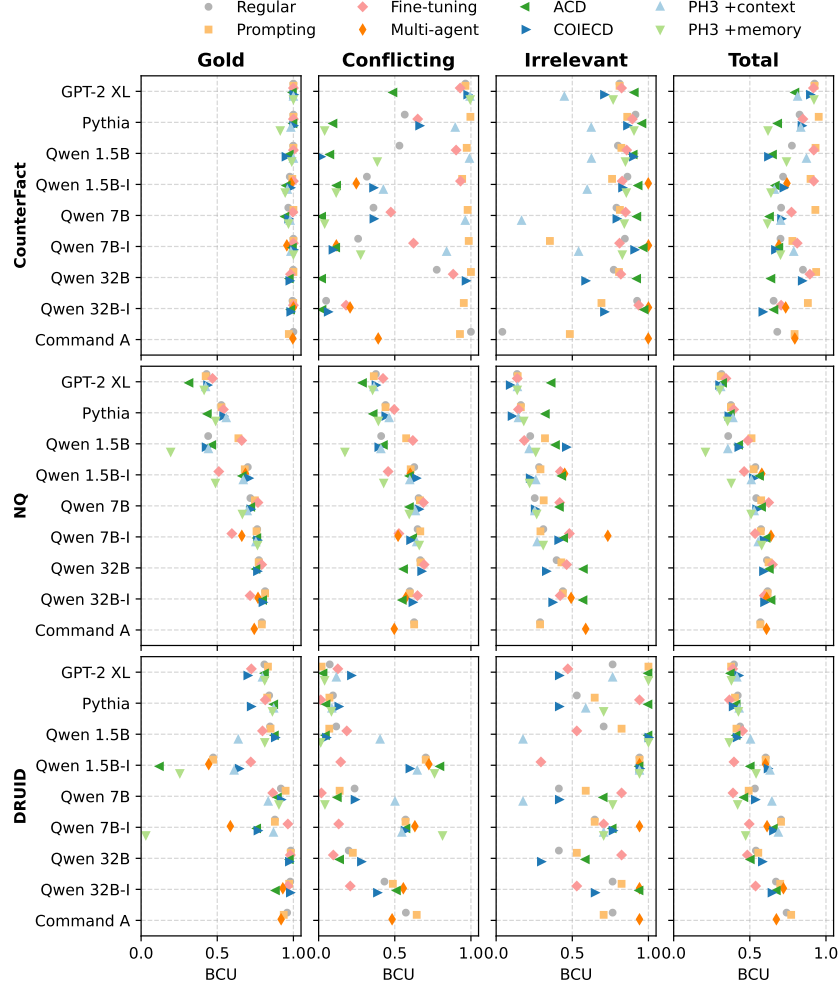


Figure 3: BCU scores for the evaluated context utilisation manipulation methods applied to the evaluated models and datasets. ‘Total’ denotes the averaged performance across all context types. A high BCU score is desirable regardless of context type.

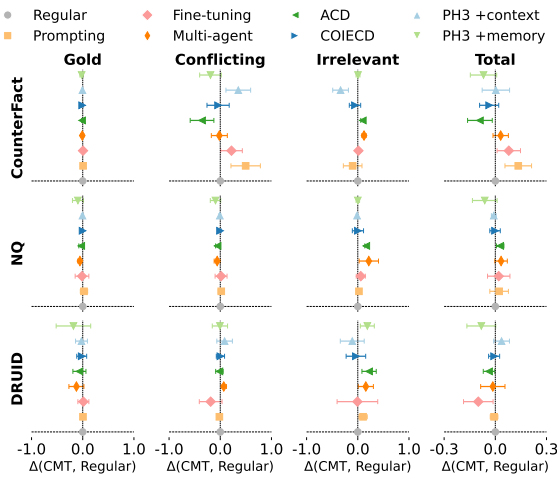


Figure 4: Model-averaged relative performance (Δ) of each CMT compared to Regular across datasets and context types. The horizontal bars represent the standard deviation.

but underperforms when applied to irrelevant contexts. Conversely, ACD, which handles irrelevant context effectively, performs worse in the conflicting context setting. Unsurprisingly, these findings highlight that the effectiveness of each CMT is closely tied to the alignment between the objective of the CMT and the type of context being provided. RAG practitioners knowing beforehand that their retrieval system is e.g. prone to return irrelevant information, may prioritise robustness over strong context utilisation and can select e.g. ACD as the CMT most suitable to their needs.

Prompting-based CMTs, such as Prompting and Multi-agent, show relatively stable performance across context types, without substantial drops in Δ . Compared to other CMTs, they offer this robustness with lower optimisation and implementation costs. Multi-agent shows clear gains in irrelevant contexts but limited efficacy in gold and conflicting

settings. This suggests that LMs are capable of identifying irrelevant contexts, but remain limited in effectively utilising relevant ones.

In realistic RAG scenarios, it will not be known beforehand what types of context will be provided to the LM. Therefore, it is important that CMTs work optimally across all context types. Our work shows that while we have CMTs that work well for relevant or irrelevant contexts *alone*, there currently are no CMTs that handle *both* relevant and irrelevant contexts well.

6.3 Impact of Model and Input Features

See Tables 6 and 7 for Spearman’s ρ between BCU and the features described in Section 5. Results are averaged across models.

Larger LMs perform better on NQ and DRUID. Corroborating our findings in Section 6.1, we observe a positive correlation with model size ($\rho \approx 0.3$) on DRUID gold contexts. Multi-agent also works significantly better with bigger LMs on DRUID gold contexts ($\rho = 0.42$). In addition, we observe a positive correlation with model size on NQ gold contexts ($\rho \in [0.20, 0.37]$). For CounterFact, we observe how model size does not correlate with performance.

Instruction-tuning is beneficial for conflicting and irrelevant DRUID contexts. We note how instruction tuning generally correlates with improved performance on conflicting and irrelevant DRUID contexts ($\rho \in [0.29, 0.77]$ depending on CMT). The conflicting DRUID contexts frequently require the LM to be able to abstain (i.e. respond with a ‘None’) when presented with insufficient contexts, which is something instruction-tuned models may be more adept at.

Conversely, instruction-tuning is clearly detrimental for conflicting CounterFact contexts ($\rho \leq -0.36$), potentially because the LMs have been more tuned to be critical of unreliable information, as opposed to following a pure causal language modelling objective.

A strong model memory corresponds to high performance on irrelevant contexts from NQ and CounterFact. We observe high correlations ($\rho \approx 0.36$) between memory strength and robustness to irrelevant contexts for Regular on CounterFact and NQ. These correlations increase when Fine-tuning, ACD or Prompting is applied. Furthermore, we observe for CounterFact how strong

Regular model memory correlates with low performance on conflicting contexts ($\rho = -0.44$). This is expected – previous work has already shown how LMs are resistant to synthesised contexts that contradict the internal model memory (Longpre et al., 2021; Xie et al., 2024a).

Answer position matters little for context utilisation. We measure low correlation values (below 0.3) across all settings for answer position in the context and Flesch reading ease score, and have thus omitted them in Table 7. Previous work has already found the Flesch reading ease score to show low correlations with LM context utilisation; our work further supports this finding (Hagström et al., 2024). Liu et al. (2024) found the answer position impactful for the utilisation of long contexts. CUB does not contain equally long contexts, which potentially explains why we do not see the same impact of answer position.

Context utilisation on gold NQ contexts is degraded on long contexts with high distractor rates. We measure weak negative correlations with context length ($\rho = -0.23$) and distractor rate ($\rho = -0.19$) with respect to Regular performance on gold NQ contexts. This is expected – long gold contexts or contexts with a high rate of distractors should be more difficult to process and utilise. We hypothesise the fairly low correlation levels are a consequence of each feature alone not being sufficiently predictive of model performance.

7 Conclusion

We introduce CUB, a benchmark that evaluates CMTs across diverse context types, datasets, and models. Under CUB, we evaluate a representative set of CMTs, covering varying context utilisation objectives and techniques. Results on CUB reveal a trade-off across most CMTs between robustness to irrelevant context and faithful utilisation of relevant context. Our analysis of features impacting context utilisation highlights the strong influence of model features, while input features have limited impact when analysed in separation. Overall, our findings highlight the need for holistic testing, as tests on synthesised datasets may show inflated performance, and the need for CMTs that can adapt to varied context conditions. Taken together, our work paves the way for the development of more effective RAG systems.

Limitations

CUB only incorporates contexts with lengths of up to that of a paragraph. It would also be relevant to evaluate CMTs in long-context settings. The long-context setting was not included in CUB, and left for future work, as it is fundamentally different from the normal context setting studied in CUB, posing new challenges for context utilisation and its evaluation, associated with a different set of CMTs (Shaham et al., 2023; Zhang et al., 2024a; Min et al., 2023; Zhang et al., 2024b).

While the dataset selection for CUB was performed to cover a wide span of task difficulty and RAG scenarios, the insights provided by CUB are limited to those derived from the underlying datasets. Moreover, all datasets are in English, leaving open the question of whether the findings generalise across languages (Chirkova et al., 2024).

Lastly, CUB does not explicitly consider datasets involving temporal dynamics, while it would be interesting to study. Time-sensitive information may lead to naturally occurring conflicts in context, adding nuance to the analysis of context utilisation (Loureiro et al., 2022; Xiong et al., 2024).

References

Stella Biderman, Hailey Schoelkopf, Quentin Gregory Anthony, Herbie Bradley, Kyle O’Brien, Eric Hallahan, Mohammad Aflah Khan, Shivanshu Purohit, USVSN Sai Prashanth, Edward Raff, et al. 2023. Pythia: A suite for analyzing large language models across training and scaling. In *International Conference on Machine Learning*, pages 2397–2430. PMLR.

Nadezhda Chirkova, David Rau, Hervé Déjean, Thibault Formal, Stéphane Clinchant, and Vassilina Nikoulina. 2024. Retrieval-augmented generation in multilingual settings. In *Proceedings of the 1st Workshop on Towards Knowledgeable Language Models (KnowLLM 2024)*, pages 177–188, Bangkok, Thailand. Association for Computational Linguistics.

Yung-Sung Chuang, Linlu Qiu, Cheng-Yu Hsieh, Ranjay Krishna, Yoon Kim, and James R. Glass. 2024. Lookback lens: Detecting and mitigating contextual hallucinations in large language models using only attention maps. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 1419–1436, Miami, Florida, USA. Association for Computational Linguistics.

Yilun Du, Shuang Li, Antonio Torralba, Joshua B. Tenenbaum, and Igor Mordatch. 2024. Improving factuality and reasoning in language models through

multiagent debate. In *Proceedings of the 41st International Conference on Machine Learning, ICML’24*. JMLR.org.

Feiteng Fang, Yuelin Bai, Shiwen Ni, Min Yang, Xiaojun Chen, and Ruifeng Xu. 2024. Enhancing noise robustness of retrieval-augmented language models with adaptive adversarial training. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 10028–10039, Bangkok, Thailand. Association for Computational Linguistics.

Shangbin Feng, Weijia Shi, Yike Wang, Wenxuan Ding, Vidhisha Balachandran, and Yulia Tsvetkov. 2024. Don’t hallucinate, abstain: Identifying LLM knowledge gaps via multi-LLM collaboration. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 14664–14690, Bangkok, Thailand. Association for Computational Linguistics.

Yunfan Gao, Yun Xiong, Xinyu Gao, Kangxiang Jia, Jinliu Pan, Yuxi Bi, Yi Dai, Jiawei Sun, Meng Wang, and Haofen Wang. 2024. Retrieval-augmented generation for large language models: A survey. *Preprint*, arXiv:2312.10997.

Mor Geva, Jasmijn Bastings, Katja Filippova, and Amir Globerson. 2023. Dissecting recall of factual associations in auto-regressive language models. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 12216–12235, Singapore. Association for Computational Linguistics.

Lovisa Hagström, Denitsa Saynova, Tobias Norlund, Moa Johansson, and Richard Johansson. 2023. The effect of scaling, retrieval augmentation and form on the factual consistency of language models. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 5457–5476, Singapore. Association for Computational Linguistics.

Lovisa Hagström, Sara Vera Marjanović, Haeun Yu, Arnav Arora, Christina Lioma, Maria Maistro, Pepa Atanasova, and Isabelle Augenstein. 2024. A reality check on context utilisation for retrieval-augmented generation. *Preprint*, arXiv:2412.17031.

Lovisa Hagström, Ercong Nie, Ruben Halifa, Helmut Schmid, Richard Johansson, and Alexander Junge. 2025. Language model re-rankers are steered by lexical similarities. *Preprint*, arXiv:2502.17036.

Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. 2021. Lora: Low-rank adaptation of large language models. *Preprint*, arXiv:2106.09685.

Qiao Jin, Bhuwan Dhingra, Zhengping Liu, William Cohen, and Xinghua Lu. 2019. PubMedQA: A dataset for biomedical research question answering. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the*

745	9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP), pages 2567–2577, Hong Kong, China. Association for Computational Linguistics.	802
746		803
747		804
748		805
749	Zhuoran Jin, Pengfei Cao, Hongbang Yuan, Yubo Chen, Jiexin Xu, Huaijun Li, Xiaojian Jiang, Kang Liu, and Jun Zhao. 2024. Cutting off the head ends the conflict: A mechanism for interpreting and mitigating knowledge conflicts in language models . In <i>Findings of the Association for Computational Linguistics: ACL 2024</i> , pages 1193–1215, Bangkok, Thailand. Association for Computational Linguistics.	806
750		807
751		808
752		809
753		810
754		811
755		812
756		813
757	Mandar Joshi, Eunsol Choi, Daniel Weld, and Luke Zettlemoyer. 2017. TriviaQA: A large scale distantly supervised challenge dataset for reading comprehension . In <i>Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 1601–1611, Vancouver, Canada. Association for Computational Linguistics.	814
758		
759		
760		
761		
762		
763		
764	Youna Kim, Hyuhng Joon Kim, Cheonbok Park, Choonghyun Park, Hyunsoo Cho, Junyeob Kim, Kang Min Yoo, Sang-goo Lee, and Taeuk Kim. 2024. Adaptive contrastive decoding in retrieval-augmented generation for handling noisy contexts . In <i>Findings of the Association for Computational Linguistics: EMNLP 2024</i> , pages 2421–2431, Miami, Florida, USA. Association for Computational Linguistics.	815
765		816
766		817
767		818
768		819
769		820
770		821
771		822
772	Tom Kwiatkowski, Jennimaria Palomaki, Olivia Redfield, Michael Collins, Ankur Parikh, Chris Alberti, Danielle Epstein, Illia Polosukhin, Jacob Devlin, Kenton Lee, Kristina Toutanova, Llion Jones, Matthew Kelcey, Ming-Wei Chang, Andrew M. Dai, Jakob Uszkoreit, Quoc Le, and Slav Petrov. 2019. Natural questions: A benchmark for question answering research . <i>Transactions of the Association for Computational Linguistics</i> , 7:452–466.	823
773		
774		
775		
776		
777		
778		
779		
780		
781	Daliang Li, Ankit Singh Rawat, Manzil Zaheer, Xin Wang, Michal Lukasik, Andreas Veit, Felix Yu, and Sanjiv Kumar. 2023. Large language models with controllable working memory . In <i>Findings of the Association for Computational Linguistics: ACL 2023</i> , pages 1774–1793, Toronto, Canada. Association for Computational Linguistics.	824
782		825
783		826
784		827
785		828
786		829
787		830
788	Nelson F. Liu, Kevin Lin, John Hewitt, Ashwin Paranjape, Michele Bevilacqua, Fabio Petroni, and Percy Liang. 2024. Lost in the middle: How language models use long contexts . <i>Transactions of the Association for Computational Linguistics</i> , 12:157–173.	831
789		832
790		833
791		834
792		
793	Pengfei Liu, Weizhe Yuan, Jinlan Fu, Zhengbao Jiang, Hiroaki Hayashi, and Graham Neubig. 2023. Pre-train, prompt, and predict: A systematic survey of prompting methods in natural language processing . <i>ACM Comput. Surv.</i> , 55(9).	835
794		836
795		837
796		838
797		839
798	Shayne Longpre, Kartik Perisetla, Anthony Chen, Nikhil Ramesh, Chris DuBois, and Sameer Singh. 2021. Entity-based knowledge conflicts in question answering . In <i>Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing</i> , pages 7052–7063, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.	840
799		841
800		842
801		
		843
		844
		845
		846
		847
		848
		849
		850
		851
		852
		853
		854
		855
		856
		857
		858
		859
		860

861	Fabio Petroni, Tim Rocktäschel, Sebastian Riedel,	Trusting your evidence: Hallucinate less with context-	918
862	Patrick Lewis, Anton Bakhtin, Yuxiang Wu, and	aware decoding. In <i>Proceedings of the 2024 Confer-</i>	919
863	Alexander Miller. 2019. Language models as knowl-	ence of the North American Chapter of the Associ-	920
864	edge bases? In <i>Proceedings of the 2019 Confer-</i>	ation for Computational Linguistics: <i>Human Lan-</i>	921
865	<i>ference on Empirical Methods in Natural Language Pro-</i>	<i>guage Technologies (Volume 2: Short Papers)</i> , pages	922
866	<i>cessing and the 9th International Joint Conference</i>	783–791, Mexico City, Mexico. Association for Com-	923
867	<i>on Natural Language Processing (EMNLP-IJCNLP)</i> ,	putational Linguistics.	924
868	pages 2463–2473, Hong Kong, China. Association		
869	for Computational Linguistics.		
870	Alec Radford, Jeffrey Wu, Rewon Child, David Luan,	Kurt Shuster, Spencer Poff, Moya Chen, Douwe Kiela,	925
871	Dario Amodei, Ilya Sutskever, et al. 2019. Language	and Jason Weston. 2021. Retrieval augmentation	926
872	models are unsupervised multitask learners. <i>OpenAI</i>	reduces hallucination in conversation . In <i>Findings</i>	927
873	<i>blog</i> , 1(8):9.	<i>of the Association for Computational Linguistics:</i>	928
		<i>EMNLP 2021</i> , pages 3784–3803, Punta Cana, Do-	929
		minican Republic. Association for Computational	930
		Linguistics.	931
874	Pranav Rajpurkar, Robin Jia, and Percy Liang. 2018.	Juraj Vladika and Florian Matthes. 2023. Scientific	932
875	Know what you don’t know: Unanswerable ques-	fact-checking: A survey of resources and approaches .	933
876	tions for SQuAD . In <i>Proceedings of the 56th Annual</i>	In <i>Findings of the Association for Computational</i>	934
877	<i>Meeting of the Association for Computational Lin-</i>	<i>Linguistics: ACL 2023</i> , pages 6215–6230, Toronto,	935
878	<i>guistics (Volume 2: Short Papers)</i> , pages 784–789,	Canada. Association for Computational Linguistics.	936
879	Melbourne, Australia. Association for Computational		
880	Linguistics.		
881	Denitsa Saynova, Lovisa Hagström, Moa Johansson,	Alexander Wan, Eric Wallace, and Dan Klein. 2024.	937
882	Richard Johansson, and Marco Kuhlmann. 2025.	What evidence do language models find convincing?	938
883	Fact recall, heuristics or pure guesswork? precise	In <i>Proceedings of the 62nd Annual Meeting of the</i>	939
884	interpretations of language models for fact comple-	<i>Association for Computational Linguistics (Volume 1:</i>	940
885	tion . <i>Preprint</i> , arXiv:2410.14405.	<i>Long Papers)</i> , pages 7468–7484, Bangkok, Thailand.	941
		Association for Computational Linguistics.	942
886	Michael Schlichtkrull, Zhijiang Guo, and Andreas Vla-	Han Wang, Archiki Prasad, Elias Stengel-Eskin, and	943
887	chos. 2023. Averitec: A dataset for real-world claim	Mohit Bansal. 2024. Adacad: Adaptively decoding	944
888	verification with evidence from the web . <i>Preprint</i> ,	to balance conflicts between contextual and paramet-	945
889	arXiv:2305.13117.	ric knowledge . <i>Preprint</i> , arXiv:2409.07394.	946
890	Uri Shaham, Maor Ivgi, Avia Efrat, Jonathan Berant,	Zhengxuan Wu, Aryaman Arora, Atticus Geiger, Zheng	947
891	and Omer Levy. 2023. ZeroSCROLLS: A zero-shot	Wang, Jing Huang, Dan Jurafsky, Christopher D.	948
892	benchmark for long text understanding . In <i>Find-</i>	Manning, and Christopher Potts. 2025. Axbench:	949
893	<i>ings of the Association for Computational Linguis-</i>	Steering llms? even simple baselines outperform	950
894	<i>tics: EMNLP 2023</i> , pages 7977–7989, Singapore.	sparse autoencoders . <i>Preprint</i> , arXiv:2501.17148.	951
895	Association for Computational Linguistics.		
896	Sagi Shaier, Lawrence Hunter, and Katharina von der	Jian Xie, Kai Zhang, Jiangjie Chen, Renze Lou, and	952
897	Wense. 2024. Desiderata for the context use of ques-	Yu Su. 2024a. Adaptive chameleon or stubborn sloth:	953
898	tion answering systems . In <i>Proceedings of the 18th</i>	Revealing the behavior of large language models in	954
899	<i>Conference of the European Chapter of the Associa-</i>	knowledge conflicts . In <i>The Twelfth International</i>	955
900	<i>tion for Computational Linguistics (Volume 1: Long</i>	<i>Conference on Learning Representations</i> .	956
901	<i>Papers)</i> , pages 777–792, St. Julian’s, Malta. Associa-		
902	tion for Computational Linguistics.	Jian Xie, Kai Zhang, Jiangjie Chen, Renze Lou, and	957
		Yu Su. 2024b. Adaptive chameleon or stubborn sloth:	958
		Revealing the behavior of large language models in	959
		knowledge conflicts . In <i>The Twelfth International</i>	960
903	Xiaoyu Shen, Rexhina Biloshmi, Dawei Zhu, Jiahuan	<i>Conference on Learning Representations</i> .	961
904	Pei, and Wei Zhang. 2024. Assessing “implicit” re-	Siheng Xiong, Ali Payani, Ramana Kompella, and Fara-	962
905	trieval robustness of large language models . In <i>Pro-</i>	marz Fekri. 2024. Large language models can learn	963
906	<i>ceedings of the 2024 Conference on Empirical Meth-</i>	temporal reasoning . In <i>Proceedings of the 62nd An-</i>	964
907	<i>ods in Natural Language Processing</i> , pages 8988–	<i>annual Meeting of the Association for Computational</i>	965
908	9003, Miami, Florida, USA. Association for Compu-	<i>Linguistics (Volume 1: Long Papers)</i> , pages 10452–	966
909	tational Linguistics.	10470, Bangkok, Thailand. Association for Compu-	967
		tational Linguistics.	968
910	Freda Shi, Xinyun Chen, Kanishka Misra, Nathan	Rongwu Xu, Zehan Qi, Zhijiang Guo, Cunxiang Wang,	969
911	Scales, David Dohan, Ed Chi, Nathanael Schärli, and	Hongru Wang, Yue Zhang, and Wei Xu. 2024.	970
912	Denny Zhou. 2023. Large language models can be	Knowledge conflicts for LLMs: A survey . In <i>Pro-</i>	971
913	easily distracted by irrelevant context. In <i>Proceed-</i>	<i>ceedings of the 2024 Conference on Empirical Meth-</i>	972
914	<i>ings of the 40th International Conference on Machine</i>	<i>ods in Natural Language Processing</i> , pages 8541–	973
915	<i>Learning, ICML’23</i> . JMLR.org.	8565, Miami, Florida, USA. Association for Compu-	974
916	Weijia Shi, Xiaochuang Han, Mike Lewis, Yulia	tational Linguistics.	975
917	Tsvetkov, Luke Zettlemoyer, and Wen-tau Yih. 2024.		

An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li, Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tianhao Li, Tingyu Xia, Xingzhang Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan, Yaqiong Liu, Zeyu Cui, Zhenru Zhang, and Zihan Qiu. 2024. Qwen2.5 technical report. *arXiv preprint arXiv:2412.15115*.

Ori Yoran, Tomer Wolfson, Ori Ram, and Jonathan Berant. 2024. [Making retrieval-augmented language models robust to irrelevant context](#). *Preprint*, arXiv:2310.01558.

Xiaowei Yuan, Zhao Yang, Yequan Wang, Shengping Liu, Jun Zhao, and Kang Liu. 2024. [Discerning and resolving knowledge conflicts through adaptive decoding with contextual information-entropy constraint](#). In *Findings of the Association for Computational Linguistics: ACL 2024*, pages 3903–3922, Bangkok, Thailand. Association for Computational Linguistics.

Huajian Zhang, Yumo Xu, and Laura Perez-Beltrachini. 2024a. [Fine-grained natural language inference based faithfulness evaluation for diverse summarisation tasks](#). In *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1701–1722, St. Julian’s, Malta. Association for Computational Linguistics.

Zhenyu Zhang, Runjin Chen, Shiwei Liu, Zhewei Yao, Olatunji Ruwase, Beidi Chen, Xiaoxia Wu, and Zhangyang Wang. 2024b. [Found in the middle: How language models use long contexts better via plug-and-play positional encoding](#). *Preprint*, arXiv:2403.04797.

Yu Zhao, Alessio Devoto, Giwon Hong, Xiaotang Du, Aryo Pradipta Gema, Hongru Wang, Xuanli He, Kam-Fai Wong, and Pasquale Minervini. 2025. [Steering knowledge selection behaviours in LLMs via SAE-based representation engineering](#). In *Proceedings of the 2025 Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pages 5117–5136, Albuquerque, New Mexico. Association for Computational Linguistics.

Zheng Zhao, Emilio Monti, Jens Lehmann, and Haytham Assem. 2024. [Enhancing contextual understanding in large language models through contrastive decoding](#). In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pages 4225–4237, Mexico City, Mexico. Association for Computational Linguistics.

A Additional results

A.1 CUB results

The exact CUB results can be found in Tables 3 and 4. CCU scores can be found in Figure 5. For the CCU scores, we note that they generally follow the same trends as the BCU scores in Figure 3; some CMTs perform better on gold, conflicting or irrelevant contexts, while none are superior when all context types are taken into consideration. The only disparate trend at odds with the BCU scores is that Fine-tuning Qwen models that have been instruction-tuned stand out by performing extra poorly with respect to CCU score. We hypothesise that this is a consequence of an increase in $P_M(t_C|Q)$ (i.e. prediction probability without context) from the fine-tuning, yielding less room for improvement in prediction confidence when context is introduced.

A.2 Analysis of inflated CMT performance on CounterFact

The inflated performance on CounterFact, observed in Figures 3 and 4, can potentially be explained by a suboptimal default prompt for CounterFact. Following previous work, the default prompt only contained the example to be completed, without any additional instructions or few-shot examples. For NQ and DRUID, the default prompt contained task instructions and few-shot examples. Furthermore, we observe how Prompting performs best on CounterFact on average, with a near perfect performance, indicating that a better default prompt may have neutralised any additional improvements from other CMTs. This raises the question of whether certain CMTs only address low context utilisation when caused by poor prompting, finding no leverage if the prompt already is adequate.

A.3 Quality check of irrelevant NQ contexts

For the CUB evaluation, we find 244 (14%) NQ samples with the context type ‘irrelevant’ for which at least 5 of the 9 evaluated LMs switch prediction to the gold answer *after* having seen the sample context. This indicates that some of the irrelevant contexts may actually be gold, as a result of quality issues with the annotation for NQ (in our sampling we assume that Wikipedia paragraphs not annotated as gold are not gold). However, we also note for some of these 244 samples that the context may simply be the heading of a Wikipedia page with the same title as the gold answer (e.g. “<H1> Scythe

Dataset		CounterFact				NQ				DRUID			
Model	Method	Gold	Conflict.	Irrel.	Tot.	Gold	Conflict.	Irrel.	Tot.	Gold	Conflict.	Irrel.	Tot.
GPT-2 XL	Regular	100.0	96.4	81.0	92.5	43.0	37.6	13.7	31.4	80.9	7.3	76.5	39.6
	Fine-tuning	100.0	92.9	82.4	91.8	46.9	42.3	13.9	34.3	72.4	12.6	47.1	38.7
	Prompting	100.0	96.4	81.0	92.5	42.4	36.2	14.2	30.9	83.3	1.9	100.0	37.7
	PH3 +context	100.0	99.4	44.8	81.4	42.3	36.4	14.0	30.9	79.6	11.6	76.5	41.5
	PH3 +memory	100.0	99.5	76.8	92.1	41.4	35.4	13.9	30.2	81.1	3.9	100.0	37.9
	COIECD	100.0	97.6	70.8	89.5	43.4	37.4	9.0	29.9	69.8	21.3	41.2	42.4
	ACD	99.6	49.1	91.0	79.9	31.8	29.1	36.4	32.4	81.3	3.2	100.0	37.6
PYTHIA 6.9B	Regular	100.0	56.5	91.5	82.7	52.7	43.9	16.2	37.6	84.1	9.4	52.9	42.1
	Fine-tuning	100.0	65.1	89.4	84.8	54.0	49.6	14.6	39.4	81.5	1.4	94.1	36.6
	Prompting	100.0	99.6	86.1	95.2	52.7	43.9	16.2	37.6	82.8	7.1	64.7	40.3
	PH3 +context	98.3	89.7	62.4	83.5	55.9	46.3	14.6	38.9	87.1	8.7	58.8	43.0
	PH3 +memory	91.4	4.0	90.5	61.9	48.9	39.2	18.1	35.4	86.2	8.4	70.6	42.5
	COIECD	99.9	66.0	86.0	84.0	53.9	43.8	10.2	35.9	72.0	13.0	41.2	38.8
	ACD	100.0	9.7	96.0	68.6	43.8	36.1	32.6	37.5	87.4	5.2	100.0	41.3
QWEN2.5 1.5B	Regular	99.9	53.1	80.0	77.6	44.0	41.1	22.4	35.8	84.7	11.6	70.6	43.6
	Fine-tuning	100.0	90.3	85.7	92.0	66.1	61.9	18.5	48.8	79.7	18.5	52.9	45.3
	Prompting	100.0	97.2	82.2	93.2	63.9	57.5	32.1	51.1	85.0	7.0	82.4	41.2
	PH3 +context	100.0	99.0	62.5	87.2	44.2	40.9	21.7	35.6	63.8	40.4	17.6	50.5
	PH3 +memory	98.9	38.5	84.9	74.1	19.4	17.3	26.0	20.9	81.2	1.4	100.0	36.5
	COIECD	94.8	1.2	89.8	61.9	42.4	39.2	45.8	42.5	87.8	4.8	100.0	41.3
	ACD	97.6	7.7	90.3	65.2	46.7	42.8	39.3	42.9	87.8	4.8	100.0	41.3
QWEN2.5 1.5B <i>Instruct</i>	Regular	97.6	31.7	86.2	71.8	70.1	62.8	28.2	53.7	47.3	70.3	94.1	60.4
	Fine-tuning	100.0	93.2	82.7	92.0	51.0	45.6	42.2	46.3	72.0	14.5	29.4	39.6
	Prompting	99.3	94.2	76.1	89.9	68.1	60.5	29.1	52.5	47.3	70.3	94.1	60.4
	Multi-agent	98.6	24.7	99.9	74.4	68.5	60.2	45.0	57.9	44.4	72.4	94.1	60.3
	PH3 +context	96.0	42.5	59.8	66.1	67.1	59.9	26.0	51.0	61.1	64.7	94.1	63.2
	PH3 +memory	94.6	11.5	85.5	63.9	48.8	42.7	22.0	37.8	25.4	76.1	94.1	54.1
	COIECD	97.8	35.8	82.7	72.1	70.5	63.9	22.1	52.1	64.1	59.6	94.1	61.7
	ACD	95.6	12.1	93.5	67.1	66.7	60.0	43.4	56.7	12.3	79.9	94.1	50.6
QWEN2.5 7B	Regular	96.6	36.0	79.0	70.5	71.7	65.6	25.3	54.2	91.8	23.6	41.2	53.3
	Fine-tuning	99.6	47.4	85.0	77.4	76.7	68.8	41.7	62.4	86.4	1.8	82.4	39.0
	Prompting	100.0	97.8	81.3	93.0	74.7	66.5	31.2	57.5	94.9	13.8	58.8	49.3
	PH3 +context	97.8	96.3	16.7	70.3	69.7	63.6	25.3	52.8	83.4	50.1	17.6	64.5
	PH3 +memory	96.8	4.0	84.2	61.6	66.5	59.5	26.6	50.8	90.5	4.1	76.5	42.0
	COIECD	96.6	36.0	79.0	70.5	71.7	65.6	25.3	54.2	91.8	23.6	41.2	53.3
	ACD	94.7	2.3	92.7	63.2	72.3	59.9	41.9	58.0	89.8	12.6	70.6	46.4
QWEN2.5 7B <i>Instruct</i>	Regular	100.0	25.9	84.5	70.1	76.2	65.0	31.0	57.4	87.8	57.1	64.7	70.5
	Fine-tuning	100.0	62.3	81.0	81.1	59.6	52.7	48.1	53.5	96.4	13.2	70.6	49.6
	Prompting	100.0	98.6	35.3	78.0	75.8	66.7	29.1	57.2	87.8	57.1	64.7	70.5
	Multi-agent	95.7	11.6	100.0	69.1	66.1	52.2	73.3	63.9	58.6	63.2	94.1	61.3
	PH3 +context	98.3	84.0	54.1	78.8	75.3	64.4	26.9	55.5	86.9	54.7	70.6	68.8
	PH3 +memory	100.0	27.6	82.8	70.1	76.4	66.1	30.9	57.8	3.1	81.4	70.6	47.3
	COIECD	99.9	9.1	90.6	66.5	76.2	60.1	40.8	59.0	76.4	56.5	76.5	65.2
	ACD	99.6	11.5	96.9	69.3	76.3	62.1	44.6	61.0	76.2	57.6	76.5	65.8
QWEN2.5 32B	Regular	99.9	77.6	77.2	84.9	77.3	66.7	39.7	61.2	98.2	19.8	41.2	54.0
	Fine-tuning	98.1	88.4	81.9	89.4	79.2	69.2	46.3	64.9	98.0	9.7	82.4	48.4
	Prompting	100.0	100.0	80.7	93.6	77.2	66.9	42.8	62.3	98.2	22.5	52.9	55.6
	COIECD	97.4	96.5	58.5	84.1	76.1	67.4	32.7	58.7	97.1	27.8	29.4	57.9
	ACD	97.6	2.3	92.6	64.1	75.7	56.1	57.6	63.1	97.6	14.1	58.8	50.6
QWEN2.5 32B <i>Instruct</i>	Regular	99.4	4.9	92.6	65.6	81.4	59.9	43.8	61.7	97.9	43.2	76.5	67.2
	Fine-tuning	100.0	18.0	93.6	70.5	71.6	64.9	42.0	59.5	96.4	20.8	52.9	53.8
	Prompting	99.9	95.3	69.1	88.1	81.4	59.9	43.8	61.7	97.2	48.7	82.4	70.0
	Multi-agent	100.0	20.6	100.0	73.5	76.8	57.2	49.2	61.1	93.1	55.6	94.1	72.1
	COIECD	98.0	6.0	70.8	58.3	79.7	61.6	36.8	59.4	97.7	38.3	64.7	64.3
	ACD	98.4	2.5	97.5	66.1	80.1	55.2	57.4	64.2	88.5	51.4	94.1	67.7
COMMAND A	Regular	100.0	100.0	4.1	68.0	79.2	62.7	28.9	56.9	95.9	57.3	76.5	74.2
	Prompting	97.0	92.8	48.4	79.4	79.2	62.7	28.9	56.9	93.6	64.4	70.6	77.2
	Multi-agent	99.6	39.1	99.9	79.6	74.3	49.7	58.8	61.0	91.9	48.2	94.1	67.4

Table 3: BCU scores on CUB. A high BCU score is desirable regardless of context type. Gold denotes relevant contexts that also contain the gold answer. Conflict. denotes ‘Conflicting’ – relevant contexts that contain a conflicting answer, dissimilar from the correct answer or model memory. Irrel. denotes irrelevant contexts. Tot. denotes the average performance across all context types. Values marked in **bold** indicate the top CMT score across LMs for each dataset and context type.

Dataset		CounterFact				NQ				DRUID			
Model	Method	Gold	Conflict.	Irrel.	Tot.	Gold	Conflict.	Irrel.	Tot.	Gold	Conflict.	Irrel.	Tot.
GPT-2 XL	Regular	100.0	2.9	69.7	57.5	43.0	8.1	20.8	24.0	80.9	69.0	64.7	74.2
	Fine-tuning	100.0	3.2	70.6	57.9	46.9	7.7	23.8	26.2	72.4	65.5	41.2	68.4
	Prompting	100.0	2.9	69.7	57.5	42.4	7.5	20.3	23.5	83.3	73.8	76.5	78.0
	PH3 +context	100.0	0.4	29.8	43.4	42.3	7.8	20.4	23.6	79.6	65.7	52.9	71.7
	PH3 +memory	100.0	0.4	65.1	55.1	41.4	7.4	20.1	23.0	81.1	72.6	76.5	76.3
	COIECD	100.0	2.3	67.7	56.7	43.4	7.1	19.4	23.3	69.8	51.0	47.1	59.1
	ACD	99.6	29.4	72.3	67.1	31.8	7.7	18.1	19.2	81.3	73.0	76.5	76.6
PYTHIA 6.9B	Regular	100.0	37.2	91.4	76.2	52.7	9.8	29.6	30.8	84.1	49.9	47.1	64.7
	Fine-tuning	100.0	26.5	91.8	72.8	54.0	5.6	26.6	28.8	81.5	74.4	70.6	77.5
	Prompting	100.0	0.5	86.1	62.2	52.7	9.8	29.6	30.8	82.8	57.1	47.1	68.3
	PH3 +context	98.3	2.5	62.1	54.3	55.9	8.4	30.0	31.5	87.1	55.2	52.9	69.0
	PH3 +memory	91.4	86.0	90.4	89.2	48.9	11.5	29.7	30.1	86.2	55.1	64.7	68.7
	COIECD	99.9	27.3	86.0	71.0	53.9	9.8	27.4	30.4	72.0	32.9	35.3	50.0
	ACD	100.0	77.6	95.9	91.2	43.8	12.1	29.7	28.6	87.4	69.2	82.4	77.2
QWEN2.5 1.5B	Regular	99.9	41.9	74.2	72.0	44.0	7.7	22.0	24.6	84.7	63.5	52.9	72.7
	Fine-tuning	100.0	5.5	77.0	60.8	66.1	18.8	42.4	42.5	79.7	60.3	58.8	68.7
	Prompting	100.0	1.6	79.7	60.4	63.9	17.0	38.5	39.8	85.0	69.8	58.8	76.4
	PH3 +context	100.0	0.7	50.1	50.3	44.2	12.6	25.5	27.5	63.8	26.9	11.8	42.9
	PH3 +memory	98.9	52.8	78.0	76.6	19.4	8.1	10.4	12.7	81.2	74.5	70.6	77.4
	COIECD	94.8	71.9	79.0	81.9	42.4	16.3	27.6	28.8	87.8	72.7	70.6	79.3
	ACD	97.6	70.8	79.4	82.6	46.7	15.5	28.0	30.1	87.8	72.7	70.6	79.3
QWEN2.5 1.5B <i>Instruct</i>	Regular	97.6	54.5	79.6	77.2	70.1	16.1	37.1	41.2	47.3	11.1	0.0	26.8
	Fine-tuning	100.0	7.0	78.0	61.7	51.0	7.6	27.8	28.8	72.0	28.5	47.1	47.5
	Prompting	99.3	5.4	74.1	59.6	68.1	15.7	38.8	41.0	47.3	11.1	0.0	26.8
	Multi-agent	98.6	68.7	83.0	83.4	68.5	16.9	36.1	40.6	44.4	10.0	0.0	24.9
	PH3 +context	96.0	35.9	58.2	63.4	67.1	15.4	34.7	39.1	61.1	18.9	0.0	37.2
	PH3 +memory	94.6	68.9	78.3	80.6	48.8	13.1	25.8	29.3	25.4	7.2	0.0	15.1
	COIECD	97.8	50.4	77.1	75.1	70.5	15.5	35.9	40.7	64.1	19.2	0.0	38.7
QWEN2.5 7B	ACD	95.6	77.7	82.1	85.1	66.7	19.0	39.0	41.6	12.3	3.6	0.0	7.4
	Regular	96.6	52.2	72.6	73.8	71.7	16.7	39.0	42.6	91.8	57.6	23.5	72.3
	Fine-tuning	99.6	45.1	77.1	73.9	76.7	18.5	50.5	48.6	86.4	74.8	70.6	79.8
	Prompting	100.0	2.4	86.2	62.9	74.7	17.9	44.6	45.8	94.9	64.2	35.3	77.4
	PH3 +context	97.8	0.2	6.0	34.7	69.7	17.0	38.7	41.9	83.4	30.5	5.9	53.4
	PH3 +memory	96.8	88.6	79.4	88.2	66.5	17.6	37.7	40.6	90.5	73.4	70.6	80.8
	COIECD	96.6	52.2	72.6	73.8	71.7	16.7	39.0	42.6	91.8	57.6	23.5	72.3
QWEN2.5 7B <i>Instruct</i>	ACD	94.7	85.5	80.4	86.9	72.3	23.9	47.2	47.8	89.8	68.1	47.1	77.5
	Regular	100.0	42.0	85.4	75.8	76.2	19.8	47.1	47.8	87.8	28.3	0.0	54.1
	Fine-tuning	100.0	34.8	88.0	74.3	59.6	8.1	35.3	34.4	96.4	65.0	64.7	78.6
	Prompting	100.0	1.9	37.5	46.5	75.8	20.3	46.0	47.4	87.8	28.3	0.0	54.1
	Multi-agent	95.7	85.5	94.0	91.7	66.1	21.4	40.9	42.9	58.6	18.5	29.4	36.0
	PH3 +context	98.3	12.5	55.6	55.5	75.3	18.5	44.1	46.0	86.9	31.5	0.0	55.5
	PH3 +memory	100.0	50.9	83.8	78.2	76.4	20.1	47.7	48.1	3.1	2.5	0.0	2.7
QWEN2.5 32B	COIECD	99.9	75.0	90.8	88.6	76.2	25.8	48.2	50.1	76.4	29.2	5.9	49.7
	ACD	99.6	85.1	94.0	92.9	76.3	25.0	49.3	50.3	76.2	29.1	5.9	49.5
	Regular	99.9	21.4	75.0	65.4	77.3	20.8	47.7	48.7	98.2	58.5	29.4	75.7
	Fine-tuning	98.1	9.8	77.2	61.7	79.2	20.3	55.9	51.9	98.0	66.6	64.7	80.3
	Prompting	100.0	0.2	80.7	60.3	77.2	19.9	50.2	49.2	98.2	57.5	41.2	75.2
	COIECD	97.4	3.2	59.7	53.4	76.1	18.8	43.9	46.3	97.1	47.4	17.6	68.9
	ACD	97.6	85.7	81.3	88.2	75.7	31.4	53.3	53.5	97.6	66.1	47.1	79.8
QWEN2.5 32B <i>Instruct</i>	Regular	99.4	81.0	93.5	91.3	81.4	28.6	52.2	54.2	97.9	41.8	29.4	66.2
	Fine-tuning	100.0	78.5	92.2	90.2	71.6	13.3	44.3	43.2	96.4	61.8	47.1	76.8
	Prompting	99.9	3.2	70.6	57.9	81.4	28.6	52.2	54.2	97.2	36.2	11.8	62.6
	Multi-agent	100.0	78.5	94.7	91.1	76.8	22.7	40.7	46.8	93.1	31.7	17.6	58.4
	COIECD	98.0	9.7	72.4	60.0	79.7	23.4	49.4	50.9	97.7	43.3	29.4	66.9
	ACD	98.4	94.7	95.4	96.2	80.1	35.3	55.4	57.0	88.5	36.0	17.6	58.8
COMMAND A	Regular	100.0	0.0	4.4	34.8	79.2	12.3	33.8	41.9	95.9	30.3	5.9	58.8
	Prompting	97.0	0.7	47.8	48.5	79.2	12.3	33.8	41.9	93.6	23.3	0.0	53.8
	Multi-agent	99.6	32.2	90.2	74.0	74.3	13.5	40.4	42.8	91.9	33.2	23.5	58.7

Table 4: Accuracy with respect to gold label on CUB. Gold denotes relevant contexts that also contain the gold answer. Conflict. denotes ‘Conflicting’ – relevant contexts that contain a conflicting answer, dissimilar from the correct answer or model memory. Irrel. denotes irrelevant contexts. Tot. denotes the average performance across all context types. Values marked in **bold** indicate the top CMT score across LMs on each dataset and context type.

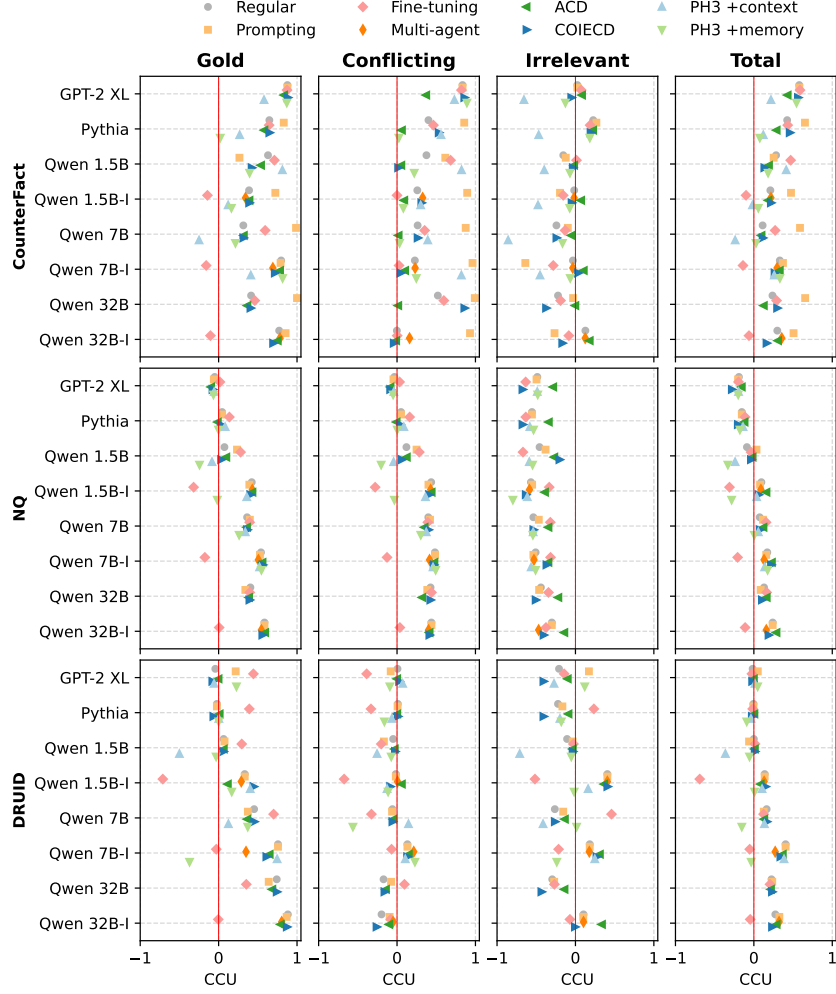


Figure 5: CCU scores for the evaluated context utilisation manipulation methods applied to the evaluated models and datasets. ‘Total’ denotes the averaged performance across all context types. A high CCU score is desirable regardless of context type. The red vertical lines indicate scores of 0.

when the gold answer is “scythe” for the query “what is the name of the weapon the grim reaper carries?”), without providing sufficient evidence with respect to the question, raising the question of whether they should be considered relevant by the model.

A.4 Performance of Relevance Judgement

For the Multi-agent technique, we investigate whether instruction-tuned LMs are capable of identifying irrelevant context when explicitly prompted to do so. According to Table 5, the Multi-agent approach demonstrates strong performance in detecting irrelevant contexts and in recognising gold contexts as relevant. Although it does not reliably maintain a closed-book response when directly generating responses (i.e. Regular), it can accurately detect irrelevance when equipped with an explicit relevance assessment setup.

The prediction accuracy of relevance assessment on conflicting contexts is consistently lower than that on other contexts. This discrepancy is particularly evident in the conflicting contexts of the CounterFact dataset. For instance, we found that LMs often generate feedback such as: “X is Y, not Z. Therefore, the context is irrelevant”. This suggests that LM interprets factual inconsistency with its internal knowledge as a signal of irrelevance, even when instructed to ignore its own memory.

One possible explanation for this behaviour lies in the nature of the CounterFact dataset itself. Contexts in CounterFact are typically composed of single-sentence facts, which may lack sufficient surrounding information to render the context trustworthy from the model’s perspective. Such behaviour is less pronounced in NQ and DRUID datasets, where the provided contexts are relatively longer and richer, offering more semantic cues that

	Gold	Conflict.	Irrel.	All
QWEN2.5 1.5B-I				
CounterFact	98.56	24.25	99.88	74.23
NQ	92.44	91.89	26.26	70.13
DRUID	93.27	96.52	17.65	94.79
QWEN2.5 7B-I				
CounterFact	99.16	10.68	99.88	69.91
NQ	80.70	76.14	59.35	72.05
DRUID	82.53	65.56	94.12	73.06
QWEN2.5 32B-I				
CounterFact	99.64	19.57	99.40	72.87
NQ	94.74	92.50	25.77	70.94
DRUID	98.66	76.25	88.24	86.05
COMMAND A				
CounterFact	100.00	99.88	99.88	99.92
NQ	94.31	91.82	37.69	74.56
DRUID	93.11	68.55	88.24	79.31

Table 5: Multi-agent: Relevance assessment accuracy

may help the LM interpret the information as contextually anchored (Xie et al., 2024b).

The performance of relevance assessment is particularly low on the NQ dataset compared to other datasets. Since irrelevant contexts of NQ dataset are sampled from the same document and may be topically or semantically similar to the question, distinguishing relevance may become more challenging.

A.5 Features Impacting Context Utilisation

See Table 6 for the correlation values between model features and context utilisation. See Table 7 for the correlation values between input features and context utilisation.

B Data Collection

B.1 CounterFact

Samples from the CounterFact dataset can be found in Table 8. The relations covered by the dataset are *capital of* (80%), *country of origin* (9%), *location of formation* (9%), *field of work* (1%) and *country of citizenship* (1%).

Rate of memorisation of CUB models We evaluate all Regular LMs on the samples from CUB CounterFact without context. The results can be found in Table 9. We observe rates above 70% for all models. As expected, the highest memorisation rate is found for Pythia. The lowest is found for

Dataset	Context	CMT	Corr.
Model size			
DRUID	Gold	Multi-agent	0.42
DRUID	Gold	ACD	0.41
NQ	Gold	PH3 +memory	0.37
DRUID	Gold	Regular	0.36
DRUID	Gold	Prompting	0.36
NQ	Conflicting	PH3 +memory	0.33
NQ	Gold	Regular	0.20
NQ	Irrelevant	Regular	0.14
NQ	Conflicting	Regular	0.09
CounterFact	Gold	Regular	0.04
CounterFact	Irrelevant	Regular	0.02
CounterFact	Conflicting	Regular	-0.01
DRUID	Conflicting	Regular	-0.08
DRUID	Irrelevant	Regular	-0.20
DRUID	Irrelevant	PH3 +memory	-0.33
CounterFact	Conflicting	Fine-tuning	-0.33
DRUID	Irrelevant	COIECD	-0.44
Instruct tuned			
DRUID	Conflicting	PH3 +memory	0.77
DRUID	Irrelevant	PH3 +context	0.65
DRUID	Conflicting	ACD	0.54
DRUID	Conflicting	Prompting	0.46
DRUID	Conflicting	Regular	0.40
DRUID	Conflicting	COIECD	0.34
DRUID	Irrelevant	Regular	0.29
NQ	Gold	Regular	0.13
CounterFact	Irrelevant	Regular	0.12
NQ	Irrelevant	Regular	0.06
NQ	Conflicting	Regular	0.05
CounterFact	Gold	Regular	0.01
DRUID	Gold	Regular	-0.19
CounterFact	Conflicting	Regular	-0.36
DRUID	Gold	ACD	-0.38
CounterFact	Conflicting	PH3 +context	-0.43
DRUID	Gold	PH3 +memory	-0.72
Strength of memory			
DRUID	Conflicting	PH3 +memory	0.54
NQ	Irrelevant	Fine-tuning	0.47
NQ	Irrelevant	ACD	0.39
CounterFact	Irrelevant	Fine-tuning	0.39
NQ	Irrelevant	Prompting	0.39
NQ	Irrelevant	COIECD	0.38
DRUID	Conflicting	ACD	0.37
NQ	Irrelevant	Regular	0.37
CounterFact	Irrelevant	Regular	0.35
DRUID	Conflicting	Prompting	0.34
CounterFact	Irrelevant	ACD	0.32
CounterFact	Irrelevant	PH3 +memory	0.31
CounterFact	Irrelevant	COIECD	0.30
DRUID	Conflicting	Regular	0.26
NQ	Gold	Regular	0.18
DRUID	Irrelevant	Regular	0.15
NQ	Conflicting	Regular	0.09
CounterFact	Gold	Regular	0.04
DRUID	Gold	Regular	0.02
CounterFact	Conflicting	ACD	-0.31
CounterFact	Conflicting	COIECD	-0.42
DRUID	Gold	PH3 +memory	-0.43
CounterFact	Conflicting	Regular	-0.44

Table 6: Spearman’s ρ between BCU and different model aspects. Correlation values for Regular or with an absolute value above 0.3 are shown. Correlation values with an absolute value below 0.3 are marked in gray. Significant correlation values (p-value < 0.05) are marked in **bold**.

GPT-2 XL, which can be expected as the model is quite small and old.

Prompt templates Following the same approach as previous work, no specific prompt template was used for the LMs evaluated on CounterFact. The

Dataset	Context	CMT	Corr.
Context length			
CounterFact	Irrelevant	Regular	0.06
CounterFact	Conflicting	Regular	0.04
CounterFact	Gold	Regular	0.02
DRUID	Conflicting	Regular	-0.02
DRUID	Irrelevant	Regular	-0.02
NQ	Irrelevant	Regular	-0.06
DRUID	Gold	Regular	-0.08
NQ	Conflicting	Regular	-0.22
NQ	Gold	Regular	-0.23
DRUID	Irrelevant	Multi-agent	-0.32
Query-context overlap			
DRUID	Gold	Regular	0.02
DRUID	Irrelevant	Regular	-0.03
NQ	Gold	Regular	-0.06
NQ	Conflicting	Regular	-0.08
NQ	Irrelevant	Regular	-0.08
DRUID	Conflicting	Regular	-0.13
DRUID	Irrelevant	Multi-agent	-0.30
Distractor rate			
CounterFact	Gold	Regular	0.00
NQ	Conflicting	Regular	-0.19
NQ	Gold	Regular	-0.19
CounterFact	Conflicting	Regular	-0.22
CounterFact	Conflicting	ACD	-0.34
CounterFact	Conflicting	Multi-agent	-0.49
Relevance judgement			
CounterFact	Conflicting	Multi-agent	0.53
CounterFact	Conflicting	Regular	0.17
NQ	Irrelevant	Regular	0.11
DRUID	Irrelevant	Regular	0.05
NQ	Gold	Regular	0.04
DRUID	Gold	Regular	0.03
NQ	Conflicting	Regular	0.02
CounterFact	Irrelevant	Regular	0.01
CounterFact	Gold	Regular	-0.01
DRUID	Conflicting	Regular	-0.15
NQ	Irrelevant	Multi-agent	-0.36
DRUID	Irrelevant	Multi-agent	-0.49

Table 7: Spearman’s ρ between BCU and different input aspects. Correlation values for Regular or with an absolute value above 0.3 are shown. Correlation values with an absolute value below 0.3 are marked in gray. Significant correlation values (p-value < 0.05) are marked in **bold**.

Prompt	Type
Fact: Athens, the capital city of Greece. Athens, the capital city of	Gold
Fact: Thomas Ong is a citizen of Pakistan. Thomas Ong is a citizen of	Conflicting
Fact: Melbourne, that is the capital of Jordan. Prince Oscar Bernadotte is a citizen of	Irrelevant

Table 8: CounterFact prompts with contexts and corresponding context types. For prompts without context, the first line (starting with “Fact:”) is simply removed.

LMs were evaluated in a simple sentence completion format as shown in Table 8.

However, since the sentence completion format is less compatible with the instruction-tuned models, we added a small prompt template for the evaluation of the instruction-tuned Qwen models on

Model	Accuracy
GPT-2 XL	71.8
Pythia	99.6
Qwen 1.5B	77.0
Qwen 1.5B-I	83.1
Qwen 7B	79.7
Qwen 7B-I	93.6
Qwen 32B	78.0
Qwen 32B-I	94.5
Command A	90.6

Table 9: Accuracy, proxying memorisation rate, on samples from CounterFact without context.

CounterFact, as follows.

Prompt without context for instruction-tuned LMs.

```
<|im_start|>system
You are Qwen, created by Alibaba Cloud. You are
a helpful assistant.<|im_end|>
<|im_start|>user
Complete the following sentence. Only answer
with the next word.
<prompt><|im_end|>
<|im_start|>assistant
```

Prompt with context for instruction-tuned LMs.

```
<|im_start|>system
You are Qwen, created by Alibaba Cloud. You are
a helpful assistant.<|im_end|>
<|im_start|>user
Complete the following sentence. Only answer
with the next word.
Fact: <context>
<prompt><|im_end|>
<|im_start|>assistant
```

B.2 NQ

We retain all samples from the development set of NQ⁴ for which a short answer of fewer than five tokens is identified in the raw HTML of the corresponding Wikipedia pages. Samples from the NQ dataset can be found in Table 10.

Sampling of conflicting contexts For a given question, context and short answer, we perform the following steps to identify substitute answers for conflicting contexts:

1. Check if the short answer is a date⁵. If so, sample a new random date in the interval [1900, 2030) and format it in the same way as the gold date.
2. If the short answer is not a date, prompt an LLM⁶ with the question and short answer to

⁴https://console.cloud.google.com/storage/browser/natural_questions/v1.0/dev

⁵Using the dateutil.parser in Python.

⁶The Cohere model command-r-plus-08-2024 from <https://docs.cohere.com/v2/docs/command-r-plus>.

Question	Short answer	Context	Type
when did the movie napoleon dynamite come out?	June 11, 2004	<Table> <Tr> <Th colspan="2"> Napoleon Dynamite </Th> </Tr> <Tr> <Td colspan="2"> Theatrical release poster </Td> </Tr> <Tr> <Th> Directed by </Th> <Td> Jared Hess </Td> </Tr> <Tr> <Th> Produced by </Th> <Td> Jeremy Coon Chris Wyatt Sean Covey Jory Weitz </Td> </Tr> <Tr> <Th> Screenplay by </Th> <Td> Jared Hess Jerusha Hess </Td> </Tr> <Tr> <Th> Based on </Th> <Td> Peluca by Jared Hess </Td> </Tr> <Tr> <Th> Starring </Th> <Td> Jon Heder Jon Gries Efrén Ramírez Tina Majorino Aaron Ruell Diedrich Bader Haylie Duff </Td> </Tr> <Tr> <Th> Music by </Th> <Td> John Swihart </Td> </Tr> <Tr> <Th> Cinematography </Th> <Td> Munn Powell </Td> </Tr> <Tr> <Th> Edited by </Th> <Td> Jeremy Coon </Td> </Tr> <Tr> <Th> Production company </Th> <Td> MTV Films Napoleon Pictures Access Films </Td> </Tr> <Tr> <Th> Distributed by </Th> <Td> Fox Searchlight Pictures (North America) Paramount Pictures (International) </Td> </Tr> <Tr> <Th> Release date </Th> <Td> January 17, 2004 (2004 - 01 - 17) (Sundance) June 11, 2004 (2004 - 06 - 11) (United States) </Td> </Tr> <Tr> <Th> Running time </Th> <Td> 95 minutes </Td> </Tr> <Tr> <Th> Country </Th> <Td> United States </Td> </Tr> <Tr> <Th> Language </Th> <Td> English </Td> </Tr> <Tr> <Th> Budget </Th> <Td> \$400,000 </Td> </Tr> <Tr> <Th> Box office </Th> <Td> \$46.1 million </Td> </Tr> </Table> <P> The Lupus Foundation of America (LFA), founded in 1967, is a national voluntary health organization based in Washington, D.C. with a network of chapters, offices and support groups located in communities throughout the United States . The Foundation is devoted to solving the mystery of lupus, one of the world's cruellest, most unpredictable and devastating diseases, while giving caring support to those who suffer from its brutal impact . Its mission is to improve the quality of life for all people affected by lupus through programs of research, education, support and advocacy . </P>	Gold
when was the lupus foundation of america founded?	1977		Conflicting
who has scored the most tries in rugby union?	Daisuke Ohata	<P> This is a list of the leading try scorers in rugby union test matches . It includes players with a minimum of 30 test tries . </P>	Irrelevant

Table 10: NQ samples and corresponding context types.

provide a substitute answer of the same format. If the proposed answer is already found in the sample context, prompt the model, for a maximum of 20 times, to generate another answer until a substitute answer not already found in the context has been generated.

The prompt used to query an LLM for a substitute answer was as follows:

Prompt for getting substitute answers.

```
## Instructions
Please provide an incorrect answer to the example below.
The incorrect answer should be incorrect in the sense that it should be significantly different from the original answer. At the same time, it should be a plausible answer to the given question.
The incorrect answer should follow the same formatting as the original answer such that it should be possible to directly replace the original answer with the incorrect answer in any context.
The incorrect answer should be a single word or a short phrase.
Only output the incorrect answer.

## Example
Question: <question>
Original answer:<target_true>
Incorrect answer:
```

In the event that the model generated a substitute answer that already could be found in the context, the previous model answer was added to the chat history together with the following new user query:

Prompt for getting another substitute answer.

Please provide another incorrect answer following the same format as the original answer. Only output the incorrect answer.

Quality of conflicting contexts A manual inspection of 200 samples found the method reliable for producing adequate conflicting contexts with an accuracy of 90% (11 samples corresponded to poor formatting, 4 were too similar to gold, and 4 were dropped due to data formatting issues or the LLM being unable to generate a substitute answer not already found in the context). In addition, we inspect the CUB results to ascertain the quality of the conflicting context sampling, see Appendix A.

We also experimented with a method based on named entities and random sampling for producing substitute answers for the conflicting contexts. In the method, the entity type of the answer to be replaced was detected and another named entity of the same type was randomly sampled from a NE dataset as the replacement. We found this method to work poorly compared to the LLM based approach. Mainly because the detected NEs lacked sufficient information for a successful sampling of replacements (e.g. “2024” and “last year” may both be labelled as time entities, while they are not interchangeable in all contexts).

Sampling of irrelevant contexts Given a query and a corresponding Wikipedia page, the NQ anno-

tators were instructed to mark the first paragraph in the Wikipedia page that contains an answer to the query. Therefore, to ensure that we only sample irrelevant contexts, we perform the sampling over all paragraphs before the gold paragraph in the given Wikipedia page.

We use the Jina Reranker v2⁷ to identify the most relevant non-gold paragraph. It is a modern LM re-ranker that has been proven to work well on NQ (Hagström et al., 2025).

Prompt templates The 2-shot prompts used to evaluate the LMs on NQ were as follows.

Prompt without context.

```
Answer the following questions.
Question: When is the first episode of House of
the Dragon released?
Answer: August 21, 2022

Question: In what country will the 2026 Winter
Olympics be held?
Answer: Italy

Question: <question>
Answer:
```

Prompt with context.

```
Answer the following questions based on the
context below.
Question: When is the first episode of House of
the Dragon released?
Context: <Table> <Tr> <Th> Season </Th> <Th>
Episodes </Th> <Th> First released </Th> <Th>
Last released </Th> </Tr> <Tr> <Td> 1 </Td>
<Td> 10 </Td> <Td> August 21, 2022 </Td> <
Td> October 23, 2022 </Td> </Tr> <Tr> <Td> 2
</Td> <Td> 8 </Td> <Td> June 16, 2024 </Td>
<Td> August 4, 2024
</Td> </Tr> </Table>
Answer: August 21, 2022

Question: Where will the 2026 Winter Olympics be
held?
Context: <P> The 2026 Winter Olympics (Italian:
Olimpiadi invernali del 2026), officially
the XXV Olympic Winter Games and commonly
known as Milano Cortina 2026, is an upcoming
international multi-sport event scheduled
to take place from 6 to 22 February 2026 at
sites across Lombardy and Northeast Italy.
</P>
Answer: Lombardy and Northeast Italy

Question: <question>
Context: <context>
Answer:
```

For the instruction-tuned Qwen models, a chat template with slightly different prompt templates was used. The 2-shot prompt templates for the instruction-tuned models were as follows.

Prompt without context for instruction-tuned LMs.

```
<|im_start|>system
You are Qwen, created by Alibaba Cloud. You are
a helpful assistant.<|im_end|>
<|im_start|>user
Answer the question. Only answer with the answer.
Examples of questions and desired answers
are given below.

# Example 1
Question: When is the first episode of House of
the Dragon released?
Answer: August 21, 2022

# Example 2
Question: In what country will the 2026 Winter
Olympics be held?
Answer: Italy

# Now, answer the following question (only with
the answer):
Question: <question>
Answer: <|im_end|>
<|im_start|>assistant
```

Prompt with context for instruction-tuned LMs.

```
<|im_start|>system
You are Qwen, created by Alibaba Cloud. You are
a helpful assistant.<|im_end|>
<|im_start|>user
Answer the question based on the provided
context. Only answer with the answer.
Examples of questions and desired answers
are given below.

# Example 1
Question: When is the first episode of House of
the Dragon released?
Context: <Table> <Tr> <Th> Season </Th> <Th>
Episodes </Th> <Th> First released </Th> <Th>
Last released </Th> </Tr> <Tr> <Td> 1 </Td>
<Td> 10 </Td> <Td> August 21, 2022 </Td> <
Td> October 23, 2022 </Td> </Tr> <Tr> <Td> 2
</Td> <Td> 8 </Td> <Td> June 16, 2024 </Td>
<Td> August 4, 2024
</Td> </Tr> </Table>
Answer: August 21, 2022

# Example 2
Question: Where will the 2026 Winter Olympics be
held?
Context: <P> The 2026 Winter Olympics (Italian:
Olimpiadi invernali del 2026), officially
the XXV Olympic Winter Games and commonly
known as Milano Cortina 2026, is an upcoming
international multi-sport event scheduled
to take place from 6 to 22 February 2026 at
sites across Lombardy and Northeast Italy.
</P>
Answer: Lombardy and Northeast Italy

# Now, answer the following question (only with
the answer):
Question: <question>
Context: <context>
Answer: <|im_end|>
<|im_start|>assistant
```

⁷jinaai/jina-reranker-v2-base-multilingual

B.3 DRUID

We map the stances of DRUID to context type using the following approach:

1. Gold: If the evidence is relevant and the stance of the evidence aligns with the claim verdict reached by the fact-check site (here considered gold). This automatically encompasses most samples with evidence that has been sampled from a fact-check site, as the stance of the evidence is likely to align with the FC verdict.
2. Conflicting: If the evidence is relevant and the stance of the evidence does not align with the claim verdict. This automatically encompasses all samples with insufficient evidence, as the original FC verdicts always are True, Half True or False.
3. Irrelevant: If the evidence is irrelevant.

Samples from the DRUID dataset can be found in Table 11. The evidence stance and fact-check verdict distributions per context type can be found in Tables 12 and 13.

Prompt templates The 2-shot prompts used for evaluating the LMs on DRUID were as follows.

Prompt without context.

```
Are the following claims True or False? Answer
None if you are not sure or cannot answer.

Claimant: Viral post
Claim: "the new coronavirus has HIV proteins
that indicate it was genetically modified in
a laboratory."
Answer: False

Claimant: Sara Daniels
Claim: "Blackpink released the single 'You me
too' in 2026."
Answer: None

Claimant: <claimant>
Claim: "<claim>"
Answer:
```

Prompt with context.

```
Are the claims True or False based on the
accompanying evidence? If you are not sure
or cannot answer, say None.

Claimant: Viral post
Claim: "the new coronavirus has HIV proteins
that indicate it was genetically modified in
a laboratory."
Evidence: "Microbiologists say the spike
proteins found in the new coronavirus are
different from the ones found in HIV. [...]
There is no evidence to suggest the
coronavirus was genetically modified."
Answer: False
```

```
Claimant: Sara Daniels
Claim: "Blackpink released the single 'You me
too' in 2026."
Evidence: "Blackpink released their album 'Born
Pink' in 2022."
Answer: None

Claimant: <claimant>
Claim: "<claim>"
Evidence: "<evidence>"
Answer:
```

For the instruction-tuned Qwen models, a chat template with slightly different prompt templates was used for compatibility. The 2-shot prompt templates for the instruction-tuned models were as follows.

Prompt without context for instruction-tuned LMs.

```
<|im_start|>system
You are Qwen, created by Alibaba Cloud. You are
a helpful assistant.<|im_end|>
<|im_start|>user
Is the claim True or False? Answer None if you
are not sure or cannot answer. Only answer
with True, False or None. Examples of claims
and desired answers are given below.

# Example 1
Claimant: Viral post
Claim: "the new coronavirus has HIV proteins
that indicate it was genetically modified in
a laboratory."
Answer: False

# Example 2
Claimant: Sara Daniels
Claim: "Blackpink released the single 'You me
too' in 2026."
Answer: None

# Now, answer for the following claim:
Claimant: <claimant>
Claim: "<claim>"
Answer (True, False or None):<|im_end|>
<|im_start|>assistant
```

Prompt with context for instruction-tuned LMs.

```
<|im_start|>system
You are Qwen, created by Alibaba Cloud. You are
a helpful assistant.<|im_end|>
<|im_start|>user
Is the claim True or False based on the
accompanying evidence? If you are not sure
or cannot answer, say None. Only answer with
True, False or None. Examples of claims,
evidence and desired answers are given below
.

# Example 1
Claimant: Viral post
Claim: "the new coronavirus has HIV proteins
that indicate it was genetically modified in
a laboratory."
Evidence: "Microbiologists say the spike
proteins found in the new coronavirus are
```

Claimant	Claim	Verdict	Evidence	Type
Viral Claim	Harvard professor Charles Lieber was arrested for manufacturing and selling the new coronavirus to China	False	Lieber was arrested on January 28 for "making false statements to the agency of the United States Government," or lying to federal authorities about his ties to China, as per the fact-check report. The channel added that prosecutors have never alleged that Lieber was involved in manufacturing and/or selling a virus to China. The full federal court complaint against Dr Lieber can be read here.</p><p>The report also clarified Lieber's links to Wuhan. The report stated, "Lieber travelled to WUT (Wuhan University of Technology) in mid-November 2011 ostensibly in order to participate in a Nano-Energy Materials Forum."</p><p>On July 29, Dr Lieber's attorney Marc Mukasey told WCVB Channel 5 that he didn't hide anything or get paid as the government alleges.</p><p>Thus, the social media claim that Harvard professor Dr Charles Lieber "made and sold" the Covid-19 virus to China is false.</p> (See attached file: List of Black Money Holders from Wiki	Gold
FACEBOOK POST	WikiLeaks has published the 1st list of black money holders in Swiss banks.	False		Conflicting
Irish Congress of Trade Unions (ICTU)	One in five school staff in Northern Ireland are assaulted at least once a week.	False	Finnegan, who died in January 2002, had also abused boys at St. Colman's College, a prestigious Catholic boys' secondary school in Newry, Northern Ireland. He taught there from 1967 to 1971 and again from 1973 to 1976, when he was appointed president of the school. He served in that post until 1987. [...] Admitted on October 9, 2014 to sample charges of indecently assaulting four boys as young as 10 at St Mary's CBS primary school in Mullingar between 1984 and 1987. Jailed for two years at Mullingar Circuit Court sitting in Tullamore. This concluded a ten-year investigation by detectives in Mullingar. [...] When Smyth returned to Kilnacrott in 1983, he again began abusing children in Belfast, including the girl who, on February 23, 1990, would meet with a social worker at the Catholic Family Welfare Society in Belfast and start all the Smyth revelations.	Irrelevant

Table 11: DRUID samples and corresponding context types.

Context	Evidence stance	Count
Gold	Refutes	1,579
	Supports	359
Conflicting	Refutes	35
	Insufficient-refutes	437
	Insufficient-contradictory	163
	Insufficient-neutral	892
	Insufficient-supports	585
Irrelevant	Supports	367
	not applicable	83

Table 12: Stance distribution per context type for DRUID.

Context	FC verdict	Count
Gold	False	1,579
	True	359
Conflicting	False	1,842
	Half True	276
	True	361
Irrelevant	False	54
	Half True	13
	True	16

Table 13: Fact-check verdict distribution per context type for DRUID.

different from the ones found in HIV. [...]
There is no evidence to suggest the coronavirus was genetically modified."
Answer: False
Example 2
Claimant: Sara Daniels
Claim: "Blackpink released the single 'You me too' in 2026."
Evidence: "Blackpink released their album 'Born Pink' in 2022."
Answer: None
Now, answer for the following claim:
Claimant: <claimant>
Claim: "<claim>"

Evidence: "<evidence>"
Answer (True, False or None):<|im_end|>
<|im_start|>assistant

C CCU metric

BCU cannot measure the difference in model behaviour when context is introduced, as it does not take model behaviour without context into consideration. To address this, we introduce CCU. Given a query Q and context C , CCU measures the change in probability for token t as follows.

$$CCU(t) = \begin{cases} \frac{P_M(t|Q,C) - P_M(t|Q)}{1 - P_M(t|Q)} & \text{if } P_M(t|Q,C) \geq P_M(t|Q), \\ \frac{P_M(t|Q,C) - P_M(t|Q)}{P_M(t|Q)} & \text{otherwise.} \end{cases} \quad (1)$$

For relevant contexts C we record $CCU(t_C)$, i.e. the scores for the token promoted by the context. For irrelevant contexts we record the $CCU(t_M)$, i.e. the scores for the top token predicted by the model when prompted without context (memory). The range of CCU is $[-1, 1]$, for which a value of -1 denotes that the model goes completely *against* the context when the context is relevant or against its memory when the context is irrelevant, and vice versa for CCU values of 1. We report the averaged CCU per context type.

By measuring the token probabilities before and after context is introduced, the CCU metric more accurately captures how the LM is impacted by

context. However, this metric excludes the Command A model, which does not provide the output logits necessary to compute CCU scores.

D Hyperparameter Search

D.1 Prompting

The tuned prompt found for each model and dataset can be found in Table 14. Different sets of prompts were experimented with depending on dataset and model type. A set of 11 to 12 prompts were produced for each of CounterFact, NQ and DRUID for the three different model types (causal LM, instruction-tuned LMs and Command A), respectively. Prompts with the same number are similar to each other across model types (e.g. Prompt #2 for Qwen2.5 on DRUID is similar to Prompt #2 for instruction-tuned Qwen2.5 on DRUID). Prompt sets across different datasets are dissimilar as they are adapted to align the instructions and few-shot examples to the given dataset. Prompt sets across different model types for the same dataset are dissimilar as small tweaks need to be applied for the instruction-tuned models that work less well in a purely causal language modelling setting, and for Command A that is a chat-based model. All prompts will be possible to view in the code repository of the paper.

D.2 PH3

The tuned attention head configurations for PH3 can be found in Table 15. The head configurations are grouped by the top number of identified attention heads to consider and to what extent we allow mixing between context and memory heads. E.g. #25 all denotes all top-25 context and memory heads detected, #3 memory denotes the top-3 memory heads, allowing for overlap with context heads, and #1 only memory denotes memory heads detected without overlap with context heads when considering the top-1 context and memory heads.

D.3 Context-aware Contrastive Decoding: COIECD

Unlike other CMTs, the hyperparameters used in COIECD, α and λ , are selected following the original paper, Yuan et al. (2024), using the gold context from the validation set of NQ dataset. This deviation is necessary, as optimising COIECD’s hyperparameters by maximising the average BCU across all context types causes the model to converge to using only the output distribution without

Dataset	Model		Prompt
CounterFact	GPT2-XL	1.5B	<i>default</i>
		6.9B	Prompt #10 (ChatGPT)
		1.5B	Prompt #1 (Jin et al. (2024))
		7B	Prompt #11 (ChatGPT)
		32B	Prompt #8 (ChatGPT)
	QWEN2.5-I	1.5B	Instruct-prompt #4 (manual)
		7B	Instruct-prompt #11 (ChatGPT)
		32B	Instruct-prompt #3 (manual)
	COMMAND A		Prompt #5 (ChatGPT)
NQ	GPT2-XL	1.5B	Prompt #2 (manual)
		6.9B	<i>default</i>
		1.5B	Prompt #7 (ChatGPT)
		7B	Prompt #6 (ChatGPT)
		32B	Prompt #5 (manual)
	QWEN2.5-I	1.5B	Prompt #5 (manual)
		7B	Prompt #3 (manual)
		32B	<i>default</i>
	COMMAND A		<i>default</i>
DRUID	GPT2-XL	1.5B	Prompt #8 (ChatGPT)
		6.9B	Prompt #2 (manual)
		1.5B	Prompt #2 (manual)
		7B	Prompt #11 (Microsoft Copilot)
		32B	Prompt #1 (manual)
	QWEN2.5-I	1.5B	<i>default</i>
		7B	<i>default</i>
		32B	Prompt #2 (manual)
	COMMAND A		Prompt #1 (manual)

Table 14: The tuned prompts for each LM. *default* denotes that the original prompt template (seen in Appendix B) worked best. “-I” denotes instruction-tuned model versions. The source of the prompt is indicated in parenthesis.

context in the decoding step. This outcome arises from the nature of COIECD, where always relying on the distribution without context results in a BCU score of 1.0 for irrelevant contexts, while also causing the model to ignore context, including gold and conflicting contexts. To prevent COIECD from collapsing into regular generation without context and to enable meaningful comparison with other CMTs, we follow the hyperparameter search from the original paper. While Yuan et al. (2024) uses the same hyperparameter values across all models, our models exhibit different tendencies during hyperparameter search. Therefore, we tune the hyperparameters separately for each model to ensure a fair comparison with other methods. We search α in the range [0.0, 2.0] and λ in the range [0.1, 1.0], and the hyperparameters for each model are in Table 16.

E Implementation Details of Fine-tuning

We fine-tune the LMs with a learning rate of $5e-5$,⁸ using warm-up. To avoid overfitting, we use early stopping based on the loss on the validation set. For QA datasets, we use the train split from SQuAD 2.0

⁸Experiments with other learning rates yielded insignificant changes in performance on the validation set.

(Rajpurkar et al., 2018), and TriviaQA (Joshi et al., 2017). For a FC dataset, we take the train split from AVeriTeC (Schlichtkrull et al., 2023). For a sentence completion dataset, we take the static partition of the DYNAMICQA (Marjanovic et al., 2024). We only create counterfactual training examples with DYNAMICQA dataset. The detailed statistics for mixing the selected datasets can be found in Table 17.

F Additional Details of Multi-agent

Algorithm 1 Multi-agent

```

1: Given: question  $q$ , context  $c$ 
2: Stage1: Relevance Assessment
3: Predict  $f_{\text{rel}} \sim \text{LM}_{\text{rel}}(f_{\text{rel}} \mid q, c)$ 
4: if  $f_{\text{rel}} = \text{Relevant}$  then
5:   Proceed to Stage 2
6: else
7:   return  $\text{LM}(a \mid q)$   $\triangleright$  Answer w/o  $c$ 
8: end if
9: Stage 2: Context-Faithfulness
10: Predict  $a_c \sim \text{LM}(a_c \mid q, c)$ 
11: Predict  $f_{\text{faith}} \sim \text{LM}_{\text{faith}}(f_{\text{faith}} \mid q, c, a_c)$ 
12: if  $f_{\text{faith}} = \text{Faithful}$  then
13:   return  $a_c$   $\triangleright$  Answer w/  $c$ 
14: else
15:   Proceed to Stage 3
16: end if
17: Stage 3: Self-Refinement
18: return  $\text{LM}(a \mid q, c, a_c, f_{\text{faith}})$   $\triangleright$  Self-Refined

```

We design the Multi-agent approach to investigate whether LMs can explicitly handle the two objectives of context utilisation: (1) being robust to irrelevant context and (2) being faithful to relevant context. Rather than directly generating an answer, an LM is guided to perform intermediate reasoning steps, each handled by a dedicated LM agent. This decomposition allows us to understand whether LMs can explicitly recognise when the context should be used and whether their answer aligns with it when it is. While self-refinement and LM agent have been used broadly in reasoning tasks (Du et al., 2024; Feng et al., 2024; Madaan et al., 2023), our motivation is grounded in examining two components of context utilisation separately. Notably, self-refinement is only applied when the context is assessed as relevant but the answer is assessed as unfaithful, reflecting our focus on improving the usage of relevant context. By

structuring the problem in this way, we aim to better understand the extent to which LMs can reason about context relevance and faithfulness.

Figure 2 and Algorithm 1 outline the Multi-agent procedure employed in our framework. Given a question and the context, the model first undergoes a relevance assessment stage, where it is explicitly instructed to determine whether the context is relevant to the question (Shen et al., 2024). If assessed as irrelevant, the model answers without the context; if relevant, it incorporates the context to generate the initial answer and proceeds to the next stage. In the context faithfulness assessment, the model is instructed to provide feedback on whether its answer faithfully reflects the provided context. If deemed faithful, the answer is retained as the final answer. If the prediction is assessed as unfaithful, the model is instructed to refine its answer using the question, context, initial answer, and feedback derived from the faithfulness assessment. This self-refinement stage encourages the model to self-correct based on its own feedback. To ensure consistency in output formatting during refinement, we incorporate two-shot demonstrations.

The templates for relevance assessment, context faithfulness, and self-refinement are presented below. Task-specific templates for each dataset are available in the released code.

Relevance Assessment (NQ)

You are a relevance assessment expert. Your task is to evaluate whether the provided context is relevant to the question.

Context: {context}
Question: {question}

If the provided context is relevant to the question, answer "Relevant", otherwise answer "Irrelevant". Do not rely on your own knowledge or judge the factual accuracy of the context.

Answer:

Context faithfulness (CounterFact and NQ)

You are a context-faithfulness expert. Your task is to evaluate whether the proposed answer faithfully uses the information in the provided context.

Context: {context}
Question: {question}
Proposed answer: {response}

Does the answer faithfully reflect the content of the context? Do not rely on your own

1710
1711
1712
1713
1714

1715
1716
1717
1718
1719
1720
1721
1722
1723
1724
1725
1726
1727
1728
1729
1730
1731
1732

knowledge or judge the factual accuracy of the context. Please explain briefly.

Feedback:

Self-refinement (NQ)

Your task is to generate the best possible final answer to the question, based on the expert feedback.

You may keep the original proposed answer if it is correct, or revise it if the feedback suggests it is incorrect or unsupported. Generate only the final answer. Do not include any explanation or repeat the prompt.

{Two demonstrations}

Context: {context}
Question: {question}
Proposed answer: {response}
Feedback on context faithfulness: {feedback}
Final answer:

G Input Features

We detect the input features described in Section 5.2 as follows:

- Context length is measured by the number of characters in the context.
- Flesch reading ease score is measured with the textstat⁹ module.
- Query-context overlap is measured as the size of the set of words that form the intersection of the set of words in the query and context, respectively, normalised by the size of the set of query words. CounterFact is excluded from this analysis as its synthetic samples yield trivial results for this feature.
- The answer position is measured as the index of the answer in the context normalised by context length. This feature is only detectable for gold and conflicting contexts for CounterFact and NQ.
- The distractor rate is measured as the number of answer entities found in the context, divided by the total number of entities in the context with an entity type that matches the answer entity type(s).¹⁰ This feature is similarly only measurable for gold and conflicting contexts from CounterFact and NQ.
- Relevance is given by the relevance agent based on Qwen 32B Instruct from the Multi-agent setup. It labels context as either ‘relevant’ or ‘irrelevant’.

⁹<https://github.com/textstat/textstat>
¹⁰Named entities are detected using spaCy and en_core_web_trf.

H Computational Resources

GPT2-XL was evaluated using one Nvidia T4 GPU. Pythia, Qwen 1.5B and Qwen 7B using one A40 GPU. Qwen 32B was evaluated using four A40 GPUs. The compute budget for all CMTs was about 14 hours per model for CounterFact, 28 hours per model for NQ and 21 hours per model for DRUID, amounting to a total of about 900 GPU hours.

The costs for the experiments with Cohere Command A amounted to a total of about 120 USD.

I Use of AI assistants

AI assistants like Copilot and ChatGPT were intermittently used to generate template code and rephrase sentences in the paper, etc. However, no complete paper sections or code scripts have been generated by an AI assistant. All generated content has been inspected and verified by the authors.

1763

1764
1765
1766
1767
1768
1769
1770
1771
1772
1773

1774

1775
1776
1777
1778
1779
1780

Model	Mode	CounterFact	NQ	DRUID
GPT2-XL	+context	#25 all L18H10, L21H10, L21H7, L22H18, L22H20, L24H6, L26H14, L26H20, L26H8, L27H15, L27H5, L28H15, L29H5, L29H9, L30H21, L30H8, L31H0, L31H3, L31H8, L32H13, L33H14, L33H18, L33H2, L33H7, L34H17, L34H20, L35H17, L35H19, L35H21, L36H17, L36H2, L37H7, L38H24, L38H7, L39H12, L39H9, L40H13, L40H23, L41H5, L41H9, L42H24, L43H15, L47H0	#1 all L28H15, L35H19	#5 only memory* L32H13, L35H19, L42H24, L43H15
	+memory	#12 memory L26H14, L26H8, L32H13, L33H14, L35H19, L38H24, L40H23, L41H5, L42H24, L43H15, L47H0, L30H8	#7 only context L27H15, L28H15, L29H9, L33H2, L34H17, L37H7	#22 all L21H10, L22H20, L24H6, L26H14, L26H20, L26H8, L27H15, L27H5, L28H15, L29H9, L30H21, L30H8, L31H0, L31H3, L31H8, L32H13, L33H14, L33H18, L33H2, L33H7, L34H17, L34H20, L35H17, L35H19, L36H17, L36H2, L37H7, L38H24, L38H7, L39H12, L39H9, L40H13, L40H23, L41H5, L42H24, L43H15, L47H0
PYTHIA 6.9B	+context	#15 memory L10H27, L14H6, L16H16, L17H28, L19H11, L19H21, L20H11, L20H18, L21H8, L27H22, L18H7, L19H28, L20H2, L20H8, L24H5	#17 only memory L10H27, L14H28, L14H6, L16H16, L17H28, L19H11, L19H21, L20H11, L20H18, L21H8, L22H12, L27H22	#10 only context L12H11, L12H13, L14H0, L15H17, L17H14, L20H2, L8H11
	+memory	#25 only context L10H1, L12H11, L12H13, L13H12, L14H0, L14H23, L15H17, L17H14, L18H10, L19H1, L19H20, L21H10, L23H25, L29H22, L8H11, L8H24	#12 only context L12H11, L12H13, L14H0, L14H23, L15H17, L17H14, L19H31, L20H2, L8H11	#17 only context L10H1, L12H11, L12H13, L13H12, L14H0, L14H23, L15H17, L17H14, L18H10, L19H1, L19H31, L8H11
QWEN2.5 1.5B	+context	#15 only memory L10H0, L10H1, L13H1, L16H1, L17H0, L18H0, L1H1, L3H0	#12 only memory L10H0, L13H1, L16H1, L17H0, L18H0, L1H1	#17 only context L14H1, L16H0, L18H1, L19H0, L19H1, L20H1, L24H1, L26H0, L26H1, L9H0
	+memory	#5 only context L15H1, L16H0, L27H0	#12 only context L14H1, L16H0, L18H1, L19H0, L24H1, L27H0	#12 only memory L10H0, L13H1, L16H1, L17H0, L18H0, L1H1
QWEN2.5 1.5B <i>Instruct</i>	+context	#7 only memory L15H0, L1H1, L21H0	#1 only context L19H1	#10 only context L14H0, L17H1, L19H1, L22H0, L26H0
	+memory	#1 only context L19H1	#12 only context* L14H0, L17H1, L19H1, L22H0, L26H0, L27H0	#5 only context L17H0, L19H1, L22H0
QWEN2.5 7B	+context	#7 memory L0H0, L17H1, L18H2, L19H0, L21H0, L22H2, L23H0	#1 only context L27H0	#3 only memory L0H0, L22H2
	+memory	#15 only context L13H0, L17H0, L18H1, L18H3, L22H0, L24H3, L25H1, L26H0, L27H0, L27H2	#5 only context L22H0, L27H0, L27H2	#12 only context L16H3, L17H0, L18H1, L18H3, L22H0, L24H3, L26H0, L27H0, L27H2
QWEN2.5 7B <i>Instruct</i>	+context	#17 only memory L11H1, L12H0, L13H3, L14H3, L16H1, L17H0, L17H3, L18H2, L1H1, L20H0, L21H2, L26H3, L3H0	#5 context L18H0, L18H3, L22H2, L23H0, L27H2	#5 only context L18H0, L18H3, L27H2
	+memory	#3 only context L18H0	#3 only context L18H0	#17 all L0H0, L11H1, L12H0, L13H3, L14H3, L15H1, L16H0, L16H1, L17H0, L17H3, L18H0, L18H1, L18H2, L18H3, L19H0, L19H3, L1H1, L20H0, L20H2, L20H3, L21H0, L21H2, L22H0, L22H2, L23H0, L26H3, L27H0, L27H2, L3H0, L8H1

Table 15: Tuned PH3 attention head configurations for each model and evaluation dataset. +context indicates heads for which pruning leads to increased context usage and vice versa for +memory. Configurations marked with * denote that they yielded degraded performance compared to the standard setting (no mechanistic intervention) on the validation set.

Model	λ	α
GPT2-XL	0.50	1.00
PYTHIA 6.9B	0.50	1.00
QWEN2.5 1.5B	1.00	0.50
QWEN2.5 1.5B INSTRUCT	0.50	1.00
QWEN2.5 7B	1.00	1.00
QWEN2.5 7B INSTRUCT	0.50	0.50
QWEN2.5 32B	0.50	1.00
QWEN2.5 32B INSTRUCT	0.50	1.50

Table 16: Selected COIECD hyperparameters λ and α for each model, evaluated on gold contexts from NQ’s validation set. For models with multiple (λ, α) pairs attaining the maximum score, we choose the setting that lies near the midpoint of the optimal region.

Dataset	Dataset weight	Context type	Context weight
SQuAD 2.0	0.4	Relevant	0.65
		Irrelevant	0.25
		Empty	0.1
TriviaQA	0.3	Relevant	0.65
		Irrelevant	0.25
		Empty	0.10
AVeriTeC	0.15	Relevant	0.65
		Irrelevant	0.25
		Empty	0.10
DYNAMICQA	0.15	Relevant	0.50
		Irrelevant	0.05
		Empty	0.05
		Counterfactual	0.40

Table 17: Sampling weight for each dataset. We first sample the number of instances for each dataset following the dataset sampling weight. Then, each context type is determined by the context sampling weight.