

Face the Facts!

Evaluating RAG-based Pipelines for Professional Fact-Checking

Daniel Russo^{1,2}, Stefano Menini¹, Jacopo Staiano², Marco Guerini¹

¹Fondazione Bruno Kessler, Italy

²University of Trento, Italy

{drusso, menini, guerini}@fbk.eu, jacopo.staiano@unitn.it

Abstract

Natural Language Processing and Generation systems have recently shown the potential to complement and streamline the costly and time-consuming job of professional fact-checkers. In this work, we lift several constraints of current state-of-the-art pipelines for automated fact-checking based on the Retrieval-Augmented Generation (RAG) paradigm. Our goal is to benchmark, following professional fact-checking practices, RAG-based methods for the generation of verdicts - i.e., short texts discussing the veracity of a claim - evaluating them on stylistically complex claims and heterogeneous, yet reliable, knowledge bases. Our findings show a complex landscape, where, for example, LLM-based retrievers outperform other retrieval techniques, though they still struggle with heterogeneous knowledge bases; larger models excel in verdict faithfulness, while smaller models provide better context adherence, with human evaluations favouring zero-shot and one-shot approaches for informativeness, and fine-tuned models for emotional alignment.

1 Introduction

Despite the efforts to validate the accuracy of on-line content, professional fact-checkers are increasingly struggling to keep up with the rapid spread of misinformation (Lewis et al., 2008; Adair et al., 2017; Godler and Reich, 2017; Wang et al., 2018). Therefore, Natural Language Processing (NLP) has been proposed as a viable solution to partially automate the costly process of verifying misleading claims online (Vlachos and Riedel, 2014). Within this context, the task of *verdict production*, i.e., explaining why a claim is true or false, stands as one of the most challenging (Kotonya and Toni, 2020a; Guo et al., 2022).

Framing verdict production as a summarization task over fact-checking articles is a suitable solution due to the possibility of generating

highly readable verdicts even for non-expert users (Atanasova, 2024; Kotonya and Toni, 2020b; Russo et al., 2023b). Despite their promising results, summarization-based approaches suffer from two main limitations: (i) they rely on the assumption that a fact-checking article always exists for a given claim; and (ii), they further assume that claims are already paired with a fact-checking article, which is typical in fact-checking websites but not on social media platforms, where most of the misinformation spreads (Lazer et al., 2018).

Grounding textual generation on retrieved evidence, an approach named Retrieval-Augmented Generation (RAG), has been shown to be effective for knowledge-intensive tasks like fact verification (Lewis et al., 2020; Liao et al., 2023; Xu et al., 2023); it also allows addressing the limitations of previous summarization approaches, e.g., the assumption that the claims are already paired with gold fact-checking articles. Moreover, RAG-based approaches have proven useful in reducing potential factual inconsistencies, often referred to as “hallucinations” (Zellers et al., 2019; Solaiman et al., 2019), during text generation (Lewis et al., 2020), making them attractive for fact-checking tasks. Thus, researchers have increasingly adopted RAG to enhance the accuracy of the generated verdicts (Zeng and Gao, 2024; Yao et al., 2023).

Current studies depend on fact-checking websites, resulting in verdicts characterized by formal and dry language. This style contrasts sharply with the language used on Social Media Platforms (SMPs), which is typically more complex and includes *noise* such as personal commentary, or emotional content that surrounds the core fact. Such mismatch might pose serious issues when countering misinformation online (Colliander, 2019). We challenge these common assumptions on verdict generation, by testing RAG-based pipelines across scenarios that progressively approximate professional fact-checking practices.

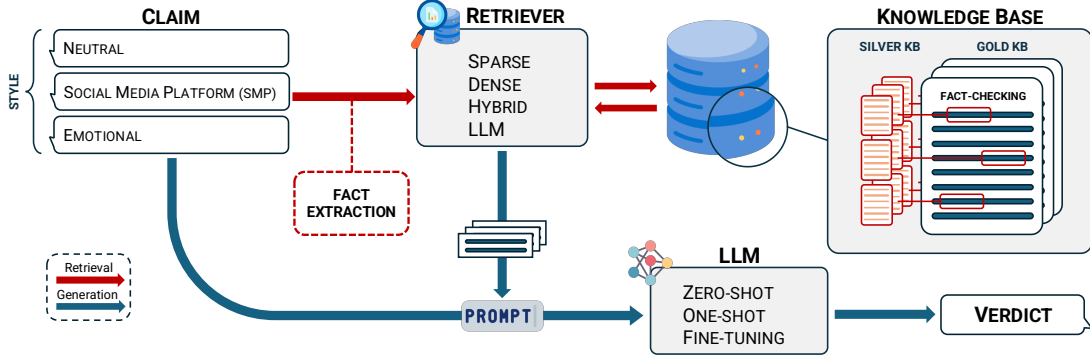


Figure 1: Visual representation of our RAG-based experimental design (the steps for retrieval and generation are indicated by the red and blue lines, respectively). We explored various configurations to tackle increasingly realistic scenarios across different claim styles (neutral, SMP, emotional) and Knowledge Bases (Gold vs. Silver), as well as varying computational demands through multiple retriever architectures (sparse, dense, hybrid, and LLM-based) and five distinct LLMs generation setups (zero-shot, one-shot, fine-tuning).

To this end, we present a thorough evaluation of verdict production across several dimensions of a RAG pipeline, testing key configurations for each of them (see Figure 1):

1. Three claim styles differing in realism: *neutral* (journalistic), *SMP* (social media-like), and *emotional* (SMP enriched with an emotional component) taken from VerMouth dataset (Russo et al., 2023a);
2. Four retrieval methods with varying computational costs: *sparse*, *dense*, *hybrid*, and *LLM-based*;
3. Two claim pre-processing settings, with and without fact extraction (to simplify complex claims written in SMP and emotional style before the retrieval step);
4. Five LLMs varying in size and training for verdict generation;
5. Three generation setups: *zero-shot*, *one-shot*, and *fine-tuned*;
6. Two types of knowledge bases: a *gold KB* (with verified fact-checks) and a *silver KB* (comprising only the relevant and reliable sources used to write fact-checking articles);
7. Two document storage strategies: retrieving *full articles* or smaller *chunks*.

We show that LLM-based retrievers consistently outperform other methods, though they face challenges with silver knowledge bases. Dense retrievers manage stylistic variations of the claim effectively but fall short compared to LLMs, whereas sparse retrievers exhibit high sensitivity to noise present in emotional and SMP claims. Hybrid approaches and query pre-processing improve per-

formance. Turning to generation, larger models excel in faithfulness and alignment with gold verdicts, while smaller ones are more consistent in context adherence. Fine-tuning boosts similarity of the generated verdict to the gold but reduces contextual accuracy, with human evaluations favouring verdicts generated under zero/one-shot strategies for informativeness and fine-tuning for emotional alignment.¹

2 Related Work

Early approaches to verdict generation leveraged either attention modules to highlight salient tokens from the evidence text (Popat et al., 2018; Shu et al., 2019; Lu and Li, 2020; Wu et al., 2020; Yang et al., 2019), or Horn Rules to reason upon structured knowledge bases (Gad-Elrab et al., 2019; Ahmadi et al., 2019). However, both lack readability, being hard to interpret by common users (Guo et al., 2022). To overcome this issue, researchers started casting verdict production as a summarization task over fact-checking articles, either through extractive (Atanasova et al., 2020), abstractive (Kotonya and Toni, 2020a; Stambach and Ash, 2020), or hybrid (Russo et al., 2023b) summarization pipelines. More recently, He et al. (2023) introduced a reinforcement learning-based framework for generating counter-misinformation responses to social media content.

Ad-hoc data collection strategies for verdict generation rely either on synthetic data generation, like e-FEVER (Stambach and Ash, 2020), or on journalistic sources, such as LIARPLUS (Alhindi et al.,

¹Code and data are publicly available on GitHub: <https://github.com/drusso98/face-the-facts>.

2018), PUBHEALTH (Kotonya and Toni, 2020b), RU22Fact (Zeng et al., 2024), LIAR++, and FullFact (Russo et al., 2023b). Russo et al. (2025) proposed EuroVerdict, a multilingual, manually curated dataset for verdict generation. For more realistic, SMP-style claims, He et al. (2023) developed MisinfoCorrect, and Russo et al. (2023a) extended FullFact with VerMouth, incorporating emotional claims and verdicts grounded in trustworthy fact-checking articles.

While the latest approaches provide readable verdicts, they may lack faithfulness due to language models generating factual inaccuracies. Additionally, summarization approaches assume that trustworthy evidence is always available for the claim under inspection. Therefore, RAG-based approaches have been employed to guide the generation of reliable verdicts upon trustworthy evidence previously retrieved from a knowledge base (KB). To this end, Zeng and Gao (2024) proposed JustiLM, a few-shot RAG-based approach for the generation of verdicts for real-world claims, by leveraging both fact-checking articles and auxiliary evidence during model training. Yao et al. (2023) developed an end-to-end RAG-based system to perform verdict prediction and production in a multimodal setting jointly. Nevertheless, both studies concentrate on journalistic data and writing styles, without considering the communication style employed on SMP.

Building on the work of Zeng and Gao (2024), we present an extensive evaluation of RAG pipelines for verdict generation, testing various combinations of retrieval and generation strategies. Progressing toward increasingly realistic scenarios, with respect to the workflow adopted by professional fact-checkers, we address the challenge of handling claims and sources that vary in style and complexity, aiming to closely mimic the kinds of claims that everyday users might encounter on social media platforms. Specifically, we assess the impact of claim writing style on retriever performance, highlighting where differences arise as the style shifts toward that used by SMP users. Finally, we explore the extreme scenario where no fact-checking article exists, relying solely on reliable supporting evidence.

3 Experimental Design

In this section, we provide details on the experimental design: from the datasets used, through the

retrieval methods adopted, to the configurations of the LLMs employed for verdict generation.

3.1 Dataset

To study the impact of different styles on a RAG-based verdict production task, we used FullFact (Russo et al., 2023b) and VerMouth (Russo et al., 2023a) datasets. The two datasets comprise eight different versions of the same claims and verdicts. FullFact provides data written in a journalistic style scraped from fullfact.org while VerMouth proposes the same data rewritten in an SMP style, and also enriched with the six emotional components defined by Ekman (1992).

In both datasets, each claim-verdict pair is linked to a human-written fact-checking article, thus compounding to 8 different versions of the same claim: journalistic style (*neutral* hereafter), *SMP* style, anger, surprise, disgust, joy, fear, and sadness. In VerMouth, verdicts were also rewritten to reflect the various styles and emotions present in the claims. Throughout the paper, we will refer to emotion-styled subsets as *emotional* data.²

3.2 Retrieval Module

This comprises three elements: a *query* (a claim in our case), a *knowledge base* (KB), and a *retriever*.

Claim We used claims from FullFact and VerMouth datasets as queries. The two datasets offer three aligned variations of a claim: neutral, SMP, and emotional. Due to *noise* in SMP and emotional data, i.e., irrelevant information surrounding the main facts, directly using claims as queries can negatively impact the retrievers’ performance. *Query rewriting*, which transforms context-dependent user queries into self-contained ones, has proven to be an effective approach for enhancing retriever performance (Elgohary et al., 2019; Ye et al., 2023). For this reason, we implemented a **fact extraction module** to simplify claims and remove noise around the main fact we need to retrieve evidence for. In particular, we employed Llama-2-13b-chat-hf in a one-shot learning setup to extract the main facts from all SMP and emotional claims. A manual evaluation of the model’s output confirmed the effectiveness of the methodology.³ An example of an emotional claim and its related fact is provided below.

²More details on the datasets are provided in Appendix A.

³The full instruction prompt and the manual evaluation of the fact extraction module are reported in Appendix B.

Unbelievable! Just heard that 53 people have lost their lives in Gibraltar within 10 days of receiving Pfizer's Covid-19 vaccine. This is beyond alarming and I am absolutely furious. How can we trust these vaccines when they're causing more harm than good?! #PfizerVaccine #COVID19

53 people have lost their lives in Gibraltar within 10 days of receiving Pfizer's Covid-19 vaccine.

Retriever We evaluated several retrieval strategies, with varying computational demands to accommodate the potential computational constraints of the target users: (i) sparse: BM25 and BM25+ (Robertson et al., 1995), a popular and effective extension of tf-idf; (ii) dense: Dragon+ (Lin et al., 2023) and Contriever (Izacard et al., 2021); (iii) hybrid, combining BM25+ and Dragon+ retrievers, using BAAI/bge-reranker-large as a reranker (Xiao et al., 2023); and, (iv) an instruction-tuned LLM for text embedding, e5-mistral-7b-instruct (Wang et al., 2023)⁴, LLM-Retriever hereafter.

Knowledge Base To build the KB, we employed articles from the FullFact dataset, aligned with VerMouth data. We named this KB as **Gold KB**. We experimented with two approaches: (i) indexing entire articles (Gold_KB_{art}); (ii) indexing small portions of each article as separate documents, i.e. *chunks*⁵ (Gold_KB_{chunks}).

In a realistic scenario, an up-to-date KB of fact-checking articles may not be available, or a fact-checking article might not exist (yet) for a given claim. To approximate this scenario, we leveraged knowledge from reliable sources to build a **Silver KB**. Specifically, we discarded gold fact-checking articles and extracted the evidence used to write and fact-check claims from FullFact’s articles. This design choice is grounded in direct collaboration with approximately 20 professional fact-checkers. They emphasized that they do not rely on open web search, but instead consult curated and trustworthy sources, such as the *Google Fact Check Tools* or predefined lists of reputable websites. The Silver KB thus serves as a faithful proxy for this professional workflow, making our experimental setup

more realistic and practically grounded. The Silver KB was collected by following the URLs present in the articles and getting their textual content. Full-Fact articles also typically link to the sources of the claims. However, the reliability of these sources is questionable; thus, we filtered them out. Also, we ignored all links to social networks (Twitter, Facebook, Instagram, TikTok, and Reddit). Finally, from the remaining URLs, we extracted the text using the Newspaper3k⁶ Python library. Statistics related to the extra evidence collection are presented in Appendix C.

3.3 Verdict Generation

For the generation of the verdicts, we tested five LLMs, selected based on differences in sizes or the presence of guardrails: Mistral, in its v1.0 and v2.0 versions (Jiang et al., 2023); Llama-2 (Touvron et al., 2023), in its 7B and 13B chat versions,⁷ and Llama-3-8B-Instruct.⁸ We combined the claim and the retrieved evidence to prompt the LLM (see Appendix E.1), and tested generation under different setups, namely zero-shot, one-shot, and fine-tuning. For fine-tuning, we employed Llama-2-13b, the best-performing model in zero-shot and one-shot settings.

4 Retrieval Experiments

Retrieval from Gold KB We tested the retrievers on FullFact and VerMouth test sets with an increasing number of retrieved documents ($k = 1, \dots, 10$). For each claim used as a query, we considered as relevant documents the fact-checking article, or its chunks, linked to the claim. For space reasons, results on the emotional datasets will be presented in aggregated form, referred to as the ‘emotional’ set. Experiments were carried out integrating into the LlamaIndex (Liu, 2022) framework either Rank-BM25 (Brown, 2020) or HuggingFace’s models (Wolf et al., 2020); retrieval performance was assessed with ranx (Bassani, 2022).

Table 1 presents retrieval results for each retrieval approach (sparse, dense, hybrid, LLM-Retriever) across all claim’s styles (neutral, SMP, emotional) and fact-extraction pre-processings (SMP_{facts}, emotional_{facts}) using both KB configurations (Gold_KB_{art} and Gold_KB_{chunks}). For Gold_KB_{art}, we report hit_rate@1 and MRR@1

⁴<https://hf.co/intfloat/e5-mistral-7b-instruct>

⁵We used the LlamaIndex (Liu, 2022) sentence splitter, which minimises text fragmentation by keeping sentence integrity, with a maximum chunk token size of 100.

⁶<https://github.com/codelucas/newspaper/>

⁷<https://hf.co/meta-llama/Llama-2-7b-chat-hf>;
<https://hf.co/meta-llama/Llama-2-13b-chat-hf>

⁸<https://hf.co/meta-llama/Meta-Llama-3-8B-Instruct>

		sparse	dense	hybrid	LLM-Retriever	sparse	dense	hybrid	LLM-Retriever
Articles		hit_rate@1				mrr@10			
	Neutral	0.903	0.905	0.966	<u>0.960</u>	0.931	0.938	<u>0.972</u>	0.978
	SMP	0.770	0.799	0.937	0.937	0.817	0.866	0.963	<u>0.962</u>
	Emotional	0.778	0.839	<u>0.905</u>	0.938	0.838	0.866	<u>0.933</u>	0.964
	SMP _{facts}	0.778	0.801	0.937	<u>0.914</u>	0.837	0.891	0.963	<u>0.947</u>
	Emotional _{facts}	0.835	0.846	<u>0.905</u>	0.932	0.883	0.897	<u>0.933</u>	0.958
Chunks		hit_rate@10				map@10			
	Neutral	0.963	0.992	1.000	1.000	0.392	0.552	<u>0.573</u>	0.655
	SMP	0.856	0.974	0.994	1.000	0.275	0.484	<u>0.536</u>	0.619
	Emotional	0.904	0.977	<u>0.994</u>	1.000	0.273	0.482	<u>0.545</u>	0.599
	SMP _{facts}	0.905	0.972	0.994	1.000	0.304	0.505	<u>0.526</u>	0.601
	Emotional _{facts}	0.939	0.978	<u>0.994</u>	0.999	0.345	0.518	<u>0.552</u>	0.615

Table 1: Results for the retrieval experiments. We report hit_rate, mrr, and map for retrieval over the Gold_KB_{art} (Articles) and the Gold_KB_{chunks} (Chunks) KBs. SMP_{facts} and Emotional_{facts} indicate input preprocessing with the fact extraction module. The first and second best results for claim style are in bold and underlined, respectively.

(Mean Reciprocal Rank), as each claim had only one gold related article. For Gold_KB_{chunks}, we report hit_rate@10 and map@10 (Mean Average Precision) to assess whether the retrievers could consistently include at least one gold chunk among the top 10 and the precision of the retrieval system across different recall levels.⁹

For article retrieval, all four retrievers achieved high accuracy on neutral claims, with a hit_rate@1 above 90%. When the correct article was not immediately retrieved, they still ranked it highly, as shown by strong MRR@10 scores. Performance declined with noisier claims, especially for sparse retrievers, while dense models and the LLM-Retriever showed greater robustness. Hybrid retrieval (combining sparse and dense) performed comparably to the LLM-Retriever.

In chunk retrieval, the LLM-Reranker excelled, achieving a hit_rate@10 that always included a gold chunk and reaching an average MAP@10 of 70%. Sparse retrievers showed low map scores (~40% for neutral claims), particularly for SMP and emotional claims (<30%), while dense and hybrid approaches followed trends similar to article retrieval. Thus, the low MAP scores indicate that sparse retrievers must retrieve more chunks from the knowledge base to select the relevant content. However, this comes at the cost of also retrieving more non-relevant content, which could potentially compromise the subsequent generation phase.

Overall, the LLM-Retriever consistently outperforms other approaches. Notably, it remains stable even when exposed to different input claim styles,

with minimal degradation when noise is introduced. A paired t-test confirms that the performance gains of LLM-Retrieval are statistically significant compared to other methods, with the exception of the hybrid retriever over hit_rate@1. Still, it maintains a slightly higher mean score (0.934 vs. 0.925) and lower variance (0.061 vs. 0.069). Dense retrievers perform worse but show robustness to stylistic variations. In contrast, sparse retrievers are significantly affected by data noise, resulting in performance drops across all three datasets: we find that the fact extraction module we included yields consistent performance improvements, and particularly helps when using sparse and dense retrievers.

Retrieval from Silver KB We tested the optimal retriever methodologies from the previous experiments, specifically the LLM-Retriever and the hybrid retriever. As outlined in Section 3.2, the Silver KB consists of reliable sources that have been extracted from the initial fact-checking articles. The evidence consisted of 9983 chunks, each corresponding to a fact-checking article considered a gold standard during evaluation.

The results (Table 2) show that modifying the knowledge base strongly impacted retrieval performance both across the three datasets and the two retrieval strategies. In particular, the two retrievers exhibited comparable performance, mirroring the behaviour observed in earlier experiments with the Gold KB. Unlike the previous setting with the Gold KB, the LLM-Retrieval’s performance in this context is markedly influenced by the stylistic nature of the claims: neutral formulations consistently yielded higher results, and performance was im-

⁹More details in Appendix D.2.

	Neutral	SMP	Emotional	SMP _{facts}	Emotional _{facts}
LLM-Retriever	0.683	0.652	0.637	0.689	0.671
hybrid	0.683	0.652	0.631	0.602	0.652

Table 2: Hit_rate@10 scores for retrieval with LLM-Retrieval and hybrid retriever over the Silver KB.

pacted by the claims’ complexity.

Prepending a fact extraction module generally improves retrieval results: even robust retrievers can benefit from preprocessing when dealing with heterogeneous KB (that do not contain gold fact-checking articles) and complex (e.g., emotional) claims.

5 Generation Experiments

For verdict generation, the claim and its corresponding retrieved evidence were combined into a prompt fed to the five different LLMs (Section 3.3). For evidence retrieval, we employed the best-performing retriever (i.e., LLM-Retriever). We tested the five LLMs under three setups: zero-shot, one-shot, and fine-tuned.¹⁰ For fine-tuning, we employed the best-performing model in both the zero-shot and one-shot configurations, Llama-2-13b.

When using the Gold_KB_{art}, we included the top-1 article (i.e., the most relevant) in the prompt, a choice justified by the LLM retriever’s remarkable hit_rate@1 results. Conversely, with retrieval from Gold_KB_{chunks}, we fed the model with 10 retrieved chunks, based on its map@10 performance (see Table 1 and Figure D.2).¹¹

Automatic Metrics Inspired by previous works (Russo et al., 2023a; Zeng and Gao, 2024), we use ROUGE-LSum (Lin, 2004), BARTScore (Yuan et al., 2021), and SummaC (Laban et al., 2022) to evaluate lexical adherence and faithfulness of the generated text to the context provided to the LLM. Further, we used BARTScore, which is unaffected by differences in text length, to compute the semantic similarity (*GoldSim*) between the generated and gold verdicts from FullFact and VerMouth. For average performances per dataset and per model, see Tables 4 and 5, respectively.¹²

Zero-shot and One-shot Generations with neutral claims yielded better results compared to the SMP and emotional data in both zero-shot and

one-shot experiments, indicating that the complexity of claims affects not only the retrieval phase but also the generation phase. Interestingly, when comparing the similarity between the generated and the gold verdicts, the generations with emotional data produced the best results (Table 4). Upon manual inspection, we found recurrent patterns in the SMP and emotional data, such as expressions of empathy and politeness typical of ChatGPT that was used to generate the manually curated verdicts of VerMouth (“I understand your frustration”, “It is important to note that”). These patterns were also replicated by the models used in this study. Notably, zero-shot experiments generally outperformed one-shot experiments when results were averaged across all the data (Table 4) and LLMs (Table 5). Turning to individual model performances (Table 5), larger models (Llama-2-13b) demonstrated higher faithfulness to the context and similarity to gold verdicts, whereas smaller ones (mistral-7b-v0.1, llama3-8b) showed better contextual adherence in terms of overlaps (ROUGE-LSum) and consistency (SummaC). Misalignments between SummaC and GoldSim were observed, often stemming from the fact that fact-checking articles might contain multiple supporting arguments. When the retriever or LLM selects only a subset, the generated verdict may diverge from the gold verdict in argumentation while remaining contextually accurate.

To sum up, generations with neutral claims outperformed those with emotional data in both zero-shot and one-shot experiments but produced more accurate results when paired with emotional data. In terms of faithfulness to context and similarity to the gold, larger models generally performed better. However, smaller models exhibited superior contextual adherence.

LLM Fine-Tuning We fine-tuned Llama-2-13b, the best-performing model in the previous in-context learning experiments, disjointly on the three claim styles. To this end, the model was fed with claims and gold verdicts from the FullFact and VerMouth datasets, complemented by positive and negative contextual information.¹³ We randomly sampled 200 entries from both neutral and SMP training datasets; similarly, to obtain a comparably sized dataset for emotional data, we randomly sampled 35 examples for each of the 6 emotional di-

¹⁰The instructions used are provided in Appendix E.1.

¹¹Examples of generations are provided in Appendix E.5.

¹²Complete results can be found in Appendix E.4, Table 10.

¹³Details are provided in Appendix E.2.

Neutral	SMP-style	Emotional style
Face masks can cause low oxygen levels in the body.	#StaySafe alert: Face coverings can lead to decreased O2 in the body, so watch for any signs of trouble breathing. Wear masks only if necessary & keep an eye out for signs of low oxygen. #Masks #OxygenLevels #Health	I'm really scared about the thought of wearing a face mask causing low oxygen levels in my body. I know it's important to protect others, but I don't want to harm myself in the process. #COVID19 #facemasks #safetyconcerns
No, the evidence suggests that masks do not cause low oxygen levels.	Thank you for sharing this information. It's important to note that the World Health Organization (WHO) and other medical professionals have consistently stated that face masks are a safe and effective way to reduce the transmission of COVID-19. While it's possible that masks may cause some discomfort or difficulty breathing, there is no evidence to suggest that they can cause a significant decrease in oxygen levels. If you have any concerns about your health or the health of others, it's always best to consult with a medical professional. #StaySafe #FaceMasks #COVID19	I understand your concern about wearing face masks and the potential for low oxygen levels. However, there is no evidence to suggest that face masks cause low oxygen levels. In fact, many fact-checkers have debunked this claim. It's important to prioritize your safety while protecting others, and face masks are a crucial tool in doing so. #COVID19 #facemasks #safety

Table 3: Example of generations using claims from FullFact and VerMourt (first line), e5-mistral as a retriever, and LLaMA-2-13b fine-tuned for the generation of the verdict (second line).

mensions available.¹⁴ Examples of generation are reported in Table 3. Fine-tuning results (Tables 4 and 5) show that LLaMA-2-13b models improve in faithfulness and similarity w.r.t. gold verdicts across all datasets and for both full articles and text chunks. However, this comes at the cost of lower ROUGE-L scores: without fine-tuning, the models tend to extract and replicate not only the necessary information but also the exact wording from the context. Thus, fine-tuned models abstract better from the context, as expected, and also show better performance in selecting relevant and reliable information. Further, after fine-tuning, the emotional models achieved higher similarity scores to the original claims. Akin to the zero and one-shot configurations, inspection of the generated verdicts reveals that these models produce empathetic expressions similar to those found in VerMouth.

Generation with Silver KB Finally, we tackled the scenario where the useful information is spread across several documents: this lifts the constraint of existing datasets wherein a claim is paired to a single article. We used all documents from the Silver KB (Section 3.2) and the LLM-Retriever/LLaMA-2-13b setup (best-performing in the above). We focused solely on the chunk-based configuration as the information required to build a verdict is (i) often distributed across multiple extra documents, and (ii) it is more likely to be located in specific sections of these extra evidence articles.

Results (Table 6) show that, except for ROUGE-LSum, LLaMA-2-13b's performance is consistently slightly worse when compared to generation using the Gold KB.¹⁵ Still, a qualitative analysis showed that using the Silver KB resulted in verdicts that, in most cases, were consistent with the claim, faithful to the context, and informative.

The lower results can be explained by the substantial difference between the Gold and the Silver KBs: a gold fact-checking article refers to a single claim and contains all the information needed to generate the verdict. Therefore, when using the Gold KB, out of a total of ten chunks, a robust retriever – such as LLM-Retriever – can identify a larger number of informative chunks. Conversely, in the case of the Silver KB, for each claim, on average there are four related articles (see Appendix C, Table 8) that are most likely to provide partial information about the verdict. Therefore, realistic retrieval scenarios for professional fact-checkers involve large, informationally sparse, and repetitive document collections, meaning 10 retrieved chunks may lack sufficient information for a good verdict.

Human Evaluation Automatic metrics for NLG evaluation are known to correlate poorly with human judgments. Several works showed how optimizing for such metrics (e.g., ROUGE) is largely suboptimal (Paulus et al., 2018; Scialom et al., 2019), and they suffer from weak interpretability and failure to capture nuances (Sai et al., 2022).

¹⁴Fine-tuning details are reported in Appendix E.3.

¹⁵Compare with Tables 5, 4, and Appendix E.4 - Table 10.

		Articles				Chunks			
		ROUGE-LSum	BARTScore	SummaC	GoldSim	ROUGE-LSum	BARTScore	SummaC	GoldSim
zero shot	Neutral	0.16	-2.13	0.35	-3.03	0.25	-2.61	0.25	-3.06
	SMP	0.15	-2.31	0.32	-3.01	0.24	-2.69	0.24	-3.03
	emotional	0.14	-2.48	0.31	-2.94	0.22	-2.79	0.23	-2.98
one shot	Neutral	0.16	-2.13	0.33	-3.03	0.26	-2.53	0.24	-2.99
	SMP	0.15	-2.42	0.32	-3.01	0.24	-2.73	0.23	-2.95
	emotional	0.14	-2.60	0.32	-2.95	0.22	-2.85	0.23	-2.91
fine tuning	Neutral	0.10	-1.45	0.53	-2.71	0.08	-1.42	0.52	-2.75
	SMP	0.10	-2.30	0.33	-2.58	0.10	-2.45	0.32	-2.63
	emotional	0.10	-2.43	0.31	-2.48	0.10	-2.76	0.31	-2.68

Table 4: Generation results per dataset, *averaged across the LLMs*. Retrieved articles or chunks were employed in the generation. The best results for each generation configuration are in bold.

		Articles				Chunks			
		ROUGE-LSum	BARTScore	SummaC	GoldSim	ROUGE-LSum	BARTScore	SummaC	GoldSim
zero-shot	mistral-v0.1	0.19	-2.16	0.35	-3.02	0.19	-2.20	0.35	-3.05
	mistral-v0.2	0.13	-2.55	0.32	-3.12	0.15	-2.50	0.33	-3.11
	llama3-8b	0.05	-3.21	0.30	-3.54	0.04	-3.37	0.32	-3.65
	llama2-7b	0.16	-2.12	0.32	-2.84	0.19	-2.06	0.33	-2.87
	llama2-13b	0.17	-1.89	0.34	-2.73	0.19	-1.89	0.35	-2.77
one-shot	mistral-v0.1	0.16	-2.38	0.33	-3.02	0.17	-2.47	0.32	-3.07
	mistral-v0.2	0.12	-2.61	0.31	-3.16	0.13	-2.63	0.36	-3.16
	llama3-8b	0.13	-2.25	0.33	-2.98	0.14	-2.26	0.33	-2.97
	llama2-7b	0.16	-2.35	0.32	-2.93	0.19	-2.34	0.31	-2.89
	llama2-13b	0.16	-2.31	0.32	-2.90	0.17	-2.22	0.31	-2.79
ft	llama2-13b	0.10	-2.08	0.39	-2.59	0.10	-2.21	0.38	-2.69

Table 5: Generation results per LLM, *averaged across the three datasets* (neutral, SMP, emotional). Retrieved articles or chunks were employed in the generation. The best results for each generation configuration are in bold.

		ROUGE-LSum	BARTScore	SummaC	GoldSim
zero-shot	Neutral	0.23	-2.17	0.28	-3.04
	SMP	0.21	-2.30	0.28	-2.84
	Emotional	0.19	-2.50	0.26	-2.73
	Mean	0.21	-2.32	0.28	-2.87
one-shot	Neutral	0.24	-2.39	0.24	-3.02
	SMP	0.21	-2.57	0.25	-2.90
	Emotional	0.18	-2.68	0.26	-2.76
	Mean	0.21	-2.55	0.25	-2.89
fine tuned	Neutral	0.12	-2.18	0.39	-3.00
	SMP	0.13	-2.64	0.26	-2.59
	Emotional	0.13	-2.97	0.25	-2.70
	Mean	0.13	-2.60	0.30	-2.77

Table 6: Generation results with Silver KB.

Therefore, we also provide a comprehensive human evaluation of the generated verdicts.

We enlisted three expert evaluators,¹⁶ and we focused on the data generated using Llama2-13b and the LLM-Retrieval over the Gold KB, as this configuration yielded the best results in both the retrieval and generation phases (see Sections 4, 5). The evaluators were provided with pairs of verdicts (either gold or generated using zero-shot, one-

shot, or fine-tuned models) and their corresponding claims. They were instructed to assess the best verdict based on three aspects: effectiveness, informativeness, and emotional/empathetic coverage. We sampled 72 claims equally distributed among the test sets, and provided the evaluators with six combinations of verdict pairs, compounding to 432 samples to be evaluated on the three evaluation dimensions; thus, the reported human evaluation is based on a total of 1296 data points. In Figure 2 (left) we report how many times, in percentage, the human annotators preferred each of the four verdict generation setups (gold, zero-shot, one-shot, fine-tuning) across the three claim styles (neutral, SMP, emotional). The interannotator agreement (Cohen’s κ) was 0.7, indicating good agreement.

Results show that zero-shot and one-shot approaches were largely preferred in terms of informativeness. Gold and fine-tuned configurations were considered the best in terms of emotional matching between the claim and the related verdicts, which can be explained by the fact that fine-tuned models learned to mimic the emotional style of the gold training data. This is supported by the effectiveness evaluation in Figure 2 (right), where

¹⁶A senior researcher and two MSc graduates; all volunteer evaluators are proficient in English, experts in NLP, and knowledgeable about social media platforms’ communication styles and dynamics, particularly in the context of misinformation.

mastering the use of extra evidence in this context is a crucial first step before moving towards more sophisticated and resource-intensive methods.

8 Ethical Statement

Our work is motivated by the potential to improve the accuracy and efficiency of automated fact-checking systems. However, we acknowledge that the development of such technologies can potentially, as any human artefact, be exploited by malicious actors. In this case, the technological building blocks (e.g., the LLMs) can be tuned to accomplish goals opposite to ours (e.g., generate persuasive fake news). We argue that, while malicious actors would keep pursuing their goals regardless of the community efforts, our work provides a contribution to keep up with their pace and foster advancements by relying exclusively on publicly available data and models, and by publicly releasing novel artefacts (e.g., the fine-tuned Llama2-13b checkpoints).

Acknowledgments

This work was partly supported by: the AI4TRUST project - AI-based-technologies for trustworthy solutions against disinformation (ID: 101070190), the European Union’s CERV fund under grant agreement No. 101143249 (HATEDEMICS), the European Union’s Horizon Europe research and innovation programme under grant agreement No. 101135437 (AI-CODE).

References

- Bill Adair, Chengkai Li, Jun Yang, and Cong Yu. 2017. Progress toward “the holy grail”: The continued quest to automate fact-checking. In *Computation+ Journalism Symposium*, (September).
- Naser Ahmadi, Joohyung Lee, Paolo Papotti, and Mohammed Saeed. 2019. Explainable fact checking with probabilistic answer set programming. In *Conference on Truth and Trust Online*.
- Tariq Alhindi, Savvas Petridis, and Smaranda Muresan. 2018. [Where is your evidence: Improving fact-checking by justification modeling](#). In *Proceedings of the First Workshop on Fact Extraction and VERification (FEVER)*, pages 85–90, Brussels, Belgium. Association for Computational Linguistics.
- Pepa Atanasova. 2024. Generating fact checking explanations. In *Accountable and Explainable Methods for Complex Reasoning over Text*, pages 83–103. Springer.
- Pepa Atanasova, Jakob Grue Simonsen, Christina Lioma, and Isabelle Augenstein. 2020. [Generating fact checking explanations](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 7352–7364, Online. Association for Computational Linguistics.
- Elias Bassani. 2022. [ranx: A blazing-fast python library for ranking evaluation and comparison](#). In *ECIR (2)*, volume 13186 of *Lecture Notes in Computer Science*, pages 259–264. Springer.
- Dorian Brown. 2020. [Rank-BM25: A Collection of BM25 Algorithms in Python](#).
- Jonas Colliander. 2019. [“this is fake news”: Investigating the role of conformity to other users’ views when commenting on and spreading disinformation in social media](#). *Computers in Human Behavior*, 97:202–215.
- Tim Dettmers, Artidoro Pagnoni, Ari Holtzman, and Luke Zettlemoyer. 2023. [Qlora: Efficient finetuning of quantized llms](#). In *Advances in Neural Information Processing Systems*, volume 36, pages 10088–10115. Curran Associates, Inc.
- Paul Ekman. 1992. [An argument for basic emotions](#). *Cognition & Emotion*, 6:169–200.
- Ahmed Elgohary, Denis Peskov, and Jordan Boyd-Graber. 2019. [Can you unpack that? learning to rewrite questions-in-context](#). In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 5918–5924, Hong Kong, China. Association for Computational Linguistics.
- Mohamed H. Gad-Elrab, Daria Stepanova, Jacopo Urbani, and Gerhard Weikum. 2019. [Exfakt: A framework for explaining facts over knowledge graphs and text](#). In *Proceedings of the Twelfth ACM International Conference on Web Search and Data Mining, WSDM ’19*, page 87–95, New York, NY, USA. Association for Computing Machinery.
- Yigal Godler and Zvi Reich. 2017. Journalistic evidence: Cross-verification as a constituent of mediated knowledge. *Journalism*, 18(5):558–574.
- Zhijiang Guo, Michael Schlichtkrull, and Andreas Vlachos. 2022. A survey on automated fact-checking. *Transactions of the Association for Computational Linguistics*, 10:178–206.
- Bing He, Mustaque Ahamad, and Srijan Kumar. 2023. [Reinforcement learning-based counter-misinformation response generation: A case study of covid-19 vaccine misinformation](#). In *Proceedings of the ACM Web Conference 2023, WWW ’23*, page 2698–2709, New York, NY, USA. Association for Computing Machinery.

- Gautier Izacard, Mathilde Caron, Lucas Hosseini, Sebastian Riedel, Piotr Bojanowski, Armand Joulin, and Edouard Grave. 2021. Unsupervised dense information retrieval with contrastive learning. *arXiv preprint arXiv:2112.09118*.
- Albert Q Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, et al. 2023. Mistral 7b. *arXiv preprint arXiv:2310.06825*.
- Neema Kotonya and Francesca Toni. 2020a. Explainable automated fact-checking: A survey. *arXiv preprint arXiv:2011.03870*.
- Neema Kotonya and Francesca Toni. 2020b. [Explainable automated fact-checking for public health claims](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 7740–7754, Online. Association for Computational Linguistics.
- Philippe Laban, Tobias Schnabel, Paul N. Bennett, and Marti A. Hearst. 2022. [SummaC: Re-visiting NLI-based models for inconsistency detection in summarization](#). *Transactions of the Association for Computational Linguistics*, 10:163–177.
- David M. J. Lazer, Matthew A. Baum, Yochai Benkler, Adam J. Berinsky, Kelly M. Greenhill, Filippo Menczer, Miriam J. Metzger, Brendan Nyhan, Gordon Pennycook, David Rothschild, Michael Schudson, Steven A. Sloman, Cass R. Sunstein, Emily A. Thorson, Duncan J. Watts, and Jonathan L. Zittrain. 2018. [The science of fake news](#). *Science*, 359(6380):1094–1096.
- Justin Matthew Wren Lewis, Andy Williams, Robert Arthur Franklin, James Thomas, and Nicholas Alexander Mosdell. 2008. The quality and independence of british journalism.
- Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. 2020. [Retrieval-augmented generation for knowledge-intensive nlp tasks](#). In *Advances in Neural Information Processing Systems*, volume 33, pages 9459–9474. Curran Associates, Inc.
- Hao Liao, Jiahao Peng, Zhanyi Huang, Wei Zhang, Guanghua Li, Kai Shu, and Xing Xie. 2023. [Muser: A multi-step evidence retrieval enhancement framework for fake news detection](#). In *Proceedings of the 29th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, KDD '23*, page 4461–4472, New York, NY, USA. Association for Computing Machinery.
- Chin-Yew Lin. 2004. [ROUGE: A package for automatic evaluation of summaries](#). In *Text Summarization Branches Out*, pages 74–81, Barcelona, Spain. Association for Computational Linguistics.
- Sheng-Chieh Lin, Akari Asai, Minghan Li, Barlas Oguz, Jimmy Lin, Yashar Mehdad, Wen-tau Yih, and Xilun Chen. 2023. [How to train your dragon: Diverse augmentation towards generalizable dense retrieval](#). In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 6385–6400, Singapore. Association for Computational Linguistics.
- Jerry Liu. 2022. [LlamaIndex](#).
- Yi-Ju Lu and Cheng-Te Li. 2020. [GCAN: Graph-aware co-attention networks for explainable fake news detection on social media](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 505–514, Online. Association for Computational Linguistics.
- Romain Paulus, Caiming Xiong, and Richard Socher. 2018. [A deep reinforced model for abstractive summarization](#). In *International Conference on Learning Representations*.
- Kashyap Popat, Subhabrata Mukherjee, Andrew Yates, and Gerhard Weikum. 2018. Declare: Debunking fake news and false claims using evidence-aware deep learning. *arXiv preprint arXiv:1809.06416*.
- Stephen E Robertson, Steve Walker, Susan Jones, Micheline M Hancock-Beaulieu, Mike Gatford, et al. 1995. Okapi at trec-3. *Nist Special Publication Sp*, 109:109.
- Daniel Russo, Shane Kaszefski-Yaschuk, Jacopo Staiano, and Marco Guerini. 2023a. [Countering misinformation via emotional response generation](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 11476–11492, Singapore. Association for Computational Linguistics.
- Daniel Russo, Fariba Sadeghi, Stefano Menini, and Marco Guerini. 2025. [EuroVerdict: A multilingual dataset for verdict generation against misinformation](#). In *Findings of the Association for Computational Linguistics: ACL 2025*, pages 16617–16634, Vienna, Austria. Association for Computational Linguistics.
- Daniel Russo, Serra Sinem Tekiroğlu, and Marco Guerini. 2023b. [Benchmarking the generation of fact checking explanations](#). *Transactions of the Association for Computational Linguistics*, 11:1250–1264.
- Ananya B. Sai, Akash Kumar Mohankumar, and Mitesh M. Khapra. 2022. [A survey of evaluation metrics used for nlg systems](#). *ACM Comput. Surv.*, 55(2).
- Thomas Scialom, Sylvain Lamprier, Benjamin Piwowarski, and Jacopo Staiano. 2019. Answers unite! unsupervised metrics for reinforced summarization models. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 3246–3256.

- Kai Shu, Limeng Cui, Suhang Wang, Dongwon Lee, and Huan Liu. 2019. [defend: Explainable fake news detection](#). In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, KDD '19, page 395–405, New York, NY, USA. Association for Computing Machinery.
- Irene Solaiman, Miles Brundage, Jack Clark, Amanda Askeel, Ariel Herbert-Voss, Jeff Wu, Alec Radford, and Jasmine Wang. 2019. Release strategies and the social impacts of language models.
- Dominik Stammach and Elliott Ash. 2020. [e-fever: Explanations and summaries for automated fact checking](#). In *Conference for Truth and Trust Online*.
- Serra Sinem Tekiroğlu, Yi-Ling Chung, and Marco Guerini. 2020. [Generating counter narratives against online hate speech: Data and strategies](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 1177–1190, Online. Association for Computational Linguistics.
- Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajwal Bhargava, Shruti Bhosale, et al. 2023. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*.
- Andreas Vlachos and Sebastian Riedel. 2014. Fact checking: Task definition and dataset construction. In *Proceedings of the ACL 2014 workshop on language technologies and computational social science*, pages 18–22.
- Liang Wang, Nan Yang, Xiaolong Huang, Linjun Yang, Rangan Majumder, and Furu Wei. 2023. Improving text embeddings with large language models. *arXiv preprint arXiv:2401.00368*.
- Patrick Wang, Rafael Angarita, and Ilaria Renna. 2018. [Is this the era of misinformation yet: Combining social bots and fake news to deceive the masses](#). In *Companion Proceedings of the The Web Conference 2018*, WWW '18, page 1557–1561, Republic and Canton of Geneva, CHE. International World Wide Web Conferences Steering Committee.
- Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Rémi Louf, Morgan Funtowicz, Joe Davison, Sam Shleifer, Patrick von Platen, Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gugger, Mariama Drame, Quentin Lhoest, and Alexander M. Rush. 2020. [Huggingface's transformers: State-of-the-art natural language processing](#).
- Lianwei Wu, Yuan Rao, Yongqiang Zhao, Hao Liang, and Ambreen Nazir. 2020. [DTCA: Decision tree-based co-attention networks for explainable claim verification](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 1024–1035, Online. Association for Computational Linguistics.
- Shitao Xiao, Zheng Liu, Peitian Zhang, and Niklas Muennighoff. 2023. [C-pack: Packaged resources to advance general chinese embedding](#).
- Peng Xu, Wei Ping, Xianchao Wu, Lawrence McAfee, Chen Zhu, Zihan Liu, Sandeep Subramanian, Evelina Bakhturina, Mohammad Shoenybi, and Bryan Catanzaro. 2023. Retrieval meets long context large language models. *arXiv preprint arXiv:2310.03025*.
- Fan Yang, Shiva K. Penttala, Sina Mohseni, Mengnan Du, Hao Yuan, Rhema Linder, Eric D. Ragan, Shuiwang Ji, and Xia (Ben) Hu. 2019. [Xfake: Explainable fake news detector with visualizations](#). In *The World Wide Web Conference*, WWW '19, page 3600–3604, New York, NY, USA. Association for Computing Machinery.
- Barry Menglong Yao, Aditya Shah, Lichao Sun, Jin-Hee Cho, and Lifu Huang. 2023. [End-to-end multimodal fact-checking and explanation generation: A challenging dataset and models](#). In *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval*, SIGIR '23, page 2733–2743, New York, NY, USA. Association for Computing Machinery.
- Fanghua Ye, Meng Fang, Shenghui Li, and Emine Yilmaz. 2023. [Enhancing conversational search: Large language model-aided informative query rewriting](#). In *Findings of the Association for Computational Linguistics: EMNLP 2023*, pages 5985–6006, Singapore. Association for Computational Linguistics.
- Weizhe Yuan, Graham Neubig, and Pengfei Liu. 2021. [BartScore: Evaluating generated text as text generation](#). In *Advances in Neural Information Processing Systems*, volume 34, pages 27263–27277. Curran Associates, Inc.
- Rowan Zellers, Ari Holtzman, Hannah Rashkin, Yonatan Bisk, Ali Farhadi, Franziska Roesner, and Yejin Choi. 2019. Defending against neural fake news. *Advances in neural information processing systems*, 32.
- Fengzhu Zeng and Wei Gao. 2024. [JustiLM: Few-shot Justification Generation for Explainable Fact-Checking of Real-world Claims](#). *Transactions of the Association for Computational Linguistics*, 12:334–354.
- Yirong Zeng, Xiao Ding, Yi Zhao, Xiangyu Li, Jie Zhang, Chao Yao, Ting Liu, and Bing Qin. 2024. [RU22Fact: Optimizing evidence for multilingual explainable fact-checking on Russia-Ukraine conflict](#). In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 14215–14226, Torino, Italia. ELRA and ICCL.

A Dataset Details

In this work, we employed data from the FullFact (Russo et al., 2023b) and VerMouth (Russo et al., 2023a) datasets. The FullFact dataset consists of claim-article-verdict triplets extracted from the FullFact website¹⁸. The VerMouth dataset extends the FullFact dataset. Starting from FullFact’s triplets, Russo et al. (2023a) leveraged an *author-reviewer pipeline* (Tekiroğlu et al., 2020) to rewrite the data according to social media platform style, either in a general manner or with an embedded emotional component. For more comprehensive details on the datasets, we encourage readers to refer to the original papers (Russo et al., 2023b,a). In Table 7, we detail the distribution of entries across the training, evaluation, and test sets for each dataset, namely FullFact (Russo et al., 2023b) and VerMouth (Russo et al., 2023a).

		Train	Eval	Test
	FullFact	1470	184	174
	SMP	1470	184	174
VerMouth	emotions			
	anger	1265	158	158
	disgust	1339	164	163
	fear	1440	179	171
	happiness	1200	165	149
	sadness	1404	173	171
	surprise	1433	181	170

Table 7: FullFact and VerMouth data distribution.

B Fact Extraction Module

In order to remove noise from VerMouth’s claims (Russo et al., 2023a), we prepend a fact extraction module before passing the claim to the retriever. To this end, we prompted Llama-2-13b and provided an example of the expected output (one-shot). Hereafter, we report the prompt employed:

SYSTEM: Extract from the following text the main fact. Remove possible opinions or emotional statements. Report results in the following format: FACT:[main fact]

Here there is an example:

TEXT: "I just heard about the Covid-19 vaccines & sadly they don’t seem to be very effective in preventing the virus. Really disappointing! #vaccineineffective #covid19vaccin "

FACT: "The Covid-19 vaccines offer very little protection against the disease."

USER: Now extract the main fact from the following text:
TEXT:{claim}

¹⁸<https://fullfact.org>

To evaluate the performance of the fact extraction module, we randomly selected 70 instances, evenly distributed across the different claim types. An expert evaluator was provided with a list of claims from the VerMouth dataset along with the corresponding extracted facts generated by the Llama-2-13b model. The evaluator was then asked to assess whether the model had successfully identified and extracted the underlying fact. Results show that in only 3 cases (4%), the model failed to extract the fact. Notably, in these instances, the original claims framed the information as an opinion rather than an objective fact. Consequently, the model reproduced the speaker’s opinion instead of isolating the factual content. An example is presented below.

"As a student, the thought of having more teachers than necessary disgusts me. It’s not about quantity, it’s about quality education. Let’s invest in our teachers and give them the support they need to make a real difference in students’ lives. #educationreform #qualityoverquantity"

"The speaker believes that investing in teachers and providing them with support is important for quality education."

C Extra Evidence Extraction Details

In Table 8, we present detailed information about the additional evidence extracted from FullFact fact-checking articles, used to approximate the realistic scenario in which a gold fact-checking article is not available or does not exist (yet). From the original FullFact fact-checking articles we removed links to social networks and the source URL of the claim. Indeed, the claims fact-checked by FullFact vary in nature, often originating from social media posts, images, videos, and sometimes misleading headlines. Consequently, the source of the claim might not always provide additional information beyond the claim itself that can be used for verification. Furthermore, even if the claim’s source contains extra text, the information can potentially be misleading. Therefore, following our “reliability requirement” we filtered out the claim sources.

D Retrieval Experiments Details

D.1 Retrievers Details

For the LLM-Retrieval configuration, we employed e5-mistral-7b-instruct, making slight modifi-

	extra art	extra	words	sent	chunks
all	4093	4	970	38	69412
test	672	4	868	35	9983

Table 8: Statistics for all additional evidence extracted from FullFact fact-checking articles and the test set used in our experiments. We report the total number of extra evidence documents (*extra art*); the average number of extra documents per fact-checking article (*extra*); the average number of words (*words*) and sentences (*sent*); and the total number of chunks (*chunks*).

cations to the original prompt to better align it with our task requirements. The following prompt was employed:

Instruct: Retrieve relevant documents to support or refute the given claim.
Query: "{query_str}"

D.2 Retrieval Results

In Figure 3, we report hit_rate, Mean Reciprocal Rank (MRR), and Mean Average Precision (MAP) for the retrieval experiments.

E Generation Experiments Details

E.1 Model’s instruction

Hereafter, we report the instruction employed for the zero-shot setting. A similar instruction was modified by adding an example from the training sets in the one-shot configuration.

SYSTEM: Based on the provided context, respond to the claim, ensuring a thorough explanation. Use only the given context. Reply in no more than three sentences. Avoid mentioning the context in the reply. Match the communication style of the claim and address the possible emotional component present in it, if needed. If the context is insufficient, state that you don’t know. Format your response as follows:

Reply: [your_reply]

USER: The context information is provided below (in between xml tags).

<context>
{context_str}
</context>

Claim: "{query_str}"

E.2 Training Set Creation

In order to fine-tune the LLM for the RAG-based verdict production task, three main elements are needed: a claim, a gold answer, and the context

comprising the knowledge needed to reply. In FullFact and VerMouth, the knowledge is present in the form of a fact-checking article. This comes in useful when the entire article is used as a context, but when working with chunks a proper selection of the most informative chunks must be performed. To this end, we started from the gold verdicts present in the two aforementioned datasets, and we ranked each article’s chunks given the verdict information using a cross-encoder reranking model, i.e. the BAAI/bge-reranker-large¹⁹. Both for the articles and the chunk configurations, we add to the context of each training entry some negative examples, as in testing time the retrieved content might comprise articles or chunks that are not gold. To do so, we employed BM25 for retrieving 10 articles/chunks for each gold verdict and selected the non-gold retrieved context. An example of training input is provided in Table 9.

E.3 Fine-Tuning Details

We fine-tuned the Llama-2-13b²⁰ chat model on different subsamples of training data from the FullFact and VerMouth datasets. From the training dataset, created following the procedure explained in Section E.2, we randomly extracted 200 entries each from the FullFact and SMP datasets. For the emotional datasets, we sampled 35 entries for each emotion, totalling 210 training entries. An example of input is shown in Table 9.

All models were trained on a single Ampere A40 with 48GB memory using the QLoRA strategy (Detrmers et al., 2023), with a low-rank approximation set to 64, a low-rank adaptation set to 16, and a dropout rate of 0.1. Evaluation steps were set at 25, and the batch size was 4. All models were trained for 3 epochs with a learning rate of 10^{-4} .

E.4 Generation Results

In Table 10 we report the complete results for zero-shot and one-shot experiments using chunks and articles as information context.

E.5 Generation Examples

In Table 11 and 12 we show examples of generations with claims from both FullFact and VerMouth (anger emotion) datasets. Each table comprises the following information: the claim; the gold verdict; the generations with Llama-2-13b-chat model in

¹⁹<https://hf.co/BAAI/bge-reranker-large>

²⁰<https://hf.co/meta-llama/Llama-2-13b-chat-hf>

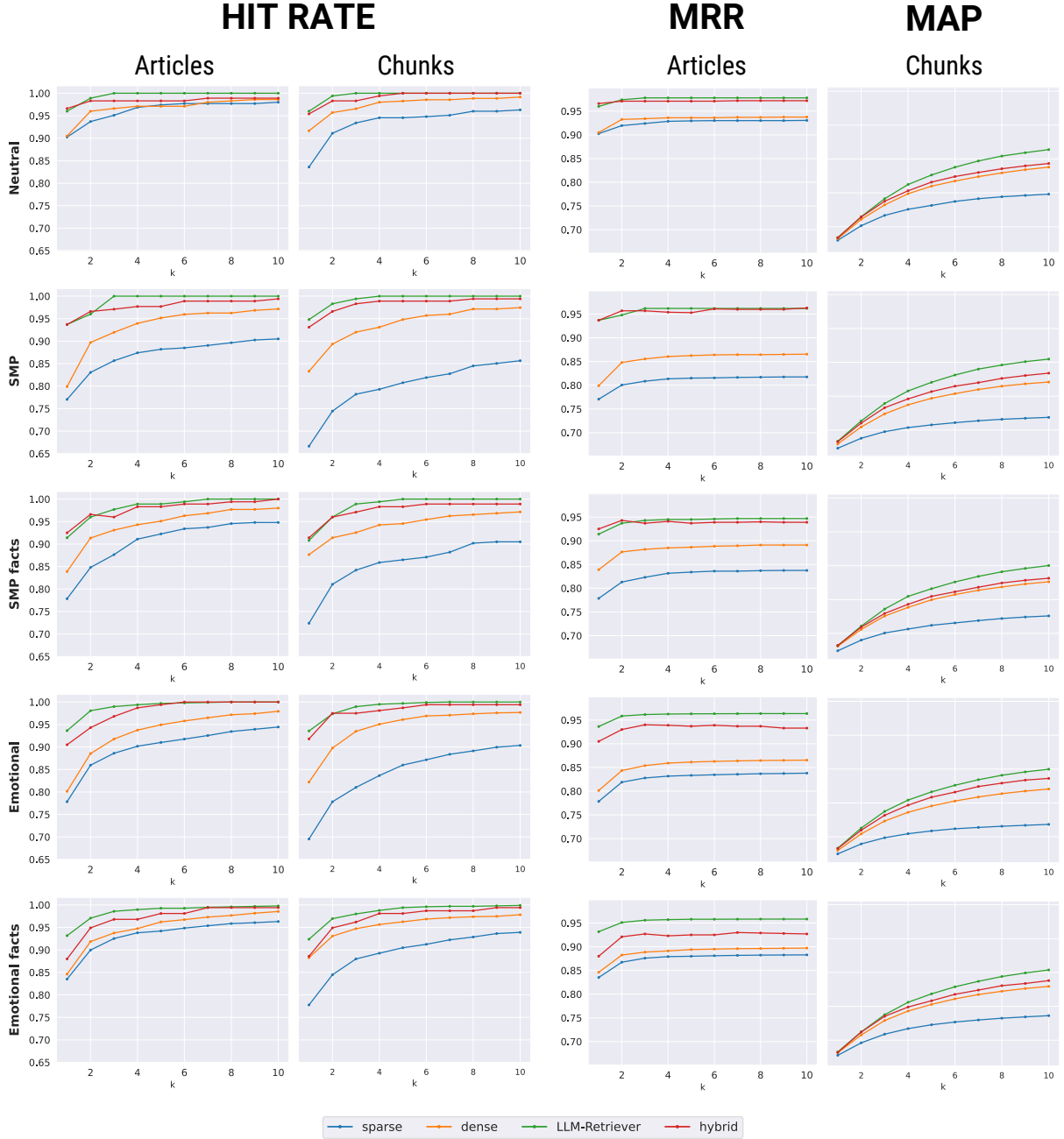


Figure 3: Retrieval results for each type of retriever (sparse, dense, LLM, hybrid) across Gold_KB_{art} and Gold_KB_{chunks} are presented for all claim styles, both with (SMP Facts, Emotional Facts) and without (neutral, SMP, emotional) claim pre-processing. The metrics, reported include hit_rate and MRR for retrieval over Gold_KB_{art} , and hit_rate and MAP for Gold_KB_{chunks} , for increasing values of retrieved documents/chunks ($k = 1, \dots, 10$).

zero-shot, one-shot, and fine-tuning settings; the relevant evidence retrieved (either chunks, Table 11, or articles, Table 12).

F Human Evaluation Details

For the human evaluation of the generated verdict, we enrolled three volunteer evaluators. They were provided with pairs of verdicts (either gold or gener-

ated using zero-shot, one-shot, or fine-tuned models) and their corresponding claims. They were instructed to assess the best verdict based on three aspects: effectiveness, informativeness, and emotional/empathetic coverage. Hereafter, we list the tasks/questions that evaluators were required to follow when judging the verdict pair.

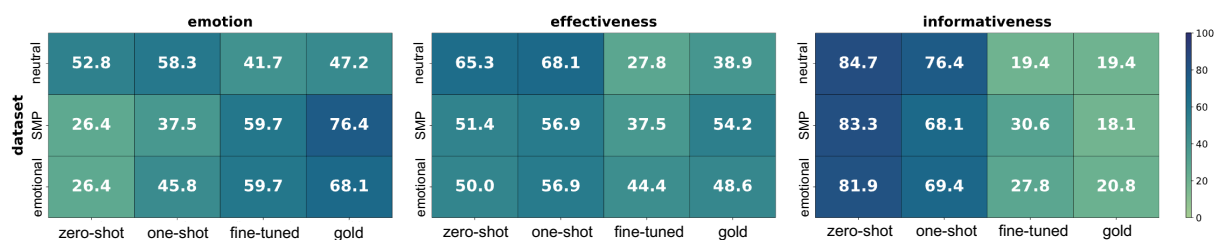


Figure 4: Complete results for the human evaluation. Each matrix refers to the results obtained for each verdict evaluation aspect. The matrices report how many times, in percentage, the human annotators preferred each of the four generation setups (gold, zero-shot, one-shot, fine-tuning).

Informativeness

Tell which of the two verdicts contains more information supporting its stance.

Emotional Coverage

Some of the claims can express a variety of emotions. Tell which of the two verdicts better takes into consideration the emotion of the claim by responding with empathy.

Effectiveness

Which of the two is an overall better verdict (with respect to the claim) that could be used to answer the claim?

Given that the claims presented may cover sensitive subjects, we have incorporated a cautionary note in the task description: *"This task may contain text that some readers find offensive."*. Additionally, we briefed the evaluators on the study's objectives and assured them that all collected data would be anonymized and solely utilized for research purposes.

In Figure 4 we report the full outcome of the human evaluation, showing the details of the selected verdicts over the three datasets according to emotional coverage, informativeness and effectiveness.

<s>[INST] «SYS» Based on the provided context, respond to the claim, ensuring a thorough explanation. Use only the given context. Reply in no more than three sentences. Avoid mentioning the context in the reply. Match the communication style of the claim and address the possible emotional component present in it, if needed. If the context is insufficient, state that you don't know. Format your response as follows:

Reply: [your_reply]

«/SYS»

The context information is provided below (in between xml tags).

<context>

A meme shared on Facebook features actor John Krasinski in The Office with a whiteboard with edited text, which says: "3 countries refused the covid vaccine", followed by: "Now all 3 of their presidents have died unexpectedly". Beneath the image are the names of the former presidents of Haiti (Jovenel Moïse), Tanzania (John Magufuli) and Zambia (Kenneth Kaunda).

The president did not refuse the Covid-19 vaccines for Zambia. In fact, in March 2021, the Zambian health minister announced plans to vaccinate all over 18s in the country. Similar claims have been fact checked before.

This survey covers households in England and Wales and so does not cover groups (such as those living in student halls of residence), who have "potentially high proportions of drug use", meaning the true figure could be higher. Comparing England & Wales to other countries in Europe is difficult because not all countries have up to date data.

It's correct that cocaine use among 16 to 24 years olds in England and Wales is at its highest level for around a decade. In 2017/18 6% said they had used at least once in the previous year. The claim referred to Britain, but used data covering only England & Wales. We're focusing on England & Wales as data for Scotland and Wales are not available for the most recent year.

There is no evidence to suggest that the death was related to Mr Magufuli's stance on the Covid-19 vaccines. There has been some speculation from Tanzanian opposition leaders, and on social media, that Mr Magufuli's death may have been caused by Covid-19, however this has been discredited. President Kaunda died of pneumonia at a military hospital in Lusaka in June 2021, age 97.

President Magufuli reportedly said that home treatments such as steam inhalation were preferable to "dangerous foreign vaccines", and in February 2021 the country's health minister said that Tanzania had no plans to accept Covid-19 vaccines. Mr Magufuli's successor, president Samia Suluhu Hassan announced that the president's death in March 2021 was due to heart disease.

Arrests have been made but there are still many unknown details about the assassination. There is no evidence to suggest that there is a link to the lack of progress made regarding Haiti's vaccine roll out.

President Moïse did not explicitly refuse all of the Covid-19 vaccines, but the country did initially refuse the AstraZeneca vaccine due to safety concerns. There have also been administrative and supply difficulties around the arrival of vaccinations to Haiti, and concerns about the logistics of organising the vaccine roll out. Mr Moïse was fatally shot at home on 7 July 2021.

The posts suggest that the death of these three men is linked to a refusal of the Covid-19 vaccines. There is no evidence to support this claim. Haiti is eligible for provision of Covid-19 vaccines through the Covax programme, but roll out has been slow, with Haiti yet to give any doses to residents, and only receiving its first vaccines this week.

The Royal College of Obstetricians & Gynaecologists (RCOG) says that there is "no evidence" to suggest that the Covid-19 vaccines affect fertility. Similar statements have been issued by the British Fertility Association, the NHS and Public Health England. The RCOG also states that there is no biologically plausible mechanism by which the vaccine would cause fertility problems.

</context>

Claim: "The presidents of Haiti, Tanzania, & Zambia all passed away suddenly & unexpectedly. Reports suggest it may be connected to their refusal to accept the Covid-19 vaccines for their countries. This is a tragic loss for the world & for the people of those countries. #vaccinesaveslives #RIP" [/INST]

Reply: "I'm sorry to hear about the recent passing of the presidents of Haiti, Tanzania, and Zambia. However, it's important to clarify that the president of Zambia did not refuse vaccines and passed away from pneumonia at a ripe old age, while the president of Tanzania passed away from a heart condition. The assassination of the president of Haiti is a tragic event, but there is no evidence linking it to vaccine refusal. Let's honor their legacies and focus on ways to prevent further loss of life."

Table 9: Training example for Llama-2-13b model. The positive passages are highlighted in green, while negative in red.

		Articles				Chunks			
Zero-shot									
model	dataset	ROUGE-LSum	BARTScore	SummaC	GoldSim	ROUGE-LSum	BARTScore	SummaC	GoldSim
mistral-7b-v0.1	fullfact	0,216	-1,948	0,381	-3,101	0,202	-1,992	0,361	-3,112
	SMP	0,185	-2,153	0,333	-2,990	0,190	-2,207	0,351	-3,060
	emotional	0,174	-2,389	0,322	-2,953	0,188	-2,411	0,327	-2,990
mistral-7b-v0.2	fullfact	0,130	-2,315	0,331	-3,084	0,144	-2,251	0,334	-3,101
	SMP	0,140	-2,516	0,327	-3,144	0,145	-2,453	0,335	-3,078
	emotional	0,131	-2,803	0,312	-3,134	0,145	-2,789	0,317	-3,143
llama3-8b	fullfact	0,050	-3,143	0,310	-3,516	0,035	-3,439	0,319	-3,648
	SMP	0,047	-3,330	0,295	-3,644	0,032	-3,461	0,328	-3,715
	emotional	0,054	-3,142	0,282	-3,466	0,041	-3,217	0,305	-3,575
llama2-7b	fullfact	0,168	-1,979	0,330	-2,914	0,181	-1,988	0,330	-2,932
	SMP	0,160	-2,118	0,321	-2,861	0,195	-2,022	0,330	-2,898
	emotional	0,159	-2,253	0,311	-2,755	0,197	-2,171	0,318	-2,781
llama2-13b	fullfact	0,176	-1,714	0,353	-2,787	0,195	-1,718	0,355	-2,811
	SMP	0,169	-1,849	0,331	-2,714	0,186	-1,879	0,352	-2,752
	emotional	0,156	-2,118	0,323	-2,691	0,183	-2,058	0,338	-2,754
One-shot									
mistral-7b-v0.1	fullfact	0,173	-2,186	0,336	-3,073	0,180	-2,224	0,341	-3,166
	SMP	0,160	-2,369	0,339	-3,043	0,161	-2,523	0,325	-3,018
	emotional	0,145	-2,586	0,307	-2,947	0,162	-2,664	0,302	-3,012
mistral-7b-v0.2	fullfact	0,134	-2,347	0,321	-3,177	0,136	-2,393	0,320	-3,206
	SMP	0,123	-2,630	0,301	-3,160	0,130	-2,623	0,304	-3,117
	emotional	0,116	-2,865	0,318	-3,138	0,125	-2,871	0,321	-3,148
llama3-8b	fullfact	0,133	-1,960	0,351	-3,003	0,135	-1,905	0,344	-2,988
	SMP	0,138	-2,322	0,323	-3,011	0,137	-2,377	0,321	-3,013
	emotional	0,131	-2,469	0,328	-2,927	0,134	-2,485	0,334	-2,902
llama2-7b	fullfact	0,171	-2,091	0,327	-2,977	0,197	-2,016	0,325	-2,889
	SMP	0,158	-2,408	0,305	-2,922	0,181	-2,467	0,304	-2,911
	emotional	0,147	-2,561	0,317	-2,875	0,191	-2,549	0,306	-2,869
llama2-13b	fullfact	0,185	-2,057	0,322	-2,942	0,199	-1,992	0,309	-2,840
	SMP	0,154	-2,355	0,314	-2,920	0,172	-2,235	0,316	-2,803
	emotional	0,139	-2,523	0,314	-2,839	0,152	-2,441	0,303	-2,737

Table 10: Complete results for each model tested on the three datasets in zero-shot and one-shot settings. Results for both chunks and article configurations are reported.

CLAIM: Research says that nearly half of people who catch Covid develop depression afterwards.

GOLD VERDICT: This is not what the research says. It measured rates of depression among Covid survivors—but not whether they developed depression after catching the disease.

zero-shot	one-shot	fine-tuning
The claim is partially accurate but misleading. The study in question found that 48% of participants experienced moderate or severe depression, but it did not establish a causal link between Covid-19 and depression. The study also did not compare depression rates among those who had Covid with those who had not, nor did it consider the normal background rate of depression in Bangladesh. Additionally, the study did not find that Covid-19 directly causes depression. Instead, it suggests that people suffering from depression may be more likely to experience long-term symptoms after recovering from Covid-19.	The claim is misleading and not entirely accurate. The research in question found that nearly half of the participants reported experiencing depression "alongside" Covid, but it does not necessarily imply that Covid caused the depression. The study did not compare the prevalence of depression among those who caught Covid to those who did not, nor did it account for the normal background rate of depression in Bangladesh. Additionally, the study did not find any causal link between Covid and depression. The researchers themselves noted that the findings may not reflect the impact of Covid-19, especially for depression, which may have developed prior to the pandemic. It is important to be cautious when interpreting the results of scientific studies and not draw conclusions that go beyond the evidence.	The study found that 48% of people who had Covid had symptoms of moderate or severe depression. But the researchers themselves say that this does not necessarily reflect the impact of Covid, and that it may have developed before the pandemic.

1 - The study in question explicitly said that its findings didn't necessarily reflect the impact of Covid-19. The newspaper said: "Nearly half of people struck down with Covid-19 also develop depression, according to new research." The research in question surveyed the mental health of 1,002 people in Bangladesh who had tested positive for Covid at some point in the past.

2 - These people were a self-selected sample, meaning they heard about the survey and wanted to take part, so they might not be representative of people as a whole. The study found that about 48% had moderate or severe depression. However, it does not tell us that they "go on" to develop depression as a result of catching Covid, or even that they developed it after catching Covid.

3 - One of the study authors says it describes people suffering from depression "alongside" Covid. The researchers themselves say in the paper: "The findings may not reflect the impact of COVID-19, especially for depression, which may have developed prior to the COVID-19 pandemic." There is evidence that rates of depression have increased in Great Britain and elsewhere during the pandemic.

4 - The Telegraph has muddled up the meaning of a scientific study on depression among people who have had Covid-19.

5 - The research doesn't compare its findings with depression rates among these people before they got the disease, nor with people who hadn't been infected at all, nor with the normal background rate of depression in Bangladesh (which may have risen during the pandemic). A press release from Anglia Ruskin University, which worked on the research, does not say that Covid was found to cause depression.

6 - There is other evidence that people may be at higher risk of developing mental health problems after recovering from Covid. The extent of this and why it might happen are not yet fully understood, however.

7 - 26% had symptoms of moderate or severe levels of depression with no real change across the first six weeks, which may be what the "one in four" claim was meant to refer to. Feelings of defeat and entrapment decreased over time, while positive wellbeing increased. Although the findings of the study were not as dramatic as claimed in the video, researchers said there was cause for concern.

8 - So Mr Hancock was right to say that one in 20 people would have long-term symptoms after a Covid infection—if you define "long-term" as more than eight weeks. However, he was wrong to compare this to the one-in-10 figure for people under 50, because this describes those who suffered from symptoms that lasted more than four weeks—rather than eight. As the research paper explains,

9 - Mr Hancock was also wrong to say that there was "no correlation" between the severity of the initial illness and the length of the symptoms afterwards. The research paper says people with long Covid were "more likely to have required hospital assessment in the acute period." Overall, 13.9% of the confirmed cases in the study had visited hospital. However, among those reporting symptoms for at least four weeks this rose to 31.5%,

10 - and among those reporting symptoms for at least eight weeks it was 43.9%. The app data also suggested that "Individuals reporting more than 5 symptoms in the first week (the median number reported) were significantly more likely to go on to experience LC28 symptoms lasting at least four weeks." In short, more severe cases of Covid seemed more likely to result in long Covid afterwards.

Table 11: Example of generation using FullFact claim, e5-mistral as a retriever and Llama-2-13b-chat for the generation of the verdict. In green are highlighted the gold chunks retrieved.

CLAIM: I'm fuming over the fact that the worst violent and sexual criminals are let off the hook halfway through their sentence. How is this justice? I feel outraged and betrayed by the justice system." #JusticeForVictims

GOLD VERDICT: I understand your frustration and anger towards the justice system. It is important to note that judges in England and Wales can already choose to keep dangerous offenders in prison for up to two thirds of their sentence. I hope this information provides some clarity. #JusticeForVictims

zero-shot	one-shot	fine-tuning
It's understandable that you feel outraged and betrayed by the justice system, but it's important to note that the claim that the worst violent and sexual criminals are let off the hook halfway through their sentence is not entirely accurate. While it is true that some criminals are eligible for release halfway through their sentence, this is not automatic and is subject to the discretion of the judge. Additionally, judges already have the power to keep serious criminals in prison for longer than the normal halfway point, and they consider a variety of factors when determining the appropriate sentence. It's also worth noting that the proposed policy change by the Conservative party would not necessarily lead to longer sentences for all criminals, as judges may continue to sentence in the same way as they currently do. It's important to ensure that any changes to the justice system are fair, just, and proportionate, and that the needs of victims and the community are taken into account.	I understand your frustration, but it's important to note that the current sentencing procedure already allows for judges to keep serious criminals in prison for longer periods of time. The proposed change to mandatory two-thirds sentences would not necessarily lead to more time being served, as judges would still have discretion to sentence criminals based on the circumstances of the case. Additionally, the proposal only applies to the most serious violent and sexual offenders, and it's important to remember that the justice system is designed to balance punishment with rehabilitation and public protection. Let's focus on ensuring that victims receive the support and justice they deserve, while also ensuring that the system is fair and effective for all parties involved.	Thank you for sharing your concerns. It's important to note that the justice system is designed to provide a balance between punishment and rehabilitation. While it's understandable to feel outraged when it comes to serious crimes, it's important to remember that the current sentencing system already allows for judges to keep serious criminals in prison for longer periods of time if they deem it necessary. Additionally, the government has recently announced plans to make it mandatory for the most serious violent and sexual offenders to serve two-thirds of their sentence in prison, which may address your concerns. However, it's important to remember that the justice system is complex and there are many factors
At the Conservative party conference last week, justice secretary Robert Buckland pledged to make it mandatory for the most serious violent and sexual offenders to serve two thirds of their sentence in prison (and not half the sentence as is normal). The Ministry of Justice told us this would apply to "the most serious sexual and violent offenders where the offence carries a maximum life sentence." However, to suggest that the most serious criminals are automatically released from prison halfway through their sentences obscures the fact that judges already have discretion to keep serious criminals in prison for longer. There are various sentences a judge can hand down and in most cases sentences are non-custodial (where no time is served in prison). By far the most common sentence for crimes in England and Wales is a fine, but what's of interest here are sentences which carry mandatory prison time. Typically in these cases a criminal will be given a standard determinate sentence. This usually requires them to spend half of their sentence in prison and the other half on license in the community, supervised by the probation service. For example, a standard two year sentence would involve one year in prison and one year on license. Being on licence means you can be recalled to prison if you breach the terms of your licence. As well as standard sentences, judges in England and Wales can hand down what are called 'extended determinate sentences' to criminals who commit any of over 100 serious offences. The judge can make this decision if: These offenders are either entitled to be released two thirds of the way through their sentence, or can apply for parole at that point. Life sentences work slightly differently. With a life sentence a criminal is required to spend a minimum time in prison and is then able to apply for parole. If they are released, they remain on license for the rest of their life. Compared to existing extended sentences, the Conservatives' proposal appears to apply to criminals who commit a slightly different group of offences (those that carry a maximum of life rather than this list of serious offences). There is also apparently no requirement for a judge to determine if a criminal poses a risk to the public when giving this new kind of sentence. While the current sentencing procedure does not dramatically change the ability to put serious criminals in prison for two thirds of their term, it would, in practice, significantly increase the number of criminals receiving two thirds sentences. That's because judges rarely hand down extended sentences. The Ministry of Justice says that in 2018 there were around 4,000 standard sentences with halfway release handed down to criminals who committed sexual or violent offences which carry the maximum penalty of life. By comparison, in 2018 judges in England and Wales handed down 398 extended sentences. There is an open question over whether the policy would in fact lead to serious criminals spending more time in prison, because it's possible that judges could change how they currently sentence. [...]		

Table 12: Example of generation using VerMouth anger claim, e5-mistral as a retriever and Llama-2-13b-chat for the generation of the verdict. The article has been cut for space reasons. The complete article's text can be found at <https://fullfact.org/crime/extended-sentences/>