

Med-V1: Small Language Models for Zero-shot and Scalable Biomedical Evidence Attribution

Qiao Jin¹, Yin Fang¹, Lauren He¹, Yifan Yang¹, Guangzhi Xiong², Zhizheng Wang¹, Nicholas Wan¹, Joey Chan¹, Donald C. Comeau¹, Robert Leaman¹, Charalampos S. Floudas³, Aidong Zhang², Michael F. Chiang^{1,4}, Yifan Peng^{5,6}, and Zhiyong Lu^{1,+}

¹Division of Intramural Research, National Library of Medicine, National Institutes of Health

²Department of Computer Science, University of Virginia

³Center for Cancer Research, National Cancer Institute, National Institutes of Health

⁴National Eye Institute, National Institutes of Health

⁵Department of Population Health Sciences, Weill Cornell Medicine

⁶Institute of AI for Digital Health, Weill Cornell Medicine

⁺Corresponding author (zhiyong.lu@nih.gov).

ABSTRACT

Assessing whether an article supports an assertion is essential for hallucination detection and claim verification. While large language models (LLMs) have the potential to automate this task, achieving strong performance requires frontier models such as GPT-5 that are prohibitively expensive to deploy at scale. To efficiently perform biomedical evidence attribution, we present Med-V1, a family of small language models with only three billion parameters. Trained on high-quality synthetic data newly developed in this study, Med-V1 substantially outperforms (+27.0% to +71.3%) its base models on five biomedical benchmarks unified into a verification format. Despite its smaller size, Med-V1 performs comparably to frontier LLMs such as GPT-5, along with high-quality explanations for its predictions. We use Med-V1 to conduct a first-of-its-kind use case study that quantifies hallucinations in LLM-generated answers under different citation instructions. Results show that the format instruction strongly affects citation validity and hallucination, with GPT-5 generating more claims but exhibiting hallucination rates similar to GPT-4o. Additionally, we present a second use case showing that Med-V1 can automatically identify high-stakes evidence misattributions in clinical practice guidelines, revealing potentially negative public health impacts that are otherwise challenging to identify at scale. Overall, Med-V1 provides an efficient and accurate lightweight alternative to frontier LLMs for practical and real-world applications in biomedical evidence attribution and verification tasks. Med-V1 is available at <https://github.com/ncbi-nlp/Med-V1>.

1 Introduction

Evidence attribution, also known as verification, is a foundational biomedical task that assesses whether an assertion is supported by a source document [1]. A common example is automatic fact-checking. Such systems usually retrieve evidence relevant to a given claim and then apply a verification model to determine the claim’s veracity based on that evidence [2–4]. The need for robust verification has grown with the increasing use of generative artificial intelligence (AI) tools such as ChatGPT, which are known to hallucinate when generating statements with citations [5–7]. As a result, their outputs must be thoroughly verified before downstream usage [8–10], especially in high-stakes domains such as healthcare [11]. While frontier large language models (LLMs) have shown strong performance across various biomedical tasks [12–19] such as question answering [20–22] and clinical trial matching [23–25], they often rely on commercial systems and are prohibitively expensive for routine, large-scale deployment. With the rapid growth of AI-generated statements, it is conceivable that verification becomes a heavily used utility task. In practice, such large-scale verification can only be realistically performed by small language models [26], which are more cost-effective and scalable. As such, the objective of this work is to develop a lightweight foundation model that can perform biomedical verification comparably with frontier LLMs, generalize well to unseen data distribution, and provide faithful explanations for its predictions.

To this end, we first introduce a novel LLM-driven pipeline that generates MedFact-Synth, a large-scale synthetic dataset comprising biomedical claims and their high-quality verification reports. Using MedFact-Synth, we subsequently fine-tune small LLMs with only three billion parameters [27, 28], resulting in Med-V1. For evaluation, we compile MedFact-Bench, a benchmark consisting of five standard and re-purposed biomedical verification datasets [29–33]. Experimental results show that Med-V1 substantially outperforms its base models, achieving improvements of 27.0% to 71.3% across the datasets. Despite its small size, Med-V1 reaches performance levels comparable to frontier LLMs such as GPT-5 [34]. Further analysis shows that

Med-V1 consistently presents high-quality explanations for its predictions.

To demonstrate the real-world utility of Med-V1, we conduct two large-scale use case studies that analyze evidence attribution in both LLM-generated answers and high-stakes clinical practice guidelines. First, we apply Med-V1 to LLM-generated answers under seven citation-instruction formats by extracting claim-citation pairs and verifying each claim against its cited article, which enables large-scale quantification of citation validity and hallucination rates across formats. This analysis shows that citation format strongly affects citation validity and hallucination, and that GPT-5 generates more claims but exhibits hallucination rates similar to its predecessor GPT-4o. Second, we apply Med-V1 to citation statements in recent clinical practice guidelines and identify a nontrivial subset whose cited sources do not support the corresponding claims; after manually validating 100 flagged cases, we find that a significant proportion of them are genuine misattributions, and many involve treatment effects that might lead to potential risks in public health.

In summary, these results show that Med-V1 achieves frontier-level verification performance while remaining compact and computationally efficient. This makes it a practical and accessible alternative to proprietary LLMs for large-scale biomedical verification, hallucination detection, and citation auditing.

2 Results

2.1 Overview of Med-V1 Training and Inference

This work contains two phases: (i) synthetic data generation and model training (Figure 1a) and (ii) Med-V1 inference for downstream tasks (Figure 1b).

Phase I: Synthetic data generation and model training. In the first phase, we construct MedFact-Synth, a large-scale synthetic training set including 1.5 million instances. Each instance contains: a synthetic claim to be verified, a source article serving as evidence, a rationale explaining the verification, and a 5-point Likert-scale verdict, ranging from strong contradiction (-2) and partial contradiction (-1) to neutral (0), partial agreement (+1), and strong agreement (+2). To build this dataset, we begin by sampling one million articles from the PubMed 2025 baseline [35]. For each article, GPT-4o-mini [36] is prompted to generate two claims: one that the article may support and one that it may refute. To collect diverse claim-article pairs for verification, we use MedCPT [37] to retrieve the top 10 most relevant PubMed articles for each claim. A panel of frontier LLMs then verifies each pair, generating both rationales and verdicts via a voting-based mechanism. Then, we utilize MedFact-Synth to fine-tune Med-V1, which is a three-billion-parameter (3B) small LLM available in two variants: Med-V1-L3B, initialized from Llama-3.2-3B-Instruct [27] and Med-V1-Q3B, initialized from Qwen2.5-3B-Instruct [28]. The quality of MedFact-Synth is examined in the next section, and technical details of the dataset construction and training process are described in the Methods section.

Phase II: Inference and downstream applications. During the inference phase, Med-V1 can be applied to various tasks where an assertion-source pair is available. Here, an assertion refers to a declarative statement whose factuality needs to be checked, and the source refers to the biomedical article or document against which the assertion is verified. Assertions may originate from various settings. For example, boolean questions can be reformulated as declarative assertions, factual claims can be collected on social media, and citation statements can be extracted from AI-generated or human-written content. When verification sources are not already provided, they can be easily retrieved using an evidence search system [38, 39]. These assertion-source pairs are then fed into Med-V1, which outputs both a structured verdict (agreement level ranging from -2 to +2) and a natural-language explanation (Figure 1b). This unified formulation allows Med-V1 to perform multiple downstream applications, including binary question answering, claim verification, and citation attribution.

In this work, we evaluate Med-V1 on standard claim verification datasets and re-purposed binary question answering datasets, under a strict zero-shot setting. We further validate the quality of Med-V1 explanation, as well as its capability to detect LLM hallucinations and identify citation misattributions in clinical guidelines.

2.2 MedFact-Synth is Large-scale and of High-quality

Figure 2a shows an example from the MedFact-Synth dataset. In this example, the claim is generated to be supported by its source article (PMID: 9035431). However, the relevant paper retrieved by MedCPT (PMID: 24935303) actually refutes the claim, as determined by the frontier LLM panel. Because claims are paired with independently retrieved articles rather than being restricted to their source papers, MedFact-Synth naturally includes a mixture of supportive, neutral, and contradictory evidence. As a result, this design leads MedFact-Synth to capture more realistic verification scenarios than settings in which claims are always aligned with their originating article.

Overall, MedFact-Synth contains 1.5 million instances with a reasonably balanced label distribution: 38.3% (strong and partial) agreement, 36.4% neutral, and 25.3% (strong and partial) contradiction (Figure 2b). This balance ensures that Med-V1 is exposed to both supportive and conflicting evidence at scale, rather than being dominated by trivially supported claims. The average lengths are 20.7 words for claims, 132.8 words for rationales, and 271.3 words for source articles (Figure 2c). Thus,

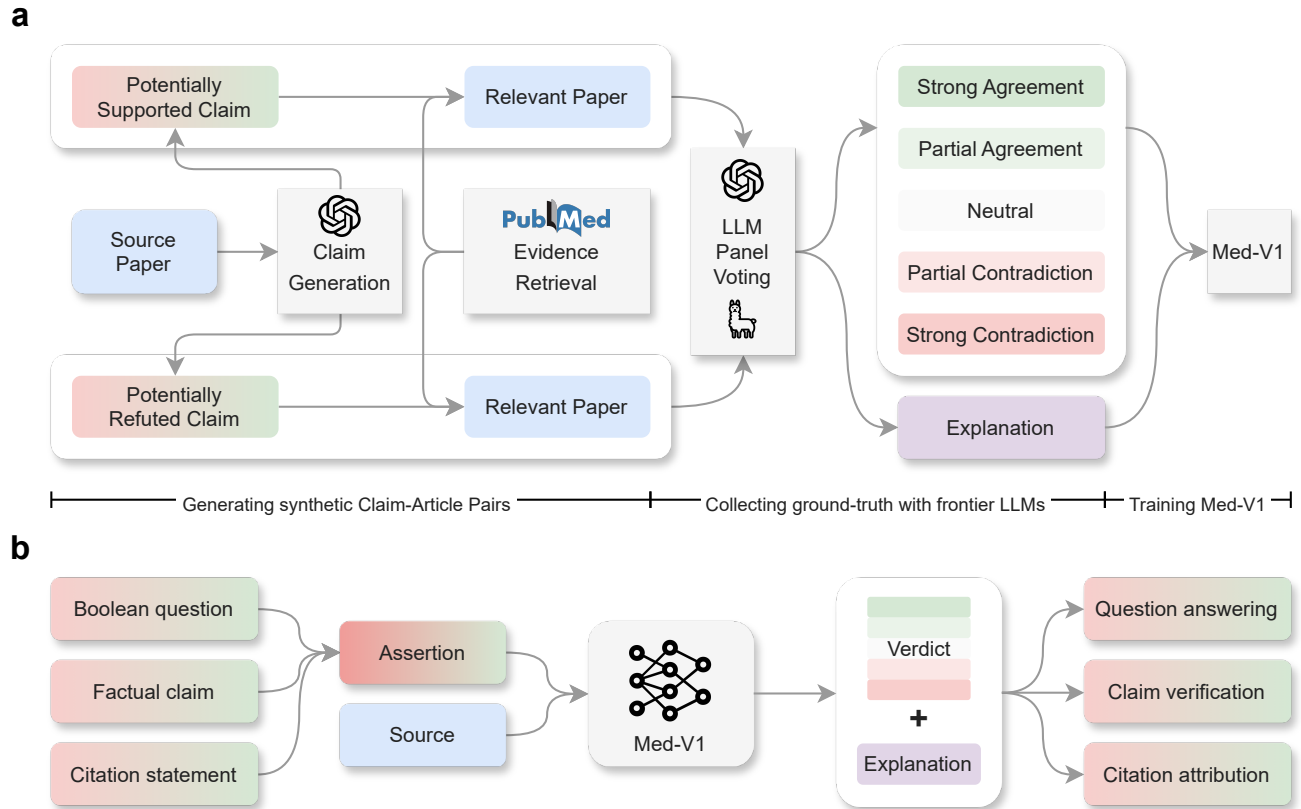


Figure 1. Overview of Med-V1 training and inference. **a:** MedFact-Synth construction and Med-V1 training. Synthetic claims are generated from source papers and then verified by a panel of LLMs using relevant papers retrieved from PubMed. The resulting verified claim-evidence pairs form the MedFact-Synth dataset, which is then used to train Med-V1 through a combination of supervised fine-tuning and reinforcement learning. **b:** Inference with Med-V1. Given an assertion and a source biomedical article, Med-V1 assesses whether the article supports the assertion. The assertions can be derived from Boolean questions, factual claims, or citation statements, corresponding to the applications of question answering, claim verification, and citation attribution, respectively. Med-V1 outputs both a 5-point Likert rating of agreement and a natural-language explanation of its verdict.

claims remain concise enough to function as standalone assertions, while rationales and articles provide sufficient context for non-trivial reasoning without requiring full-length papers.

To assess the quality of MedFact-Synth, we randomly sample 100 instances and recruit three annotators (Z.W., N.W., J.C.) with biomedical training to rate the claim verdict using the same 5-point Likert scale as in MedFact-Synth. Annotators were instructed to follow the original label definitions and to base their judgments solely on the title and abstract of each source article. After an initial round of independent annotations, a second round of discussion was conducted to reach a consensus verdict for each instance, which we treat as the ground-truth label.

Figure 2d shows the confusion matrices comparing each annotator and the MedFact-Synth labels against these consensus ground-truth labels. Overall, the synthetic MedFact-Synth achieves the highest annotation accuracy of 0.74, much higher than the individual annotators' accuracies ranging from 0.54 to 0.69. Most discrepancies between MedFact-Synth and the consensus labels differ by only one point on the 5-point scale, indicating that the synthetic labels rarely produce severe errors, such as labeling a strongly supported claim as strongly contradicted.

Taken together, these results confirm that MedFact-Synth is large-scale, diverse, and high-quality, and that the LLM panel can provide labels that are at least comparable to, and often more consistent than, those of individual human annotators.

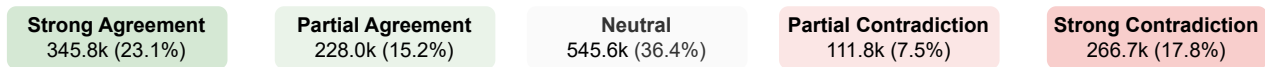
2.3 Med-V1 Closes the Performance Gap between Frontier and Lightweight LLMs

To quantitatively measure the performance of Med-V1, we introduce MedFact-Bench. It contains five biomedical verification datasets: SciFact [29], HealthVer [31], MedAESQA [30], and two repurposed datasets, PubMedQA-Fact (derived from

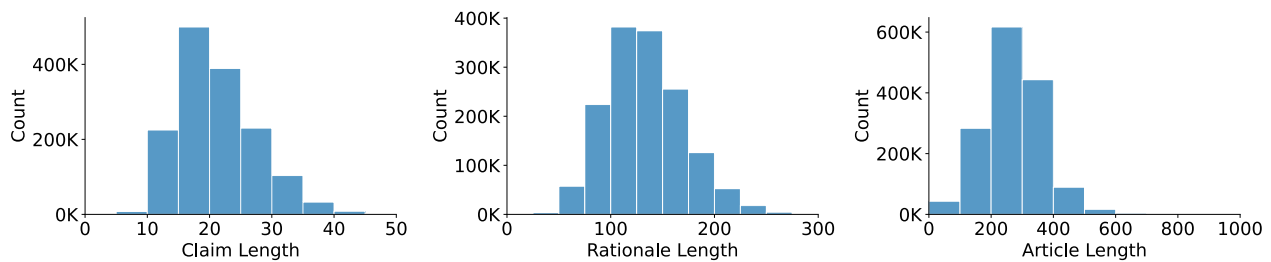
a

Synthetic Claim (generated from PMID: 9035431) Aprotinin is more effective than epsilon-aminocaproic acid in reducing blood loss during orthotopic liver transplantation.	Verification Rationale The article directly compares the effectiveness of aprotinin and epsilon-aminocaproic acid (EACA) in reducing blood loss during liver transplantation. According to the results, EACA had the lowest intraoperative blood loss, which contradicts the claim that aprotinin is more effective. The claim states that aprotinin is more effective than EACA in reducing blood loss, but the article suggests the opposite. Therefore, the article clearly and directly refutes the claim, indicating a strong contradiction.
Source (PMID: 24935303, retrieved by MedCPT) Aminocaproic Acid (amicar) as an alternative to aprotinin (trasylo) in liver transplantation. INTRODUCTION: This study compared clinical outcomes for a large number of liver transplant patients receiving intraoperative epsilon-aminocaproic acid (EACA), aprotinin, [...] EACA had the lowest intraoperative blood loss and required the fewest transfusions of plasma. [...]	
Verdict: Strong Contradiction (-2)	

b



c



d

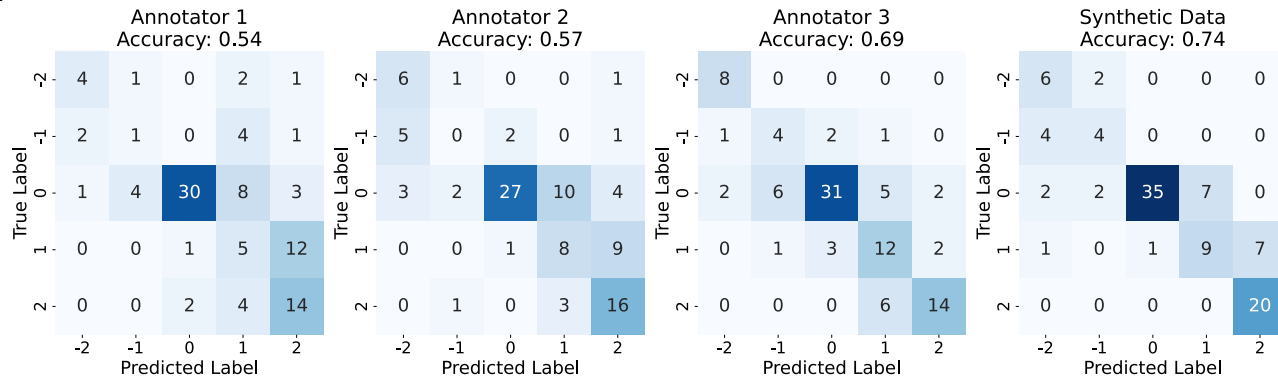


Figure 2. Generation and evaluation of MedFact-Synth. **a:** An example from MedFact-Synth. **b:** Distribution of veracity labels in MedFact-Synth. **c:** Word-count distribution of the claims, rationales, and articles in MedFact-Synth. **d:** Confusion matrices comparing each annotation method with the true labels, which are the ground-truth derived from annotator consensus. Annotations in MedFact-Synth (Synthetic Data) achieve higher accuracy than those of human annotators.

PubMedQA [32]) and BioASQ-Fact (derived from BioASQ [33]). Together, MedFact-Bench contains 14,274 claim–article pairs annotated with three veracity labels: support, not enough information (NEI), and contradict. As these datasets do not distinguish between partial and strong labels, we map Med-V1’s prediction into the same three-way taxonomy by merging strong and partial support into support, strong and partial contradiction into contradict, and neutral to NEI.

Figure 3 reports the performance of different LLMs on MedFact-Bench. Notably, a wide variety of frontier LLMs, from Llama-3.3-70B-Instruct to GPT-4o and GPT-5, achieve tightly clustered results across nearly all tasks, with average accuracies ranging from 0.717 to 0.736. This narrow distribution suggests that frontier model performance on MedFact-Bench has largely plateaued, and that an accuracy of around 0.73 may represent a practical “saturation” level given the noise and label ambiguity in current benchmarks. In contrast, much smaller 3B-parameter models such as Llama-3.2-3B-Instruct and Qwen2.5-3B-Instruct

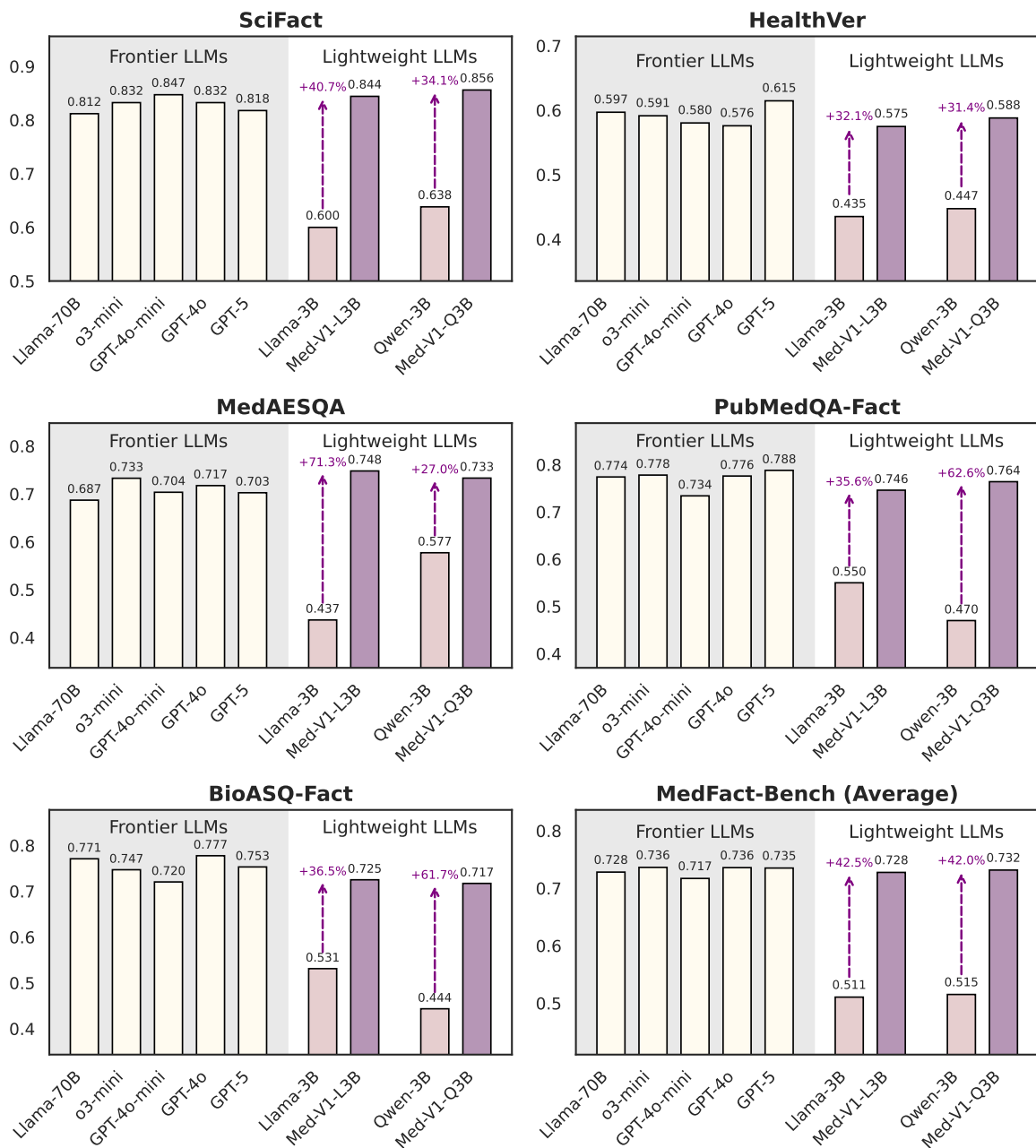


Figure 3. Zero-shot accuracies of different LLMs on MedFact-Bench. Performance is reported for each component dataset, and the (macro-)average accuracy is the main evaluation metric of MedFact-Bench. Frontier LLMs include large-scale open LLMs (e.g., 70B parameters) and the latest proprietary LLMs. Lightweight LLMs are 3B-parameter models. Med-V1-L3B is fine-tuned from Llama-3.2-3B-Instruct (Llama-3B), and Med-V1-Q3B is fine-tuned from Qwen2.5-3B-Instruct (Qwen-3B). Relative improvements over the respective base models are indicated. Both Med-V1 variants show significant performance improvements over their backbones. Llama-70B denotes Llama-3.3-70B-Instruct.

perform substantially worse in this setting, with average accuracies of only 0.511 and 0.515, respectively. These gaps highlight the challenge of achieving reliable verification with lightweight models in the absence of domain-specific fine-tuning.

With training on MedFact-Synth, Med-V1 significantly improves over the base 3B-parameter models. On average, Med-V1-L3B achieves 42.5% higher accuracy than Llama-3B (0.728 vs. 0.511), and Med-V1-Q3B achieves 42.0% higher accuracy than Qwen-3B (0.732 vs. 0.515). These gains are consistent across all datasets and for both the Llama- and Qwen-based

variants (27.0% to 71.3% relative performance improvements), demonstrating that the training pipeline is effective regardless of the underlying architecture. Overall, these results show that Med-V1 substantially narrows the performance gap between lightweight and frontier LLMs on biomedical verification, even though it is trained entirely on synthetic data and evaluated strictly in a zero-shot setting.

2.4 Error Analysis Reveals High-quality Reasoning Despite Dataset Noise

To better understand Med-V1’s errors, especially its weaker performance on certain datasets, we conduct an error analysis across the five MedFact-Bench components. For each model variant (Med-V1-L3B and Med-V1-Q3B), we randomly sample 20 misclassified predictions from each dataset (100 total examples per model variant) and manually examine the source article, claim, ground-truth label, and model output. Based on this analysis, we categorize errors into four types: model prediction also acceptable (Type A), bad claim (Type B), model knowledge error (Type C), and model reasoning error (Type D).

Type A and B errors do not reflect true model failures. Type A errors consist of cases where, upon close inspection, the model’s reasoning and answer are acceptable. For these cases, the resulting discrepancies between the ground truth and model error occurs either because the provided ground truth label does not accurately represent the claim-source pair, or because the model weighs evidence appropriately to conclude partial agreement or disagreement when the ground truth is neutral. The latter scenario occurs most frequently in the MedAESQA dataset. For instance, one MedAESQA claim states that “Trauma-focused cognitive behavioral therapy (TFCBT), prolonged exposure therapy, eye movement desensitization and reprocessing (EMDR), and stress management are effective in reducing PTSD symptoms.” While the source addresses the efficacy of EMDR and TFCBT, it does not mention prolonged exposure therapy or stress management. The model labels this claim-source pair as partial agreement, whereas the ground truth is neutral. However, because the source supports a subset of the information in the claim, the model’s classification as partial agreement is defensible. Type B errors address claims that are of poor quality. For example, the claim is ambiguously phrased, not a true standalone assertion, or contains fabricated or non-interpretable biomedical concepts. Overall, Type B errors are less prevalent than others, but the frequency varies by dataset. BioASQ-Fact and MedAESQA exhibit the most Type A labels across both Med-V1-L3B and Med-V1-Q3B cases (Type A; $n = 17, 16$ for Med-V1-L3B and $n = 17, 14$ for Med-V1-Q3B, respectively). In total, Types A and B account for the majority of examined errors (71% of errors of the Med-V1-L3B sample and 66% of errors of the Med-V1-Q3B sample), indicating that dataset quality could limit the accuracy of claim verification.

The remaining two categories represent true model errors. Type C (knowledge errors) refers to inaccurate or incomplete factual knowledge and assumptions, such as misinterpreting the source, overlooking key evidence, or introducing unsupported details during reasoning. These are most common in SciFact ($n = 6, 5$) and PubMedQA-Fact ($n = 4, 5$). In contrast, Type D (reasoning errors) represents cases where the model correctly interprets the evidence but derives a self-contradictory or logically invalid conclusion. A notable pattern is that the step-by-step reasoning appears coherent, but ultimately reaches a verdict that contradicts its own evidence. PubMedQA-Fact and SciFact also exhibit the highest numbers of Type D errors.

Thus, although Med-V1 achieves lower accuracies appear generally lower on some datasets such as HealthVer (Figure 3), many observed “errors” likely arise from dataset noise, rather than true failures of model reasoning. Both Med-V1 variants produced similar error distributions, reinforcing this interpretation. These results highlight the importance of ongoing efforts to refine biomedical verification benchmarks. At the same time, the remaining Type C (knowledge) and Type D (reasoning) errors provide concrete targets for future model improvements.

Table 1. Distribution of different error types for Med-V1 across MedFact-Bench. SF: SciFact; HV: HealthVer; MQA: MedAESQA; PQA: PubMedQA-Fact; BA: BioASQ-Fact. Type A: incorrect ground truth; Type B: bad claim; Type C: model knowledge error; Type D: model reasoning error. Only Types C and D errors reflect true model failures

Error	Med-V1-L3B						Med-V1-Q3B					
	SF	HV	MQA	PQA	BA	Sum	SF	HV	MQA	PQA	BA	Sum
Type A	11	14	16	6	17	64	9	11	14	11	17	62
Type B	0	3	2	2	0	7	0	4	0	0	0	4
Sum	11	17	18	8	17	71	9	14	14	9	18	66
Type C	6	0	1	4	1	12	5	2	3	5	2	17
Type D	3	3	1	8	2	17	6	3	3	4	1	17
Sum	9	3	2	12	3	29	11	5	6	9	3	34

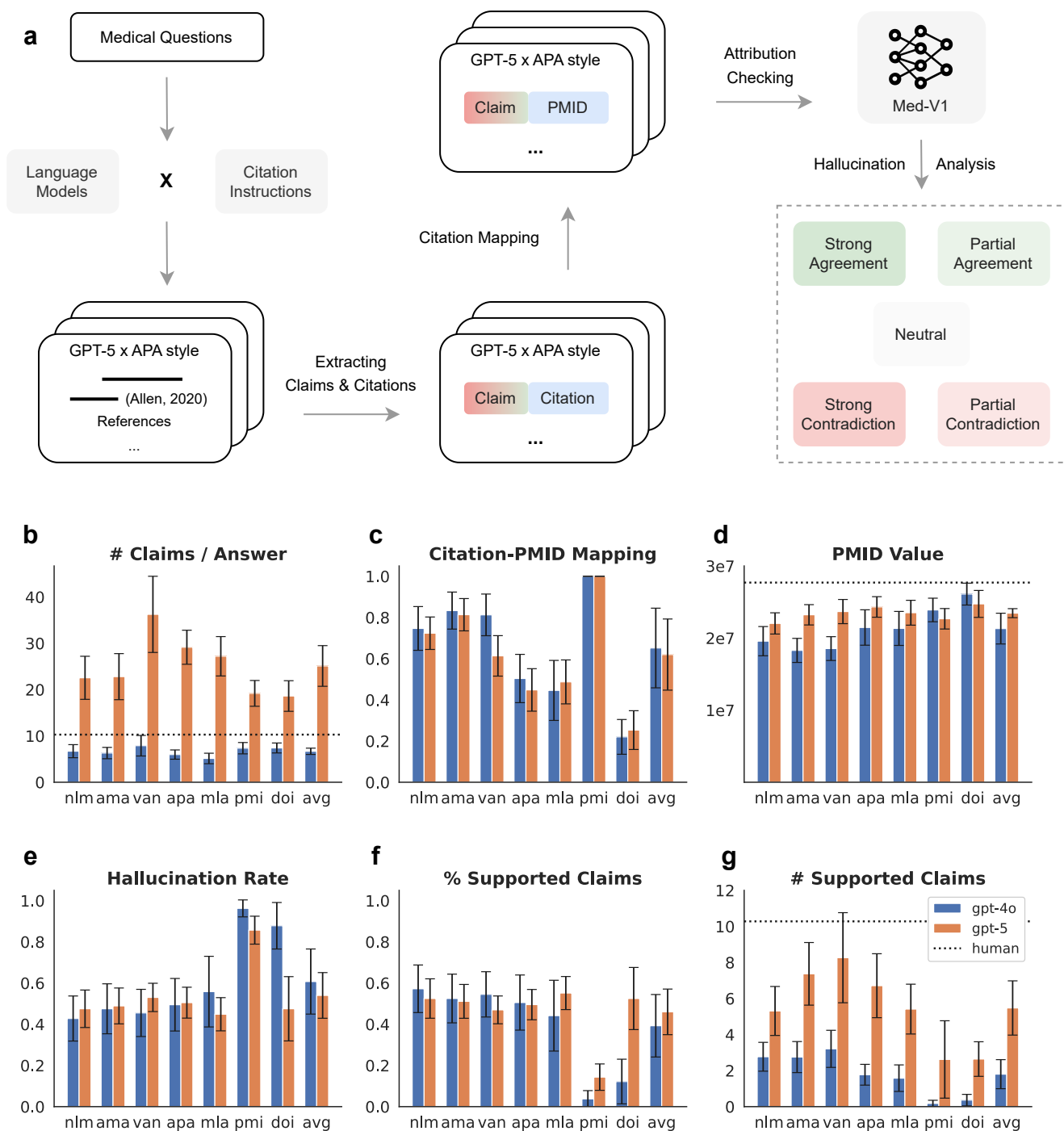


Figure 4. Detecting LLM Hallucination with Med-V1. **a:** Overview of this use case study. We use Med-V1 to analyze the hallucination rates of different LLMs and citation instructions. **b:** average number of claims (citation statements) per LLM-generated answer. **c:** The proportions of the generated citations that can be mapped to a PubMed ID (PMID). **d:** The average PMID values, which reflect their recency, generated by different methods. **e:** The hallucination rates of different methods. **f:** The proportion of supported claims generated by different methods. **g:** The number of supported claims generated by different methods. Statistics of the human-generated answers are shown in dotted horizontal lines. Avg: the averaged metric for each model. Error bars reflect 95% confidence intervals estimated from 2,000 bootstrap iterations.

2.5 Use Case 1: Detecting LLM Hallucinations with Med-V1

On the MedAESQA dataset that verifies AI-generated claims against their citations, both Med-V1 variants outperform GPT-5 in accuracy (Med-V1-L3B: 0.748, Med-V1-Q3B: 0.733, GPT-5: 0.703; Figure 3). Empowered by this capability of Med-V1, we perform a first-of-its-kind use case study to analyze the hallucination rates of different LLMs under various citation instructions. As shown in Figure 4a, we prompt two LLMs (GPT-4o and GPT-5) to answer the 40 medical questions from MedAESQA using seven different citation style instructions, including NLM (National Library of Medicine), AMA (American Medical Association), Vancouver, APA (American Psychological Association), MLA (Modern Language Association), and directly citing the PMID or DOI. Claim-citation pairs from LLM-generated answers are then extracted using GPT-5, and the citation texts are mapped to PMIDs using the single citation matcher of PubMed [40]. Titles and abstracts of the mapped PMIDs are retrieved from the PubMed database, and they are fed into Med-V1-L3B for attribution checking together with the claims.

On average, human experts answer with 10.3 claims for each question (Figure 4b). In comparison, GPT-4o generates fewer claims (5.1 to 7.4) while GPT-5 generates many more claims (18.6 to 36.3). While GPT-4o generates similar numbers of claims under different citation styles, GPT-5 generates the most claims when using the Vancouver style. Overall, citations generated by GPT-4o and GPT-5 have similar rates of successful PMID mapping (Figure 4c). Except for directly citing PMIDs that gets perfect mapping, NLM, AMA, and Vancouver-style citations (61.4% to 83.3%) are more likely to be mapped successfully to PMIDs than APA and MLA citations (44.6% to 50.4%). Most of the LLM-generated DOIs are invalid and cannot be mapped to any PMID. Overall, the successful mapping rate of GPT-4o-generated citations is similar to that of GPT-5. Additionally, since PMIDs are indexed in a temporal order, we use their absolute values as a proxy for publication recency (Figure 4d). Human-generated citations are the most up-to-date, with an average PMID value of 27696376, which corresponds to an article published in 2016.

Among those citations with successful PMID mapping, we define hallucinations as claims that are not partially or fully supported by the PubMed source, corresponding to the Med-V1 predicted scores of +1 (partial agreement) and +2 (strong agreement), respectively. Among NLM, AMA, Vancouver, APA, and MLA citation styles, similar hallucination rates are observed with between GPT-4o (42.8% to 55.8%) and GPT-5 (44.9% to 53.0%), as shown in Figure 4e. As expected, both models show high hallucination rates (96.3% for GPT-4o and 85.7% for GPT-5) when instructed to directly cite PMIDs. Interestingly, while GPT-4o also hallucinates over 80% when citing DOI, GPT-5 shows a lower hallucination rate (47.5%), which is consistent with other non-ID-based citation formats. This suggests that GPT-5 has improved memorization of DOI compared to GPT-4o. Figure 4f shows the proportion of supported claims, which inversely reflects these discussed trends in hallucination rates. Finally, while GPT-5 still generates more claims (2.6 to 8.3) that can be supported by the source than GPT-4o (0.2 to 3.2), both remain below the human baseline of 10.3 claims (Figure 4g). For both models, the highest number of supported claims occurs under the Vancouver citation format, and the fewest occur when directly citing PMIDs or DOIs.

In summary, GPT-5 produces far more claims and citations than GPT-4o (especially under the Vancouver style), but both models show similar PMID-mapping patterns across citation styles and generally rely on less recent sources than humans. Among successfully mapped citations, their hallucination rates are comparable for standard formats, whereas direct PMID citations cause extreme hallucinations, and DOI citations remain problematic for GPT-4o. GPT-5, however, shows improved DOI performance, with hallucination rates similar to other non-ID-based styles, suggesting improved DOI memorization. Despite these improvements, both models still generate fewer supported claims than human experts overall.

2.6 Use Case 2: Identifying High-Stakes Misattributions with Med-V1

To illustrate the practical value of Med-V1 in real-world settings, we apply it at scale to clinical practice guidelines, which depend heavily on accurate citation and might impact high-stakes medical decision-making. We collect all publicly available guideline articles published between 2015 and 2025 in PubMed Central, resulting in 6,152 articles. From these, we extract decontextualized statements associated with a single citation, resulting in a total of 57k statement-source pairs.

For each pair, we apply Med-V1-L3B to classify if the cited article supports the citation statement. As shown in Figure 5a, most citations (54.2%) are predicted to be strongly (24.4%) or partially (29.8%) supported by the source, which is expected as guideline recommendations are typically grounded in evidence that the authors regard as supportive. Importantly, a smaller portion is flagged as partial contradiction (1.6k pairs, 3%) or strong contradiction (1.1k pairs, 2%). In such cases, a guideline statement might cite a source that does not support the claim or even directly contradicts it. To estimate the real prevalence of these misattributions, we randomly sample 100 flagged contradictions for review, including 50 partial contradiction and 50 strong contradiction cases. The human reviewer evaluates whether each flagged case is a valid misattribution, or assigns it one of four other categories: a poor quality claim, not enough information arising from the use of titles and abstracts as the source instead of full text, a model error, or other explanation.

In our sample, 28 cases are validated as misattributions – the citation statement in the guideline is actually contradicted by the cited article (Figure 5b). To characterize their potential impact, we categorize claims into six domains: effectiveness, risk/etiology, prevalence, diagnostic, prognostic, and expert opinion/policy. The most frequently misattributed domain

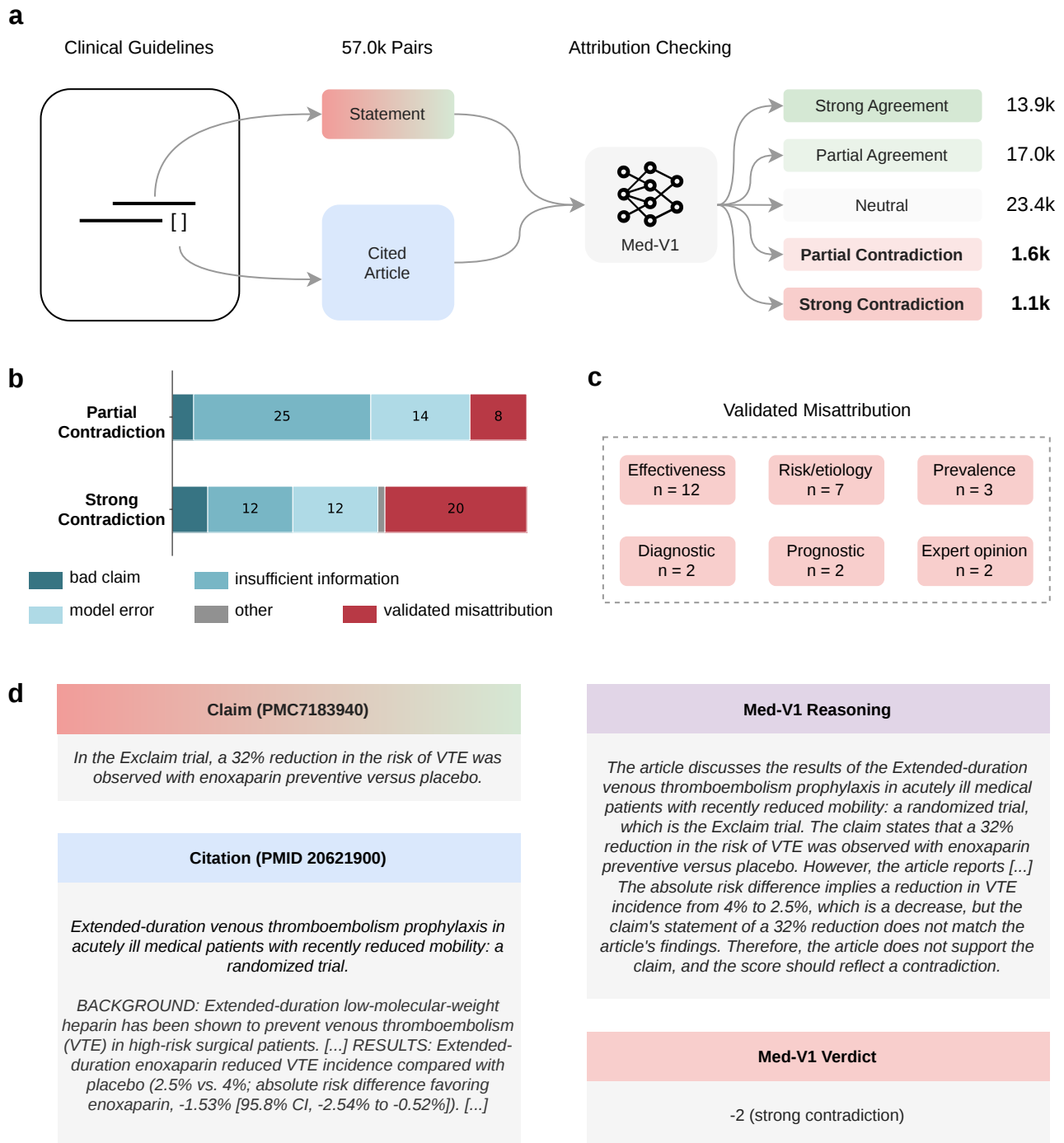


Figure 5. Identifying High-stakes Misattributions with Med-V1. **a:** Overview of this use case study. We extract citation statements and their source articles from clinical guidelines, and automatically check their attribution validity with Med-V1. **b:** Distribution of the manual validation of 50 partial contradiction and 50 strong contradiction samples. **c:** Topic distribution of the manually validated misattributions. **d:** An example of the manually-validated misattribution identified by Med-V1. In this example, the claim of a “32%” reduction (from the clinical practice guideline PMC7183940) contradicts the results presented in its cited source.

(accounting for more than 40% of the validated misattributions) is the effectiveness of a treatment. These errors are often made when reporting percentages or statistical significance (Figure 5c). For example, one claim mentions “a 32% reduction in the risk of VTE was observed with enoxaparin preventive versus placebo.” However, this is clearly refuted by the cited source, which reports a 2.5% vs 4% absolute difference, corresponding to a 37.5% relative risk reduction. This example illustrates that errors in calculations can result in a misleading perceived effect (Figure 5d). The second most common misattributed category is risk/etiology, in which the claim discusses the development of adverse events due to an intervention or exposure factor. For instance, one claim states “There are several compelling evidences that childhood OB is the most common risk factor for both cardiovascular (CV) and metabolic diseases,” while the corresponding citation mentions that when body mass index (BMI) is adjusted for, there is a weak negative association between childhood BMI and metabolic variables. This example highlights how misattributions can arise from misrepresenting association findings. These two categories of misattributions are particularly consequential, as they can directly inform decisions regarding treatment and prevention.

It is important to highlight that a claim categorized as a validated misattribution does not necessarily indicate that the claim itself is universally incorrect, rather that it is not supported by its cited source. Overall, this large-scale analysis demonstrates that Med-V1 can efficiently audit citation accuracy in biomedical documents with potential public health impact. This approach provides a practical way to identify inconsistencies that would be extremely difficult to detect manually at the guideline scale.

3 Discussion

In this work, we introduce Med-V1, a family of 3B-parameter biomedical verification models trained on the large-scale MedFact-Synth corpus to perform zero-shot evidence attribution with both structured verdicts and natural-language explanations. Across MedFact-Bench, Med-V1 substantially improves over lightweight base models and narrows the gap to much larger frontier LLMs, while remaining far more efficient to deploy. Analyses of MedFact-Synth indicate that the synthetic supervision is diverse and consistent, and the error analysis suggests that a substantial fraction of benchmark “errors” reflect dataset ambiguity or label noise rather than systematic failures of reasoning. These results support the view that targeted post-training on high-quality synthetic supervision can unlock strong verification ability in compact open models.

Biomedical verification faces a persistent supervision bottleneck because high-quality claim–evidence labeling is expensive and difficult to scale across medical subdomains. Traditional biomedical claim verification systems are typically supervised classifiers trained on manually annotated datasets [41, 42]. This strategy limits their adaptability to new tasks and prevents them from providing explicit reasoning beyond a final label. Recent work on synthetic verification data [43, 44] helps alleviate data scarcity, but largely retains the black-box paradigm by producing labels without rationales. As a result, such models are not often used in modern, LLM-centered pipelines [7, 11], where users instead prompt large proprietary or open-source frontier LLMs in a zero-shot manner to evaluate claims or citations. Although effective, these models are expensive to run at scale and pose practical challenges in biomedical settings. Med-V1 is designed to bridge this gap by offering frontier-level verification performance in a lightweight 3B-parameter model that supports fully zero-shot use while still producing structured verdicts and explicit, interpretable reasoning.

The two use cases illustrate how a lightweight verifier can serve as reusable infrastructure for auditing. The hallucination analysis shows that citation instructions strongly influence citation mappability and downstream support rates, implying that citation format is not merely a presentation choice but can alter factual reliability. The guideline audit shows that automated verification can surface potentially high-stakes misattributions at a scale that is otherwise infeasible, enabling risk-based triage in which domain experts review a focused set of flagged contradictions. These applications do not require perfect verification accuracy to be valuable, because operational utility often depends on high precision among the most suspicious cases rather than uniform performance across all instances.

Our study has several limitations. First, Med-V1 is trained only on titles and abstracts from PubMed, primarily because full-text articles are limited and subject to licensing constraints. That said, Med-V1 can potentially apply to full-texts by coupling with a retrieval module that identifies relevant sections to be used. Although abstracts summarize the main findings, many clinically important verification tasks require methodological details or results that are discussed only in the full text, making extension to long-form evidence an important direction for future work. Second, we evaluate Med-V1 solely in a zero-shot setting, without integrating it into end-to-end systems. In practice, deployed verification models will require retrieval components to identify relevant evidence for a given assertion, and future work should explore how Med-V1 performs when coupled with such retrieval modules. Finally, our current formulation focuses on claim verification in isolation and does not address other key dimensions of evidence-based decision-making. Real-world fact-checking also requires assessing evidence quality, distinguishing background information from actionable recommendations, and considering factors such as study design, generalizability, and risk of bias.

Taken together, our findings suggest that high-quality synthetic supervision, paired with targeted post-training, can serve as a practical alternative to expensive frontier LLMs for many biomedical verification tasks. Beyond Med-V1 itself, our framework for generating reasoning-aware synthetic corpora and using them to align small open models may generalize to other domains

that require transparent and scalable verification, such as clinical document summarization. As the volume of AI-generated biomedical text continues to grow, lightweight verification models like Med-V1 may play an increasingly important role as reusable infrastructure for auditing the factual consistency and citation integrity of automated systems.

4 Methods

4.1 Constructing MedFact-Synth with synthetic data generation

Figure 6 summarizes the pipeline used to construct MedFact-Synth, our large-scale synthetic corpus for supervising biomedical verification. The pipeline operates in three main stages: (i) sampling biomedical articles and generating synthetic claims from them, (ii) using dense retrieval to match each claim to potentially relevant articles, and (iii) assigning consensus veracity labels and rationales using a panel of LLMs.

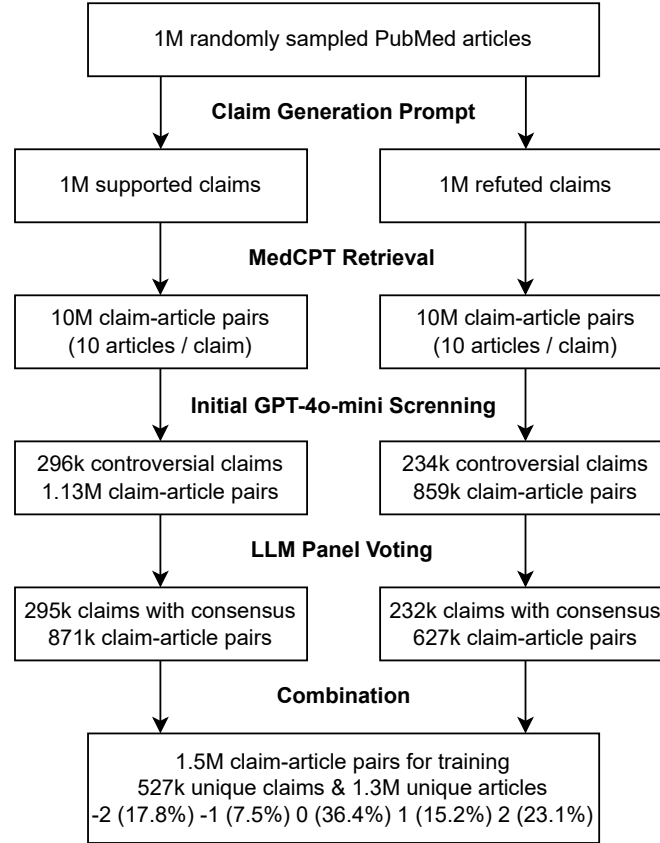


Figure 6. Overview of the proposed pipeline for synthetic data generation. One million PubMed articles are randomly sampled and used to generate two standalone claims per article: one potentially supported and one potentially refuted by the source article. Each claim is paired with ten candidate articles retrieved using MedCPT, resulting in approximately 20 million claim–article pairs. An initial screening stage using GPT-4o-mini filters for controversial claim–article pairs that receive mixed verdict signals. These pairs are then evaluated by a panel of frontier LLMs, and a consensus label is assigned when sufficient agreement is reached. Finally, supported- and refuted-claim pipelines are merged to construct MedFact-Synth, yielding 1.5 million high-quality claim–article pairs with five-point Likert-scale veracity labels and natural-language rationales, covering a balanced distribution of agreement, neutrality, and contradiction.

Stage 1: Sampling articles and generating claims. We begin from a local copy of the PubMed 2025 baseline and randomly sample one million articles with non-empty titles and abstracts. For each sampled article, we prompt GPT-4o-mini twice: once to generate a claim that can be supported by the article (by Prompt 1), and once to generate a claim that can be refuted by it (by Prompt 2). In both cases, the model is instructed to produce standalone, interpretable statements that do not rely on the original text for context. For refuted claims, we additionally discourage trivial negations (e.g., simply inserting “not”

into an otherwise supported claim). This procedure yields two synthetic claims per article, each tied to the source paper but phrased for independent verification.

Prompt 1: Supported Claim Generation Prompt

System Prompt

You are a biomedical annotation expert, and your task is to generate a claim that can be supported by the provided article. The claim should be interpretable on its own, without relying on the article. Do not generate anything else than the claim.

User Prompt

Here is the article:

Title: {title}

Abstract: {abstract}

Prompt 2: Refuted Claim Generation Prompt

System Prompt

You are a biomedical annotation expert, and your task is to generate a claim that can be refuted by the provided article. The claim should be interpretable by its own, without relying on the article. Avoid using simple negative words such as not and no. Do not generate anything else than the claim.

User Prompt

Here is the article:

Title: {title}

Abstract: {abstract}

Stage 2: Matching claims to candidate evidence. Rather than always pairing each claim with its source article, we aim to mimic realistic verification settings in which a claim must be checked against independently retrieved literature. To achieve this, we use a dense retrieval system based on MedCPT [37] to encode each synthetic claim into an embedding and search a large index of PubMed article embeddings. For every claim, we retrieve the top k most similar articles (with $k = 10$ in our experiments), together with their titles and abstracts. This step decouples claims from their originating papers and naturally produces a diverse set of claim–article pairs spanning strong agreement, partial agreement, neutrality, and contradiction.

Stage 3: Multi-LLM reasoning and consensus labeling. In the final stage, we assign high-quality veracity labels and rationales to the retrieved claim–article pairs using a panel of LLMs. Each pair is first evaluated by GPT-4o-mini with the Med-V1 Verification Prompt (Prompt 3), which instructs the model to (i) reason step by step about how the article relates to the claim, and (ii) assign a discrete score on a five-point Likert scale: strong contradiction (−2), partial contradiction (−1), neutral/unrelated (0), partial agreement (+1), or strong agreement (+2). This initial pass yields both scores and rationales for tens of millions of claim–article pairs.

To focus labeling resources on the most informative examples, we then identify “controversial” claims whose retrieved articles receive a mixture of supportive, contradicting, and neutral scores. For these challenging pairs, we collect additional evaluations from two stronger LLMs (Llama-3.3-70B-Instruct and o3-mini), resulting in three independent scores and associated rationales per pair. We then derive a consensus “ground-truth” score only when at least two of the three models agree, and the third differs by at most one point on the five-point scale. For pairs that meet this criterion, we randomly select one rationale from one of the agreeing models and store it alongside the final label, the claim, and the article abstract.

Applying these filters produces MedFact-Synth, a high-quality synthetic dataset containing approximately 1.5 million claim–article pairs, covering hundreds of thousands of unique claims and over one million unique articles. Each instance includes (i) a standalone biomedical claim, (ii) a title–abstract pair serving as the verification source, (iii) a structured veracity label on the five-point Likert scale, and (iv) a natural language explanation generated by the LLM panel. This reasoning-aware supervision is subsequently used to post-train Med-V1.

Prompt 3: Med-V1 Verification Prompt

System Prompt

You are a fact-checking expert trained in evidence-based medicine. Your task is to evaluate how strongly an **article** agrees or disagrees with a **claim**. The **article** is retrieved from a search engine using the **claim** as the query.

Use the following five-point scale:

- **** -2 Strong Contradiction **** – The article clearly and directly refutes the claim.
- **** -1 Partial Contradiction **** – The article provides mixed or indirect evidence against the claim.
- **** 0 Neutral / Unrelated **** – The article does not address the claim, offers insufficient information, or is irrelevant to the claim.
- **** 1 Partial Agreement **** – The article offers some indirect or tentative support for the claim.
- **** 2 Strong Agreement **** – The article explicitly and strongly supports the claim.

Note that the **article** might not describe the exact same subjects, interventions, or measurements as the **claim**. In this case, please note the difference and assign a score of 0.

Output in two parts only and do not output anything else:

<think>[your detailed, step-by-step explanation for scoring]</think>

<score>[the integer score only, i.e., -2, -1, 0, 1, or 2]</score>

User Prompt

Article:

{source-article}

Claim:

{claim}

4.2 Training Med-V1

We use a two-stage post-training procedure that first applies supervised fine-tuning (SFT) and then reinforcement learning (RL) [45]. Using this framework, we train two variants of Med-V1: Med-V1-L3B, initialized from Llama-3.2-3B-Instruct [27], and Med-V1-Q3B, initialized from Qwen2.5-3B-Instruct [28]. Both variants follow the same training schedule and hyperparameters. Their only difference is the backbone model.

Supervised fine-tuning. In the SFT stage, we formulate verification as a standard conditional generation task. Given a claim and a source article, the model is trained to produce the corresponding reasoning trace and veracity score provided by MedFact-Synth. We optimize a language modeling objective using the AdamW optimizer with a learning rate of 2×10^{-5} , an effective batch size of 128, and train for 10 epochs on approximately 1.5 million instances. All SFT experiments are conducted on Nvidia H200 (140GB) GPUs with distributed data-parallel training. In this setup, training Med-V1 on $8 \times$ H200 GPUs requires on the order of tens of GPU-days.

Reinforcement learning. After SFT, we apply RL to further refine the model’s alignment with human preferences. Specifically, we design a rule-based reward function that evaluates each prediction in two steps. First, it checks whether the output adheres to the format required in the Med-V1 verification: one <think>...</think> block followed by one <score>...</score> block. Any output that violates this format receives a reward of -1. For well-formatted predictions, the function then extracts the model’s predicted score and compares it with the ground-truth label. Outputs with invalid or uninterpretable scores are again assigned a reward of -1. Valid scores are rewarded according to $0.5 \times (2 - |y_{\text{pred}} - y_{\text{true}}|)$, which assigns higher values to predictions whose 5-point Likert scores are closer to the ground truth.

We then update the LLMs (policy model) using Group Relative Policy Optimization (GRPO) [46], which estimates advantages without a separate critic model. For each prompt, the policy generates multiple candidate responses, and each is scored by the rule-based reward function. The mean reward of all responses for a given prompt serves as the baseline, and each response is compared with it to compute its advantage. Similarly, we perform GRPO training on 1.5M MedFact-Synth instances. For each prompt, the policy generates 5 rollouts, all evaluated by the rule-based reward function. The actor model is initialized from the SFT checkpoint and trained with a maximum prompt length of 3,072 tokens and a maximum response length of 1,024 tokens. We use AdamW with a learning rate of 1×10^{-3} and a global batch size of 1,440. KL regularization follows the low-variance formulation with a coefficient 10^{-3} to constrain drift from the SFT model. Training runs for 3 epochs on a single node with $8 \times$ Nvidia H200 (140GB) GPUs, for a total of 3,120 steps.

4.3 Developing MedFact-Bench

To evaluate Med-V1, we curate MedFact-Bench, a benchmark comprising five biomedical verification datasets: SciFact [29], HealthVer [31], MedAESQA [30], PubMedQA-Fact (re-purposed PubMedQA [32]), and BioASQ-Fact (re-purposed BioASQ [33]). Each instance in MedFact-Bench consists of a claim–source pair annotated with one of three labels: support, not enough information (NEI), or contradict. Because Med-V1 produces a 5-point Likert score that reflects the degree of agreement between the claim and the source, we map the model’s predictions into this three-way taxonomy: strong/partial agreement → support, strong/partial contradiction → contradict, and neutral → NEI. We report macro-average accuracy across the five datasets as the primary metric, as it equally weights each dataset and reduces the influence of varying dataset sizes.

Claim verification datasets. SciFact [29] and HealthVer [31] are established biomedical claim verification datasets that provide diverse claim–source pairs annotated with support, contradict, or NEI labels. SciFact contains expert-curated scientific claims derived from citation statements and paired with abstracts that include the relevant evidence. HealthVer comprises real-world medical claims verified against biomedical abstracts. For consistency, we use the pre-processed versions of these datasets [47], containing 340 SciFact instances and 903 HealthVer instances.

Citation attribution dataset. MedAESQA [30] is an expert-annotated benchmark designed to evaluate citation attribution in LLM-generated health answers. Each example contains an LLM-generated answer with one or more citation markers, and annotators label whether the cited PubMed article supports the associated statement. We preprocess MedAESQA by flattening each answer into individual statement–PMID pairs, removing bracketed citation markers from the statement text, and normalizing sentence formatting. We then map each cited PMID to the title and abstract of the corresponding PubMed record in our local PubMed 2025 baseline, and discard pairs with missing or invalid PMIDs. MedAESQA provides fine-grained labels including supporting, contradicting, neutral, and not relevant. To include MedAESQA in MedFact-Bench under a unified three-way claim-verification taxonomy, we map neutral and not relevant into NEI. This yields a subset of 9,106 verification instances.

Re-purposed Boolean QA datasets. In addition to these established verification datasets, MedFact-Bench includes two re-purposed Boolean QA datasets, which we refer to as PubMedQA-Fact and BioASQ-Fact. We convert each QA instance into a verification example using a two-step process. First, each yes/no question is rewritten as a declarative claim under the assumption that the answer is “yes.” Second, we map the original QA answers to verification labels: yes → support, maybe → NEI, and no → contradict. This repurposing transforms large-scale QA datasets into fact-checking benchmarks compatible with the unified MedFact-Bench framework. PubMedQA-Fact is derived from PubMedQA [32], where each question corresponds to a PubMed title and the associated source is the abstract with its conclusion removed (the “reasoning-required” setting). In MedFact-Bench, we include the re-purposed PubMedQA test set of 500 instances. BioASQ-Fact is constructed from the BioASQ yes/no QA task [33], whose questions and answers are curated by biomedical experts and each supported by PubMed abstracts. In MedFact-Bench, we include all 3,425 yes/no questions from the BioASQ-13B training set.

Prompt 4: Question Conversion Prompt

System Prompt

You are a helpful assistant. Your task is to convert a yes/no question into a declarative statement. The statement should be a claim that is true if the answer to the question is “yes”. Do not output anything else than the converted statement.

User Prompt

Convert the following question into a claim, assuming the answer is “yes”:

Question: {question}

4.4 Error Analysis

We conduct an error analysis using the Med-V1 predictions of MedFact-Bench. From each of the component datasets of MedFact-Bench, we randomly sample 20 cases in which Med-V1 produced a label different from the ground truth, resulting in 100 cases for in-depth evaluation for each model variant. Our objective is to understand and characterize the underlying reasons for the discrepancies between ground truth and model output. For each case, we review the claim, source abstract, and Med-V1-generated reasoning. We document the source of error and derive a set of four error categories to classify these cases. Because We use a hierarchical procedure to make sure that only one error category is assigned to each case.

We first assess the claim itself, and if it is of poor quality, then the case is immediately labeled as a Type B (bad claim) error, and no further analysis is conducted. An example is the ambiguous claim “a popular treatment to tamp down the immune system in severely ill patients may help a few, but could harm many others,” where “popular treatment”, “a few”, and “others” lack the necessary context to draw a conclusion.

For cases with coherent claims, we review the source and draw a conclusion about the claim’s validity. We compare this conclusion to the ground truth label, and if the label is contradictory but the model’s prediction matches our classification, we assign the case as Type A (model prediction also acceptable). For instance, the claim “there is software for automated analysis of FISH images” is labeled as supported in the ground truth, but the source is about creating FISH images and does not discuss image analysis techniques, so the correct label should be “not enough information”.

If the claim and ground truth are consistent with each other, then we next examine Med-V1’s reasoning to identify the source of error. For the special case in which the ground truth is neutral, but Med-V1 correctly identifies that some components of the claim are directly supported or refuted by the source while others are not mentioned, thereby concluding a partial agreement or disagreement, we assign the case Type A. For true model errors, we assign Type C (knowledge-based) or Type D (reasoning-based). For the claim “Non-HDL-cholesterol is a better predictor of long-term outcome in patients after acute myocardial infarction compared to LDL-cholesterol,” Med-V1 makes a knowledge error through misinterpreting the source’s evidence regarding LDL-cholesterol, resulting in a contradictory label when the claim is actually supported. In contrast, a reasoning-based error arises for the claim “Febrile seizures reduce the threshold for development of epilepsy,” for which the model incorrectly concludes that the source supports the claim, even though it also acknowledges that the source does not directly confirm the claim. This method is applied across both Med-V1-L3B and Med-V1-Q3B outputs. We tabulate the error counts across a total of 200 cases and review the distribution of categories by type and dataset.

4.5 Detecting LLM Hallucinations with Med-V1

We conduct a use case study to quantify hallucinations in citation-grounded LLM answers under different citation instructions. We use the first 40 medical questions from MedAESQA [30] and prompt two LLMs (GPT-4o and GPT-5) to generate answers under seven citation-instruction formats: NLM, AMA, Vancouver, APA, MLA, and directly citing PMIDs or DOIs.

Prompt 5: Claim Extraction Prompt

System Prompt

You are an expert in biomedical literature and citation analysis. Your task is to extract every factual claim and its corresponding full citation from the provided text.

Instructions:

1. Identify every sentence or clause that makes a factual claim supported by a citation.
2. For each claim, identify the inline citation marker (e.g., “[1]”, “(Smith, 2023)”, “I”, “(PMID: 12345)”).
3. Resolve this inline citation to its **full reference entry** from the bibliography/reference list at the end of the text.
 - If the text uses numeric citations (e.g., AMA, Vancouver, NLM), match the number to the numbered reference list.
 - If the text uses author-date/page citations (e.g., APA, MLA), match the author/date or author/page to the alphabetical reference list.
 - If the text uses only inline identifiers (e.g., PMID, DOI) and has no reference list, use the full inline citation string itself (e.g., “PMID: 12345”).
4. Output the results as a strict JSON list of objects (**only JSON**; no additional text).
5. If a claim has multiple citations, repeat the same claim multiple times, once per citation, each time with a different citation.

JSON Format (example structure):

```
[
  {
    "claim": "The exact text of the factual claim.",
    "citation": "The full text of the corresponding reference entry (e.g., '1. Author AA. Title. Journal. Year...')."
  }
]
```

User Prompt

{model_answer}

To evaluate citation attribution, we apply a two-step extraction and processing workflow. First, GPT-5 is instructed with Prompt 5 to extract all factual claims from each generated answer and to resolve each in-text citation marker to the corresponding full reference entry when a reference list is present. This produces a set of claim–citation pairs per answer. Second, we parse

and normalize the extracted citations. For all styles, we remove leading numbering and trailing punctuation. For PMID- and DOI-based instructions, we further extract the bare identifier string. For non-identifier citation styles (e.g., NLM and AMA), we retain the full citation text after normalization.

We then map citations to PubMed articles. For non-identifier citation styles, we map each citation string to PMIDs using the PubMed single citation matcher [40], retaining the top-ranked match returned by the system. For DOI-based citations, we resolve DOIs to PMIDs using the National Center for Biotechnology Information (NCBI) ID Converter API [48] and standardize them into PMID-based mappings. For successfully mapped PMIDs, we retrieve titles and abstracts from a local PubMed 2025 baseline and use them as the verification source. For each extracted claim and its mapped source article, we apply Med-V1-L3B to predict a 5-point Likert score from strong contradiction (−2) and partial contradiction (−1) to neutral (0), partial agreement (+1), and strong agreement (+2). We define a claim as supported when the Med-V1 score is +1 or +2. Conditioned on successful PMID mapping, we define a hallucination when the claim is not partially or fully supported by its mapped source, corresponding to Med-V1 scores of −2, −1, or 0.

We report the following metrics by model and citation instruction: (i) the average number of extracted claims per answer, (ii) the PMID mapping rate, (iii) the average mapped PMID value as a proxy for recency, (iv) the hallucination rate among successfully mapped citations, (v) the proportion of supported claims, and (vi) the number of supported claims. Human expert statistics are computed from the expert answers in MedAESQA using the same claim extraction pipeline.

4.6 Identifying High-Stakes Misattributions with Med-V1

For this case study, we first compile a corpus of clinical guidelines by retrieving all the free full-text articles with the publication type of “practice guideline” and publication year between 2015 and 2025 (6,152 PMIDs, on Sep 16, 2025). Then, we extract the citation statement sentences from these articles using the BioC-preprocessed version of PubMed Central [49]. These articles have annotations indicating exactly where each citation appears in the text. Since some citation sentences might not be verification-worthy (e.g., citations to an experimental method, or a contextualized citation sentence that has ambiguous pronouns), we further apply GPT-4o-mini with Prompt 6 to filter out such claims, which leads to about 57k pairs of claim-citation pairs. We apply Med-V1 on this corpus and take the subset of cases in which the model determines there is disagreement between the source and claim. Because the source is obtained from the guideline claim’s citations, any instance of disagreement is proposed to be a misattribution. To ensure a balanced sample across Med-V1’s two levels of disagreement (partial contradiction and strong contradiction), we first stratify the misattribution cases by output label, then draw a random sample of 50 cases from each group, totaling to 100 cases for analysis.

Prompt 6: Claim Verification-worthy Check Prompt

System Prompt

You are a biomedical expert, and your task is to classify if a biomedical claim can be fact-checked.
Please respond with “yes” if the claim meets the requirement, and “no” otherwise. Only output “yes” or “no”.

User Prompt

Claim: “{claim}”

Our annotation procedure is similar to the one used for the error analysis. For each case, we evaluate, in order, the guideline claim, source (i.e. abstract of the article cited by the guideline), and Med-V1’s reasoning. As in the error analysis, we first classify cases with bad claims. Proceeding to the source, we recognize that a claim may be supported by evidence in the full text that is not present in the abstract. However, because Med-V1 uses titles and abstracts as input, we introduce a “not enough information” category to convey that this is due to lack of abstract detail, rather than a model error. When information in the abstract clearly does not support the claim, we label the case as a validated misattribution. On the other hand, if the abstract supports the claim, then Med-V1’s output is a model error. We do not further subcategorize these errors, as this is addressed in the error analysis section.

We perform a second-stage analysis on cases that are classified as validated misattributions, focusing on their potential clinical risk. We reuse the topic classification of clinical questions [50] to categorize claims into different domains/topics. Our classification results form the basis for further assessment of the potential risk associated with each misattribution.

Data Availability

MedFact-Synth is available at <https://huggingface.co/datasets/ncbi/MedFact-Synth>. MedFact-Bench is available at <https://huggingface.co/datasets/ncbi/MedFact-Bench>.

PubMed 2025 baseline can be downloaded from <https://ftp.ncbi.nlm.nih.gov/pubmed/baseline/>. Pre-processed SciFact and HealthVer can be downloaded from <https://github.com/dwadden/multivers>. MedAESQA can be downloaded from <https://github.com/deepaknlp/MedAESQA>. PubMedQA can be downloaded from <https://github.com/pubmedqa/pubmedqa>. BioASQ can be downloaded from <http://participants-area.bioasq.org/>.

Code Availability

Our code is available at <https://github.com/ncbi-nlp/Med-V1>. The Med-V1-L3B model is available at <https://huggingface.co/ncbi/Med-V1-L3B>, and the Med-V1-Q3B model is available at <https://huggingface.co/ncbi/Med-V1-Q3B>.

Llama-3.2-3B-Instruct is available at <https://huggingface.co/meta-llama/Llama-3.2-3B-Instruct>. Qwen2.5-3B-Instruct is available at <https://huggingface.co/Qwen/Qwen2.5-3B-Instruct>. We access GPT-5, GPT-4, GPT-4o, o3-mini, and Llama-3.3-70B-Instruct through the application programming interface (API) provided by Microsoft Azure <https://azure.microsoft.com/>, and generation temperature is set to 0 wherever applicable for deterministic outputs.

Author Contributions

Q.J. conceived the study. Q.J. and Z.L. designed the experiments. Q.J., Y.F., Y.Y., and D.C. conducted data collection, model training, and evaluation. Z.W., N.W., J.C., R.L., and C.F. performed data annotations and error analysis. L.H., N.W., C.F., G.X., A.Z., M.C., and Y.P. analyzed case studies. M.C., Y.P., and Z.L. supervised the project. All authors contributed to writing the manuscript and approved the submitted version.

Acknowledgements

This research was supported by the Intramural Research Program of the National Institutes of Health (NIH). The contributions of the NIH author(s) are considered Works of the United States Government. This research was also partially supported by the NIH Pathway to Independence Award K99LM014903 (Q.J.), as well as R01LM014344 (Y.P.) and R01LM014573 (Y.P.). The findings and conclusions presented in this paper are those of the author(s) and do not necessarily reflect the views of the NIH or the U.S. Department of Health and Human Services.

Competing interests

None declared.

References

1. Vladika, J. & Matthes, F. Scientific fact-checking: A survey of resources and approaches. In Rogers, A., Boyd-Graber, J. & Okazaki, N. (eds.) *Findings of the Association for Computational Linguistics: ACL 2023*, 6215–6230, DOI: [10.18653/v1/2023.findings-acl.387](https://doi.org/10.18653/v1/2023.findings-acl.387) (Association for Computational Linguistics, Toronto, Canada, 2023).
2. Guo, Z., Schlichtkrull, M. & Vlachos, A. A survey on automated fact-checking. *Transactions Assoc. for Comput. Linguist.* **10**, 178–206, DOI: [10.1162/tac1_a_00454](https://doi.org/10.1162/tac1_a_00454) (2022). https://direct.mit.edu/tac1/article-pdf/doi/10.1162/tac1_a_00454/1987018/tac1_a_00454.pdf.
3. Wadden, D. *et al.* SciFact-open: Towards open-domain scientific claim verification. In Goldberg, Y., Kozareva, Z. & Zhang, Y. (eds.) *Findings of the Association for Computational Linguistics: EMNLP 2022*, 4719–4734, DOI: [10.18653/v1/2022.findings-emnlp.347](https://doi.org/10.18653/v1/2022.findings-emnlp.347) (Association for Computational Linguistics, Abu Dhabi, United Arab Emirates, 2022).
4. Petroni, F. *et al.* Improving wikipedia verifiability with ai. *Nat. Mach. Intell.* **5**, 1142–1148 (2023).
5. Zuccon, G., Koopman, B. & Shaik, R. Chatgpt hallucinates when attributing answers. In *Proceedings of the Annual International ACM SIGIR Conference on Research and Development in Information Retrieval in the Asia Pacific Region, SIGIR-AP '23*, 46–51, DOI: [10.1145/3624918.3625329](https://doi.org/10.1145/3624918.3625329) (Association for Computing Machinery, New York, NY, USA, 2023).
6. Liu, N., Zhang, T. & Liang, P. Evaluating verifiability in generative search engines. In Bouamor, H., Pino, J. & Bali, K. (eds.) *Findings of the Association for Computational Linguistics: EMNLP 2023*, 7001–7025, DOI: [10.18653/v1/2023.findings-emnlp.467](https://doi.org/10.18653/v1/2023.findings-emnlp.467) (Association for Computational Linguistics, Singapore, 2023).
7. Wu, K. *et al.* An automated framework for assessing how well llms cite relevant medical references. *Nat. Commun.* **16**, 3615 (2025).
8. Jin, Q., Leaman, R. & Lu, Z. Retrieve, summarize, and verify: how will chatgpt affect information seeking from the medical literature? *J. Am. Soc. Nephrol.* **34**, 1302–1304 (2023).
9. Augenstein, I. *et al.* Factuality challenges in the era of large language models and opportunities for fact-checking. *Nat. Mach. Intell.* **6**, 852–863 (2024).
10. Asai, A. *et al.* Openscholar: Synthesizing scientific literature with retrieval-augmented lms. *arXiv preprint arXiv:2411.14199* (2024).
11. Wang, X. *et al.* MedCite: Can language models generate verifiable text for medicine? In Che, W., Nabende, J., Shutova, E. & Pilehvar, M. T. (eds.) *Findings of the Association for Computational Linguistics: ACL 2025*, 18891–18913, DOI: [10.18653/v1/2025.findings-acl.967](https://doi.org/10.18653/v1/2025.findings-acl.967) (Association for Computational Linguistics, Vienna, Austria, 2025).
12. Thirunavukarasu, A. J. *et al.* Large language models in medicine. *Nat. medicine* **29**, 1930–1940 (2023).
13. Clusmann, J. *et al.* The future landscape of large language models in medicine. *Commun. medicine* **3**, 141 (2023).
14. Tian, S. *et al.* Opportunities and challenges for chatgpt and large language models in biomedicine and health. *Briefings Bioinforma.* **25** (2023).
15. Wang, B. *et al.* Pre-trained language models in biomedical domain: A systematic survey. *ACM Comput. Surv.* **56**, 1–52 (2023).
16. Shah, N. H., Entwistle, D. & Pfeffer, M. A. Creation and adoption of large language models in medicine. *Jama* **330**, 866–869 (2023).
17. He, Y. *et al.* Foundation model for advancing healthcare: Challenges, opportunities and future directions. *IEEE Rev. Biomed. Eng.* (2024).
18. Omiye, J. A., Gui, H., Rezaei, S. J., Zou, J. & Daneshjou, R. Large language models in medicine: the potentials and pitfalls: a narrative review. *Annals internal medicine* **177**, 210–220 (2024).
19. Liu, F. *et al.* Application of large language models in medicine. *Nat. Rev. Bioeng.* 1–20 (2025).
20. Singhal, K. *et al.* Large language models encode clinical knowledge. *Nature* **620**, 172–180 (2023).

21. Nori, H. *et al.* Can generalist foundation models outcompete special-purpose tuning? case study in medicine. *arXiv preprint arXiv:2311.16452* (2023).
22. Chen, S. *et al.* The effect of using a large language model to respond to patient messages. *The Lancet Digit. Heal.* **6**, e379–e381 (2024).
23. Wong, C. *et al.* Scaling clinical trial matching using large language models: A case study in oncology. In *Machine Learning for Healthcare Conference*, 846–862 (PMLR, 2023).
24. Jin, Q. *et al.* Matching patients to clinical trials with large language models. *Nat. communications* **15**, 9074 (2024).
25. Wornow, M. *et al.* Zero-shot clinical trial patient matching with llms. *NEJM AI* **2**, A1cs2400360 (2025).
26. Wang, F. *et al.* A survey on small language models in the era of large language models: Architecture, capabilities, and trustworthiness. In *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V. 2*, 6173–6183 (2025).
27. Grattafiori, A. *et al.* The llama 3 herd of models (2024). [2407.21783](#).
28. Qwen *et al.* Qwen2.5 technical report (2025). [2412.15115](#).
29. Wadden, D. *et al.* Fact or fiction: Verifying scientific claims. In Webber, B., Cohn, T., He, Y. & Liu, Y. (eds.) *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 7534–7550, DOI: [10.18653/v1/2020.emnlp-main.609](#) (Association for Computational Linguistics, Online, 2020).
30. Gupta, D., Bartels, D. & Demner-Fushman, D. A dataset of medical questions paired with automatically generated answers and evidence-supported references. *Sci. Data* **12**, 1035 (2025).
31. Sarrouiti, M., Ben Abacha, A., Mrabet, Y. & Demner-Fushman, D. Evidence-based fact-checking of health-related claims. In Moens, M.-F., Huang, X., Specia, L. & Yih, S. W.-t. (eds.) *Findings of the Association for Computational Linguistics: EMNLP 2021*, 3499–3512, DOI: [10.18653/v1/2021.findings-emnlp.297](#) (Association for Computational Linguistics, Punta Cana, Dominican Republic, 2021).
32. Jin, Q., Dhingra, B., Liu, Z., Cohen, W. & Lu, X. PubMedQA: A dataset for biomedical research question answering. In Inui, K., Jiang, J., Ng, V. & Wan, X. (eds.) *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, 2567–2577, DOI: [10.18653/v1/D19-1259](#) (Association for Computational Linguistics, Hong Kong, China, 2019).
33. Krithara, A., Nentidis, A., Bougiatiotis, K. & Paliouras, G. Bioasq-qa: A manually curated corpus for biomedical question answering. *Sci. Data* **10**, 170 (2023).
34. OpenAI. GPT-5 System Card. Tech. Rep., OpenAI (2025). PDF.
35. Sayers, E. W. *et al.* Database resources of the national center for biotechnology information in 2025. *Nucleic acids research* **53**, D20–D29 (2025).
36. Hurst, A. *et al.* Gpt-4o system card. *arXiv preprint arXiv:2410.21276* (2024).
37. Jin, Q. *et al.* Medcpt: Contrastive pre-trained transformers with large-scale pubmed search logs for zero-shot biomedical information retrieval. *Bioinformatics* **39**, btad651 (2023).
38. Jin, Q., Leaman, R. & Lu, Z. Pubmed and beyond: biomedical literature search in the age of artificial intelligence. *EBioMedicine* **100** (2024).
39. Hersh, W. Search still matters: information retrieval in the era of generative ai. *J. Am. Med. Informatics Assoc.* **31**, 2159–2161 (2024).
40. Fiorini, N., Leaman, R., Lipman, D. J. & Lu, Z. How user intelligence is improving pubmed. *Nat. biotechnology* **36**, 937–945 (2018).
41. Pradeep, R., Ma, X., Nogueira, R. & Lin, J. Scientific claim verification with VerT5erini. In Holderness, E. *et al.* (eds.) *Proceedings of the 12th International Workshop on Health Text Mining and Information Analysis*, 94–103 (Association for Computational Linguistics, online, 2021).

42. Wadden, D. *et al.* MultiVerS: Improving scientific claim verification with weak supervision and full-document context. In Carpuat, M., de Marneffe, M.-C. & Meza Ruiz, I. V. (eds.) *Findings of the Association for Computational Linguistics: NAACL 2022*, 61–76, DOI: [10.18653/v1/2022.findings-naacl.6](https://doi.org/10.18653/v1/2022.findings-naacl.6) (Association for Computational Linguistics, Seattle, United States, 2022).
43. Wright, D. *et al.* Generating scientific claims for zero-shot scientific fact checking. In Muresan, S., Nakov, P. & Villavicencio, A. (eds.) *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 2448–2460, DOI: [10.18653/v1/2022.acl-long.175](https://doi.org/10.18653/v1/2022.acl-long.175) (Association for Computational Linguistics, Dublin, Ireland, 2022).
44. Zhang, J., Qian, J., Zhou, Y. & Peng, Y. Enhancing health fact-checking with llm-generated synthetic data (2025). [2508.20525](https://arxiv.org/abs/2508.20525).
45. Guo, D. *et al.* Deepseek-r1 incentivizes reasoning in llms through reinforcement learning. *Nature* **645**, 633–638 (2025).
46. Shao, Z. *et al.* Deepseekmath: Pushing the limits of mathematical reasoning in open language models. *arXiv preprint arXiv:2402.03300* (2024).
47. Wadden, D. *et al.* MultiVerS: Improving scientific claim verification with weak supervision and full-document context. In Carpuat, M., de Marneffe, M.-C. & Meza Ruiz, I. V. (eds.) *Findings of the Association for Computational Linguistics: NAACL 2022*, 61–76, DOI: [10.18653/v1/2022.findings-naacl.6](https://doi.org/10.18653/v1/2022.findings-naacl.6) (Association for Computational Linguistics, Seattle, United States, 2022).
48. Sayers, E. W. *et al.* Database resources of the national center for biotechnology information. *Nucleic acids research* **49**, D10–D17 (2021).
49. Comeau, D. C., Wei, C.-H., Islamaj Doğan, R. & Lu, Z. Pmc text mining subset in bioc: about three million full-text articles and growing. *Bioinformatics* **35**, 3533–3535 (2019).
50. Munn, Z., Stern, C., Aromataris, E., Lockwood, C. & Jordan, Z. What kind of systematic review should i conduct? a proposed typology and guidance for systematic reviewers in the medical and health sciences. *BMC medical research methodology* **18**, 5 (2018).