

Review

Educational Materials for *Helicobacter pylori* Infection: A Comparative Evaluation of Large Language Models Versus Human Experts

Giulia Ortu ¹, Elettra Merola ¹ , Giovanni Mario Pes ¹  and Maria Pina Dore ^{1,2,*} 

¹ Dipartimento di Medicina, Chirurgia e Farmacia, University of Sassari, Clinica Medica, Viale San Pietro 8, 07100 Sassari, Italy; giguortu@gmail.com (G.O.); emerola@uniss.it (E.M.); gmpes@uniss.it (G.M.P.)

² Department of Medicine, Baylor College of Medicine, One Baylor Plaza Blvd, Houston, TX 77030, USA

* Correspondence: mpdore@uniss.it; Tel.: +39-079-229886

Abstract

Helicobacter pylori infects about half of the global population and is a major cause of peptic ulcer disease and gastric cancer. Improving patient education can increase screening participation, enhance treatment adherence, and help reduce gastric cancer incidence. Recently, large language models (LLMs) such as ChatGPT, Gemini, and DeepSeek-R1 have been explored as tools for producing patient-facing educational materials; however, their performance compared to expert gastroenterologists remains under evaluation. This narrative review analyzed seven peer-reviewed studies (2024–2025) assessing LLMs' ability to answer *H. pylori*-related questions or generate educational content, evaluated against physician- and patient-rated benchmarks across six domains: accuracy, completeness, readability, comprehension, safety, and user satisfaction. LLMs demonstrated high accuracy, with mean accuracies typically ranging from approximately 77% to 95% across different models and studies, and with most models achieving values above 90%, comparable to or exceeding that of general gastroenterologists and approaching senior specialist levels. However, their responses were often judged as incomplete, described as “correct but insufficient.” Readability exceeded the recommended sixth-grade level, though comprehension remained acceptable. Occasional inaccuracies in treatment advice raised safety concerns. Experts and medical trainees rated LLM outputs positively, while patients found them less clear and helpful. Overall, LLMs demonstrate strong potential to provide accurate and scalable *H. pylori* education for patients; however, heterogeneity between LLM versions (e.g., GPT-3.5, GPT-4, GPT-4o, and various proprietary or open-source architectures) and prompting strategies results in variable performance across studies. Enhancing completeness, simplifying language, and ensuring clinical safety are key to their effective integration into gastroenterology patient education.

Keywords: *Helicobacter pylori*; large language models (LLMs); patient education; gastroenterology; artificial intelligence in healthcare



Academic Editor: Isidoros Perikos

Received: 14 October 2025

Revised: 25 November 2025

Accepted: 26 November 2025

Published: 28 November 2025

Citation: Ortu, G.; Merola, E.; Pes, G.M.; Dore, M.P. Educational Materials for *Helicobacter pylori* Infection: A Comparative Evaluation of Large Language Models Versus Human Experts. *AI* **2025**, *6*, 311. <https://doi.org/10.3390/ai6120311>

Copyright: © 2025 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

Helicobacter pylori infection can lead to gastritis, peptic ulcer, mucosa-associated lymphoid tissue (MALT) lymphoma, and is a well-established risk factor for gastric adenocarcinoma [1]. Given that *H. pylori* contributed to approximately 4.8% of global cancer incidence in 2018 (excluding non-melanoma skin cancers), controlling this infection is a public health priority. Early detection and eradication of *H. pylori* have been shown to

reduce the incidence of gastric cancer [2]. The success of eradication programs partly depends on public awareness and engagement. Unfortunately, knowledge of *H. pylori* in the general population is often poor, and individuals with low awareness are less likely to undergo testing or adhere to treatment. This gap highlights the importance of effective patient education in improving disease outcomes [3].

Patient education initiatives for *H. pylori* aim to clearly and understandably convey information on transmission, risks, testing, and treatment to encourage informed decision-making and adherence [4]. Traditional patient education materials (pamphlets, websites) require significant efforts by experts to create content that is accurate, comprehensive, and pitched at the right literacy level [5]. In recent years, advances in artificial intelligence (AI) have led to the development of new tools for generating health information. Large language models (LLMs) like OpenAI's ChatGPT can produce human-like text and have been explored in various medical applications (e.g., drafting medical notes and assisting clinical decision support) [6]. In gastroenterology, there is growing interest in using LLMs to address patient questions and provide guidance on conditions such as inflammatory bowel disease and *H. pylori* infection. LLMs offer the potential for an on-demand, scalable approach to disseminate health information, but their reliability and quality must be rigorously evaluated before clinical integration [7].

Early studies evaluated LLM-generated information on *H. pylori*. They examined content quality on multiple fronts, including factual accuracy, completeness of information, readability for the average patient, and safety (absence of misleading or harmful advice). They also gauge how well patients understand the AI-provided information and how satisfied users are with the answers, compared to traditional expert-derived content.

Importantly, successive editions of ChatGPT (e.g., GPT-3.5, GPT-4, GPT-4o) and other LLMs differ in training data, architecture, and alignment procedures, which can materially influence the accuracy, completeness, and style of their responses. Evaluating individual model versions, rather than treating "ChatGPT" as a single entity, is therefore essential for a nuanced understanding of performance.

Several recent narrative and umbrella reviews have examined the broader use of LLMs in healthcare, including gastroenterology research and clinical practice [6,7]. However, to our knowledge, no prior review has focused explicitly on LLM-generated patient-educational materials for *H. pylori* infection. Given the high prevalence of *H. pylori*, its established role in gastric carcinogenesis, and the availability of multiple empirical studies directly comparing LLM-generated and gastroenterologist-written materials in this domain, *H. pylori* provides a timely and clinically relevant case study to explore the opportunities and limitations of AI-based patient education.

In this context, "educational materials" in the present article are defined as written or text-based resources intended primarily for patients or laypersons, rather than for the education of medical students, residents, or other healthcare professionals.

This review provides a summary analysis of LLM-generated *H. pylori* educational content in comparison with content written by gastroenterologists. We focus on six key quality domains: (i) accuracy (correctness of information); (ii) completeness (breadth and depth of content); (iii) readability (reading level and ease of understanding); (iv) patient comprehension (how well target audiences understand the material); (v) safety (freedom from harmful or misleading information); and (vi) user satisfaction. By synthesizing results from recent studies, we aim to determine whether LLMs can serve as reliable patient education tools in *H. pylori* infection and identify any necessary improvements or safeguards to enhance their usefulness in clinical practice.

The present article is conceived as a narrative, non-systematic review, providing an integrative synthesis of a rapidly emerging evidence base rather than a systematic meta-analysis.

2. Methods

2.1. Study Design and Scope

The primary objective was to synthesize recent evidence on the quality of LLM-generated educational materials on *H. pylori* infection, compared with content produced by human gastroenterology experts. Our focus was explicitly on patient-facing information rather than on educational resources for medical trainees. Given the limited number of available studies and their methodological heterogeneity, the review does not follow the PRISMA guidelines and does not conduct a formal meta-analysis. Instead, we provide a structured qualitative synthesis complemented by descriptive summaries.

2.2. Search Strategy

We searched PubMed/MEDLINE and Google Scholar for English-language, peer-reviewed publications from January 2023 to January 2025. The following keyword combinations were used: (“*Helicobacter pylori*” OR “*H. pylori*”) AND (“large language model” OR “LLM” OR “ChatGPT” OR “artificial intelligence”) AND (“patient education” OR “educational materials” OR “information” OR “counseling”). The reference lists of relevant articles were also manually screened to identify additional studies.

2.3. Eligibility Criteria and Study Selection

Studies were eligible for inclusion if they met the following criteria: (i) empirical research articles or letters reporting original data; (ii) evaluation of at least one LLM (e.g., ChatGPT-3.5, ChatGPT-4, ChatGPT-4o, Gemini, ERNIE Bot, DeepSeek) in answering *H. pylori*-related questions or generating *H. pylori* educational materials; (iii) comparison of LLM-generated content with information provided by clinicians (e.g., gastroenterologists) or established references; and (iv) assessment of at least one of the six predefined domains: accuracy, completeness, readability, comprehension, safety, or user satisfaction.

Exclusion criteria were: non-empirical articles (e.g., commentaries, editorials without original data), conference abstracts without complete data, technical AI papers without patient-facing content, and studies not specific to *H. pylori*.

Two authors (G.O., M.P.D.) independently screened titles and abstracts for relevance, followed by full-text assessment of potentially eligible articles. Discrepancies were resolved by discussion among the authors. The final selection comprised seven studies (six original research articles and one letter), which are summarized in Table 1. Because of differences in study design, outcome definitions, scoring scales, and prompting strategies, a formal PRISMA flow diagram and meta-analytic pooling were not undertaken.

Table 1. Studies included in the review.

Study	Type of Study	AI Models Analyzed	Raters/Respondents
“Assessing Accuracy of ChatGPT on Addressing <i>Helicobacter pylori</i> Infection-Related Questions: A National Survey and Comparative Study” (Hu et al., 2024) [8]	National survey and comparative study	ChatGPT-3.5 and ChatGPT-4	Gastroenterologists (national sample; experience level variably reported)

Table 1. Cont.

Study	Type of Study	AI Models Analyzed	Raters/Respondents
“Comparative analysis of large language models in medical counseling: A focus on <i>Helicobacter pylori</i> infection” (Kong et al., 2024) [9]	Comparative analysis	ChatGPT-4, ChatGPT-3.5, ERNIE Bot 4.0	Gastroenterologists and physicians as expert raters
“Exploring the capacities of ChatGPT: A comprehensive evaluation of its accuracy and repeatability in addressing <i>Helicobacter pylori</i> -related queries” (Lai et al., 2024) [10]	Observational study	ChatGPT-3.5	Gastroenterologists (expert panel)
“Artificial Intelligence-Generated Patient Education Materials for <i>Helicobacter pylori</i> Infection: A Comparative Analysis” (Zeng et al., 2024) [11]	Comparative analysis	Bing Copilot, Claude 3 Opus, Gemini Pro, ChatGPT-4, ERNIE Bot 4.0	Gastroenterologists and non-expert patients
“Assessing the Capabilities of Novel Open-Source Artificial Intelligence—DeepSeek in <i>Helicobacter pylori</i> -Related Queries” (Du et al., 2025) [12]	Letter to the editor (comparative analysis)	DeepSeek (V3 and R1), ChatGPT-4o and OpenAI-o1	Gastroenterologists (expert raters)
“Potential application of ChatGPT in <i>Helicobacter pylori</i> disease relevant queries” (Gao et al., 2024) [13]	Evaluation study	ChatGPT-4	Gastroenterologists, medical students, and laypersons
“Comparative evaluation of the accuracy and reliability of ChatGPT versions in providing information on <i>Helicobacter pylori</i> infection” (Ye et al., 2025) [14]	Comparative evaluation study	ChatGPT-3.5, ChatGPT-4, ChatGPT-4o	Gastroenterologists (expert raters)

2.4. Data Extraction and Synthesis

For each included study, we extracted information on: study design; LLMs and versions evaluated; comparators; type of raters or respondents (e.g., board-certified gastroenterologists, trainees, patients, laypersons); language(s); and outcomes related to accuracy, completeness, readability, comprehension, safety, and user satisfaction. Particular attention was paid to the prompts used to query the LLMs (e.g., instructions regarding reading level or language), as these may influence performance. Given the heterogeneity of metrics (e.g., different Likert scales, percentages of correct answers, various readability indices), results were synthesized qualitatively (Table 1).

These studies collectively examined multiple LLM platforms (including various versions of OpenAI’s ChatGPT, such as GPT-3.5, GPT-4, GPT-4.5, GPT4o, and other models like ERNIE Bot and DeepSeek, Gemini, etc.). They often included multiple languages (typically English and Chinese). For each study, we extracted data pertaining to the six predefined domains of interest: accuracy, completeness, readability, patient comprehension, safety, and user satisfaction (for both patients and providers). Due to heterogeneity in study designs and measurement scales, a meta-analytic quantitative synthesis was not feasible; instead, results were integrated qualitatively. Key findings from each domain

were compared and summarized, emphasizing consistent patterns or notable discrepancies between LLM- and expert-generated content. We ensured that the source data from these studies supported any interpretative claims. All data presented are reported as originally stated in the studies, with representative examples and statistical results cited to illustrate each point. No additional experimental data were generated for this review.

3. Results

3.1. Accuracy

Across all studies, LLM-generated content on *H. pylori* demonstrated high factual accuracy, often approaching or matching that of expert gastroenterologists (Supplementary Table S1). In a national-based survey, ChatGPT (both GPT-3.5 and GPT-4 versions) answered *H. pylori*-related clinical questions correctly 92% of the time (median accuracy), outperforming the average accuracy of 1279 gastroenterologists ($\approx 80\%$) on the same questions. Notably, ChatGPT's accuracy was comparable to that of senior subspecialists for many topics. Several independent evaluations corroborated GPT-4 as the most accurate LLM [8]. For example, one study that scored answers on a 5-point scale found that ChatGPT-4o had the highest average accuracy ($\sim 4.7/5$), significantly above older versions [9]. Similarly, the Chinese LLM DeepSeek-R1, in a letter-based study by Du et al., achieved 95.2% accuracy, outperforming ChatGPT-4o and OpenAI-o1 [10]. LLM accuracy did vary by content area. Certain knowledge domains (e.g., basic facts, indications for testing) were handled quite accurately by ChatGPT, whereas nuanced clinical management questions were more challenging. In Lai et al., 61.9% of ChatGPT's answers were rated completely correct, and an additional 33.3% accurate but inadequate, leaving only $\sim 5\%$ of answers outright incorrect or containing errors. The few errors primarily involved outdated treatment recommendations or misinterpretation of complex scenarios [11]. In addition, variations were observed depending on the language chosen for the study. In one investigation analyzing three models (ChatGPT-4, ChatGPT-3.5, and ERNIE Bot 4.0) across English and Chinese tasks, both languages achieved satisfactory performance ($\sim 90\%$), yet a discrepancy was observed (91.1% in English vs. 88.9% in Chinese). Despite this difference, no statistically significant variation was found among the LLMs. Notably, ChatGPT-3.5 recommended serological testing for post-treatment follow-up, which is in discordance with current clinical guidelines [12]. Similarly, Zeng et al. assessed Patient Educational Materials (PEMs) generated in both English and Chinese by five LLMs (ChatGPT-4, ChatGPT-3.5, Claude 3 Opus, Gemini Pro, ERNIE Bot) and by a physician. All responses were considered acceptable across both LLMs and physicians, except for the Chinese outputs produced by Claude 3 Opus [13]. Consistent with Kong et al., English-language outputs demonstrated overall superior performance [12]. Overall, these data indicate that current LLMs can deliver predominantly accurate *H. pylori* information, echoing findings that ChatGPT performs at or above the level of practicing clinicians on knowledge-based queries. However, "accuracy" here assumes static factual queries; dynamic clinical decision accuracy (e.g., choosing optimal therapy) remains an area for caution, as discussed under the Section 3.5. Variability in prompting strategies across studies may partly explain differences not only in accuracy but more prominently in completeness, as discussed below. Figure 1 provides a descriptive bar chart summarizing the approximate range of accuracy values for the main models evaluated to facilitate visual comparison across heterogeneous metrics (Figure 1).

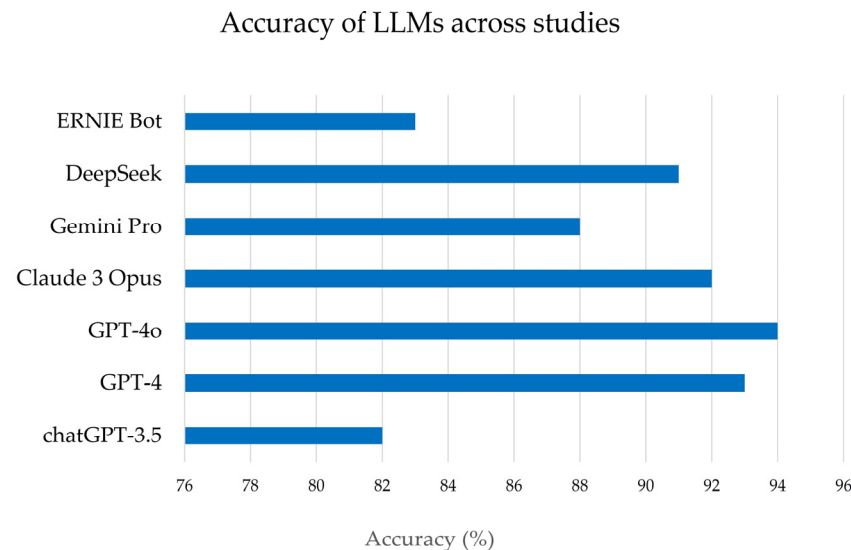


Figure 1. Descriptive comparison of accuracy values across the large language models (LLMs) evaluated in the included studies. Accuracy is expressed as the approximate percentage of correct responses provided by each model based on published assessments.

3.2. Completeness

In contrast to accuracy, the completeness of LLM-generated educational content was consistently identified as a weakness (Supplementary Table S2). Across studies, prompts varied in terms of specificity, requested level of detail, and explicit instructions to adapt language to a sixth-grade reading level. Such differences in prompt design likely contributed to variation in both the length and completeness of LLM responses. Completeness refers to whether the content covers all relevant aspects of the topic with sufficient depth and context. Multiple studies have found that LLM responses often omit specific details or caveats that expert-written answers would typically include. Kong et al. observed that while 90% of LLM answers met the accuracy thresholds, only 45.6% were judged to be sufficiently complete (using a $\geq 2/3$ Likert completeness criterion) [9]. Similarly, expert raters in Zeng et al. [11] reported that LLM-generated patient education materials were frequently missing some content elements, rating most LLM outputs as “unsatisfactory” in completeness by gastroenterologist standards. For instance, in the Chinese-language materials, four of five AI-generated brochures had mean completeness scores < 2 (on a 3-point scale) when evaluated by specialists, indicating that important information was lacking. Interestingly, the physician-written material was also not perfectly comprehensive in every case. In the English versions, completeness scores for the doctors’ and AI’s brochures were statistically comparable, suggesting that even experts might condense information. Patients tended to perceive the information as more complete than the experts did. In this study, lay patients received slightly higher completeness scores on average than gastroenterologists for the same AI outputs [11]. This discrepancy suggests that non-expert readers may not recognize the nuances that are missing. Nonetheless, from a clinical perspective, specific LLM answers were too superficial. Common gaps included a lack of detail on *H. pylori* transmission prevention, incomplete explanations of diagnostic steps, and insufficient emphasis on follow-up and antibiotic resistance issues [8,9,13]. In Gao et al., experts still gave ChatGPT-4 high completeness marks ($\sim 2.8/3$), likely because the prompts in that study were designed around guideline topics [13]. Yet, even there, the consensus was that more exhaustive coverage would be beneficial. Overall, ensuring adequacy of content remains a challenge. LLMs may require more effective prompting or iterative querying to elicit all key information for patients. Improving completeness is critical because patients need not just correct facts, but a whole picture of their condition

and care. Figure 2 provides descriptive bar plots of completeness scores across models, derived from the original Likert-based ratings.

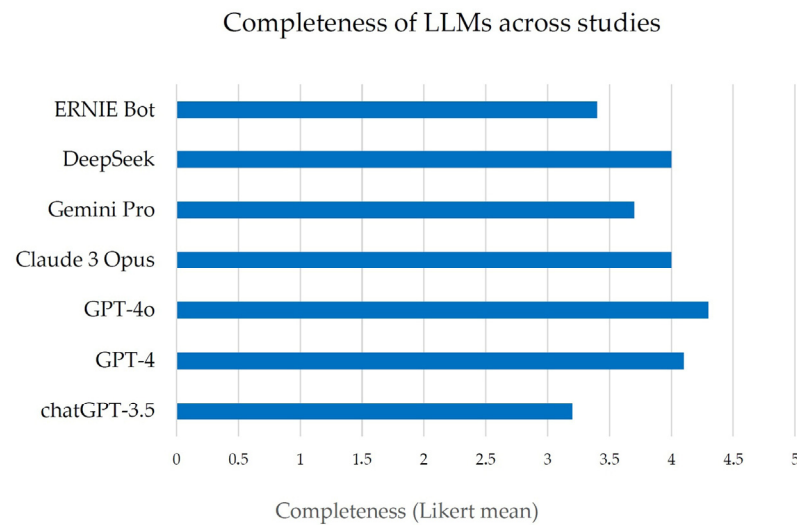


Figure 2. Descriptive comparison of completeness scores for large language models (LLMs) across the included studies. Completeness reflects the mean Likert-scale ratings assigned by expert reviewers to evaluate the extent to which each response covered all relevant clinical information. Variability in scoring systems and prompting strategies among studies limits direct comparability; the figure provides a qualitative visualization of general trends.

3.3. Readability

The readability of the educational content, specifically the reading grade level and textual clarity, was another key point of comparison (Supplementary Table S3). An ideal patient handout should be written at roughly a 6th-grade reading level (as per American Medical Association and other health literacy guidelines) to be easily understood by the majority of adults. Both LLM- and expert-generated materials often failed to meet this benchmark [11]. Zeng et al. explicitly tested for reading level and found that none of the patient education documents (neither AI-generated nor physician-written) achieved the 6th-grade level; all were more complex than recommended. This held even though the prompt to the LLMs had specifically requested a “sixth-grade reading level” [11]. Quantitatively, in Ye et al., the readability of responses generated by ChatGPT 3.5, 4, and 4o was evaluated through expert review and quantitative analysis. Outputs were assessed based on word count and standard readability metrics, including the Flesch Reading Ease (FRE) and Flesch-Kincaid Grade Level (FKGL). ChatGPT 4o produced the longest responses, followed by ChatGPT 4 and 3.5. While ChatGPT 4 showed numerically “superior” (FRE~25), the differences among models were not statistically significant [14]. Similarly, an open-source model, DeepSeek, produced particularly verbose and complex answers that were harder to read. Evaluators noted DeepSeek’s responses “required significant simplification for layperson comprehension” [12]. In general, while LLMs employ a conversational tone that avoids overly technical jargon, they still often produce sentences and vocabulary that exceed the level of a middle-school reader. This highlights the need to simplify the language further or to use health literacy tools. No matter how accurate an educational passage is, if many patients find it linguistically challenging, its utility is diminished. Both AI developers and medical communicators should prioritize readability optimization to ensure that *H. pylori* educational content is accessible to patients with varying literacy levels. “You do not really understand something unless you can explain it to your grandmother,” as the apocryphal aphorism sound.

3.4. Patient Comprehension

Readability metrics offer an objective measure of text complexity, but actual patient comprehension remains the ultimate test of practical education (Supplementary Table S4). Several studies directly assessed how well people (patients or non-experts) understood information provided by LLMs, compared with experts. Overall, initial evidence suggests that properly written content, whether generated by AI or physicians, can be understood by patients at least in a controlled setting; however, comprehension can vary depending on the audience's background knowledge. In Zeng et al. [11], 50 patients were asked to rate the comprehensibility of *H. pylori* brochures (without knowing which source they came from). All materials, including those generated by LLMs, were rated as satisfactorily understandable (≥ 2 on a 3-point ease-of-understanding scale). There was virtually no difference in median comprehension scores between AIs and human-written brochures in that study, indicating that, from a patient's perspective, the clarity was adequate [11]. These findings are consistent with other reports. Kong et al. documented high performance in this domain, with ChatGPT 3.5, ChatGPT 4, and ERNIE Bot achieving 100% of responses rated as sufficiently comprehensible in both English and Chinese. However, it is essential to note that in this study, evaluators were not patients or laypersons, but rather physicians [9]. Similarly, Lai et al. confirmed the practical applicability of ChatGPT 3.5 by stating that it "can provide correct answers to the majority of *H. pylori*-related queries." Moreover, they highlighted that "It exhibited good reproducibility and delivered responses that were easily comprehensible to patients." Notably, this study did not apply a dedicated evaluation scale [10]. Gao et al., however, highlighted a critical nuance: when "ordinary people" (14 laypersons) were asked to evaluate ChatGPT-4's answers, their scores for comprehension were significantly lower than the scores given by medical experts reviewing the same answers. In that study, experts rated ChatGPT's responses nearly perfect for comprehension (mean~2.95/3), whereas non-medical participants gave lower ratings (mean~2.08/3) for how easy the answers were to understand. Medical students' ratings fell between the experts' views, closer to the experts' views. Consistently, the authors noted that "Some individuals thought that these answers were too obscure and lacked significance" and further stated that "those who have medical knowledge, who have a certain knowledge base, can more easily understand the answers" [13]. This suggests that individuals with some medical knowledge found the AI explanations clear. Still, those without such a background had more difficulty. The likely reason is that although the text is grammatically clear, truly grasping concepts like "antibiotic resistance" or "urea breath test" requires some baseline knowledge. Another qualitative finding was that patients with lower levels of education or those with low health literacy might struggle with key terms or longer answers. Du et al., in a study comparing the Chinese LLM DeepSeek (versions R1 and V3) with ChatGPT-4o and OpenAI o1, noted that despite improvements in the completeness of information in Chinese AI, outputs remain "overly complex, limiting usability for non-expert audiences" [12]. Thus, while patient comprehension of LLM-generated content can be high in many cases (especially when patients are relatively educated or the content is simplified), there is a risk that a subset of patients will misinterpret or fail to absorb the information fully. Ensuring comprehension may involve supplementing AI-generated text with illustrations, utilizing interactive Q&A to clarify misunderstandings, or tailoring explanations based on patient feedback. In future trials, it will be essential to measure knowledge gain or recall after patients use AI-provided education to quantify comprehension outcomes directly.

3.5. Safety

Any tool providing medical information must be evaluated for safety, i.e., the absence of dangerous misinformation or advice that could harm patients if followed. The studies

reviewed did not identify overtly dangerous instructions in response to *H. pylori* queries; nonetheless, they did report subtle inaccuracies and misleading information that could pose risks if left uncorrected. Zeng et al. had gastroenterologists conduct a structured safety review of each educational material, evaluating the likelihood and severity of potential harm arising from content errors. While the materials produced by both physicians and several LLMs (e.g., Bing, Gemini, ERNIE Bot) were generally judged as safe and unlikely to cause serious patient harm, some outputs raised more significant concerns. In particular, Ernie Bot was considered the safest among the LLMs, with 100% of responses classified as “No harm”, followed by Gemini Pro, which identified “Potentially mild to moderate harm” in only 20% of its Chinese responses. In contrast, Claude 3 Opus exhibited the most significant risk, with 60% of its English responses classified as “Potentially mild to moderate harm” and 20% of its Chinese responses classified as “Definitely mild to moderate harm.” These findings were attributed mainly to inaccuracies and insufficient precision in the generated content, underscoring the importance of systematic safety evaluation across different models [11]. Lai et al. reported that 16.6% of ChatGPT’s answers in the treatment domain contained a mix of correct and outdated information, such as recommending antibiotic regimens no longer considered first-line due to resistance patterns [10]. Similarly, Ye et al. highlighted notable errors, including ChatGPT-3.5 suggesting symptom relief as evidence of eradication and ChatGPT-4/4o recommending amoxicillin-containing regimens despite penicillin allergy [14]. Such inaccuracies, while not always overtly harmful, could nonetheless be dangerous and underscore the importance of up-to-date, guideline-consistent information in patient care. All authors stressed the importance of human oversight. Therefore, LLMs should be used with caution and ideally have their medical content reviewed or augmented by clinicians to catch subtle mistakes. As LLM deployment expands, implementing safety checks (for example, integrating medical knowledge bases or citing sources in answers) will be key to maintaining a high safety profile.

User Satisfaction: Ultimately, user satisfaction with the educational content is a crucial outcome, as it may influence whether patients trust and utilize the information. Because LLMs can provide information in a conversational format, one hypothesis is that patients might find this format engaging. Empirical data on satisfaction are still limited, but early indications are generally positive among healthcare professionals and mixed among patients. In Gao et al., expert gastroenterologists rated their satisfaction with ChatGPT-4’s answers at 4.55 out of 5 on average, indicating that the specialists were delighted with the quality of the AI-generated responses. These experts also rated the content’s usefulness highly (mean~2.83/3), suggesting they felt the information provided would be helpful to patients. In the same study, medical students also gave positive evaluations of ChatGPT’s answers, aligning with the notion that the content was educationally valuable. However, among ordinary laypeople, satisfaction-related metrics were more tempered. Founding that non-medical participants gave significantly lower scores for the “usefulness” of ChatGPT’s answers compared to experts [13]. This could reflect differences in expectations or understanding; if parts of the answer were not fully grasped, a layperson might not feel satisfied. Kong et al. did not formally measure patient satisfaction; however, in their discussion, they emphasized that real-world patient satisfaction remains to be studied and may depend on factors such as a person’s educational background and health context [9]. No study to date has reported on long-term satisfaction (for instance, whether patients would choose an AI tool again or recommend it to others), as most were one-time evaluations. It is also worth noting that none of the studies offered patients a choice between an AI-generated and a doctor-written brochure to determine which they preferred; such comparisons in the future could be enlightening. In summary, initial satisfaction levels with LLM-provided *H. pylori* information appear high when judged by content experts and reasonably good, but not

uniformly excellent, when judged by lay users. Bridging this gap by improving content tailoring and addressing comprehension issues could enhance patient satisfaction. After all, an accurate brochure is only valuable if the patient feels it answered their concerns helpfully. As LLMs become more user-aware (through improved prompt engineering or interactive clarification), we expect user satisfaction to improve; however, it will be vital to continue capturing patient feedback in deployments.

4. Discussion

This comparative analysis suggests that LLMs have significant potential to complement gastroenterologists in providing patient education about *H. pylori* infection; however, critical challenges must be addressed before LLM-generated content can be adopted in practice. Accuracy emerged as a clear strength of modern LLMs. The reviewed studies uniformly show that these models can deliver factually correct answers to *H. pylori* questions at a level comparable to clinicians. This high accuracy is consistent with reports that most of the AIs under study have passed medical exams, suggesting that, on average, patients querying an advanced LLM are likely to receive generally reliable information. Such a capability could be invaluable in settings where physicians are not readily available to answer every question. For instance, patients often have numerous concerns about *H. pylori* (ranging from transmission to diet to treatment side effects) that they may not fully address during a brief clinic visit. An LLM-based tool could provide immediate, accurate answers as a supplement to the physician's advice, potentially improving patient understanding and reducing anxiety. However, accuracy alone is not sufficient. The completeness of information is where LLM responses currently fall short in comparison to a thorough consultation or a well-crafted pamphlet by an expert. In practice, an incomplete answer can be as problematic as an incorrect one when critical guidance is omitted. The observation that LLMs frequently provide only partial answers (e.g., lacking detail on follow-up testing or omitting certain risk factors) underscores the risk that patients using these tools might encounter knowledge gaps. A patient might recognize that *H. pylori* causes ulcers and can be treated with antibiotics (information that an LLM is likely to provide), but not realize that family members should also be tested, or that antibiotic resistance could affect therapy, nuances that an expert would emphasize [13]. One strategy to improve completeness is better prompting: clinicians or developers could design structured prompts that ensure all key topics are covered. Another approach is an interactive Q&A, where the AI can prompt the user to learn about related issues (for example, "Would you like to hear about how to prevent reinfection?"). Until such solutions are implemented, it may be advisable for any AI-derived content to undergo review by a healthcare provider to identify and address any issues before the material is provided to patients [9].

Regarding readability and comprehension, our review underscores that current AI-generated content is not yet optimized for all patient populations. The fact that none of the evaluated materials met the target reading level for 6th graders is a call to action for both AI technology and health communication practices. Even the gastroenterologist-written materials were linguistically too advanced, a known challenge in patient education; explaining medical concepts in elementary language is difficult [11]. LLMs, with proper training or constraints, might be able to simplify language more consistently than busy clinicians. Future LLM development could focus on a "patient-friendly mode" that prioritizes shorter sentences, familiar words, and clear definitions of medical terms. Additionally, the multilingual capabilities of LLMs are a considerable asset: these models can instantly produce content in multiple languages, a task that would require significant human resources and time. This can help bridge language barriers and reach patient groups who speak different languages. Ensuring readability in each language (not just direct translation) is essen-

tial [9,11]. Our findings on patient comprehension, especially the gap between experts and laypeople in perceiving clarity, suggest that involving actual patients in the development and testing of LLM-based tools is vital. By observing where non-experts get confused, developers can tweak the AI's explanations. The AI might also incorporate visual aids and analogies to enhance understanding.

In terms of safety, although most LLMs generally adhered to clinical guidelines, some studies did report a certain level of “harmful” responses. This was primarily due to outdated or incomplete information, such as treatment recommendations that are no longer considered adequate (for example, using an antibiotic regimen that has become ineffective), which can lead to treatment failure. Nonetheless, in several studies, ChatGPT consistently included a recommendation to consult a physician alongside the information provided, which may help mitigate potential risks. Thus, minimizing the dissemination of outdated or partial content remains crucial [8,10,13]. One solution is to continually update LLM knowledge bases with current clinical guidelines, although models like GPT-4 are not easily updatable in real time [13]. Future systems may incorporate live data or retrieve information from trusted databases. Another safety measure is transparency: if the AI provides citations or sources (as some LLM-based medical assistants are starting to do), patients and providers can verify the information against reputable references. Ultimately, we envision that LLM-generated patient education will not operate in isolation but rather under a framework of human-AI collaboration: clinicians could supervise the content, or the AI could triage questions and draft answers that a clinician then reviews for accuracy and safety. Such a model would harness the efficiency of AI while prioritizing patient well-being.

User satisfaction and engagement are essential for the practical success of any patient-facing tool. Early positive feedback from physicians and trainees suggests that, if the content is of high quality, healthcare professionals are willing to trust and even recommend these AI resources. This is important, as doctors could serve as facilitators. A doctor might guide a patient to use a vetted chatbot for follow-up questions at home. However, lukewarm responses from some lay users indicate a need to improve the user experience. Patients will be satisfied not just with correct answers, but with the feeling that their concerns were addressed [9]. LLMs can adopt a conversational style that can be friendly and empathetic, but they currently lack the true personalization and emotional intelligence of a human provider. Future development could incorporate more adaptive responses, in which the AI asks the user whether the answer was helpful or if they have other concerns, thereby mimicking dialog with a doctor. Moreover, some patients might distrust information from an “algorithm.” Building trust will require showing that the AI's information is endorsed or co-developed by medical experts and that it has been tested in real patient populations with good outcomes. Over time, as patients become more accustomed to digital health tools, their satisfaction is likely to increase, provided the information is reliable and comprehensible. An often-mentioned benefit is the 24/7 availability of LLM-based assistance; patients can get answers at any time, which could improve satisfaction in the context of anxiety (e.g., a patient worrying at night about their *H. pylori* test results might consult the AI for immediate information on what to expect).

Despite ongoing technological advances, the core evaluation domains: accuracy, completeness, readability, comprehension, safety, and user satisfaction, remain essential benchmarks. Methodologically, a recurrent limitation across studies is the predominant use of Likert-type scales to assess accuracy, completeness, and comprehensibility. Although convenient, these instruments are inherently subjective; variation in rater interpretation, ceiling effects, and inconsistent scale formats undermine comparability across studies. Moreover, several investigations did not specify whether raters were blinded to the source

of the material (LLM vs. clinician) or report inter-rater reliability. Small sample sizes, single-center designs, and limited information on the representativeness of patient cohorts further restrict generalizability. Collectively, these limitations highlight the need for more rigorous, blinded, and standardized evaluation protocols.

Additional heterogeneity stems from prompt design. Some studies used detailed instructions regarding reading level, content scope, or language, whereas others relied on broad, open-ended queries. Because prompts can substantially shape LLM outputs, affecting factual content, length, structure, and clarity, comparisons between models tested with different prompts must be interpreted cautiously. Future research should fully report and, ideally, standardize prompts to isolate model-specific differences better.

Similar concerns have been raised in other specialties, for example, in dermatology, where AI-generated case-based questions lacked depth and educational nuance compared with expert-written material, highlighting the importance of careful oversight when generating any form of educational content [15].

Finally, although existing studies assessed informational quality and user perceptions, none evaluated downstream clinical outcomes such as treatment adherence, completion of eradication regimens, follow-up testing, or long-term rates of peptic ulcer recurrence and *H. pylori*-associated gastric cancer. While more accurate and comprehensible educational materials may, in theory, improve adherence and clinical outcomes, these benefits remain largely speculative. Prospective, real-world studies are required to determine whether appropriately supervised, workflow-integrated LLM-based educational interventions yield measurable improvements in patient knowledge, decision-making, adherence, and health outcomes.

5. Conclusions

In this narrative review, we synthesized evidence from 7 recent studies evaluating LLM-generated educational materials for *H. pylori* infection compared with content produced by gastroenterology experts. Overall, contemporary LLMs demonstrated high accuracy, often approaching that of practicing specialists, and generally acceptable patient comprehensibility, indicating that these tools can effectively address many common patient questions. By generating standardized, evidence-based information across languages and regions, LLMs have the potential to markedly expand access to reliable educational resources.

Nevertheless, important limitations persist. LLM-generated responses were frequently judged to be only partially complete, written above recommended patient reading levels, and occasionally inconsistent with guideline-based management, particularly regarding treatment. Such issues were exacerbated by methodological heterogeneity across studies, including small sample sizes, subjective Likert-scale ratings, variable prompting strategies, and insufficient blinding of evaluators. These factors temper the overall strength of the available evidence.

Taken together, current data suggest that LLMs may serve as valuable adjuncts to traditional gastroenterology care [16], empowering patients with greater knowledge, reinforcing physicians' counseling, and potentially improving management of *H. pylori* infection. However, fully realizing this promise will require deliberate attention to completeness, readability, safety, and continuous expert oversight to maintain alignment with evolving clinical guidelines.

Looking forward, high-quality prospective studies are essential to determine whether LLM-based patient education leads to measurable clinical benefits, such as improved knowledge retention, reduced decisional conflict, enhanced treatment adherence, and

better real-world outcomes. Establishing these effects will be crucial for defining the clinical role of this rapidly advancing technology.

Supplementary Materials: The following supporting information can be downloaded at <https://www.mdpi.com/article/10.3390/ai6120311/s1>; Table S1: Evaluation of Accuracy. Comparison between educational materials on *Helicobacter pylori* generated by large language models (LLMs), such as ChatGPT, and those produced by human gastroenterology experts; Table S2: Evaluation of Completeness. Comparison between educational materials on *Helicobacter pylori* generated by large language models (LLMs), such as ChatGPT, and those produced by human gastroenterology experts. Table S3: Evaluation of Readability. Comparison between educational materials on *Helicobacter pylori* generated by large language models (LLMs), such as ChatGPT, and those produced by human gastroenterology experts; Table S4: Evaluation of Comprehensibility. Comparison between educational materials on *Helicobacter pylori* generated by large language models (LLMs), such as ChatGPT, and those produced by human gastroenterology experts.

Author Contributions: Conceptualization, G.O. and M.P.D.; methodology, M.P.D. and G.M.P.; software, G.M.P.; validation, G.M.P., E.M. and G.M.P.; formal analysis, G.O. and G.M.P.; investigation, M.P.D.; resources, M.P.D.; data curation, M.P.D.; writing—original draft preparation, M.P.D.; writing—review and editing, M.P.D., E.M. and G.M.P.; visualization, G.M.P.; supervision, M.P.D.; project administration, M.P.D.; funding acquisition, M.P.D. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: Not applicable.

Conflicts of Interest: The authors declare no conflicts of interest.

References

1. Bray, F.; Laversanne, M.; Sung, H.; Ferlay, J.; Siegel, R.L.; Soerjomataram, I.; Jemal, A. Global cancer statistics 2022: GLOBOCAN estimates of incidence and mortality worldwide for 36 cancers in 185 countries. *CA Cancer J. Clin.* **2024**, *74*, 229–263. [[CrossRef](#)] [[PubMed](#)]
2. de Martel, C.; Georges, D.; Bray, F.; Ferlay, J.; Clifford, G.M. Global burden of cancer attributable to infections in 2018: A worldwide incidence analysis. *Lancet Glob. Health* **2020**, *8*, e180–e190. [[CrossRef](#)] [[PubMed](#)]
3. Zha, J.; Li, Y.Y.; Qu, J.Y.; Yang, X.X.; Han, Z.X.; Zuo, X. Effects of enhanced education for patients with the *Helicobacter pylori* infection: A systematic review and meta-analysis. *Helicobacter* **2022**, *27*, e12880. [[CrossRef](#)] [[PubMed](#)]
4. Hafiz, T.A.; D'Sa, J.L.; Zamzam, S.; Visbal Dionaldo, M.L.; Aldawood, E.; Madkhali, N.; Mubarak, M.A. The Effectiveness of an Educational Intervention on *Helicobacter pylori* for University Students: A Quasi-Experimental Study. *J. Multidiscip. Healthc.* **2023**, *16*, 1979–1988. [[CrossRef](#)] [[PubMed](#)]
5. Weiss, B.D. *Health Literacy and Patient Safety: Help Patients Understand. Manual for Clinicians*, 2nd ed.; American Medical Association Foundation: Chicago, IL, USA, 2007.
6. Iqbal, U.; Tanweer, A.; Rahmanti, A.R.; Greenfield, D.; Lee, L.T.; Li, Y.J. Impact of large language model (ChatGPT) in healthcare: An umbrella review and evidence synthesis. *J. Biomed. Sci.* **2025**, *32*, 45. [[CrossRef](#)] [[PubMed](#)]
7. Berry, P.; Dhanakshirur, R.R.; Khanna, S. Utilizing large language models for gastroenterology research: A conceptual framework. *Therap. Adv. Gastroenterol.* **2025**, *18*, 17562848251328577. [[CrossRef](#)] [[PubMed](#)]
8. Hu, Y.; Lai, Y.; Liao, F.; Shu, X.; Zhu, Y.; Du, Y.Q.; Lu, N.H.; National Clinical Research Center for Digestive Diseases. Assessing Accuracy of ChatGPT on Addressing *Helicobacter pylori* Infection-Related Questions: A National Survey and Comparative Study. *Helicobacter* **2024**, *29*, e13116. [[CrossRef](#)] [[PubMed](#)]
9. Kong, Q.Z.; Ju, K.P.; Wan, M.; Liu, J.; Wu, X.Q.; Li, Y.Y.; Zuo, X.L.; Li, Y.Q. Comparative analysis of large language models in medical counseling: A focus on *Helicobacter pylori* infection. *Helicobacter* **2024**, *29*, e13055. [[CrossRef](#)] [[PubMed](#)]
10. Lai, Y.; Liao, F.; Zhao, J.; Zhu, C.; Hu, Y.; Li, Z. Exploring the capacities of ChatGPT: A comprehensive evaluation of its accuracy and repeatability in addressing *Helicobacter pylori*-related queries. *Helicobacter* **2024**, *29*, e13078. [[CrossRef](#)] [[PubMed](#)]

11. Zeng, S.; Kong, Q.; Wu, X.; Ma, T.; Wang, L.; Xu, L.; Kou, G.; Zhang, M.; Yang, X.; Zuo, X.; et al. Artificial Intelligence-Generated Patient Education Materials for *Helicobacter pylori* Infection: A Comparative Analysis. *Helicobacter* **2024**, *29*, e13115. [[CrossRef](#)] [[PubMed](#)]
12. Du, R.C.; Zhu, Y.C.; Xiao, Y.T.; Yang, B.N.; Lai, Y.K.; Zhou, Z.X.; Deng, H.; Shu, X.; Lu, N.H.; Zhu, Y.; et al. Assessing the Capabilities of Novel Open-Source Artificial Intelligence-DeepSeek in *Helicobacter pylori*-Related Queries. *Helicobacter* **2025**, *30*, e70045. [[CrossRef](#)] [[PubMed](#)]
13. Gao, Z.; Ge, J.; Xu, R.; Chen, X.; Cai, Z. Potential application of ChatGPT in *Helicobacter pylori* disease relevant queries. *Front. Med.* **2024**, *11*, 1489117. [[CrossRef](#)] [[PubMed](#)]
14. Ye, Y.; Zheng, E.D.; Lan, Q.L.; Wu, L.C.; Sun, H.Y.; Xu, B.B.; Wang, Y.; Teng, M.M. Comparative evaluation of the accuracy and reliability of ChatGPT versions in providing information on *Helicobacter pylori* infection. *Front. Public Health* **2025**, *13*, 1566982. [[CrossRef](#)] [[PubMed](#)]
15. Karampinis, E.; Bozi Tzetzzi, D.A.; Pappa, G.; Koumaki, D.; Sgouros, D.; Vakirlis, E.; Liakou, A.; Papakonstantis, M.; Papadakis, M.; Mantzaris, D.; et al. Use of a Large Language Model as a Dermatology Case Narrator: Exploring the Dynamics of a Chatbot as an Educational Tool in Dermatology. *JMIR Dermatol.* **2025**, *8*, e72058. [[CrossRef](#)] [[PubMed](#)]
16. Dore, M.P.; Merola, E.; Pes, G.M. Advances and future perspectives in the pharmacological treatment of *Helicobacter pylori* infection: Taking advantage from artificial intelligence. *Clin. Res. Hepatol. Gastroenterol.* **2025**, *49*, 102689. [[CrossRef](#)] [[PubMed](#)]

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.