

STARK: Benchmarking LLM Retrieval on Textual and Relational Knowledge Bases

Shirley Wu^{*§}, Shiyu Zhao^{*§}, Michihiro Yasunaga[§], Kexin Huang[§], Kaidi Cao[§], Qian Huang[§],
Vassilis N. Ioannidis[†], Karthik Subbian[†], James Zou^{‡§}, Jure Leskovec^{‡§}
[§]Department of Computer Science, Stanford University [†]Amazon

Abstract

Answering real-world user queries, such as product search, often requires accurate retrieval of information from semi-structured knowledge bases or databases that involve blend of unstructured (*e.g.*, textual descriptions of products) and structured (*e.g.*, entity relations of products) information. However, previous works have mostly studied textual and relational retrieval tasks as separate topics. To address the gap, we develop STARK, a large-scale Semi-structure retrieval benchmark on Textual and Relational Knowledge Bases. We design a novel pipeline to synthesize natural and realistic user queries that integrate diverse relational information and complex textual properties, as well as their ground-truth answers. Moreover, we rigorously conduct human evaluation to validate the quality of our benchmark, which covers a variety of practical applications, including product recommendations, academic paper searches, and precision medicine inquiries. Our benchmark serves as a comprehensive testbed for evaluating the performance of retrieval systems, with an emphasis on retrieval approaches driven by large language models (LLMs). Our experiments suggest that the STARK datasets present significant challenges to the current retrieval and LLM systems, indicating the demand for building more capable retrieval systems that can handle both textual and relational aspects.

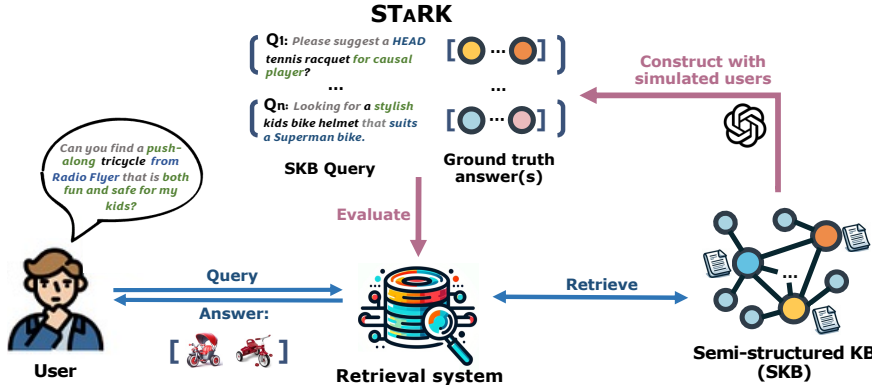


Figure 1: STARK features queries on semi-structured knowledge bases (SKBs) with textual and relational knowledge, with node entities as ground truth answers. It simulates user interactions with a SKB to evaluate LLM retrieval systems’ performance in providing accurate responses.

1 Introduction

Natural-language queries are the primary forms in human society for acquiring information, involving diverse knowledge in the real world [14, 18, 22]. For example, users on e-commerce web search

^{*‡}Equal contribution. Correspondence: {shirwu, jamesz, jure}@cs.stanford.edu.

The work of Vassilis N. Ioannidis and Karthik Subbian does not relate to their position at Amazon.

	Example query	Title of ground truth node(s)
STARK-AMAZON	<i>Looking for durable Dart World brand dart flights that resist easy tearing. Any recommendations?</i>	<Amazon Standard Flights> <Dart World Broken Glass Flight> (12 more)
	<i>What are recommended scuba diving weights for experienced divers that would fit well with my Gorilla PRO XL waterproof bag?</i>	<Sea Pearls Vinyl Coated Lace Thru Weight>
STARK-MAG	<i>Search publications by Hao-Sheng Zeng on non-Markovian dynamics.</i>	<Distribution of non-Markovian intervals...> <Comparison between non-Markovian...>
	<i>What are some nanofluid heat transfer research papers published by scholars from Philadelphia University?</i>	<A Numerical Study on Convection Around A Square Cylinder using AL2O3-H2O Nanofluid>
STARK-PRIME	<i>Could you provide a list of investigational drugs that interact with genes or proteins active in the epididymal region?</i>	<(S)-3-phenyllactic Acid>, <Anisomycin>, <Puromycin>
	<i>Search for diseases without known treatments and induce pruritus in pregnant women, potentially associated with Autoimmune.</i>	<Intrahepatic Cholestasis>
	<i>Please find pathways involving the POLR3D gene within nucleoplasm.</i>	<RNA Polymerase III Chain Elongation>
	<i>Which gene or protein associated with lichen amyloidosis can bind interleukin-31 to activate the PI3K/AKT and MAPK pathways?</i>	<OSMR>, <IL31RA>

Table 1: Example queries from STARK, which involve semi-structured information, where the relational and textual aspects are highlighted.

can express complex information needs by combining free-form elements or constraints, such as “Can you help me find a push-along tricycle from Radio Flyer that’s both fun and safe for my kid?”. Here, the underlying knowledge can be represented in semi-structured formats [29, 32, 40], referred to as semi-structured knowledge bases (SKBs), which integrate unstructured data, such as textual descriptions and natural language expressions (e.g., descriptions about a tricycle), with structured data, like entity interactions on knowledge graph or database (e.g., a tricycle with “belongs to” relation with brand Radio Flyer). This allows semi-structured knowledge bases to represent comprehensive knowledge in specific applications, making them indispensable in domains such as e-commerce [12], social media [26], and biomedicine [5, 15]. In this work, we focus on semi-structured knowledge bases (SKBs) with texts and relation information, which are two prevalent modalities on the existing knowledge bases.

In this context, as shown in Figure 1, information retrieval serves as a predominant step by identifying relevant data from the semi-structured knowledge bases, such as specific tricycles that match the request, which are then processed to generate accurate answers to user queries [34]. The retrieval accuracy on semi-structured knowledge bases is crucial for enhancing user experience, supporting informed decision-making, and preventing hallucination.

Limitations of existing works. However, previous works and benchmarks have primarily focused on either purely textual queries on unstructured knowledge [11, 17, 21, 24, 43] or structured SQL or knowledge graph queries [1, 36, 45, 47], which are inadequate to study the complexities of retrieval tasks involving semi-structured knowledge bases. Another line of works [9, 20, 37, 41, 42, 45] have explored the middle ground via multi-hop reasoning within single or across multiple documents. However, they are either limited in the span of knowledge or lack of diverse textual properties and explicit relational information. Recently, large language models (LLMs) have demonstrated significant potential on information retrieval tasks [11, 25, 33, 51]. Nevertheless, it remains an open question how effectively LLMs can be applied to the specific challenge of retrieval from SKBs. It is also important to note that existing works focus mainly on general knowledge, e.g., from wikipedia. In fact, the knowledge may also come from private sources, necessitating information retrieval systems adaptable to private semi-structured knowledge bases. Therefore, there is a gap in our understanding of how current LLM retrieval systems handle the complex interplay between textual and relational requirements in queries that optionally involve private knowledge.

Challenges. To address this gap, our goal is to develop a benchmark for LLM retrieval on SKBs. However, the lack of real queries on SKBs has become an obstacle. Although retrieval datasets with synthetic queries is a valuable supplement, accurately simulating user queries on SKBs is particularly difficult. This difficulty arises from the interdependence of textual and relational information, which can lead to challenges in precisely filtering through millions of candidates to construct the ground truth answers. Additionally, ensuring that queries are relevant to real-world applications and resembles real-world scenarios adds further complexity to the benchmarking process.

Our contribution: STARK. As shown in Figure 1, we present a large-scale Semi-structure retrieval benchmark on Textual and Relational Knowledge Bases (STARK). We propose a novel pipeline that simulates user queries and constructs precise ground truth answers using three SKBs built from

extensive texts and millions of entity relations from public sources. We validate the quality of queries in our benchmark through detailed analysis and human evaluation, focusing on their naturalness, diversity, and practicality. With STARK, we delve deeper into retrieval tasks on SKBs, evaluate the capability of current retrieval systems, and providing insights for future advancement. We point out the key features of STARK as follows:

- **Natural-sounding queries on semi-structured knowledge:** As illustrated in Table 1, the queries in our benchmark are crafted to incorporate rich relational information and complex textual properties. Additionally, these queries closely mirror the types of questions users would naturally ask in real-life scenarios, *e.g.*, with flexible query formats and possibly with extra contexts.
- **Context-specific reasoning:** The queries entail reasoning capabilities specific to the context. This includes the ability to infer customer interests, understand specialized field descriptions, and deduce relationships involving multiple subjects mentioned within the query. For example, the context “*I had a dozen 2.5-inch Brybelly air hockey pucks, so I’m trying to find matching strikers.*” entails the user’s interest in looking for complementary products. Such reasoning capabilities are crucial for accurately interpreting and responding to the nuanced requirements of each query.
- **Diverse domains:** Our benchmark spans a wide range of content and domains, reflecting the diverse nature of real-world information needs. It covers areas on product recommendation, academic paper search, and precision medicine inquiry. Therefore, STARK provides a comprehensive evaluation of retrieval systems across various contexts and knowledge domains.

Moreover, we perform extensive experiments on each benchmark dataset for the prevailing retrieval systems. We highlight research challenges especially on the capability to handle textual and relational requirements and the latency of retrieval systems on large-scale SKBs that involve million-scale entities or relations. Finally, we provide insights into future directions to build more capable retrieval systems. The benchmark datasets and code are available on <https://github.com/snap-stanford/stark>.

2 Related Work

We identify three primary categories on retrieval-oriented question answering (QA) datasets based on the knowledge sources they utilize: unstructured, structured, and semi-structured knowledge bases. We discuss the differences and limitations in the previous works that we seek to address.

Unstructured QA Datasets. The line of work explores retrieving answers from a single document [31] or multiple document sources [9, 20, 39, 41, 44]. For example, SQuAD [31] focuses on extracting answers from a single document, testing model’s ability to understand and retrieve information within a specific context. For multi-document retrieval, HotpotQA [44] demands reasoning over multiple documents and includes sentence-level supporting facts, enhancing the complexity of the retrieval task. TriviaQA [20] combines questions from trivia games with evidence documents, evaluating both retrieval and reasoning over multiple sources. Moreover, some works [23, 28] leverage generated results by search engines as ground truth answers or information sources for retrieval. However, unstructured QA datasets or knowledge bases often lack the depth of relational reasoning commonly required in answering complex user queries.

Structured QA Datasets. Within this category, the datasets challenge models to retrieve answer from structured knowledge bases, either knowledge graphs [2–4, 10, 13, 38, 48] or tabular data [49, 50]. For example, ComplexWebQuestions [38] challenge model’s abilities to decompose complex queries and interpret queries using entities and relationships defined in knowledge graphs. GraphQA [13] textualizes graphs by describing node attributes and edges with natural language, offering language model a way to handle KBQA. For datasets with tabular data, WikiSQL [50] and Spider [49] both address the task of converting natural language questions to SQL queries, with WikiSQL concentrating on single-table databases and Spider extending to complex multi-tables. While these datasets primarily focus on relational information, the lack of textual information limits questions within predefined relationships and entities, which constrains the breadth of available information, resulting in less comprehensive and nuanced responses compared to the extensive insights drawn from abundant textual data.

Semi-Structured QA Datasets. Another line of work focuses on semi-structured QA with tabular and textual data. WikiTableQuestions [30] emphasizes comprehension of table structure and

Table 2: Data statistics of our constructed semi-structured knowledge bases

	#entity types	#relation types	avg. degree	#entities	#relations	#tokens
STARK-AMAZON	2	3	3.0	1,032,407	3,886,603	592,067,882
STARK-MAG	4	4	10.6	1,872,968	19,919,698	212,602,571
STARK-PRIME	10	18	62.6	129,375	8,100,498	31,844,769

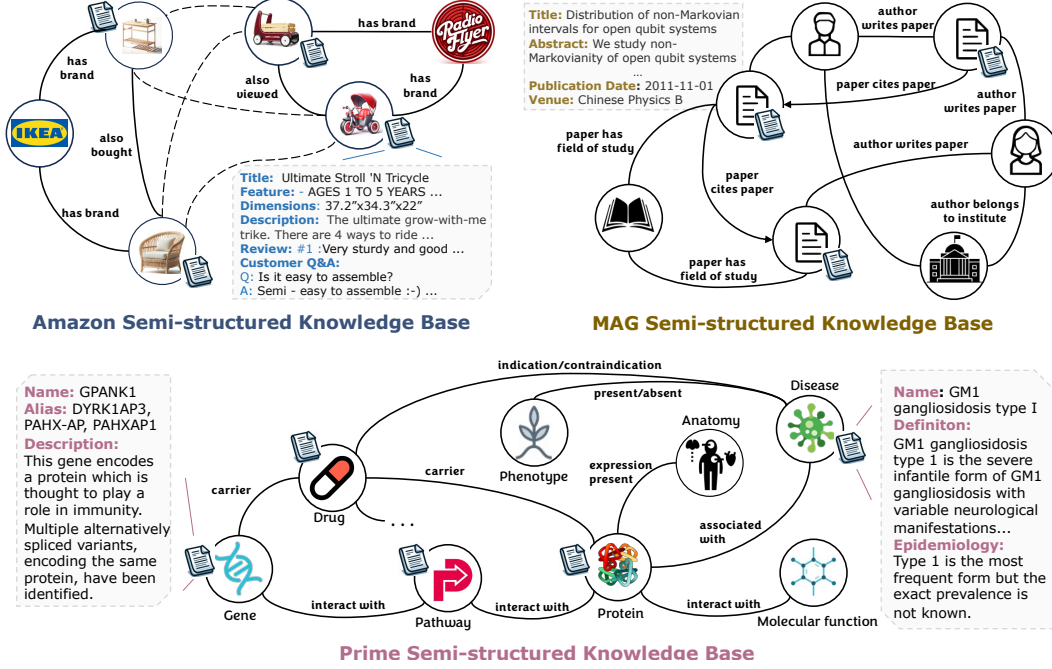


Figure 2: Demonstration of Semi-structured Knowledge Bases, where each knowledge base combines both textual and relational information in a complex way, making the retrieval tasks challenging.

textual content, reflecting the semi-structured nature of many real-world datasets. TabFact [6] and HybridQA [7] combine textual and tabular data, requiring models to validate statements or answer questions using both data types. TabMCQ [19] introduces multiple-choice questions based on tables, adding complexity in understanding and reasoning over semi-structured data. However, different from this work, these existing datasets mainly focus on tables as structured sources, which may not incorporate rich relational information naturally exists between entities. Additionally, previous attempts to combine text and tabular data through external web links or additional text columns often result in data structures that are hard to navigate. Therefore, there is an emerging need to benchmark retrieval tasks on textual and relational knowledge base, offering more insights beyond the scope of structured tabular sources.

3 Benchmarking Retrieval Tasks over Textual and Relational Knowledge

We develop three large-scale retrieval datasets on three semi-structure knowledge bases (SKBs). In Section 3.1, we introduce the schema and scale of the SKBs, which integrate entities with rich relational and textual information. In Section 3.2, we introduce the retrieval datasets and the key characteristics of the queries. Then, we highlight our novel dataset construction pipeline in Section 3.3, which simulates user queries involving textual and relational knowledge and automatically generates the ground truth answers in the retrieval datasets. Finally, we conduct data distribution analysis and human evaluation on the queries in Section 3.4.

Table 3: Dataset statistics on STARK.

	#queries	#queries w/ multiple answers	avg. #answers	train / val / test
STARK-AMAZON	9,100	7,082	17.99	0.65 / 0.17 / 0.18
STARK-MAG	13,323	6,872	2.78	0.60 / 0.20 / 0.20
STARK-PRIME	11,204	4,188	2.56	0.55 / 0.20 / 0.25

3.1 Semi-structured Knowledge Bases

As shown in Figure 2, we construct three SKBs from the relational structure between entities and the textual information associated with a subset of the entities. We present the statistics of the relational structure in Table 2 and introduce each SKB as follows:

Amazon Semi-structured Knowledge Base. This knowledge base is derived from the Sports and Outdoors category from Amazon Product Reviews [12] and Amazon Question and Answer Data [27]. The knowledge base features two entity types: product and brand, and three types of relational information: `also_bought`, `also_viewed` between product entities, and `has_brand` relation between product and brand entities. In total, it comprises around 1.03M entities (product entities: 0.96M, brand entities: 0.07M) and 3.9M relations (`also_bought`: 1.7M, `also_viewed`: 1.3M, `has_a_brand`: 0.9M). We obtain the textual information by combining the meta data from Amazon Product Reviews with the customer Q&A records from Amazon Question and Answer Data. This provides a rich amount of textual data, including product titles, descriptions, prices, customer reviews, and related customer Q&A for each product. For brand entities, we extract brand titles as the textual attribute. Especially, Amazon SKB features an extensive textual data, which is largely contributed from customer reviews and Q&A.

MAG Semi-structured Knowledge Base. This knowledge base is constructed from obgn-MAG [16], obgn-papers100M [16], and Microsoft Academic Graph (version 2019-03-22) [35, 40]. The knowledge base contains around 1.9M entities under four entity types (author: 1.1M, paper: 0.7M, institution: 9K, field_of_study: 0.06M) and 20M relations under four relation types (author_writes_paper: 7.9M, paper_has_field_of_study: 7.2M, paper_cites_paper: 4.9M, author_affiliated_with_institution: 1.0M). We construct the knowledge base by selecting papers¹ shared between obgn-MAG and obgn-papers100M. The relational structure comes from the extracted subgraph in obgn-mag based on the selected papers. The textual information such as titles and abstracts is sourced from obgn-papers100M. Additionally, we augment the textual data by integrating details from the Microsoft Academic Graph database, providing extra information like paper venue, author and institution names. This SKB demonstrates a large number of entities and relations associate with paper nodes, especially on citation and authorship relation types.

Prime Semi-structured Knowledge Base. We leverage the existing knowledge graph PrimeKG [5] which contains ten entity types including disease, drug, gene/protein, and eighteen relation types, such as `associated_with`, `synergistic_interaction`, `indication`, `expression_present`. The entity count in our knowledge base is approximately 129K, with around 8M relations. The details on entity numbers and relation tuples are available in Appendix A.2. Compared to the Amazon and MAG SKBs, Prime SKB is denser and features a greater variety of relation types. While PrimeKG already provides text information on disease and drug entities, including drug mechanisms and disease descriptions, we additionally integrate the textual details from multiple databases for gene/protein and pathway entities such as genomic position, gene activity summary and pathway orthologous event.

We include the comprehensive public sources we used to construct the SKBs in Table 8 of Appendix A.

3.2 Retrieval Tasks on Semi-structured Knowledge Bases

We develop three novel retrieval datasets. Specifically, the task is to retrieve node entities over the SKBs given any input queries. To highlight, the queries feature the interplay and fusion between relational and textual knowledge. Moreover, to elevate their applicability in practical scenarios, these

¹ We filter out non-English language papers as we only considers single-lingual queries.

queries are designed to mimic real-life query patterns in terms of the natural-sounding property and flexible formats. Each query may have single answer or multiple answers, formulated as a subset of entities from the knowledge base, which requires the model to accurately identify the ground truth entities out of potentially millions in response to each query. In Table 1, we demonstrate example queries from these benchmarks. We detail our retrieval benchmarks’ scale in Table 3 and introduce the specifics of each retrieval dataset as follows:

STARK-AMAZON. This dataset comprises 9,100 customer-focused queries aimed at product searches, with a notable 68% of these queries yielding multiple answers. The dataset prioritizes customer-oriented criteria, highlighting textual elements such as product quality, functionality, and style. Additionally, it accentuates relational aspects, including brand and product connections (*e.g.*, complementary or substitute items). The queries are framed in detailed, conversation-like formats, enriching the context and enhancing the dataset’s relevance to real-world scenarios.

STARK-MAG. This dataset comprises 13,323 paper queries, with roughly half yielding multiple answers. Beyond the single-hop relational requirements, STARK-MAG also emphasizes the fusion between the textual requirements with multi-hop queries. For example, “Are there any papers from King’s College London” highlights the metapath (*institution* $\rightarrow$ *author* $\rightarrow$ *paper*) on the relational structure. We designed three single-hop and four multi-hop relational query templates, *cf.* Appendix A.1. The textual aspects of the queries diversifies, focusing on different elements such as the paper’s topic, methodology, or the advancements it introduces, sourced mainly from the abstracts.

STARK-PRIME. This dataset, comprising 11,204 queries across all ten entity types, incorporates single-hop and multi-hop relational information similar to STARK-MAG. We developed 28 multi-hop query templates, detailed in Appendix A.2, to cover various relation types and ensure their practical relevance. For instance, the query “What is the drug that targets the genes or proteins expressed in <anatomy>?” serves applications in precision medicine and pharmacogenomics, aiding researchers and healthcare professionals in identifying drugs that act on genes or proteins associated with specific anatomical areas and enabling more targeted treatments. For entities like drug, disease, gene/protein, and pathway, the queries are a hybrid of relational requirements and textual requirements drawn from documents. A notable feature is the deliberate simulation of three distinct roles – medical scientist, doctor, and patient – for certain queries related to drugs and diseases. This is intended to diversify the language used across various user types and to evaluate the robustness of the retrieval system under varied linguistic scenarios. For entities such as effect/phenotype, the queries rely solely on relational data due to limited textual information in the knowledge base.

We divide each dataset into training, validation, and testing subsets, with the ratios detailed in Table 3. For the STARK-AMAZON dataset, we randomly allocate a subset of queries with 20 or fewer answers to the validation and testing sets, while the remaining queries are assigned to the training set. For the STARK-MAG and STARK-PRIME, we use random splits to divide the data.

3.3 Benchmark Construction

Here we detail the novel pipeline to synthesize the retrieval datasets present in Section 3.2. The core idea is to entangle relational information and textual properties, while constructing ground truth answers in an accurate and efficient way. As shown in Figure 3, the construction of the retrieval datasets generally involves four steps, and the specific processes vary depending on the characteristics of each dataset. These steps are as follows:

- **Sample Relational Requirements:** Initially, we sample a relation template, such as “(a product) belongs to <brand>” and instantiate it into a relational requirement, *e.g.*, “belongs to Radio Flyer” shown in the first step of Figure 3. Each relational requirement yields a set of candidate entities that meet the requirement, *i.e.*, products belonging to Radio Flyer. Therefore, we obtain the relational requirements and the corresponding candidates. The relation templates for the three SKBs are available in Appendix A, including 28 relation templates for the Prime SKB. Each relation template in the SKBs serves a distinct and practical purpose. For instance, the template “What is the drug that targets the genes or proteins which are expressed in <anatomy>?” is particularly valuable for precision medicine inquiries. It aids in pinpointing medications aimed at specific genetic or protein targets within defined anatomical areas, thereby contributing to the development of targeted treatments.

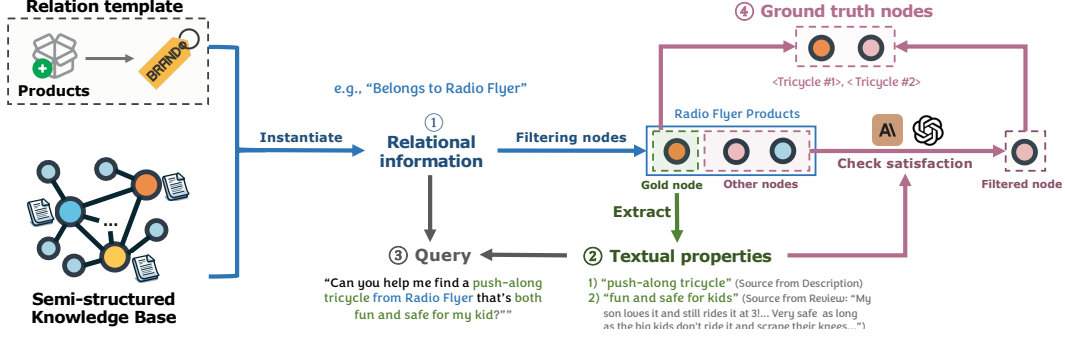


Figure 3: The process of constructing semi-structured retrieval datasets involves four main steps: 1) Sample Relational Requirement: Based on relational templates, sample a relational requirement on a SKB. 2) Extract Textual Properties: From a node that meets the relational requirement, extract relevant textual properties. 3) Combine Information: Merge the relational information and textual properties to form a natural-sounding query. 4) Construct Ground Truth Nodes: Check if nodes satisfy the textual properties using multiple language models to establish ground truth nodes.

- **Extracting Textual Properties:** Subsequently, we select one of the candidate nodes that satisfies the relational requirement, which is referred to as the gold answer. We extract textual properties that align with the interests of specific roles from the textual information of the gold answer. The roles could be customers, researchers, or medical scientists, differing across the SKBs. For example, in the second step of Figure 3, we extract the phrases of functionality and user experience from a Radio Flyer product, which come from different sources in the product document. We utilize GPT-3.5-turbo-16k for the STARK-AMAZON and Claude-v2 for the other two datasets to conduct this step, where the prompts used are included in Appendix B.
- **Combining Textual and Relational Information:** With the textual and relational requirements, we simulate queries by fusing them using LLMs. We conduct two-stage fusion using two different LLMs, *i.e.*, Claude-v2 and GPT-4-Turbo, to avoid bias that might arise from relying on a single LLM and ensure a more diverse set of simulated queries. Specifically, this first-stage integration is guided by various criteria, such as ensuring a natural-sounding query, adhering to the style of ArXiv searches, or aligning with specific roles. Furthermore, the second stage instructs the LLM to enrich the context and rephrase the language, thereby posing a more demanding reasoning challenge in comprehending the requirements of the query. In our example, by combining relational information “belongs to Radio Flyer” and textual information “push-along tricycle” and “fun and safe for kids”, we arrive at the final query in the third step.
- **Filtering Additional Answers:** In the final step, we assess whether each remaining candidate, excluding the gold answer used for extracting textual properties, meets the additional textual requirement. For the verification, we employ three Claude models. Only candidates that pass the verification across all models are included in the final ground truth answer set. Therefore, this stringent verification ensures the accuracy of the nodes in the ground truth answers. During this process, we also evaluate the accuracy of the gold nodes that pass the verification criteria. The accuracy rates STARK-AMAZON, STARK-MAG, and STARK-PRIME are 86.6%, 98.9%, and 92.3%, respectively, demonstrating the effectiveness of our filtering approach in maintaining high-quality ground truth answers.

This dataset construction pipeline is automatic, efficient, and broadly applicable to the SKBs within our context. We include the complete prompts of the above steps in Appendix B.

3.4 Benchmark Data Distribution Analysis and Human Evaluation

We study the data distribution to understand the characteristics of our benchmark datasets and conduct human evaluation on the benchmark quality. We study the data distribution in three dimensions:

- **Query and Answer Length.** We analyze the query length distribution (measured in the number of words) and the number of ground truth answers for each query. The query length reflects the amount of context information provided by users, while the number of ground truth answers

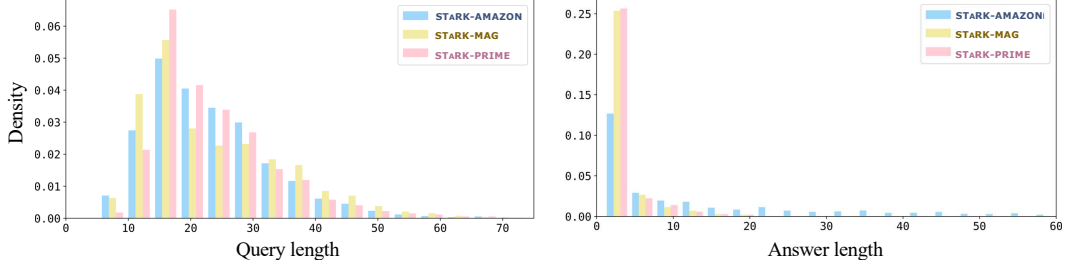
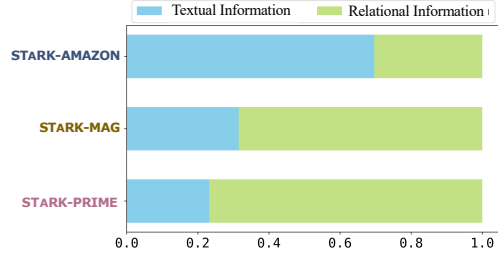


Figure 4: Distribution of query and answer lengths on STARK datasets.

Table 4: Query diversity measurement on STARK.

	Shannon Entropy	Type-Token Ratio
STARK-AMAZON	10.39	0.179
STARK-MAG	10.25	0.180
STARK-PRIME	9.63	0.143
Reference ¹	10.44	0.261

Figure 5: Average relative composition of relational vs. textual information.



indicates the level of inclusiveness or ambiguity of queries within the specific SKB. As illustrated in Figure 4, all three datasets exhibit similar query length distributions, with most queries containing around 16 words. Queries with up to 50 words typically involve mentions of other entities, such as product or paper titles, or provide more detailed context, such as specific symptoms descriptions. Interestingly, the answer length distribution for STARK-AMAZON shows a more significant long-tail pattern, with approximately 22% of answers exceeding 30 entities, whereas the answer lengths for STARK-PRIME and STARK-MAG are all within 20 entities. On average, as shown in Table 3, STARK-AMAZON presents an average answer length of 5.32, indicative of the e-commerce recommendation domain where general customer inquiries frequently result in diverse recommendations. While STARK-PRIME has the smallest average answer length, partially attributed to the smaller size of the Prime SKB compared to the other two SKBs.

- **Query Diversity.** A diverse set of queries pose challenges for broader applicability to meet varying user demands. Specifically, we measure query diversity using Shannon Entropy, which quantifies the uncertainty in the word distribution across all queries, and the Type-Token Ratio (TTR), which calculates the proportion of unique words to the total number of words. Higher values of Shannon Entropy and TTR indicate greater lexical diversity, reflecting a wider range of topics and linguistic expressions in the query set. As shown in Table 4, we observe high values of Shannon Entropy on all of the datasets. For reference, we compute the metrics for the wikipedia page of Barack Obama¹. While TTR is more sensitive to the document length, the TTR values consistently demonstrate a steady presence of unique words relative to the large length in our query document.
- **Ratio of Relational vs. Textual Information.** A key feature of our benchmark dataset is the composition of textual and relational information. Therefore, it is crucial to understand the proportionality between these two types of information. We calculate the ratio of the length of relational requirements to the length of textual requirements for each query and then average these ratios across each dataset. Intuitively, the ratio serves as an approximation of the relative distribution based on the length of requirements, and does not directly reflect the importance of each type of information in determining the final answers. As shown in Figure 5, the ratios vary across different datasets, highlighting a differing emphasis on textual versus relational information. This distribution variation presents additional challenges for retrieval systems, requiring them to adapt to the dataset characteristic and specific balance of information.

¹https://en.wikipedia.org/wiki/Barack_Obama

Table 5: Positive/Non-negative rates (%) from human evaluation.

	Naturalness	Diversity	Practicality
STARK-AMAZON	73.6 / 89.5	68.4 / 89.5	89.5 / 94.7
STARK-MAG	94.7 / 100	73.7 / 84.2	68.4 / 84.2
STARK-PRIME	67.8 / 92.8	71.4 / 82.1	71.4 / 89.3
Average	78.7 / 94.1	71.0 / 85.3	76.4 / 89.4

Human evaluation. Moreover, we qualitatively assess sampled queries from our benchmark to determine if they are natural-sounding (*resembling natural conversation*), diverse (*covering a wide range of question structures and complexity levels*), and practical (*relevant and useful in real-life situations*), with participation from 63 individuals. The evaluation results are converted from a 5-point Likert-like scale to a positive/tie/negative scale for reporting. We report the positive and non-negative rates in Table 5. On average across the three datasets, 94.1%, 85.3%, and 89.4% of participants rated neutral or above in terms of naturalness, diversity, and practicality, respectively. The results justify the quality of our benchmark and highlight its potential for diverse and realistic retrieval tasks.

4 Experiments

In this section, we aim to explore the following questions:

- **Q1:** How well do current retrieval approaches perform on our benchmark datasets?
- **Q2:** What challenges are revealed by the experiments, and more importantly, what are the potential future directions for improvement?

4.1 Retrieval Models

Vector Similarity Search (VSS). VSS embeds both the query and the concatenated textual and relational information of each candidate entity. Then, the similarity score is computed based on cosine similarity between the query and candidate embeddings. We use the `text-embedding-ada-002` model from OpenAI for generating the embeddings.

Multi-Vector Similarity Search (Multi-VSS). Multi-VSS represents candidate entities with multiple vectors, capturing detailed features for complex retrieval tasks. Here, we chunk and separately embed the textual and relational information of each candidate entity. The final score is aggregated from the similarities between the chunks and the query.

Dense Retriever. Dense Retriever finetunes a query encoder and a document encoder separately using the query answer pairs from the training dataset. We optimize the encodes using contrastive loss. Specifically, the positive pairs are constructed from the training questions and their ground truth answers, while we construct 20 hard negative answers for each query from the top false positive predictions from VSS. We use `roberta-base` as the base model and finetune the encodes over the entire training split for each dataset.

QAGNN [46]. QAGNN constructs a graph where nodes include entities found in the question or answer choices, incorporating their neighboring nodes as intermediate context. It extracts a subgraph from a larger knowledge graph, enabling a comprehensive understanding by leveraging the relational information from the knowledge graph. The approach integrates this with semantic embeddings from a language model, jointly modeling both relational and semantic information to enhance the question-answering modeling.

VSS + LLM Reranker [8, 52]. This method improves the precision of the top- v results from VSS by reranking them using language models, taking advantage of their advanced language understanding capabilities. For this purpose, we employ two different language models: GPT-4-turbo (`gpt-4-1106-preview`) and Claude-v2. We set the value of $v = 20$ for the reranking process to balance between precision and computational efficiency. Specifically, we construct a prompt that instructs the language models to assign a score between 0 and 1 to a node, based on its combined textual and relational information, with certain criteria provided for rating the node. We find that this

Table 6: Main experimental results.

		Dense Retriever (roberta)	QAGNN (roberta)	VSS (ada-002)	Multi-VSS (ada-002)	VSS+Claude2 Reranker	VSS+GPT4 Reranker
STARK-AMAZON	Hit@1	0.1529	0.2656	0.3916	0.4007	0.4266	0.4479
	Hit@5	0.4793	0.5001	0.6273	0.6498	0.6746	0.7117
	Recall@20	0.4449	0.5205	0.5329	0.5512	0.5376	0.5535
	MRR	0.3020	0.3775	0.5035	0.5155	0.5329	0.5569
STARK-MAG	Hit@1	0.1051	0.1288	0.2908	0.2592	0.3202	0.4090
	Hit@5	0.3523	0.3901	0.4961	0.5043	0.5334	0.5818
	Recall@20	0.4211	0.4697	0.4836	0.5080	0.4834	0.4860
	MRR	0.2134	0.2912	0.3862	0.3694	0.4129	0.4900
STARK-PRIME	Hit@1	0.0446	0.0885	0.1263	0.1510	0.1611	0.1828
	Hit@5	0.2185	0.2135	0.3149	0.3356	0.3582	0.3728
	Recall@20	0.3013	0.2963	0.3600	0.3805	0.3598	0.3405
	MRR	0.1238	0.1473	0.2141	0.2349	0.2466	0.2655

approach, which involves directly scoring nodes, performs better than asking for strict satisfaction of conditions [52].

4.2 Evaluation Metrics

We use the following metrics for a holistically evaluating the model performance.

Hit@ k . The Hit@ k metric assesses whether the correct item is among the top- k results from the model. We used $k = 1$ and $k = 5$ for evaluation. At $k = 1$, it evaluates the accuracy of the top recommendation; at $k = 5$, it examines the model’s precision in a wider recommendation set. It’s a straightforward yet effective way to measure the immediate relevance of answers.

Recall@ k . Recall@ k measures the proportion of relevant items in the top- k results. In our study, $k = 20$ is used, as the answer length of all of the queries in our benchmarks are equal or smaller than 20. This metric offers insight into the model’s ability to identify all relevant items, particularly in scenarios where missing any could be critical.

Mean Reciprocal Rank (MRR). MRR is a statistic for evaluating the average effectiveness of a predictive model. It calculates the reciprocal of the rank at which the first relevant item appears in the list of predictions. MRR is particularly useful for understanding the model’s performance at presenting the correct item at an early stage, without the need for a predefined k . This metric emphasizes the importance of the rank of the first correct answer, which is crucial in many practical applications where the first correct answer is often the most impactful.

4.3 Analysis

Main results. In Table 6, we report the experiment results. Firstly, we notice a significant performance gap between models trained or fine-tuned on our training datasets, such as Dense Retriever and QAGNN, and models like VSS and Multi-VSS that use embeddings generated from text-embedding-ada-002. For the dense retriever, training the encoders proves challenging when handling both textual and relational information as text, leading to training challenges. The difficulty in training a dense retriever is exacerbated by the stark differences in length between long documents and short relational pairs, making it tough for the model to learn a new knowledge base’s structure from scratch during fine-tuning. Without the advantage of pre-trained embeddings, it also struggles to align with the textual requirements. For QAGNN, it faces its own training challenges due to its high computational demands during training. Finally, the superiority of VSS and Multi-VSS demonstrates the superior pre-trained embeddings obtained from a large corpus.

Further, the Multi-VSS strategy, which breaks down the documents into chunks and embeds them into embeddings separately to capture more fine-grained context, generally outperforms VSS. However, it falls short in the Hit@1 metric on the STARK-MAG dataset, likely due to this dataset’s property for understanding more extended contexts or potential interaction or reasoning between different chunks.

Significantly, the methods combining VSS with an LLM reranker consistently outperform other approaches across most metrics. Notably, reranking the top 20 predictions with an LLM leads

Table 7: Latency (s) of the retrieval systems on STARK.

	Dense Retriever	QAGNN	VSS	Multi-VSS	VSS+Claude	VSS+GPT4
STARK-AMAZON	2.34	2.32	5.71	4.87	27.24	24.76
STARK-MAG	0.94	1.35	2.25	3.14	22.60	23.43
STARK-PRIME	0.92	1.29	0.54	0.90	29.14	26.97
Average	1.40	1.65	2.83	2.97	26.33	25.05

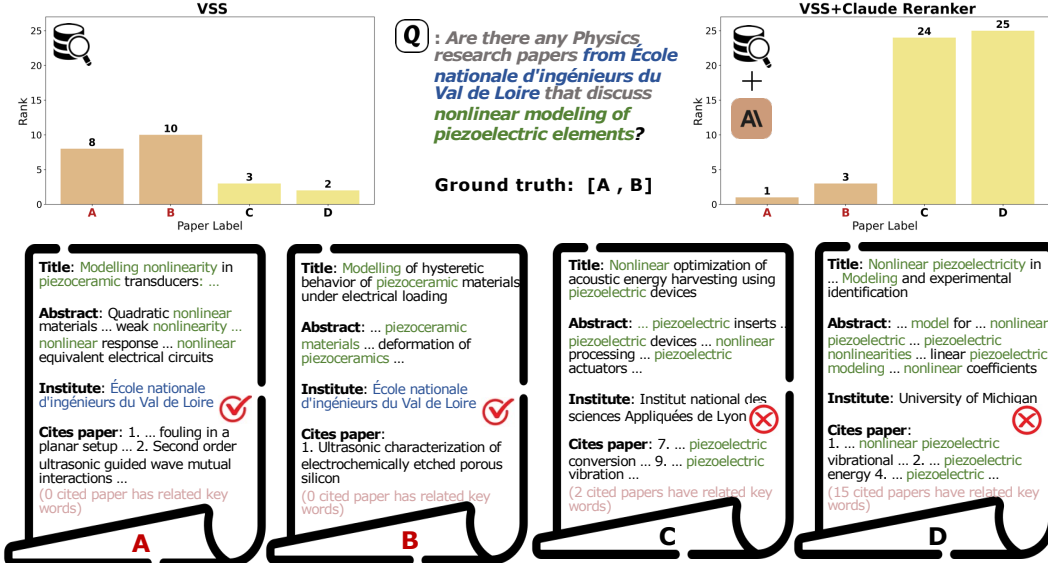


Figure 6: A case study on STARK-MAG, where VSS mistakenly ranks non-ground truth papers C and D higher due to the repeated key words in the relational information "cites paper." After reranking with Claude, it correctly prioritizes the ground truth papers A and B. This correction is attributed to the more accurate reasoning and analysis of the combined textual and relational information.

to significant improvements in Hit@ k and MRR metrics compared to using VSS alone. This improvement underscores the limitations of traditional embedding methods, which lack the advanced reasoning capabilities provided by LLMs. Additionally, using GPT-4 as a reranker tends to yield better results than using Claude, further demonstrating the effectiveness of better reasoning and context understanding in improving retrieval performance.

However, even for VSS+GPT4 Reranker, we observe insufficient performances, *e.g.*, Hit@1 is only around 18% on STARK-PRIME and 41% on STARK-MAG, meaning that for a significant portion of queries, the top-ranked answer is not the correct one. Moreover, Recall@20 metric values are lower than 60% across all of the datasets, which indicates the ranking results may not be comprehensive, potentially missing relevant answers in specific applications. Also, we notice that the MRR scores generally follow a similar trend as the Hit@1 metric. However, the MRR values remain relatively low, particularly for STARK-PRIME, indicating large space for further optimization in the ranking process.

Retrieval Latency. Another aspect vital for the practical application of retrieval systems is their latency. As shown in Table 7, we test the latency of the retrieval methods on a single NVIDIA A100-SXM4-80GB. Specifically, the Dense Retriever and QAGNN models exhibit lower average latency, making them more suitable for time-sensitive applications. In contrast, the VSS and Multi-VSS models show moderate latency, but when combined with LLM rerankers, such as Claude and GPT-4, the latency significantly increases. Therefore, in the context of semi-structured retrieval tasks, the balance between enhanced accuracy and increased latency is crucial. Effective retrieval methods should consider such trade-offs, particularly when addressing complex queries involving reasoning over both textual and relational information.

Case study. Finally, we highlight the importance of reasoning ability in achieving good retrieval results on our benchmark datasets. In Figure 6, we conduct a case study to compare the performance of VSS and VSS+Claude Reranker. Specifically, we examine a user query that requests papers from an institution on a particular topic. In this scenario, VSS fails to adequately address the relational requirement due to directly embedding the entire documents without detailed analysis. As a result, papers containing frequently repeated keywords, such as “nonlinear modeling” and “piezoelectric elements”, are mistakenly assigned high scores. However, the result improves significantly after the reranking with an LLM, where the LLM is equipped to reason the relationship between the query and the information contained in each paper, making the scores more accurately reflect the papers’ relevance to the query. This case study underscores the need for reasoning ability in grasping the complexities of queries. As embedding methods may capture some aspects of the query, they often fail to capture the detailed interaction between textual and relational information. Thus, integrating LLMs with advanced reasoning into the retrieval process can be crucial for the retrieval accuracy.

5 Conclusion

We introduce STARK, the first benchmark designed to thoroughly evaluate the capability of retrieval systems driven by LLMs in handling semi-structured knowledge bases (SKBs). This benchmark features a diverse set of queries that are semi-structured and natural-sounding, requiring context-specific reasoning across various domains, thereby setting a new standard for assessing retrieval systems in the context of SKBs. We utilize public datasets to construct three SKBs and develop an automated, general pipeline to simulate user queries that mimic real-world scenarios. Our extensive experiments on STARK reveal significant challenges for current models in handling both textual and relational information effectively and flexibly. Overall, STARK offers valuable opportunities for future research to advance the field of complex and multimodal retrieval systems, where reducing retrieval latency and incorporating strong reasoning ability into the retrieval process are identified as two prospective future directions.

6 Acknowledgement

We thank group members in Jure Leskovec and James Zou’s lab for providing valuable suggestions. We also acknowledge the support of Amazon, DARPA under Nos. N660011924033 (MCS); NSF under Nos. OAC-1835598 (CINES), CCF-1918940 (Expeditions), DMS-2327709 (IHBEM); Stanford Data Applications Initiative, Wu Tsai Neurosciences Institute, Chan Zuckerberg Initiative, Genentech, GSK, Hitachi, Juniper Networks, KDDI, and UCB.

REFERENCES

- [1] Jonathan Berant, Andrew Chou, Roy Frostig, and Percy Liang. 2013. Semantic Parsing on Freebase from Question-Answer Pairs. In *EMNLP*.
- [2] Antoine Bordes, Sumit Chopra, and Jason Weston. 2014. Question Answering with Subgraph Embeddings. In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, Alessandro Moschitti, Bo Pang, and Walter Daelemans (Eds.). Association for Computational Linguistics, Doha, Qatar.
- [3] Antoine Bordes, Nicolas Usunier, Sumit Chopra, and Jason Weston. 2015. Large-scale Simple Question Answering with Memory Networks. *CoRR* abs/1506.02075 (2015).
- [4] Shulin Cao, Jiaxin Shi, Liangming Pan, Lunyu Nie, Yutong Xiang, Lei Hou, Juanzi Li, Bin He, and Hanwang Zhang. 2022. KQA Pro: A Dataset with Explicit Compositional Programs for Complex Question Answering over Knowledge Base. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*.
- [5] Payal Chandak, Kexin Huang, and Marinka Zitnik. 2023. Building a knowledge graph to enable precision medicine. *Scientific Data* (2023).
- [6] Wenhu Chen, Hongmin Wang, Jianshu Chen, Yunkai Zhang, Hong Wang, Shiyang Li, Xiyu Zhou, and William Yang Wang. 2020. TabFact: A Large-scale Dataset for Table-based Fact Verification. In *International Conference on Learning Representations*.
- [7] Wenhu Chen, Hanwen Zha, Zhiyu Chen, Wenhan Xiong, Hong Wang, and William Yang Wang. 2020. HybridQA: A Dataset of Multi-Hop Question Answering over Tabular and Textual Data. In *Findings of the Association for Computational Linguistics: EMNLP 2020*. Association for Computational Linguistics.
- [8] Yew Ken Chia, Pengfei Hong, Lidong Bing, and Soujanya Poria. 2023. INSTRUCTEVAL: Towards Holistic Evaluation of Instruction-Tuned Large Language Models. *CoRR* abs/2306.04757 (2023).
- [9] Matthew Dunn, Levent Sagun, Mike Higgins, V. Ugur Guney, Volkan Cirik, and Kyunghyun Cho. 2017. SearchQA: A New QA Dataset Augmented with Context from a Search Engine. arXiv:1704.05179 [cs.CL]
- [10] Tiantian Gao, Paul Fodor, and Michael Kifer. 2019. Querying Knowledge via Multi-Hop English Questions. *Theory and Practice of Logic Programming* (2019).
- [11] Kelvin Guu, Kenton Lee, Zora Tung, Panupong Pasupat, and Mingwei Chang. 2020. Retrieval augmented language model pre-training. In *ICML*. PMLR.
- [12] Ruining He and Julian J. McAuley. 2016. Ups and Downs: Modeling the Visual Evolution of Fashion Trends with One-Class Collaborative Filtering. In *WWW*. ACM.
- [13] Xiaoxin He, Yijun Tian, Yifei Sun, Nitesh V. Chawla, Thomas Laurent, Yann LeCun, Xavier Bresson, and Bryan Hooi. 2024. G-Retriever: Retrieval-Augmented Generation for Textual Graph Understanding and Question Answering. arXiv:2402.07630 [cs.LG]
- [14] Lynette Hirschman and Robert Gaizauskas. 2001. Natural language question answering: the view from here. *natural language engineering* (2001).
- [15] Weihua Hu, Matthias Fey, Marinka Zitnik, Yuxiao Dong, Hongyu Ren, Bowen Liu, Michele Catasta, and Jure Leskovec. 2020. Open Graph Benchmark: Datasets for Machine Learning on Graphs. In *NeurIPS*.
- [16] Weihua Hu, Matthias Fey, Marinka Zitnik, Yuxiao Dong, Hongyu Ren, Bowen Liu, Michele Catasta, and Jure Leskovec. 2021. Open Graph Benchmark: Datasets for Machine Learning on Graphs. arXiv:2005.00687 [cs.LG]
- [17] Gautier Izacard and Edouard Grave. 2021. Leveraging Passage Retrieval with Generative Models for Open Domain Question Answering. In *EACL*.

- [18] Hasan M. Jamil. 2017. Knowledge Rich Natural Language Queries over Structured Biological Databases. In *Proceedings of the 8th ACM International Conference on Bioinformatics, Computational Biology, and Health Informatics*.
- [19] Sujay Kumar Jauhar, Peter D. Turney, and Eduard H. Hovy. 2016. TabMCQ: A Dataset of General Knowledge Tables and Multiple-choice Questions. *CoRR* abs/1602.03960 (2016).
- [20] Mandar Joshi, Eunsol Choi, Daniel S. Weld, and Luke Zettlemoyer. 2017. TriviaQA: A Large Scale Distantly Supervised Challenge Dataset for Reading Comprehension. In *ACL*.
- [21] Vladimir Karpukhin, Barlas Oguz, Sewon Min, Patrick S. H. Lewis, Ledell Wu, Sergey Edunov, Danqi Chen, and Wen-tau Yih. 2020. Dense Passage Retrieval for Open-Domain Question Answering. In *EMNLP*.
- [22] Esther Kaufmann and Abraham Bernstein. 2010. Evaluating the usability of natural language query languages and interfaces to Semantic Web knowledge bases. *Journal of Web Semantics* (2010).
- [23] Tom Kwiatkowski, Jennimaria Palomaki, Olivia Redfield, Michael Collins, Ankur Parikh, Chris Alberti, Danielle Epstein, Illia Polosukhin, Jacob Devlin, Kenton Lee, Kristina Toutanova, Llion Jones, Matthew Kelcey, Ming-Wei Chang, Andrew M. Dai, Jakob Uszkoreit, Quoc Le, and Slav Petrov. 2019. Natural Questions: A Benchmark for Question Answering Research. *Transactions of the Association for Computational Linguistics* (2019).
- [24] Kenton Lee, Ming-Wei Chang, and Kristina Toutanova. 2019. Latent Retrieval for Weakly Supervised Open Domain Question Answering. In *ACL*.
- [25] Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. 2021. Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks. arXiv:2005.11401 [cs.CL]
- [26] Svetlana Mansmann, Nafees Ur Rehman, Andreas Weiler, and Marc H. Scholl. 2014. Discovering OLAP dimensions in semi-structured data. *Inf. Syst.* (2014).
- [27] Julian J. McAuley, Christopher Targett, Qinfeng Shi, and Anton van den Hengel. 2015. Image-Based Recommendations on Styles and Substitutes. In *SIGIR*. ACM.
- [28] Tri Nguyen, Mir Rosenberg, Xia Song, Jianfeng Gao, Saurabh Tiwary, Rangan Majumder, and Li Deng. 2016. MS MARCO: A Human Generated Machine Reading Comprehension Dataset. In *Proceedings of the Workshop on Cognitive Computation: Integrating neural and symbolic approaches 2016 co-located with the 30th Annual Conference on Neural Information Processing Systems (NIPS 2016), Barcelona, Spain, December 9, 2016 (CEUR Workshop Proceedings)*.
- [29] Barlas Oguz, Xilun Chen, Vladimir Karpukhin, Stan Peshterliev, Dmytro Okhonko, Michael Sejr Schlichtkrull, Sonal Gupta, Yashar Mehdad, and Scott Yih. 2022. UniK-QA: Unified Representations of Structured and Unstructured Knowledge for Open-Domain Question Answering. In *ACL Findings*.
- [30] Panupong Pasupat and Percy Liang. 2015. Compositional Semantic Parsing on Semi-Structured Tables. Association for Computational Linguistics.
- [31] Pranav Rajpurkar, Jian Zhang, Konstantin Lopyrev, and Percy Liang. 2016. SQuAD: 100,000+ Questions for Machine Comprehension of Text. In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics, Austin, Texas.
- [32] Pum-Mo Ryu, Myung-Gil Jang, and Hyunki Kim. 2014. Open domain question answering using Wikipedia-based knowledge model. *Inf. Process. Manag.* (2014).
- [33] Weijia Shi, Sewon Min, Michihiro Yasunaga, Minjoon Seo, Rich James, Mike Lewis, Luke Zettlemoyer, and Wen-tau Yih. 2023. REPLUG: Retrieval-Augmented Black-Box Language Models. 2301.12652 (2023).

- [34] Amit Singhal. 2001. Modern Information Retrieval: A Brief Overview. *IEEE Data Eng. Bull.* (2001).
- [35] Arnab Sinha, Zhihong Shen, Yang Song, Hao Ma, Darrin Eide, Bo-June Paul Hsu, and Kuansan Wang. 2015. An Overview of Microsoft Academic Service (MAS) and Applications. In *WWW*.
- [36] Kai Sun, Yifan Ethan Xu, Hanwen Zha, Yue Liu, and Xin Luna Dong. 2023. Head-to-Tail: How Knowledgeable are Large Language Models (LLM)? A.K.A. Will LLMs Replace Knowledge Graphs? (2023).
- [37] Alon Talmor and Jonathan Berant. 2018. The Web as a Knowledge-Base for Answering Complex Questions. In *NAACL-HLT*.
- [38] Alon Talmor and Jonathan Berant. 2018. The Web as a Knowledge-Base for Answering Complex Questions. Association for Computational Linguistics, New Orleans, Louisiana.
- [39] Adam Trischler, Tong Wang, Xingdi Yuan, Justin Harris, Alessandro Sordoni, Philip Bachman, and Kaheer Suleman. 2017. NewsQA: A Machine Comprehension Dataset. In *Proceedings of the 2nd Workshop on Representation Learning for NLP*. Association for Computational Linguistics.
- [40] Kuansan Wang, Zhihong Shen, Chiyuan Huang, Chieh-Han Wu, Yuxiao Dong, and Anshul Kanakia. 2020. Microsoft Academic Graph: When experts are not enough. *Quant. Sci. Stud.* (2020).
- [41] Johannes Welbl, Pontus Stenetorp, and Sebastian Riedel. 2018. Constructing Datasets for Multi-hop Reading Comprehension Across Documents. *Trans. Assoc. Comput. Linguistics* (2018).
- [42] Johannes Welbl, Pontus Stenetorp, and Sebastian Riedel. 2018. Constructing Datasets for Multi-hop Reading Comprehension Across Documents. *Trans. Assoc. Comput. Linguistics* (2018).
- [43] Wei Yang, Yuqing Xie, Aileen Lin, Xingyu Li, Luchen Tan, Kun Xiong, Ming Li, and Jimmy Lin. 2019. End-to-End Open-Domain Question Answering with BERTserini. In *NAACL-HLT*.
- [44] Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Bengio, William Cohen, Ruslan Salakhutdinov, and Christopher D. Manning. 2018. HotpotQA: A Dataset for Diverse, Explainable Multi-hop Question Answering. Association for Computational Linguistics.
- [45] Zhilin Yang, Peng Qi, Saizheng Zhang, Yoshua Bengio, William W Cohen, Ruslan Salakhutdinov, and Christopher D Manning. 2018. HotpotQA: A dataset for diverse, explainable multi-hop question answering. *EMNLP* (2018).
- [46] Michihiro Yasunaga, Hongyu Ren, Antoine Bosselut, Percy Liang, and Jure Leskovec. 2021. QA-GNN: Reasoning with Language Models and Knowledge Graphs for Question Answering.
- [47] Wen-tau Yih, Ming-Wei Chang, Xiaodong He, and Jianfeng Gao. 2015. Semantic Parsing via Staged Query Graph Generation: Question Answering with Knowledge Base. In *ACL*.
- [48] Wen-tau Yih, Matthew Richardson, Chris Meek, Ming-Wei Chang, and Jina Suh. 2016. The Value of Semantic Parse Labeling for Knowledge Base Question Answering. Association for Computational Linguistics, Berlin, Germany.
- [49] Tao Yu, Rui Zhang, Kai Yang, Michihiro Yasunaga, Dongxu Wang, Zifan Li, James Ma, Irene Li, Qingning Yao, Shanelle Roman, Zilin Zhang, and Dragomir Radev. 2018. Spider: A Large-Scale Human-Labeled Dataset for Complex and Cross-Domain Semantic Parsing and Text-to-SQL Task. Association for Computational Linguistics, Brussels, Belgium.
- [50] Victor Zhong, Caiming Xiong, and Richard Socher. 2017. Seq2SQL: Generating Structured Queries from Natural Language using Reinforcement Learning. *CoRR* abs/1709.00103 (2017).
- [51] Yutao Zhu, Huaying Yuan, Shuting Wang, Jiongnan Liu, Wenhan Liu, Chenlong Deng, Zhicheng Dou, and Ji-Rong Wen. 2023. Large language models for information retrieval: A survey. *arXiv:2308.07107* (2023).

- [52] Honglei Zhuang, Zhen Qin, Kai Hui, Junru Wu, Le Yan, Xuanhui Wang, and Michael Bendersky. 2023. Beyond Yes and No: Improving Zero-Shot LLM Rankers via Scoring Fine-Grained Relevance Labels. *CoRR* abs/2310.14122 (2023).

A Benchmark details

	relational structure	textual information
STARK-AMAZON	Amazon Product Reviews	Amazon Product Reviews
STARK-MAG	ogbn-mag	Amazon Question and Answer Data
STARK-PRIME	PrimeKG	ogbn-papers100M, Microsoft Academic Graph disease: Orphanet; drug: DrugBank; pathway: Reactome; gene: Ensembl, NCBI Entrez, Uniprot, UCSC, CPDB

Table 8: Sources of relational structure and textual information of the benchmarks

A.1 STARK-MAG

Relational query templates:

metapath	multi-hop query template
(author $\rightarrow$ paper)	"Can you list the papers authored by <author>?"
(paper $\rightarrow$ paper)	"Which papers have been cited by the paper <paper>?"
(field_of_study $\rightarrow$ paper)	"Can you provide a list of papers in the field of <field_of_study>?"
(institution $\rightarrow$ author $\rightarrow$ paper)	"What papers have been published by researchers from <institution>?"
(paper $\rightarrow$ author $\rightarrow$ paper)	"What papers have been published by researchers that are coauthors of <paper>?"
(paper $\rightarrow$ author $\rightarrow$ paper $\leftarrow$ field_of_study $\leftarrow$ paper)	"Can you find papers that share a coauthor with <paper> and are also in the same field of study?"
(institution $\rightarrow$ author $\rightarrow$ paper $\leftarrow$ field_of_study)	"Are there any papers associated with <institution> and are in the field of <field_of_study>?"

For example, the metapath (field_of_study $\rightarrow$ paper) requires an initial field_of_study entity to be filled in the corresponding query template. For multi-hop metapaths, the last metapath (institution $\rightarrow$ author $\rightarrow$ paper $\leftarrow$ field_of_study) requires an institution entity and a field_of_study entity to initialize the query.

A.2 STARK-PRIME

Number of entities in each type.

```
#disease:          17,080
#gene/protein:     27,671
#molecular_function: 11,169
#drug:             7,957
#pathway:          2,516
#anatomy:          14,035
#effect/phenotype: 15,311
#biological_process: 28,642
#cellular_component: 4,176
#exposure:         818
```

Relational tuples.

```
[(anatomy, expression absent, gene/protein),
 (anatomy, expression present, gene/protein),
 (anatomy, parent-child, anatomy),
 (biological_process, interacts with, exposure),
 (biological_process, interacts with, gene/protein),
```

```

(biological_process, parent-child, biological_process),
(cellular_component, interacts with, exposure),
(cellular_component, interacts with, gene/protein),
(cellular_component, parent-child, cellular_component),
(disease, associated with, gene/protein),
(disease, contraindication, drug),
(disease, indication, drug),
(disease, linked to, exposure),
(disease, off-label use, drug),
(disease, parent-child, disease),
(disease, phenotype absent, effect/phenotype),
(disease, phenotype present, effect/phenotype),
(drug, carrier, gene/protein),
(drug, contraindication, disease),
(drug, enzyme, gene/protein),
(drug, indication, disease),
(drug, off-label use, disease),
(drug, side effect, effect/phenotype),
(drug, synergistic interaction, drug),
(drug, target, gene/protein),
(drug, transporter, gene/protein),
(effect/phenotype, associated with, gene/protein),
(effect/phenotype, parent-child, effect/phenotype),
(effect/phenotype, phenotype absent, disease),
(effect/phenotype, phenotype present, disease),
(effect/phenotype, side effect, drug),
(exposure, interacts with, biological_process),
(exposure, interacts with, cellular_component),
(exposure, interacts with, gene/protein),
(exposure, interacts with, molecular_function),
(exposure, linked to, disease),
(exposure, parent-child, exposure),
(gene/protein, associated with, disease),
(gene/protein, associated with, effect/phenotype),
(gene/protein, carrier, drug),
(gene/protein, enzyme, drug),
(gene/protein, expression absent, anatomy),
(gene/protein, expression present, anatomy),
(gene/protein, interacts with, biological_process),
(gene/protein, interacts with, cellular_component),
(gene/protein, interacts with, exposure),
(gene/protein, interacts with, molecular_function),
(gene/protein, interacts with, pathway),
(gene/protein, ppi, gene/protein),
(gene/protein, target, drug),
(gene/protein, transporter, drug),
(molecular_function, interacts with, exposure),
(molecular_function, interacts with, gene/protein),
(molecular_function, parent-child, molecular_function),
(pathway, interacts with, gene/protein),
(pathway, parent-child, pathway)]

```

Relational query templates.

```

{
(effect/phenotype → [phenotype absent] → disease ← [!indication] ← drug):
    "Find diseases with zero indication drug and are associated with <effect/phenotype>",
(drug → [contraindication] → disease ← [associated with] ← gene/protein):
    "Identify diseases associated with <gene/protein> and are contraindicated with <drug>",
(anatomy → [expression present] → gene/protein ← [expression absent] ← anatomy):
    "What gene or protein is expressed in <anatomy1> while is absent in <anatomy2>?",
(anatomy → [expression absent] → gene/protein ← [expression absent] ← anatomy):
}

```

```

    "What gene/protein is absent in both <anatomy1> and <anatomy2>?",
    (drug → [carrier] → gene/protein ← [carrier] ← drug):
    "Which target genes are shared carriers between <drug1> and <drug2>?",
    (anatomy → [expression present] → gene/protein → [target] → drug):
    "What is the drug that targets the genes or proteins which are expressed in <anatomy>?",
    (drug → [side effect] → effect/phenotype → [side effect] → drug):
    "What drug has common side effects as <drug>?",
    (drug → [carrier] → gene/protein → [carrier] → drug):
    "What is the drug that has common gene/protein carrier with <drug>?",
    (anatomy → [expression present] → gene/protein → enzyme → drug):
    "What is the drug that some genes or proteins act as an enzyme upon, where the genes or proteins are expressed in <anatomy>?",
    (cellular_component → [interacts with] → gene/protein → [carrier] → drug):
    "What is the drug carried by genes or proteins that interact with <cellular_component>?",
    (molecular_function → [interacts with] → gene/protein → [target] → drug):
    "What drug targets the genes or proteins that interact with <molecular_function>?",
    (effect/phenotype → [side effect] → drug → [synergistic interaction] → drug):
    "What drug has a synergistic interaction with the drug that has <effect/phenotype> as a side effect?",
    (disease → [indication] → drug → [contraindication] → disease):
    "What disease is a contraindication for the drugs indicated for <disease>?",
    (disease → [parent-child] → disease → [phenotype present] → effect/phenotype):
    "What effect or phenotype is present in the sub type of <disease>?",
    (gene/protein → [transporter] → drug → [side effect] → effect/phenotype):
    "What effect or phenotype is a [side effect] of the drug transported by <gene/protein>?",
    (drug → [transporter] → gene/protein → [interacts with] → exposure):
    "What exposure may affect <drug>'s efficacy by acting on its transporter genes?",
    (pathway → [interacts with] → gene/protein → [ppi] → gene/protein):
    "What gene/protein interacts with the gene/protein that related to <pathway>?",
    (drug → [synergistic interaction] → drug → [transporter] → gene/protein):
    "What gene or protein transports the drugs that have a synergistic interaction with <drug>?",
    (biological_process → [interacts with] → gene/protein → [interacts with] → biological_process):
    "What biological process has the common interactino pattern with gene or proteins as <biological_process>?",
    (effect/phenotype → [associated with] → gene/protein → [interacts with] → biological_process):
    "What biological process interacts with the gene/protein associated with <effect/phenotype>?",
    (drug → [transporter] → gene/protein → [expression present] → anatomy):
    "What anatomy expressed by the gene/protein that affect the transporter of <drug>?",
    (drug → [target] → gene/protein → [interacts with] → cellular_component):
    "What cellular component interacts with genes or proteins targeted by <drug>?",
    (biological_process → [interacts with] → gene/protein → [expression absent] → anatomy):
    "What anatomy does not express the genes or proteins that interacts with <biological_process>?",
    (effect/phenotype → [associated with] → gene/protein → [expression absent] → anatomy):
    "What anatomy does not express the genes or proteins associated with <effect/phenotype>?",
    (drug → [indication] → disease → [indication] → drug) & (drug → [synergistic interaction] → drug):
    "Find drugs that has a synergistic interaction with <drug> and both are indicated for the same disease.",
    (pathway → [interacts with] → gene/protein → [interacts with] → pathway) & (pathway → [parent-child] → pathway):
    "Find pathway that is related with <pathway> and both can [interacts with] the same gene/protein.",
    (gene/protein → [associated with] → disease → [associated with] → gene/protein) & (gene/protein → [ppi] → gene/protein):
    "Find gene/protein that can interact with <gene/protein> and both are associated with the same disease.",
    (gene/protein → [associated with] → effect/phenotype → [associated with] → gene/protein) & (gene/protein → [ppi] → gene/protein):
    "Find gene/protein that can interact with <gene/protein> and both are associated with the same effect/phenotype."
}

```

where $[\cdot]$ denotes the relation type.

B Prompts

B.1 Extracting textual requirements

Prompt for STARK-AMAZON: Textual requirement extraction

You are an intelligent assistant that extracts diverse positive requirements
 $\hookrightarrow$ and negative perspectives for an Amazon product. I will give you the
 $\hookrightarrow$ following information:

- product: <product name>
- dimensions: <product dimensions>
- weight: <product weight>
- description: <product description>
- features:
 - #1: <feature #1>
 - ...

```

- reviews:
  #1:
    summary: <review summary>
    text: <full review text>
  #2:
    ...
- Q&A:
  #1:
    question: <product-related question>
    answer: <answer to product-related question>
  #2:
    ...

```

Based on the given product information, you need to (1) identify the product's
 ↳ generic category, (2) list all of the negative perspectives and their
 ↳ sources, and (2) extract up to five hard and five soft requirements
 ↳ relevant to customers' interests along with their sources. (1) For
 ↳ example, the product's generic category can be "a chess book" or "a
 ↳ phone case for iphone 6", do not use the product name directly. (2)
 ↳ Negative perspectives are those that the product doesn't fulfill, which
 ↳ come from the negative reviews or Q&A. (3) For the requirements, you
 ↳ should only focus on the product's advantages and positive perspects.
 ↳ Hard requirements mean that product must fulfil, such as size and
 ↳ functionality. Soft requirements are not as strictly defined but still
 ↳ desirable, such as a product is easy-to-use. For (2) and (3), each
 ↳ source is a composite of the key and index (if applicable) separated by
 ↳ "-", such as "description", "Q&A-#1". You should provide the response
 ↳ in a specific format as follows where "item" refers to the product's
 ↳ generic category, e.g., "a chess book".

Response format:

```

{
  "item": <the product's generic category> ,
  "negative": [[<source of negative perspective>, <negative perspective
    ↳ description>]],
  "hard": [[<source of hard requirement>, <hard requirement description>],
    ↳ ...],
  "soft": [[<source of soft requirement>, <soft requirement description>], ...]
}

```

Here is an example of the response:

```

{
  "item": "a camping chair",
  "negative": [[<"reviews-#3", "the chair is not sturdy enough">, [<"Q&A-#1", "
    ↳ wrong color">]],
  "hard": [[<"description", "has a breathable mesh back">, [<"description", "the
    ↳ arm is adjustable">, [<"dimensions", "more than 35 inches long">, [<"
    ↳ features-#7", "with a arm rest cup holder">, [<"Q&A-#4", "need to come
    ↳ with a carrying bag">]],
  "soft": [[<"description", "suitable for outdoors">, [<"features-#9", "compact
    ↳ and save space">, [<"reviews-#6", "light and portable">]]
}

```

This is the information of the product that you need to write response for:

```

Information:
<product_doc>
Response:

```

Prompt for STARK-MAG: Textual requirement extraction

You are a helpful assistant that helps me extract one short requirement (no
→ more than 10 words) about a paper from the paper information that
→ researchers might be interested in. The requirement can be about the
→ paper content, publication date, publication venue, etc. The
→ requirement should be general and not too specific. I will give you the
→ paper information, and you should return a short phrase about the
→ paper, starting with 'the paper...'. This is the paper information:
<doc_info>
Please only return the short and general requirement without additional
→ comments.

Prompt for STARK-PRIME: Textual requirement extraction

You are a helpful assistant that helps me extract <n_properties> from a given <
→ target> information that a <role> may be interested in.
<role_instruction>
Each property should be no more than 10 words and start with "the <target>".
→ You should also include the source of each property as indicated in the
→ paragraph names of the information, e.g., "details.mayo_symptoms", "
→ details.summary", etc. You should return a list of properties and their
→ sources following the format:
[["<short_property1>", "<source1>"], ["<short_property2>", "<source2>"], ...]
This is the information:
<doc_info>
Please provide only the list with <n_properties> in your response. Response:

According to the role assigned to simulate the query content, the <role_instruction> as shown below is filled in accordingly.

role	role instruction
Doctor	Doctors typically ask questions aimed at diagnosing and treating. Their questions tend to be direct and practical, focusing on aspects involving side effects, symptoms, and complications etc.
Medical scientist	Medical scientists often ask questions that reflect the complexity and depth of the scientific inquiry in the medical field. Their questions tend to be detailed and specific, focusing on aspects such as: etiology and pathophysiology, genetic factors, association with pathway, protein, or molecular function.
Patient	Patients typically don't know the professional medical terminology. Their questions tend to be straightforward, focusing on practical concerns on the symptoms, effects, and inheritance etc., instead of the detailed mechanisms, which may also include more context.

B.2 Combining relational and textual requirements

Prompt for STARK-AMAZON: Fuse relational and textual requirements

You are an intelligent assistant that generates queries about an Amazon item. I
→ will provide you with the item name, requirements, and its negative
→ customer reviews. Your task is to create a natural-sounding customer

↪ query that leads to the item as the answer, using the requirements that
↪ are non-conflicting with the negative reviews, and provide the indices
↪ of the requirements used. For example:

Information:

- item: a soccer rebounder
- requirements:
#1: needs a heavy-duty 1-inch to 3-inch steel tube frame
#2: should be adjustable for practicing different skills
#3: should be durable
#4: usually be viewed together with <SKLZ Star-Kick Hands Free Solo Soccer
↪ Trainer>
- negative reviews:
#1: it was broken after a few uses

Response:

```
{
  "index": [1, 2, 4],
  "query": "Please recommend a soccer rebounder with a steel frame, about 2
    ↪ inches thick, that can adapt to different skill levels. We had a
    ↪ blast using the <SKLZ Star-Kick Hands Free Solo Soccer Trainer> with
    ↪ my family, and I'm on the lookout for something similar."
}
```

As the negative review indicates that the soccer rebounder lacks durability,
↪ your query should only incorporate requirements #1, #2, and #4 while
↪ excluding #3. A requirement should only be excluded if it conflicts
↪ with negative feedback or is unlikely to align with customers'
↪ interests. For relational requirements about another <product>, do not
↪ directly use "usually bought/viewed together with <product>" in the
↪ query. You must deduce the item's relationship with <product> into
↪ substitute or complement, and create various user scenarios, such as
↪ the item should be compatible or used with <product> (for complements)
↪ or match in style with <product> (for substitute), to make the queries
↪ sound natural. Except for <product>, you should change the description
↪ but convey similar meanings. The query structure is completely flexible.
↪ Here is the information to generate the requirement indices and a
↪ natural-sounding query:

Information:

<product_req_and_neg_comments>

Response:

Prompt for STARK-MAG: Fuse relational and textual requirements

You are a helpful assistant that helps me generate a new query by incorporating
↪ an additional requirement into a given query, and form a coherent and
↪ natural-sounding question.

This is the existing query:

<query>

This is the additional requirement:

<additional_textual_requirement>

You should be creative in combining the existing query and requirement, and

↪ flexible in structuring the new query, adding context as needed. Please
↪ return the new query without additional comments:

The prompt of a second-time rewrite by GPT-4 Turbo:

You are a helpful assistant that helps make a researcher's query about a paper
↪ more natural-sounding, akin to the language used in ArXiv web searches.

↪ You should change the description but convey similar meanings. The
↪ query structure is completely flexible. The original query:
"<query>"
Please only output the new query without additional comments:

Prompt for STARK-PRIME: Fuse relational and textual requirements

You are a helpful assistant that helps me generate a natural-sounding and
↪ coherent query as if you were a <role>. The query should be created
↪ based on a list of requirements for searching <plural_target> in a
↪ database. I will provide you with the requirements in the following
↪ format:
[<requirement1>, <requirement2>, ...]
You should create the query based solely on the given requirements. Moreover,
↪ you should craft the query from the perspective of a <role>.
<role_instruction>
For example, a query from a <role> could be
"<example_query>"
You can be flexible in structuring the query and adding additional context.
↪ Ensure that the query uses different descriptions than the original
↪ property descriptions while retaining similar meanings. The query
↪ should sound concise and natural. These are the requirements:
<requirements>
Please create the query based on the given requirements and provide only the
↪ query without additional comments. Your response:

The prompt of a second-time rewrite by GPT-4 Turbo:

You are a helpful assistant that helps me rewrite a query that searches for <
↪ plural_target> from the perspective of a <role>. You should maintain
↪ the requirements from the original query and the characteristics of the
↪ <role>, while being creative and flexible in structuring the query.
↪ Ensure the revised query is concise and natural-sounding. Original
↪ query: "<query>". Please output only the rewritten query:

B.3 Filtering additional answers

Prompt for STARK-AMAZON: Filtering additional answers

Filter products by general category

You are an intelligent assistant that identifies whether an Amazon product
↪ belongs to a given category. I will give you the product information.
↪ You should only answer yes / no in the response. For examples, the
↪ product <SKLZ Star-Kick Hands Free Solo Soccer Trainer> belongs to the
↪ category "soccer trainer" and the product <Test your Opening,
↪ Middlegame and Endgame Play - VOLUME 2> belongs to the category "a
↪ chess opening book", while <Baby Girls One-piece Shiny Athletic Leotard
↪ Ballet Tutu with Bow> doesn't belong to category "an adult tutu".

Information:

- product title: <<product_title>>
- product description: <product_description>

Does the product belong to "<target_category>"? Response (yes/no):

Filter products by requirements

You are a helpful assistant that helps me check whether an Amazon product

- ↪ satisfies the given requirements. I will provide you with the product
- ↪ information, which may include the product description, features,
- ↪ reviews, and Q&A from customers. Your task is to assess whether the
- ↪ product meets each requirement based on the provided information. If
- ↪ there is no information that supports the requirement, your response
- ↪ for that requirement is "NA". If there is relevant information that
- ↪ supports the requirement, your response for that requirement is the
- ↪ information source that fulfills the requirement. Each information
- ↪ source is a composite of the key and index (if applicable), separated
- ↪ by "-", such as "description", "features-#3", "Q&A-#1", "reviews-#2".
- ↪ If there are multiple sources,

Response:

```
{
  1: "NA" or [the information sources that satisfy the requirement #1],
  2: "NA" or [the information sources that satisfy the requirement #2],
  ...
}
```

Here is the product information:

<product_doc>

The requirements are as follows:

<customer_requirements>

Response:

Prompt for STARK-MAG: Filtering additional answers

You are a helpful assistant that helps me verify whether a given <

- ↪ target_node_type> is subject to a requirement. I will provide you with
- ↪ the <target_node_type> information and the requirement, and you should
- ↪ return only a 'True' or 'False' value, indicating whether the <
- ↪ target_node_type> meets the requirement.

This is the <target_node_type> information:

<doc_info>

This is the requirement:

<additional_textual_requirement>

Please return only the boolean value without additional comments:

Prompt for STARK-PRIME: Filtering additional answers

You are a helpful assistant tasked with verifying whether a given <target>

- ↪ satisfies each of the provided requirements. I will give you the
- ↪ requirements in the following format:

{1: <requirement1>, 2: <requirement2>, ...}

When evidence in the <target> information confirms a requirement is met, cite

- ↪ the source, for example, 'details.mayo_symptoms', 'details.summary'.
- ↪ no direct evidence exists, indicate this with 'NA'. The output in JSON
- ↪ format should be as follows:

{1: 'NA' or <source1>, 2: 'NA' or <source2>, ...}

This is the <target> information:

<doc_info>

These are the requirements:

<requirements>

Please provide only the JSON in your response. Response:"