

Entity Linking for Retrieval-Augmented Factoid Question Answering with Large Language Models

Anonymous ACL submission

Abstract

Factoid questions seek real world factual knowledge. We explore the application of retrieval augmented question answering with large language models in the context of factoid questions. Our novel approach, *Entity Retrieval*, employs entity linking to identify relevant documents, offering an alternative to dense retrieval techniques. Our findings establish that retrieval-augmented methods are particularly effective for smaller large language models, and that *Entity Retrieval* outperforms other retrieval methods while demanding reduced time and resources¹.

1 Introduction

Factoid questions seek factual information about the real world, and typically have answers that are concise single words or short phrases. These answers often reference or directly stem from a knowledge base entity (Ranjan and Balabantaray, 2016). The literature on factoid question answering is replete with a variety of well-studied methodologies, including rule-based (Shitu et al., 2020), pattern-based (Pala Er and Cicekli, 2013), and neural network-based approaches (Lukovnikov et al., 2017; Mohammed et al., 2018; Lukovnikov et al., 2019).

In recent years, Large Language Models (LLMs; OpenAI, 2023; Touvron et al., 2023; Jiang et al., 2024) have significantly transformed the field of natural language processing. Retrieval-augmented generation (RAG; Lewis et al., 2020b; Izacard and Grave, 2021; Singh et al., 2021) has emerged as a prevalent approach to question answering with LLMs. RAG systems typically utilize the retriever-reader architecture (Chen et al., 2017), where the retriever can be sparse (Peng et al., 2023), dense (Karpukhin et al., 2020), or a hybrid of the two

¹We have included our source code implementation of the project, along with the generated model answers, in the Software section of our ARR submission.

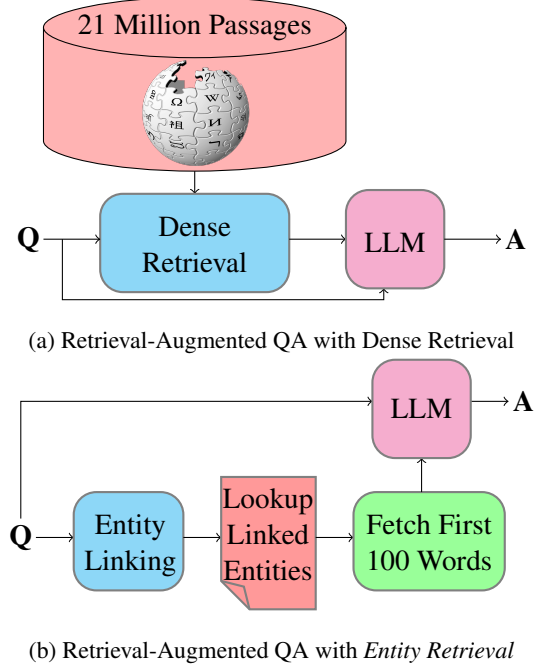


Figure 1: *Entity Retrieval* simplifies the process of obtaining augmentation documents by replacing the need to search through large indexed passages with a straightforward lookup.

(Glass et al., 2022). The reader, a generative language model (e.g., BART; Lewis et al., 2020a, T5; Raffel et al., 2020, GPT-3; Brown et al., 2020), then conditions the generated output based on the documents retrieved by the retriever. Recent RAG methodologies leverage the in-context learning capabilities of LLMs to integrate the retrieved documents into the prompt (Shi et al., 2023; Peng et al., 2023; Yu et al., 2023).

Following the prevalent belief that factoid question answering is nearing resolution (Petrochuk and Zettlemoyer, 2018), there has been a shift in the community towards reading comprehension (Joshi et al., 2017), multi-hop (Yang et al., 2018), and commonsense (Talmor et al., 2019) question answering. This shift has relegated factoid question answering to a less central position. Consequently,

a reliable evaluation of the performance of common retrieval-augmented question answering methods for factoid question types is currently lacking.

In this paper, we examine the suitability of retrieval-augmented question answering methods for factoid questions. In addition, following the established methodologies in statistical factoid question answering systems (Aghaebrahimian and Jurčiček, 2016; Li et al., 2021; Adebisi et al., 2022), we study the role of entity linking as a critical component in retrieval-augmented factoid question answering with LLMs. Our contributions in this paper are as follows:

- (1) We examine the performance of retrieval- and non-retrieval-augmented LLMs for answering factoid questions.
- (2) We propose *Entity Retrieval*, a simple yet effective approach to employ entity linking in retrieval-augmented factoid question answering with LLMs.

Figure 1 presents a schematic comparison between *Entity Retrieval* and the common dense retrieval approach (e.g. Karpukhin et al., 2020) in identifying retrieval documents to enhance question answering with LLMs.

2 Retrieval-Augmentation Techniques

Retrieval-augmentation (Lewis et al., 2020b) is a method of converting closed-book question answering² (Roberts et al., 2020) into extractive question answering (Abney et al., 2000; Rajpurkar et al., 2016), where the answers can be directly extracted from the retrieved documents. Despite the abundance of effective retrieval-augmentation techniques for question answering in existing literature, this section will concentrate on a select few methods utilized to study the factoid question answering capabilities of LLMs in this paper.

Dense Passage Retrieval (DPR; Karpukhin et al., 2020) leverages a bi-encoder architecture, wherein the initial encoder processes the question and the subsequent encoder handles the passages to be retrieved. The similarity scores between the two encoded representations are computed using a dot product. The encoded representations of the second encoder are fixed and indexed in FAISS (Johnson et al., 2019; Douze et al., 2024), while the first encoder is optimized to maximize the dot-product

scores based on positive and negative examples. The performance of DPR solidifies its position as a superior retriever compared to BM25-based (Robertson et al., 1994) sparse retrieval methods for question answering.

REPLUG (Shi et al., 2023) views LLM as a black-box, encoding each of the k most relevant retrieved documents along with the input query to generate k probability distributions over the forthcoming token. These distributions are then weighted averaged, considering the similarity of each retrieved document to the original input query. The strength of REPLUG lies in its ability to infuse the knowledge from the retrieved documents while generating the answer. This makes REPLUG a compelling candidate for studying retrieval-augmented question answering.

3 Entity Linking for Question Answering

While quite powerful, most retrieval-augmented systems are notably time and resource-intensive, necessitating the storage of extensive lookup indices and the need to attend to all retrieved documents to generate a response.

Entity linking has been an integral component of statistical factoid question answering systems (Aghaebrahimian and Jurčiček, 2016, *inter alia*). Additionally, the extensively studied field of Knowledge Base Question Answering (Cui et al., 2017, *inter alia*) has underscored the significance of entity information from knowledge bases in question answering (Salnikov et al., 2023).

A traditional neural question answering pipeline may contain entity detection, entity linking, relation prediction, and evidence integration (Mohammed et al., 2018; Lukovnikov et al., 2019), where entity detection can employ LSTM-based (Hochreiter and Schmidhuber, 1997) or BERT-based (Devlin et al., 2019) encoders. Inspired by this body of work, we investigate the relevance of entity linking as an alternative strategy to dense retrieval methods for augmenting factoid question answering with LLMs. We propose *Entity Retrieval*, a method employing a simple heuristic for implementing entity linking-based document retrieval. *Entity Retrieval* leverages entity linking to identify entities within the question and retrieves corresponding knowledge base articles, providing the first 100 words of each article as the retrieved documents (see Figure 1.b).

²Closed-book QA focuses on answering questions without additional context during inference.

4 Experiments

4.1 Setup

We focus on Wikipedia as the knowledge base and utilize the pre-existing Wikipedia passages and the dense retrieval model available in the `wiki_dpr`³ repository from huggingface. `wiki_dpr` follows established practices (Chen et al., 2017; Karpukhin et al., 2020) and segments the articles into non-overlapping text blocks of 100 words, resulting in 21,015,300 passages. These passages are processed with a pre-trained DPR context encoder, generating fixed embedding vectors stored in a FAISS index (Douze et al., 2024). Factoid questions are encoded using the DPR question encoder, and the top k relevant passages to the encoded question are retrieved from the FAISS index. We use the exact FAISS index storage, single-nq DPR question encoder, and retrieve the top 4 documents for each question, in our experiments. As well for better time efficiency, following Ram et al. (2023), we treat document retrieval as a pre-processing step, caching the most relevant passages for each question before conducting the question answering experiments.

For entity linking in *Entity Retrieval*, we select SPEL (Shavarani and Sarkar, 2023) mainly due to its near-perfect linking precision. Architecturally, SPEL comprises an entity knowledge fine-tuned RoBERTa (Liu et al., 2019) model as the encoder and a classification layer atop the encoder which maps the encoded representations to the space of predicted entities. SPEL models entity linking as structured prediction which enables it to be fast and minimal resource demanding. In this study, we employ the fine-tuned SPEL-large model with an entity vocabulary of 500K, enabling identification of entity mentions referencing the 500K most hyperlinked Wikipedia pages.

Given the proven effectiveness of utilizing initial sentences from Wikipedia pages for entities in tasks such as document classification (Shavarani and Sekine, 2020) and question answering (Choi et al., 2018), we propose employing the first 100 words of Wikipedia articles corresponding to the identified entities in questions as retrieved documents for *Entity Retrieval* settings. We consider two such settings: (1) using SPEL for question annotation and utilizing its suggested linked entities to retrieve Wikipedia articles, (2) using gold entity link annotations for dataset questions to retrieve

the Wikipedia articles.

For LLMs, we consider the open weight LLaMA 2 (Touvron et al., 2023) model in all three available sizes (7B, 13B, and 70B). However, due to hardware constraints — limited to 2 RTX A6000s with 49GB GPU memory each — we utilize the 8-bit quantized version of the 70B model. In all our experiments with LLaMA 2, we prevent it from generating sequences longer than 10 subwords. Additionally, we evaluate GPT 4 (0613 version) from OpenAI (2023).

As a public implementation of REPLUG is not available, we implement it with the haystack⁴ library, employing our cached DPR passages for each question to autoregressively generate answers.

To verify the capacity of LLMs in utilizing the retrieved documents without additional fine-tuning or further in-context examples, we do not use any training question-answer pairs in the prompts of our models. Aside from a simple instruction for answering the question, in the closed-book setting, the prompt solely comprises the question, while in the DPR and REPLUG settings, it includes the retrieved documents from the DPR cache along with the question. Similarly, for the *Entity Retrieval* settings, the prompt consists of the first 100 words of the Wikipedia pages corresponding to the identified or gold entities in the question. We follow Ram et al. (2023) for question normalization and prompt formulation.

4.2 Data

We use the following datasets in our experiments:

FactoidQA (Smith et al., 2008) contains 2203 hand crafted factoid question-answer pairs derived from Wikipedia articles, with each pair accompanied by its corresponding Wikipedia source article included in the dataset. We use OpenQA-eval (Kamalloo et al., 2023) scripts to evaluate model performance, reporting exact match (EM) and F1 scores by comparing expected answers to model responses for FactoidQA questions.

StrategyQA (Geva et al., 2021) is a complex boolean question answering dataset, constructed by presenting individual terms from Wikipedia to annotators. Its questions contain references to more than one Wikipedia entity, and necessitate implicit reasoning for binary responses. The dataset comprises 5111 answered questions which are split into two subsets: train and train_filtered subsets

³https://huggingface.co/datasets/wiki_dpr, created on a Wikipedia dump from December 20, 2018.

⁴<https://haystack.deepset.ai>

Setting	LLaMA 2						GPT 4	
	7B		13B		70B-8bQ			
	EM	F1	EM	F1	EM	F1	EM	F1
Closed-book	30.5	38.3	33.7	42.9	37.0	45.9	42.4	55.1
DPR	33.6	42.7	37.1	45.5	35.6	42.0	35.9	47.1
REPLUG	15.8	22.0	27.7	33.4	22.0	25.3	-	-
<i>Entity Retrieval</i> w/ SPEL	38.1	47.3	40.5	49.2	40.6	49.0	37.3	48.0
<i>Entity Retrieval</i> w/ oracle entities [†]	37.2	46.9	40.2	49.3	41.4	49.7	38.5	48.2

Table 1: FactoidQA evaluation results. EM refers to the exact match between predicted and expected answers, disregarding punctuation and articles (a, an, the). [†]*Entity Retrieval* with oracle results are not directly comparable to other approaches, as they leverage gold annotated entity links from the dataset.

Setting		7B		13B		70B-8bQ	
		Acc	Inv #	Acc	Inv #	Acc	Inv #
train	Closed-book	51.8	215	51.4	302	60.3	191
	DPR	52.8	212	52.5	280	52.8	336
	<i>Entity Retrieval</i> w/ SPEL	53.8	175	52.7	200	58.0	152
train_filtered	Closed-book	59.9	286	61.6	337	67.0	232
	DPR	59.8	274	63.7	296	62.7	407
	<i>Entity Retrieval</i> w/ SPEL	60.1	233	64.9	190	66.6	206

Table 2: StrategyQA evaluation with LLaMA 2 results.

containing 2290 and 2821 questions, respectively. For evaluation, we present accuracy scores by comparing model responses to the expected boolean answers in the dataset. As well, to assess model comprehension of the task, we count the number of invalid answers that deviate from Yes or No and report this count in a distinct column labeled “Inv #” for each experiment.

4.3 Results and Analysis

We generate answers to FactoidQA questions for the following settings: (1) closed-book, (2) DPR, (3) REPLUG, (4) *Entity Retrieval* with SPEL-identified entities, and (5) *Entity Retrieval* with oracle entity annotations from the dataset. Table 1 summarizes our evaluation results.

Our experimental results prove that *Entity Retrieval* is a formidable contender among retrieval-augmented techniques for factoid question answering, particularly exhibiting enhanced efficacy with smaller LLMs. The outcomes from our GPT-4⁵

⁵Despite undisclosed specifications of GPT-4 models, extrapolating from the known size of GPT-3 (175B parameters; Brown et al., 2020), it is plausible to estimate GPT-4 to surpass 200 billion parameters, with speculations suggesting over 1 trillion parameters implemented as a mixture of experts model.

experiments substantiate this assertion, revealing a consistent decline in performance across all investigated retrieval-augmentation techniques.

Furthermore, comparing the evaluation results using SPEL identified entities and the oracle entities for *Entity Retrieval*, we realize that despite SPEL’s constrained entity lexicon comprising the 500K most hyperlinked entities, its performance remains notably competitive. While acknowledging this observation, we defer a comprehensive evaluation of alternative entity linking methods beyond SPEL to future investigations.

Table 2 presents our evaluation results for the StrategyQA dataset. Notably, *Entity Retrieval* with oracle annotations is excluded due to the absence of oracle entity links for questions in StrategyQA, while the exclusion of REPLUG is attributed to its comparatively inferior performance relative to DPR in FactoidQA experiments. Our results affirm our previous inference that retrieval-augmentation is not beneficial with sufficiently large models. However, despite the complex reasoning demanded by this dataset, *Entity Retrieval* achieves comparable results to other retrieval-augmented methods, while offering better hardware efficiency. Additionally, invalid count values indicate that *Entity Retrieval* is capable of aiding the model in understanding the boolean nature of expected responses without relying on dense retrieval from millions of passages.

5 Conclusion

We highlight the disproportionate benefit of retrieval augmentation for smaller LLMs in the context of factoid question answering, and introduce *Entity Retrieval* as a promising entity linking-based alternative to dense retrieval for augmenting factoid questions in prompting LLMs.

Limitations and Ethical Considerations

We have not exhaustively explored all potential entity linking methods, which may yield insights enhancing the proposed *Entity Retrieval* approach.

Additionally, due to space constraints and a desire to expedite community engagement, we have not incorporated additional datasets (e.g. 30MFQA; Serban et al., 2016), most of which are annotated with Freebase (Bollacker et al., 2008) and have fallen into disuse following Freebase’s discontinuation. We intend to revitalize such neglected factoid question answering datasets, and we posit that revitalizing these datasets could facilitate the development of a benchmark dataset akin to MMLU (Hendrycks et al., 2021), enabling robust evaluations of newly released LLMs in terms of their factual knowledge capabilities.

Our research is on English only, and we acknowledge that factoid question answering in other languages is also relevant and important. We hope to extend our work to cover multiple languages in the future. We inherit the biases that exist in the data used in this project, and we do not explicitly de-bias the data. We are providing our code to the research community and we trust that those who use the model will do so ethically and responsibly.

References

- Steven Abney, Michael Collins, and Amit Singhal. 2000. [Answer extraction](#). In *Sixth Applied Natural Language Processing Conference*, pages 296–301, Seattle, Washington, USA. Association for Computational Linguistics.
- Emmanuel Adebisi, Bolanle Adefowoke Ojokoh, and Folasade Olubusola Isinkaye. 2022. [An open domain factoid qa framework with improved validation techniques](#). *International Journal of Information Science and Management (IJISM)*, 20(1).
- Ahmad Aghaebrahimian and Filip Jurčiček. 2016. [Open-domain factoid question answering via knowledge graph search](#). In *Proceedings of the Workshop on Human-Computer Question Answering*, pages 22–28, San Diego, California. Association for Computational Linguistics.
- Kurt Bollacker, Colin Evans, Praveen Paritosh, Tim Sturge, and Jamie Taylor. 2008. Freebase: a collaboratively created graph database for structuring human knowledge. In *Proceedings of the 2008 ACM SIGMOD international conference on Management of data*, pages 1247–1250.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind

- Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. [Language models are few-shot learners](#). *Advances in neural information processing systems*, 33:1877–1901.
- Danqi Chen, Adam Fisch, Jason Weston, and Antoine Bordes. 2017. [Reading Wikipedia to answer open-domain questions](#). In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1870–1879, Vancouver, Canada. Association for Computational Linguistics.
- Eunsol Choi, He He, Mohit Iyyer, Mark Yatskar, Wentaoh Yih, Yejin Choi, Percy Liang, and Luke Zettlemoyer. 2018. [QuAC: Question answering in context](#). In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, pages 2174–2184, Brussels, Belgium. Association for Computational Linguistics.
- Wanyun Cui, Yanghua Xiao, Haixun Wang, Yangqiu Song, Seung-won Hwang, and Wei Wang. 2017. [Kbqa: Learning question answering over qa corpora and knowledge bases](#). *Proceedings of the VLDB Endowment*, 10(5).
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: Pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.
- Matthijs Douze, Alexandr Guzhva, Chengqi Deng, Jeff Johnson, Gergely Szilvasy, Pierre-Emmanuel Mazaré, Maria Lomeli, Lucas Hosseini, and Hervé Jégou. 2024. [The faiss library](#). *arXiv preprint arXiv:2401.08281*.
- Mor Geva, Daniel Khashabi, Elad Segal, Tushar Khot, Dan Roth, and Jonathan Berant. 2021. [Did aristotle use a laptop? a question answering benchmark with implicit reasoning strategies](#). *Transactions of the Association for Computational Linguistics*, 9:346–361.
- Michael Glass, Gaetano Rossiello, Md Faisal Mahbub Chowdhury, Ankita Naik, Pengshan Cai, and Alfio Gliozzo. 2022. [Re2G: Retrieve, rerank, generate](#). In *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 2701–2715, Seattle, United States. Association for Computational Linguistics.
- Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. 2021. [Measuring massive multitask language understanding](#). In *International Conference on Learning Representations*.

412	Sepp Hochreiter and Jürgen Schmidhuber. 1997. Long	470
413	short-term memory. <i>Neural computation</i> , 9(8):1735–	471
414	1780.	472
415	Gautier Izacard and Edouard Grave. 2021. Leveraging	473
416	passage retrieval with generative models for open do-	474
417	main question answering . In <i>Proceedings of the 16th</i>	
418	<i>Conference of the European Chapter of the Associ-</i>	
419	<i>ation for Computational Linguistics: Main Volume</i> ,	
420	pages 874–880, Online. Association for Computa-	
421	tional Linguistics.	
422	Albert Q Jiang, Alexandre Sablayrolles, Antoine	
423	Roux, Arthur Mensch, Blanche Savary, Chris Bam-	
424	ford, Devendra Singh Chaplot, Diego de las Casas,	
425	Emma Bou Hanna, Florian Bressand, et al. 2024.	
426	Mixtral of experts . <i>arXiv preprint arXiv:2401.04088</i> .	
427	Jeff Johnson, Matthijs Douze, and Hervé Jégou. 2019.	
428	Billion-scale similarity search with gpus . <i>IEEE</i>	
429	<i>Transactions on Big Data</i> , 7(3):535–547.	
430	Mandar Joshi, Eunsol Choi, Daniel Weld, and Luke	
431	Zettlemoyer. 2017. TriviaQA: A large scale distant	
432	supervised challenge dataset for reading comprehen-	
433	sion . In <i>Proceedings of the 55th Annual Meeting of</i>	
434	<i>the Association for Computational Linguistics (Vol-</i>	
435	<i>ume 1: Long Papers)</i> , pages 1601–1611, Vancouver,	
436	Canada. Association for Computational Linguistics.	
437	Ehsan Kamalloo, Nouha Dziri, Charles Clarke, and	
438	Davood Rafiei. 2023. Evaluating open-domain ques-	
439	tion answering in the era of large language models .	
440	In <i>Proceedings of the 61st Annual Meeting of the</i>	
441	<i>Association for Computational Linguistics (Volume</i>	
442	<i>1: Long Papers)</i> , pages 5591–5606, Toronto, Canada.	
443	Association for Computational Linguistics.	
444	Vladimir Karpukhin, Barlas Oguz, Sewon Min, Patrick	
445	Lewis, Ledell Wu, Sergey Edunov, Danqi Chen, and	
446	Wen-tau Yih. 2020. Dense passage retrieval for open-	
447	domain question answering . In <i>Proceedings of the</i>	
448	<i>2020 Conference on Empirical Methods in Natural</i>	
449	<i>Language Processing (EMNLP)</i> , pages 6769–6781,	
450	Online. Association for Computational Linguistics.	
451	Mike Lewis, Yinhan Liu, Naman Goyal, Marjan	
452	Ghazvininejad, Abdelrahman Mohamed, Omer Levy,	
453	Veselin Stoyanov, and Luke Zettlemoyer. 2020a.	
454	BART: Denoising sequence-to-sequence pre-training	
455	for natural language generation, translation, and com-	
456	prehension . In <i>Proceedings of the 58th Annual Meet-</i>	
457	<i>ing of the Association for Computational Linguistics</i> ,	
458	pages 7871–7880, Online. Association for Computa-	
459	tional Linguistics.	
460	Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio	
461	Petroni, Vladimir Karpukhin, Naman Goyal, Hein-	
462	rich Küttler, Mike Lewis, Wen-tau Yih, Tim Rock-	
463	täschel, et al. 2020b. Retrieval-augmented generation	
464	for knowledge-intensive nlp tasks . <i>Advances in Neu-</i>	
465	<i>ral Information Processing Systems</i> , 33:9459–9474.	
466	Lin Li, Mengjing Zhang, Zhaohui Chao, and Jianwen	
467	Xiang. 2021. Using context information to enhance	
468	simple question answering . <i>World Wide Web</i> , 24:249–	
469	277.	
	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-	
	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,	
	Luke Zettlemoyer, and Veselin Stoyanov. 2019.	
	Roberta: A robustly optimized bert pretraining ap-	
	proach . <i>ArXiv</i> , abs/1907.11692.	
	Denis Lukovnikov, Asja Fischer, and Jens Lehmann.	
	2019. Pretrained transformers for simple question	
	answering over knowledge graphs . In <i>The Semantic</i>	
	<i>Web–ISWC 2019: 18th International Semantic Web</i>	
	<i>Conference, Auckland, New Zealand, October 26–</i>	
	<i>30, 2019, Proceedings, Part I 18</i> , pages 470–486.	
	Springer.	
	Denis Lukovnikov, Asja Fischer, Jens Lehmann, and	
	Sören Auer. 2017. Neural network-based question	
	answering over knowledge graphs on word and char-	
	acter level . In <i>Proceedings of the 26th international</i>	
	<i>conference on World Wide Web</i> , pages 1211–1220.	
	Salman Mohammed, Peng Shi, and Jimmy Lin. 2018.	
	Strong baselines for simple question answering over	
	knowledge graphs with and without neural networks .	
	In <i>Proceedings of the 2018 Conference of the North</i>	
	<i>American Chapter of the Association for Computa-</i>	
	<i>tional Linguistics: Human Language Technologies,</i>	
	<i>Volume 2 (Short Papers)</i> , pages 291–296, New Or-	
	leans, Louisiana. Association for Computational Lin-	
	guistics.	
	OpenAI. 2023. Gpt-4 technical report .	
	Nagehan Pala Er and Ilyas Cicekli. 2013. A factoid	
	question answering system using answer pattern	
	matching . In <i>Proceedings of the Sixth International</i>	
	<i>Joint Conference on Natural Language Processing</i> ,	
	pages 854–858, Nagoya, Japan. Asian Federation of	
	Natural Language Processing.	
	Baolin Peng, Michel Galley, Pengcheng He, Hao Cheng,	
	Yujia Xie, Yu Hu, Qiuyuan Huang, Lars Liden, Zhou	
	Yu, Weizhu Chen, et al. 2023. Check your facts and	
	try again: Improving large language models with	
	external knowledge and automated feedback . <i>arXiv</i>	
	<i>preprint arXiv:2302.12813</i> .	
	Michael Petrochuk and Luke Zettlemoyer. 2018. Sim-	
	pleQuestions nearly solved: A new upperbound and	
	baseline approach . In <i>Proceedings of the 2018 Con-</i>	
	<i>ference on Empirical Methods in Natural Language</i>	
	<i>Processing</i> , pages 554–558, Brussels, Belgium. As-	
	sociation for Computational Linguistics.	
	Colin Raffel, Noam Shazeer, Adam Roberts, Katherine	
	Lee, Sharan Narang, Michael Matena, Yanqi Zhou,	
	Wei Li, and Peter J Liu. 2020. Exploring the limits	
	of transfer learning with a unified text-to-text trans-	
	former . <i>The Journal of Machine Learning Research</i> ,	
	21(1):5485–5551.	
	Pranav Rajpurkar, Jian Zhang, Konstantin Lopyrev, and	
	Percy Liang. 2016. SQuAD: 100,000+ questions for	
	machine comprehension of text . In <i>Proceedings of</i>	
	<i>the 2016 Conference on Empirical Methods in Natu-</i>	
	<i>ral Language Processing</i> , pages 2383–2392, Austin,	
	Texas. Association for Computational Linguistics.	

527	Ori Ram, Yoav Levine, Itay Dalmedigos, Dor Muhlgay, Amnon Shashua, Kevin Leyton-Brown, and Yoav Shoham. 2023. In-context retrieval-augmented language models . <i>Transactions of the Association for Computational Linguistics</i> , 11:1316–1331.	585
528		586
529		587
530		588
531		589
532	Prakash Ranjan and Rakesh Chandra Balabantaray. 2016. Question answering system for factoid based question . In <i>2016 2nd International Conference on Contemporary Computing and Informatics (IC3I)</i> , pages 221–224. IEEE.	590
533		591
534		592
535		593
536		594
537	Adam Roberts, Colin Raffel, and Noam Shazeer. 2020. How much knowledge can you pack into the parameters of a language model? In <i>Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)</i> , pages 5418–5426, Online. Association for Computational Linguistics.	595
538		596
539		597
540		598
541		599
542		600
543	Stephen E. Robertson, Steve Walker, Susan Jones, Micheline Hancock-Beaulieu, and Mike Gatford. 1994. Okapi at trec-3 . In <i>Text Retrieval Conference</i> .	601
544		602
545		603
546	Mikhail Salnikov, Hai Le, Prateek Rajput, Irina Nikishina, Pavel Braslavski, Valentin Malykh, and Alexander Panchenko. 2023. Large language models meet knowledge graphs to answer factoid questions . In <i>Proceedings of the 37th Pacific Asia Conference on Language, Information and Computation</i> , pages 635–644, Hong Kong, China. Association for Computational Linguistics.	604
547		605
548		606
549		607
550		608
551		609
552		
553		
554	Iulian Vlad Serban, Alberto García-Durán, Caglar Gulcehre, Sungjin Ahn, Sarath Chandar, Aaron Courville, and Yoshua Bengio. 2016. Generating factoid questions with recurrent neural networks: The 30M factoid question-answer corpus . In <i>Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 588–598, Berlin, Germany. Association for Computational Linguistics.	610
555		611
556		612
557		613
558		614
559		615
560		616
561		617
562		
563	Hassan S. Shavarani and Anoop Sarkar. 2023. SpEL: Structured prediction for entity linking . In <i>Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing</i> , pages 11123–11137, Singapore. Association for Computational Linguistics.	618
564		619
565		620
566		621
567		
568		
569	Hassan S. Shavarani and Satoshi Sekine. 2020. Multi-class multilingual classification of Wikipedia articles using extended named entity tag set . In <i>Proceedings of the Twelfth Language Resources and Evaluation Conference</i> , pages 1197–1201, Marseille, France. European Language Resources Association.	
570		
571		
572		
573		
574		
575	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 . <i>arXiv preprint arXiv:2301.12652</i> .	
576		
577		
578		
579		
580	Tanzim Tamanna Shitu, Nazia Zaman, and KM Azharul Hasan. 2020. Domain specific factoid question answering by regular expression generation . In <i>2020 IEEE Region 10 Symposium (TENSYP)</i> , pages 1464–1467. IEEE.	
581		
582		
583		
584		
	Devendra Singh, Siva Reddy, Will Hamilton, Chris Dyer, and Dani Yogatama. 2021. End-to-end training of multi-document reader and retriever for open-domain question answering . <i>Advances in Neural Information Processing Systems</i> , 34:25968–25981.	
	Noah A Smith, Michael Heilman, and Rebecca Hwa. 2008. Question generation as a competitive undergraduate course project . In <i>Proceedings of the NSF Workshop on the Question Generation Shared Task and Evaluation Challenge</i> , volume 9.	
	Alon Talmor, Jonathan Herzig, Nicholas Lourie, and Jonathan Berant. 2019. CommonsenseQA: A question answering challenge targeting commonsense knowledge . In <i>Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)</i> , pages 4149–4158, Minneapolis, Minnesota. Association for Computational Linguistics.	
	Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. 2023. Llama 2: Open foundation and fine-tuned chat models . <i>arXiv preprint arXiv:2307.09288</i> .	
	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 . In <i>Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing</i> , pages 2369–2380, Brussels, Belgium. Association for Computational Linguistics.	
	Wenhao Yu, Zhihan Zhang, Zhenwen Liang, Meng Jiang, and Ashish Sabharwal. 2023. Improving language models via plug-and-play retrieval feedback . <i>arXiv preprint arXiv:2305.14002</i> .	