



GeAR: Generation Augmented Retrieval

Anonymous ACL submission

Abstract

Document retrieval techniques form the foundation for the development of large-scale information systems. The prevailing methodology is to construct a bi-encoder and compute the semantic similarity. However, such scalar similarity is difficult to reflect enough information and impedes our comprehension of the retrieval results. In addition, this computational process mainly emphasizes the global semantics and ignores the fine-grained semantic relationship between the query and the complex text in the document. In this paper, we propose a new method called **Generation Augmented Retrieval (GeAR)** that incorporates well-designed fusion and decoding modules. This enables GeAR to generate the relevant text from documents based on the fused representation of the query and the document, thus learning to "focus on" the fine-grained information. Also when used as a retriever, GeAR does not add any computational burden over bi-encoders. To support the training of the new framework, we have introduced a pipeline to efficiently synthesize high-quality data by utilizing large language models. GeAR exhibits competitive retrieval and localization performance across diverse scenarios and datasets. Moreover, the qualitative analysis and the results generated by GeAR provide novel insights into the interpretation of retrieval results.

1 Introduction

Document retrieval serve as the foundational technology behind large-scale information systems, playing a crucial role in applications such as web search, open-domain question answering (QA) (Chen et al., 2017; Karpukhin et al., 2020), and retrieval-augmented generation (RAG) (Lewis et al., 2020; Liu et al., 2024a; Gao et al., 2024). The predominant approach in passage retrieval is to construct a bi-encoder model. In this architecture, queries and documents are encoded separately, transforming each into vector representations that

enable computation of their semantic similarity in a high-dimensional space.

However, this similarity calculation process faces several challenges. First, the complex semantic relationship between query and document is mapped to a scalar similarity, which cannot reflect enough information and is difficult to understand (Brito and Iser, 2023). Second, when dealing with long documents, such as those with 256, 512, or even more tokens, identifying the section most relevant to the query and contributing most to the similarity is highly desirable but challenging to achieve (Luo et al., 2024; Günther et al., 2024). Moreover, many NLP tasks, such as sentence selection, search result highlighting, needle in a haystack (Liu et al., 2024b; An et al., 2024; Wang et al., 2024), and fine-grained citations (Gao et al., 2023; Zhang et al., 2024), require a deep and fine-grained understanding of the text.

Given this need for fine-grained understanding, the bi-encoder that simply aligns the entire document to the query seems insufficient, as its conventional contrastive loss mainly emphasizes global semantics (Khattab and Zaharia, 2020). To complement this core localization capability of the retriever, we propose a novel and challenging fundamental question: Can we enhance and integrate the information localization capability of existing retrievers without sacrificing their inherent retrieval capabilities?

To address these challenges, we proposed a novel approach **GeAR (Generation-Augmented Retrieval)**. Specifically, we construct the data into (query-document-information) triples, still using contrastive learning to optimize the similarity between the query and the document. At the same time, we design a text decoder to generate the relevant fine-grained information in the document given the query and document to enhance the retrieval and localization capabilities. Although the concept is simple, there are many challenges. First,

it is difficult to find sufficient data to support our solution to this problem in previous research work. Second, the training objectives of retrieval and generation tasks, model architectures, and more design details, as well as how to effectively train the models, have not been fully explored. To this end, we explored a complete pipeline from data synthesis, structure design, to model training. Overall, our contributions are summarized as follows:

- We proposed GeAR, which enhances the model’s ability to understand and locate text in a fine-grained manner by jointly modeling natural language understanding and natural language generation. At the same time, the inference process is very flexible to handle different tasks.
- We abstract a new retrieval task that takes into account the problems present in the current retrieval scenario. To solve this task and to support model training, we built a pipeline to synthesize a large amount of high quality data using LLM.
- Through extensive experiments, GeAR has shown competitive performance in retrieval tasks and fine-grained information localization tasks. At the same time, GeAR can also generate relevant information based on the query and document to help us understand the retrieval results, bringing a new perspective to the traditional retrieval process.

2 Related Work

2.1 Embedding-based Retrieval

Embedding-based retrieval has emerged as a cornerstone of modern information retrieval systems, enabling efficient semantic search through dense vector representations. Early approaches like Word2Vec (Mikolov et al., 2013) and GloVe (Pennington et al., 2014) demonstrated the potential of learning distributed word representations, while more recent transformer-based models such as BERT (Devlin et al., 2019) have pushed the boundaries of contextual embeddings. Bi-encoder architectures (Reimers and Gurevych, 2019) have become particularly popular for retrieval tasks (Huang et al., 2013). Recent advances include contrastive learning objectives (Karpukhin et al., 2020; Wang et al., 2022; Li et al., 2023;

Gao et al., 2021), hard negative mining strategies (Xiong et al., 2021), and knowledge distillation techniques (Hofstätter et al., 2021) to improve embedding quality while maintaining computational efficiency. Muennighoff et al. (2024) explored how to generate text and provide excellent semantic representation by distinguishing task instructions.

Multimodal information retrieval also relies on high-quality semantic representations, where the embedding space serves to bridge different modalities, including text, images, and video. Vision language models such as CLIP (Radford et al., 2021), ALBEF (Li et al., 2021), and BLIP (Li et al., 2022) have demonstrated remarkable zero-shot capabilities by learning joint embeddings derived from large scale image-text pairs. These advances made cross-modal retrieval tasks such as text-to-image search and image-to-text retrieval possible.

2.2 Information Localization

Information localization in massive corpora and contents has become a key research direction for improving response accuracy and factual basis. The classic methods used RNN or BERT to compute token representations and trained a classifier for information extraction (Seo, 2016; Wang, 2016; Chen et al., 2017; Xu et al., 2019). The heuristic hierarchical approach involves further chunking the document and then calculating the semantic similarity with the query on the chunked sentences or units for localization. However, finer chunking also results in increased computational complexity and semantic incoherence (Yang et al., 2016; Liu et al., 2021; Arivazhagan et al., 2023). With the development of generative models, there have been many recent efforts to enhance the model’s ability to find a needle in a haystack (Liu et al., 2024b; An et al., 2024; Wang et al., 2024), that is, to locate key information such as sentences in long texts. Another type of similar task is to have the model add reference information to the original text when generating responses (Gao et al., 2023; Zhang et al., 2024). Coincidentally, there have been some recent works focusing on improving the regional level understanding ability of multimodal large language models (MLLMs) (Chen et al., 2024). Despite these advances, we have found that there is currently little focus on fine-grained information localization during the retrieval stage.

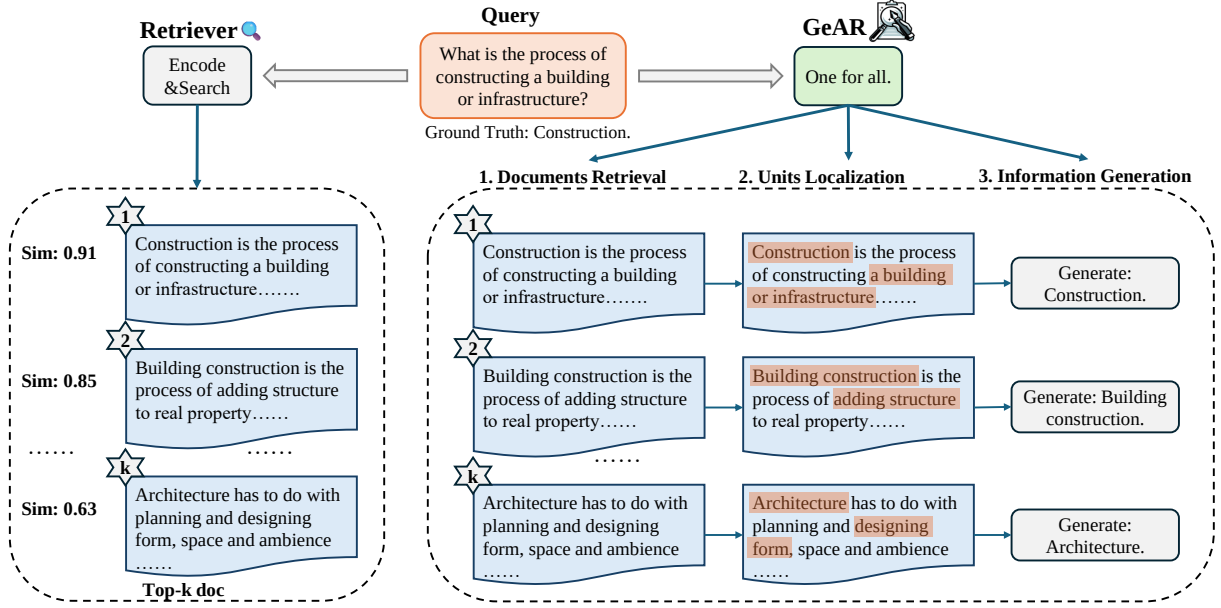


Figure 1: Comparison of functionality between classical retriever and GeAR. GeAR is designed to handle both document retrieval and fine-grained unit localization simultaneously, while also generating auxiliary information for reference.

3 Generation Augmented Retrieval

3.1 Preliminaries

In this work, we formalize the retrieval task with localization as follows: Let a document corpus as $\mathbb{D}$, which contains N documents $\{d_1, \dots, d_i, \dots, d_N\}$. Each of these documents d_i contains a number of fine-grained information units $\{u_1, \dots, u_{l_i}\}$, such as sentences, where l_i is the units number of d_i . Our goal is to find a retrieval method $f(\cdot)$, which can retrieve the relevant document d from $\mathbb{D}$, as well as the fine-grained unit u from d given query q :

$$f(q, \mathbb{D}) \rightarrow \{d\} \quad (1)$$

$$f(q, d) \rightarrow \{u\} \quad (2)$$

In this work, we explicitly define the process as two tasks, (1) the document retrieval task and (2) the fine-grained unit localization task, as Figure 1 showing. It can be seen that the triples of query, document, and unit, represented by the symbols (q, d, u) , are fundamental to the definition and resolution of this task.

3.2 Data construction

In this work, we focus on two common retrieval scenarios: (1) Question Answer Retrieval (QAR) and (2) Relevant Information Retrieval (RIR). In the following sections, we will introduce how the

data are constructed and how they correspond to the triples (q, d, u) mentioned above.

Question Answer Retrieval In this scenario, the query q is in the form of a question, and the goal is to retrieve reference documents d that support answering the question and fine-grained sentences u that contain the answer.

Relevant Information Retrieval In this scenario, the query q is in the form of a few phrases or keywords, the objective is to retrieve the documents d that correspond to the query and the fine-grained sentences u in the documents that are most relevant to the query. The scenario is very close to what users normally do when they search for information on the web. The challenge is that we have difficulty in finding suitable data in the current public dataset to drive our problem solving. Therefore, we constructed a pipeline to synthesize high quality data using a large language model. Specifically, we selected high quality Wikipedia documents (Foundation), from which we will sample sentences of appropriate length and whose subject is not a pronoun as u . Then we will leverage LLM to rewrite them as queries q . After de-duplication and relevance filtering, we get promising **5.8M** triples. Kindly refer to Appendix A for details on complete data processing procedure.

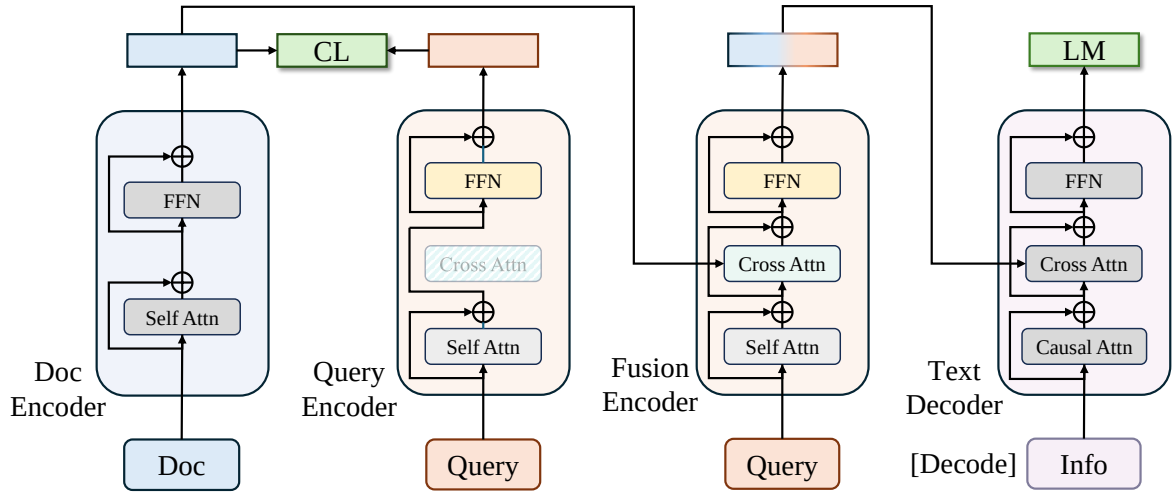


Figure 2: **GeAR**. It consists of a bi-encoder, a fusion encoder, and a text decoder. It contains two training objectives, CL represents contrastive learning loss, which aims to optimize the similarity between documents and queries. LM represents the language modeling loss for generating relevant information given documents and queries.

3.3 Model Structure

This section presents the architecture of GeAR. It is our intention that the model not only has powerful retrieval capability, but also has the ability to locate key information in documents. Inspired by advances in multimodal representation learning (Li et al., 2021, 2022; He et al., 2020), we revisit the task from a modal alignment perspective. Documents and queries can be considered as two modalities. We facilitate semantic alignment between documents and queries via a bi-encoder, and enable the model to learn to attend to fine-grained query-related information in the document via a fusion encoder and a generation task. The overview of the GeAR structure is illustrated in Figure 2.

Bi-Encoder In the same setup as the classical retrieval approach, we initialize two encoders $E_d(\cdot)$ for documents and $E_q(\cdot)$ for queries. We use mean pooling to obtain the text embedding.

Fusion Encoder The fusion encoder share most of the parameters with query encoder, but have an extra learnable cross attention module. In this part, the document embeddings from $E_d(\cdot)$ are fused with the query embeddings through cross attention at each layer of the fusion encoder.

Text Decoder The text decoder receives the fusion embeddings and generates fine-grained information¹ in the document based on the given

query and document. It uses a unidirectional causal attention instead of a bidirectional self-attention. A specific [Decode] token is added to identify the beginning of the sequence. The subsequent autoregressive decoding process will interact with the generated tokens and fusion embeddings to generate text.

3.4 Training Objectives

In this section, we present the training objectives of GeAR. We make the model capable of both retrieval as well as fine-grained semantic understanding and localization through joint natural language understanding and natural language generation modeling.

Contrastive Learning Loss (CL) We use bi-encoder to encode the queries and documents, and optimize the semantic similarity between them through contrastive learning loss (CL). In addition, we followed the practice in MoCo (He et al., 2020) and BLIP (Li et al., 2022), where a momentum Bi-Encoder is introduced to encode momentum embeddings and provide richer supervised signals as soft labels.

Language Modeling Loss (LM) The introduction of LM loss is key to enhancing the information localization capability of GeAR. LM activates the text decoder, which enables the model to generate relevant information using the fusion embeddings of document and query. It guides the model to learn the fine-grained semantic fusion of query and document. LM optimizes the cross en-

¹Note that in the generation task of the QAR scenario, the ground truth is the answer itself, not the sentence u . But in the RIR scenario and the localization task, we all used the sentence u .

trophy loss over the entire vocabulary, maximizing the likelihood of the ground truth text. The overall loss of GeAR is the sum of $\mathcal{L}_{CL}$ and $\mathcal{L}_{LM}$:

$$\mathcal{L}_{GeAR} = \mathcal{L}_{CL} + \mathcal{L}_{LM} \quad (3)$$

3.5 Inference

GeAR’s inference process is diverse and flexible. In this section, we introduce various usages of GeAR to accomplish different tasks.

Documents Retrieval For this task, we can use the bi-encoder part of GeAR to compute the similarity between query and document like the previous classic retrieval method, without introducing any additional parameters and computation cost.

Fine-Grained Units Localization The fusion encoder in GeAR calculates the fusion embedding of query and document through cross attention. We use the cross attention weights of the query on the tokens in the document to locate the units that the query pays the most attention to in the document.

Information Generation For this task, we feed the fusion embedding to the text decoder and enable autoregressive decoding. In GeAR, information generation is actually an auxiliary task, and we will present the generative performance of the model in experiments, both in terms of quantitative metrics and qualitative analysis.

4 Experiments

In this section, we first introduce the experimental setup, and then we show the overall performance of each task and more detailed analysis experiments.

4.1 Setup

Datasets For Question Answer Retrieval, we sampled 30M data from PAQ (Lewis et al., 2021) datasets to train GeAR, and sampled 1M documents and 20k queries as test set. We also evaluate the performance on another 3 QA datasets: SQuAD (Rajpurkar et al., 2016), NQ (Kwiatkowski et al., 2019), and TriviaQA (Joshi et al., 2017). These test datasets are all held out to observe the generalizability of compared methods. For Relevant Information Retrieval, we leverage the synthesized 5.8M data, of which 95% is used for training and 5% is reserved for the test set. Specific dataset statistics are in Appendix B.

Training Details To better observe the effectiveness of GeAR, we use "BERT-base-uncased" (Devlin et al., 2019) to initialize the encoders in GeAR. We trained the model for 10

epochs using a batch size of 48 (QAR) / 16 (RIR) on 16 AMD MI200 GPUs with 64GB memory. We use the AdamW (Loshchilov, 2017) optimizer with a weight decay of 0.05. The full hyperparameters and training settings are detailed in Appendix C.

Baselines We compare our approach to two classes of baseline methods, one class of text representation models that have been adequately trained on a large corpus, including SBERT (Reimers and Gurevych, 2019), specifically "all-mpnet-base" (Song et al., 2020), E5 (Wang et al., 2022), BGE (Xiao et al., 2024), and GTE (Li et al., 2023). We use both base-level models for this comparison. The other category consists different training pipelines that unify the training data and starting points, including SBERT (Reimers and Gurevych, 2019) and BGE (Xiao et al., 2024). We retrained them all using the "bert-base-uncased" to initialize and aligned the training data, referred to as SBERT_{RT} and BGE_{RT} in the following.

4.2 Overall performance

In this section, we present the overall performance on Documents Retrieval, Units Localization, and Information Generation.

Documents Retrieval Firstly, we report the comparison with existing methods on documents retrieval task in Table 1. The results demonstrate that GeAR delivers competitive performance across multiple datasets, even with limited training data. As a reference, the pre-trained SBERT model used 1.17B sentence pairs, with partial overlap between its training data and our evaluation data. To ensure a fair comparison, we retrained SBERT² and BGE³ using their open source training pipelines, aligned training data and initialization settings. As shown in the retrained model section in Table 1, GeAR achieves superior performance, underscoring the effectiveness of our training approach.

Units Localization Next, we evaluate the performance of each method on the units localization task. In the evaluation process, we provide the query and the document (q, d) to the model and observe whether it is able to find the corresponding fine-grained unit u . For the retrieval model, we split the documents into sentences and compute their similarity to the query independently, selecting the top-k sentences. In contrast, GeAR locates units based on the cross attention weights for each

²<https://huggingface.co/sentence-transformers/all-mpnet-base-v2>.

³<https://github.com/FlagOpen/FlagEmbedding>.

	SQuAD		NQ		TriviaQA		PAQ		RIR	
	R@5	M@5	R@5	M@5	R@5	M@5	R@5	M@5	R@5	M@5
<i>Pre-trained retrieval model</i>										
SBERT	0.812	0.667	0.754	0.576	0.677	0.413	0.808	0.701	0.376	0.297
E5	0.803	0.674	0.760	0.581	0.645	0.390	0.816	0.716	0.484	0.396
BGE	0.829	0.701	0.674	0.502	0.690	0.422	0.752	0.647	0.451	0.367
GTE	0.866	0.744	0.767	0.587	0.726	0.443	0.836	0.736	0.528	0.435
<i>Retrained retrieval model</i>										
SBERT _{RT}	0.742	0.585	0.739	0.550	0.577	0.342	0.859	0.742	0.739	0.631
BGE _{RT}	0.841	0.701	0.751	0.553	0.640	0.384	0.901	0.802	0.953	0.881
GeAR	0.883	0.762	0.747	0.567	0.660	0.398	0.940	0.855	0.961	0.903
GeAR _{w/o$\mathcal{L}_{LM}$}	0.889	0.776	0.755	0.565	0.660	0.399	0.955	0.877	0.963	0.907

Table 1: Comparison of documents retrieval performance on different datasets, where R@k stands for Recall@k, M@k stands for MAP@k.

	SQuAD		NQ		TriviaQA		PAQ		RIR	
	R@1	M@1	R@1	M@1	R@1	M@1	R@1	M@1	R@3	M@3
<i>Pre-trained retrieval model</i>										
SBERT	0.739	0.800	0.558	0.652	0.359	0.583	0.498	0.561	0.891	0.874
E5	0.783	0.847	0.590	0.683	0.379	0.613	0.573	0.640	0.891	0.878
BGE	0.768	0.830	0.570	0.663	0.362	0.589	0.565	0.630	0.894	0.881
GTE	0.758	0.820	0.548	0.639	0.352	0.572	0.525	0.590	0.895	0.886
<i>Retrained retrieval model</i>										
SBERT _{RT}	0.516	0.568	0.445	0.523	0.281	0.472	0.363	0.418	0.899	0.991
BGE _{RT}	0.455	0.538	0.601	0.656	0.288	0.475	0.409	0.466	0.897	0.888
GeAR	0.810	0.874	0.765	0.871	0.515	0.808	0.885	0.965	0.954	0.897

Table 2: Comparison of units localization performance on different datasets, where R@k stands for Recall@k, M@k stands for MAP@k.

sentence given the document and the query, as described in Section 3.5. The results are reported in Table 2. We found that GeAR came out ahead on all metrics, and that GeAR did not require further chunking and encoding of the document. It is observed that SBERT_{RT} and BGE_{RT} exhibit suboptimal performance, as their training objective focus solely on optimizing the overall semantic similarity between the document and the query, neglecting the fine-grained semantic relationships. In contrast, GeAR benefits from the joint end-to-end training of retrieval and generation tasks, enabling it not only to retrieve documents closely aligned with the query but also to effectively attend to fine-grained information within the document.

Information Generation Although genera-

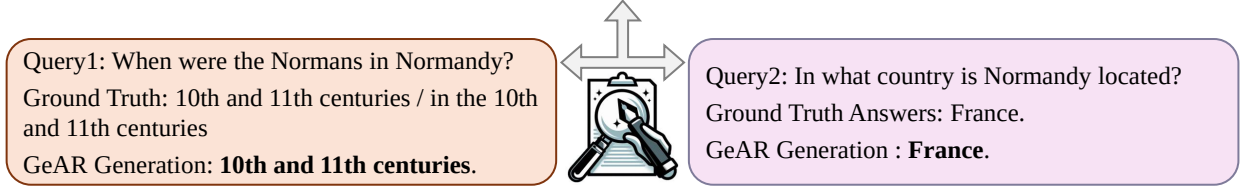
tion serves only as an auxiliary task in GeAR, we are nonetheless interested in evaluating its generation performance. Table 3 reports the Exact Match (EM) and F1 scores on the QA datasets, and the Rouge (Lin, 2004) scores on the RIR dataset. Notably, GeAR achieves strong performance on test sets with distributions similar to the training data, such as PAQ and RIR, and performs reasonably well on other test sets. Additionally, Figure 3 illustrates examples of GeAR’s ability to generate correct answers and relevant information, demonstrating its satisfactory generation capabilities.

4.3 Analysis

Visualization of Information Localization Figure 3 illustrates the information localization and

Document

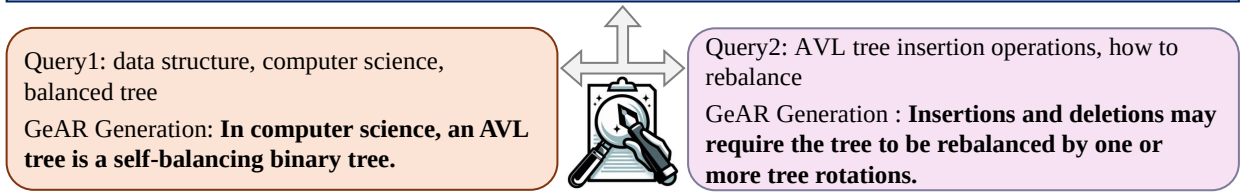
The Normans (Norman: Nourmands; French: Normands; Latin: Normanni) were the people who in the 10th and 11th centuries gave their name to Normandy, a region in France. They were descended from Norse ("Norman" comes from "Norseman") raiders and pirates from Denmark, Iceland and Norway who, under their leader Rollo, agreed to swear fealty to King Charles III of West Francia. Through generations of assimilation and mixing with the native Frankish and Roman-Gaulish populations, their descendants would gradually merge with the Carolingian-based cultures of West Francia. The distinct cultural and ethnic identity of the Normans emerged initially in the first half of the 10th century, and it continued to evolve over the succeeding centuries.



(a) Information localization and generation results of GeAR in Question Answer Retrieval scenario.

Document

[In computer science, an AVL tree is a self-balancing binary search tree.] In an AVL tree, the heights of the two child subtrees of any node differ by at most one; if at any time they differ by more than one, rebalancing is done to restore this property. [Insertions and deletions may require the tree to be rebalanced by one or more tree rotations.] The AVL tree is named after its two Soviet inventors, Georgy Adelson-Velsky and Evgenii Landis, who published it in their 1962 paper "An algorithm for the organization of information".



(b) Information localization and generation results of GeAR in Related Information Retrieval scenario. The sentences in brackets of corresponding colors are the ground truth of the query.

Figure 3: Visualization of information localization of GeAR . In the two scenarios of Question Answer retrieval and Related Information Retrieval, we propose two different queries for one document and highlight the top 10 tokens with the highest cross attention weights for the corresponding queries. The tokens with orange background are for query1 , and the tokens with purple background are for query2 . We also show the generated results of GeAR.

generation results of GeAR across different scenarios. We provide two distinct queries for one document and highlight the top 10 tokens with the highest cross attention weights corresponding to each queries. In Figure 3(a), the two queries focus on time and location, respectively. GeAR not only gave the correct answers to both queries but also dynamically adjusts its query-specific focus: it assigns higher attention weights to time-related tokens in response to the first query and prioritizes tokens related to countries and regions in response to the second query. In Figure 3(b), GeAR will focus on the definition of the AVL tree itself, and the insertion, deletion, rotation and rebalancing of the AVL tree, and generate corresponding sentences. It can be seen that the added generation task has brought improvements to the model in terms of per-

formance and qualitative effects, making it accurate in localization and generation.

Correlation of Generation and Localization

In this section, we analyze the relationship between the generation and localization tasks. As illustrated in Figure 4(a) and 4(b), we plot the performance coordinates from epoch 1 to epoch 10 during training, where the horizontal axis represents the generation performance and the vertical axis represents the localization performance. The results reveal a strong correlation between the two tasks. This observation demonstrates that the generation task, designed as a proxy, effectively enhances the model's ability to extract fine-grained information from documents. These findings highlight the synergistic relationship between generation and localization.

Localization performance of different layers

SQuAD		NQ		TriviaQA		PAQ		RIR	
EM	F1	EM	F1	EM	F1	EM	F1	Rouge-1	Rouge-L
46.6	65.2	66.1	61.0	47.4	60.0	88.1	92.4	87.4	87.1

Table 3: Generation performance of GeAR on different tasks.

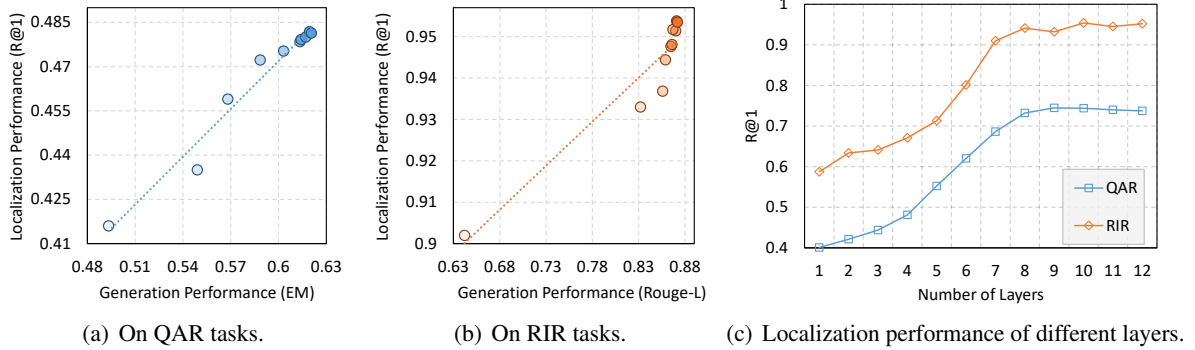


Figure 4: Plots of generation and localization performance on (a) QAR tasks and (b) RIR tasks as training progresses. (c) shown the localization performance at different layers.

In GeAR, the fusion encoder and decoder interact through the cross attention module at each layer. To investigate the relationship between localization performance and model depth, we plot the localization performance using cross attention weights across different layers in Figure 4(c). The results indicate that high-level token embeddings perform well, as they capture rich semantic information through deeper layers of the network. Interestingly, we observe that the highest layer does not yield the best localization performance. Instead, peak performance is achieved in the last 3 to 4 layers⁴. This phenomenon may arise because the representations in the highest layer are optimized to serve the final task rather than intermediate localization. Similar findings have been reported in prior studies involving encoder-only and decoder-only models (Jawahar et al., 2019; Skean et al., 2024).

The Affect of Language Modeling Objectives

In this work, we utilize only the information corresponding to the query as supervision and incorporate a language modeling objective. It enables the model to achieve stronger capabilities in both information localization and generation, without requiring additional loss functions or complex module designs. However, as a trade-off, we observe a slight decrease in retrieval performance when compared to using only the contrastive learning

objective for the retrieval task, as shown in the last two rows of Table 1. How to further design the balance between the two training objectives from the perspective of multi-task learning so that they benefit from each other is a point that can still be explored in the future.

5 Conclusion

In this work, to address the challenges of unexplainable and coarse-grained results inherent in current bi-encoder retrieval methods, we propose a direct and effective modeling method: **Generation Augmented Retrieval (GeAR)**. GeAR enhances fine-grained information localization and generation capabilities by incorporating a decoder and a lightweight cross-attention layer, while maintaining the efficiency of a bi-encoder. Experimental results across multiple retrieval tasks and two different scenarios demonstrate that GeAR achieves competitive performance. Furthermore, its ability to accurately and reasonably localize information makes it particularly promising for downstream tasks such as web search, semantic understanding, and retrieval-augmented generation (RAG). We hope this work offers valuable insights into the gradual unification of natural language understanding and generation paradigms, paving the way for more versatile and explainable retrieval systems in the future.

⁴In this work, we utilized the 10th layer for evaluation.

Limitations

Due to constraints in computational resources and associated costs, the synthesized data used in our experiments is not as comprehensive as that found in traditional retrieval scenarios. While the results demonstrate the efficacy of GeAR, applying it to more diverse and semantically rich retrieval scenarios remains an important direction for future exploration.

Additionally, the context length of limited to 512 tokens, consistent with the chunk lengths commonly used in retrieval tasks. However, recent advancements in extending the context length of retrieval models, such as those proposed in (Zhu et al., 2024), suggest exciting opportunities to overcome this limitation. Extending GeAR’s context length could further enhance its capabilities in handling long-form retrieval tasks, which we plan to investigate in future work.

We hope that the above discussions can inspire further investigation within the research community, encouraging advancements that address these limitations and contribute to the broader progress of NLP research.

References

- Shengnan An, Zexiong Ma, Zeqi Lin, Nanning Zheng, and Jian-Guang Lou. 2024. [Make your llm fully utilize the context](#). In *NeurIPS 2024*.
- Manoj Ghuhana Arivazhagan, Lan Liu, Peng Qi, Xinchu Chen, William Yang Wang, and Zhiheng Huang. 2023. [Hybrid hierarchical retrieval for open-domain question answering](#). In *Findings of the Association for Computational Linguistics: ACL 2023*, pages 10680–10689, Toronto, Canada. Association for Computational Linguistics.
- Eduardo Brito and Henri Iser. 2023. [Maxsime: Explaining transformer-based semantic similarity via contextualized best matching token pairs](#). In *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR ’23*, page 2154–2158, New York, NY, USA. Association for Computing Machinery.
- 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.
- Hong-You Chen, Zhengfeng Lai, Haotian Zhang, Xinze Wang, Marcin Eichner, Keen You, Meng Cao, Bowen

- Zhang, Yinfei Yang, and Zhe Gan. 2024. [Contrastive localized language-image pre-training](#). *arXiv preprint arXiv:2410.02746*.
- 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](#).
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. 2024. [The llama 3 herd of models](#). *arXiv preprint arXiv:2407.21783*.
- Wikimedia Foundation. [Wikimedia downloads](#).
- Tianyu Gao, Xingcheng Yao, and Danqi Chen. 2021. [SimCSE: Simple contrastive learning of sentence embeddings](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 6894–6910, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Tianyu Gao, Howard Yen, Jiatong Yu, and Danqi Chen. 2023. [Enabling large language models to generate text with citations](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 6465–6488, Singapore. Association for Computational Linguistics.
- Yunfan Gao, Yun Xiong, Xinyu Gao, Kangxiang Jia, Jinliu Pan, Yuxi Bi, Yi Dai, Jiawei Sun, Meng Wang, and Haofen Wang. 2024. [Retrieval-augmented generation for large language models: A survey](#).
- Michael Günther, Isabelle Mohr, Daniel James Williams, Bo Wang, and Han Xiao. 2024. [Late chunking: contextual chunk embeddings using long-context embedding models](#). *arXiv preprint arXiv:2409.04701*.
- Kaiming He, Haoqi Fan, Yuxin Wu, Saining Xie, and Ross Girshick. 2020. [Momentum contrast for unsupervised visual representation learning](#). In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*.
- Sebastian Hofstätter, Sheng-Chieh Lin, Jheng-Hong Yang, Jimmy Lin, and Allan Hanbury. 2021. [Efficiently teaching an effective dense retriever with balanced topic aware sampling](#). In *Proceedings of the 44th International ACM SIGIR Conference on Research and Development in Information Retrieval*, pages 113–122.

- Po-Sen Huang, Xiaodong He, Jianfeng Gao, Li Deng, Alex Acero, and Larry Heck. 2013. Learning deep structured semantic models for web search using clickthrough data. In *Proceedings of the 22nd ACM international conference on Information & Knowledge Management*, pages 2333–2338.
- Ganesh Jawahar, Benoît Sagot, and Djamé Seddah. 2019. [What does BERT learn about the structure of language?](#) In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 3651–3657, Florence, Italy. Association for Computational Linguistics.
- Jeff Johnson, Matthijs Douze, and Hervé Jégou. 2019. Billion-scale similarity search with GPUs. *IEEE Transactions on Big Data*, 7(3):535–547.
- Mandar Joshi, Eunsol Choi, Daniel Weld, and Luke Zettlemoyer. 2017. [TriviaQA: A large scale distantly supervised challenge dataset for reading comprehension](#). In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1601–1611, Vancouver, Canada. Association for Computational Linguistics.
- Vladimir Karpukhin, Barlas Oguz, Sewon Min, Patrick Lewis, Ledell Wu, Sergey Edunov, Danqi Chen, and Wen-tau Yih. 2020. [Dense passage retrieval for open-domain question answering](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 6769–6781, Online. Association for Computational Linguistics.
- Omar Khattab and Matei Zaharia. 2020. Colbert: Efficient and effective passage search via contextualized late interaction over bert. In *Proceedings of the 43rd International ACM SIGIR conference on research and development in Information Retrieval*, pages 39–48.
- 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*, 7:452–466.
- Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph E. Gonzalez, Hao Zhang, and Ion Stoica. 2023. Efficient memory management for large language model serving with pagedattention. In *Proceedings of the ACM SIGOPS 29th Symposium on Operating Systems Principles*.
- Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. 2020. [Retrieval-augmented generation for knowledge-intensive nlp tasks](#). In *Advances in Neural Information Processing Systems*, volume 33, pages 9459–9474. Curran Associates, Inc.
- Patrick Lewis, Yuxiang Wu, Linqing Liu, Pasquale Minervini, Heinrich Küttler, Aleksandra Piktus, Pontus Stenetorp, and Sebastian Riedel. 2021. [PAQ: 65 million probably-asked questions and what you can do with them](#). *Transactions of the Association for Computational Linguistics*, 9:1098–1115.
- Junnan Li, Dongxu Li, Caiming Xiong, and Steven Hoi. 2022. Blip: Bootstrapping language-image pre-training for unified vision-language understanding and generation. In *International conference on machine learning*, pages 12888–12900. PMLR.
- Junnan Li, Ramprasaath Selvaraju, Akhilesh Gotmare, Shafiq Joty, Caiming Xiong, and Steven Chu Hong Hoi. 2021. Align before fuse: Vision and language representation learning with momentum distillation. *Advances in neural information processing systems*, 34:9694–9705.
- Zehan Li, Xin Zhang, Yanzhao Zhang, Dingkun Long, Pengjun Xie, and Meishan Zhang. 2023. Towards general text embeddings with multi-stage contrastive learning. *arXiv preprint arXiv:2308.03281*.
- Chin-Yew Lin. 2004. [ROUGE: A package for automatic evaluation of summaries](#). In *Text Summarization Branches Out*, pages 74–81, Barcelona, Spain. Association for Computational Linguistics.
- Haoyu Liu, Jianfeng Liu, Shaohan Huang, Yuefeng Zhan, Hao Sun, Weiwei Deng, Furu Wei, and Qi Zhang. 2024a. [se²: Sequential example selection for in-context learning](#). In *Findings of the Association for Computational Linguistics ACL 2024*, pages 5262–5284, Bangkok, Thailand and virtual meeting. Association for Computational Linguistics.
- Nelson F. Liu, Kevin Lin, John Hewitt, Ashwin Paranjape, Michele Bevilacqua, Fabio Petroni, and Percy Liang. 2024b. [Lost in the middle: How language models use long contexts](#). *Transactions of the Association for Computational Linguistics*, 12:157–173.
- Ye Liu, Kazuma Hashimoto, Yingbo Zhou, Semih Yavuz, Caiming Xiong, and Philip Yu. 2021. [Dense hierarchical retrieval for open-domain question answering](#). In *Findings of the Association for Computational Linguistics: EMNLP 2021*, pages 188–200, Punta Cana, Dominican Republic. Association for Computational Linguistics.
- I Loshchilov. 2017. Decoupled weight decay regularization. *arXiv preprint arXiv:1711.05101*.
- Kun Luo, Zheng Liu, Shitao Xiao, and Kang Liu. 2024. Bge landmark embedding: A chunking-free embedding method for retrieval augmented long-context large language models. *arXiv preprint arXiv:2402.11573*.

725	Tomas Mikolov, Ilya Sutskever, Kai Chen, Greg S Corrado, and Jeff Dean. 2013. Distributed representations of words and phrases and their compositionality. <i>Advances in neural information processing systems</i> , 26.	<i>Conference on Empirical Methods in Natural Language Processing</i> , pages 5627–5646, Miami, Florida, USA. Association for Computational Linguistics.	781
726			782
727			783
728			
729			
730	Niklas Muennighoff, Hongjin Su, Liang Wang, Nan Yang, Furu Wei, Tao Yu, Amanpreet Singh, and Douwe Kiela. 2024. Generative representational instruction tuning .	S Wang. 2016. Machine comprehension using match-lstm and answer pointer. <i>arXiv preprint arXiv:1608.07905</i> .	784
731			785
732		Shitao Xiao, Zheng Liu, Peitian Zhang, Niklas Muennighoff, Defu Lian, and Jian-Yun Nie. 2024. C-pack: Packed resources for general chinese embeddings . In <i>Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval, SIGIR '24</i> , page 641–649, New York, NY, USA. Association for Computing Machinery.	786
733			787
734	Jeffrey Pennington, Richard Socher, and Christopher D Manning. 2014. Glove: Global vectors for word representation. In <i>Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)</i> , pages 1532–1543.		788
735			789
736			790
737			791
738			792
739	Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. 2021. Learning transferable visual models from natural language supervision. In <i>International conference on machine learning</i> , pages 8748–8763. PMLR.	Lee Xiong, Chenyan Xiong, Ye Li, Kwok-Fung Tang, Jialin Liu, Paul Bennett, Junaid Ahmed, and Arnold Overwijk. 2021. Approximate nearest neighbor negative contrastive learning for dense text retrieval . In <i>International Conference on Learning Representations (ICLR)</i> .	793
740			794
741			795
742			796
743			797
744			798
745	Pranav Rajpurkar, Jian Zhang, Konstantin Lopyrev, and Percy Liang. 2016. SQuAD: 100,000+ questions for machine comprehension of text . In <i>Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing</i> , pages 2383–2392, Austin, Texas. Association for Computational Linguistics.	Hu Xu, Bing Liu, Lei Shu, and Philip Yu. 2019. BERT post-training for review reading comprehension and aspect-based sentiment analysis . 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 2324–2335, Minneapolis, Minnesota. Association for Computational Linguistics.	801
746			802
747			803
748			804
749			805
750			806
751	Nils Reimers and Iryna Gurevych. 2019. Sentence-BERT: Sentence embeddings using Siamese BERT-networks . In <i>Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)</i> , pages 3982–3992, Hong Kong, China. Association for Computational Linguistics.	Zichao Yang, Diyi Yang, Chris Dyer, Xiaodong He, Alex Smola, and Eduard Hovy. 2016. Hierarchical attention networks for document classification . In <i>Proceedings of the 2016 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies</i> , pages 1480–1489, San Diego, California. Association for Computational Linguistics.	807
752			808
753			809
754			810
755			811
756			812
757			813
758			814
759	M Seo. 2016. Bidirectional attention flow for machine comprehension. <i>arXiv preprint arXiv:1611.01603</i> .	Jiajie Zhang, Yushi Bai, Xin Lv, Wanjuan Gu, Danqing Liu, Minhao Zou, Shulin Cao, Lei Hou, Yuxiao Dong, Ling Feng, et al. 2024. Longcite: Enabling llms to generate fine-grained citations in long-context qa. <i>arXiv preprint arXiv:2409.02897</i> .	815
760			816
761	Oscar SKEAN, Md Rifat Arefin, Yann LeCun, and Ravid Shwartz-Ziv. 2024. Does representation matter? exploring intermediate layers in large language models. <i>arXiv preprint arXiv:2412.09563</i> .		817
762			818
763			819
764			820
765	Kaitao Song, Xu Tan, Tao Qin, Jianfeng Lu, and Tie-Yan Liu. 2020. Mpnnet: Masked and permuted pre-training for language understanding . In <i>Advances in Neural Information Processing Systems</i> , volume 33, pages 16857–16867. Curran Associates, Inc.	Dawei Zhu, Liang Wang, Nan Yang, Yifan Song, Wenhao Wu, Furu Wei, and Sujian Li. 2024. LongEmbed: Extending embedding models for long context retrieval . In <i>Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing</i> , pages 802–816, Miami, Florida, USA. Association for Computational Linguistics.	821
766			822
767			823
768			824
769			825
770	Liang Wang, Nan Yang, Xiaolong Huang, Binxing Jiao, Linjun Yang, Daxin Jiang, Rangan Majumder, and Furu Wei. 2022. Text embeddings by weakly-supervised contrastive pre-training. <i>arXiv preprint arXiv:2212.03533</i> .		826
771			827
772			828
773			
774			
775	Minzheng Wang, Longze Chen, Fu Cheng, Shengyi Liao, Xinghua Zhang, Bingli Wu, Haiyang Yu, Nan Xu, Lei Zhang, Run Luo, Yunshui Li, Min Yang, Fei Huang, and Yongbin Li. 2024. Leave no document behind: Benchmarking long-context LLMs with extended multi-doc QA . In <i>Proceedings of the 2024</i>		
776			
777			
778			
779			
780			

Appendices

A Data Construction

We present here the practice of synthesizing data for Relevant Information Retrieval scenarios.

Pre-processing Firstly, we choose high-quality documents from Wikipedia (Foundation). We process the documents sentence by sentence, removing sentences with repetitive line breaks and phrases, until the document processing is complete or the token count reaches 500 (<512). We remove the documents that are too short, with a sentence count less than 3 or a token count of less than 200. Second, we filter the candidate sentences in the document that can be rewritten: we filter all the sentences that have a token count between 8 and 20 and whose first word and subject are not pronouns (the set of pronouns includes "this", "these", "it", "that", "those", "they", "he", "she", "we", "you", "I"). If the number of sentences filtered is less than 3, we discard the document.

LLM Rewriting We randomly select 3 sentences in the document and use vLLM (Kwon et al., 2023) and "Llama-3.1-70B-Instruct" (Dubey et al., 2024) to rewrite them into queries, the prompt is: "".

Post-processing We de-duplicate the keywords in the rewritten query and then reorder them. To ensure the relevance of the query to the document, we perform a round of filtering using BGE (Xiao et al., 2024) to retain the data with a similarity of 0.5 or more between the rewritten query and the document. In this way we obtain a reasonable triad of queries, documents, and units (sentences).

For the construction of Relevant Information Retrieval data, we have also tried to collect paired sentences and make LLM expand one of them into a document. However, we found that other sentences in the LLM expansion were less informative than the original sentence, for example, being some descriptive statements were generated around the original sentence. This pattern tends to cause the model to learn to locate the central sentence, or the most informative sentence, in the expanded document, leading model to ignore the query. So please be aware of this if you plan to try this way of constructing your data.

Hyperparameter	Assignment
Computing Infrastructure	16 MI200-64GB GPUs
Number of epochs	10
Batch size per GPU	48 / 16
Maximum sequence length	512
Optimizer	AdamW
AdamW epsilon	1e-8
AdamW beta weights	0.9, 0.999
Learning rate scheduler	Cosine lr schedule
Initialization learning rate	1e-5
Minimum learning rate	1e-6
Weight decay	0.05
Warmup steps	1000
Warmup learning rate	1e-6

Table 4: Hyperparameter settings

B Overview of datasets

We describe here in detail the datasets used for training and evaluation.

B.1 Training

For Question Answer Retrieval, we sampled 30M data from PAQ (Lewis et al., 2021) datasets to train GeAR. For Relevant Information Retrieval, we used the 95% of the synthetic data for training. The specific statistics are shown in Table 5.

Scenario	Data Number
QAR	30,000,000
RIR	5,676,877

Table 5: Training data statistics.

B.2 Evaluation

In the evaluation stage, we introduce the specific information of the evaluation data by task.

Documents Retrieval First, for the document retrieval task, the queries come from the test set in the respective dataset, and the candidate documents are all documents within the entirety of the dataset, including the SQuAD (Rajpurkar et al., 2016), NQ (Kwiatkowski et al., 2019), TriviaQA (Joshi et al., 2017), and RIR datasets. It is difficult to encode all the documents of the PAQ dataset because the dataset is too large. So for the PAQ dataset, we sampled 1M documents and 20k queries, all of which have no intersection with the training data. The evaluation data statistics for the document retrieval task are shown in Table 6.

Scenario	Dataset	Documents Number	Queries Number
QA	Squad	20,239	5,928
	NQ	64,501	2,889
	TriviaQA	104,160	14,000
	PAQ	932,601	20,000
RIR	RIR	2,315,413	145,562

Table 6: The evaluation data statistics for the document retrieval task.

Scenario	Dataset	Data Number
QA	Squad	5,928
	NQ	2,889
	TriviaQA	14,000
	PAQ	20,000
RIR	RIR	10,000

Table 7: The evaluation data statistics for the units localization and information generation tasks.

Units Localization and Information Generation For these two tasks, we directly use the test set data corresponding to the respective datasets. Therefore, their number is consistent with the number of queries in Table 6. For the RIR dataset, we sample 10k records as the test set. The evaluation data statistics for the units localization and information generation tasks are shown in Table 7.

C HyperParameters and Implementation Details

We run model training on 16 AMD MI200 GPUs with 64GB memory and evaluation on 8 NVIDIA Tesla V100 GPUs with 32GB memory. The learning rate is warmed-up from $1e-6$ to $1e-5$ in the first 1000 steps, and then following a cosine scheduler, where the minimum learning rate is $1e-6$. The momentum parameter for updating momentum encoder is set as 0.995, the queue size is set as 57600. We linearly ramp-up the soft labels weight from 0 to 0.4 within the first 2 epoch. The overall hyperparameters are detailed in Table 4. We use FAISS (Douze et al., 2024; Johnson et al., 2019) to store and search for vectors. The 2 encoders and 1 decoders in GeAR are the same size as "bert-base" (Devlin et al., 2019), the total number of parameters of GeAR is about 330M. The training time for QAR scenario is about 5 days, for RIR scenario is about 3 days.