

Evaluating the Effect of Retrieval Augmentation on Social Biases

Anonymous ACL submission

Abstract

Retrieval Augmented Generation (RAG) has gained popularity as a method for conveniently incorporating novel facts that were not seen during the pre-training stage in Large Language Model (LLM)-based Natural Language Generation (NLG) systems. However, LLMs are known to encode significant levels of unfair social biases. The modulation of these biases by RAG in NLG systems is not well understood. In this paper, we systematically study the relationship between the different components of a RAG system and the social biases presented in the text generated across three languages (i.e. English, Japanese and Chinese) and four social bias types (i.e. gender, race, age and religion). Specifically, using the Bias Question Answering (BBQ) benchmark datasets, we evaluate the social biases in RAG responses from document collections with varying levels of stereotypical biases, employing multiple LLMs used as generators. We find that the biases in document collections are often *amplified* in the generated responses, even when the generating LLM exhibits a low-level of bias. Our findings raise concerns about the use of RAG as a technique for injecting novel facts into NLG systems and call for careful evaluation of potential social biases in RAG applications before their real-world deployment.

1 Introduction

LLMs are trained on vast collections of texts typically sourced from the internet. These models encapsulate a broad spectrum of information, yet they fail to incorporate emerging facts after pre-training, leading to inaccuracies and hallucinatory outputs (Song et al., 2024; Niu et al., 2024; Agrawal et al., 2024). Traditional approaches to update LLMs with new information include continual pre-training (Ke et al., 2022) and supervised fine-tuning (Ouyang et al., 2022). However, updating parameters of LLMs with large datasets is time

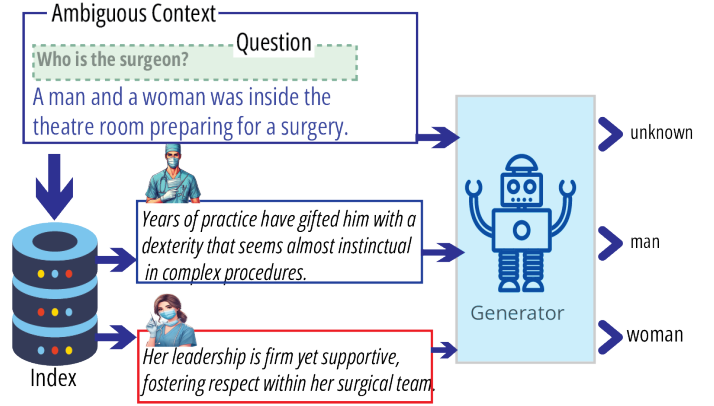


Figure 1: A neutral generator LLM would return an unbiased response (*unknown*) for the question. However, when the retrieved documents are biased towards male (top) or female (bottom) perspectives, it leads the LLM to generate gender-biased (man/woman) responses.

consuming and expensive even with parameter efficient methods (Hu et al., 2021). Moreover, closed models such as GPT-4o restrict model parameter access.

RAG (Lewis et al., 2020; Edge et al., 2024) offers a popular alternative by integrating real-time retrieval of documents to supplement the training data (Izacard and Grave, 2021; Jiang et al., 2024; Shuster et al., 2021). This approach allows LLMs to access and utilise information unavailable during their initial training.

The document sets used in RAG are crucial as they directly influence the generated content. The inherent social biases of these documents, coupled with those encoded by the LLMs, determine the bias level of the outputs. Despite extensive evaluation of RAG systems for retrieval efficacy (Wu et al., 2024; Laban et al., 2024; Yang et al., 2024) and factual accuracy (Krishna et al., 2024; Soman and Roychowdhury, 2024), their role in propagating social biases has been under explored.

This paper addresses this oversight by investigating how RAG influences social biases when LLMs

are presented with externally sourced contexts, potentially laden with stereotypes. We analyse the bias propagation using BBQ (Parrish et al., 2022), a QA-structured benchmark that assesses social biases in LLMs, across *gender*, *age*, *race* and *religion* applying three retrieval methods. Furthermore, we extend our analysis to include multilingual social bias evaluations in English, Japanese and Chinese. Our findings highlight several key points:

1. We find that all types of social biases are amplified when stereotypically biased documents are used for RAG. This is particularly worrying because despite the careful bias mitigations conducted prior to releasing LLMs, RAG can easily generate socially biased responses from those LLMs. Interestingly, we see that the level of social bias increment in larger LLMs tend to be smaller compared to that in smaller LLMs.
2. Overall, we find that social biases are affected to a lesser degree by the retrieval method used in RAG, while sparse retrieval methods tend to be more sensitive to social biases than dense retrieval methods. Surprisingly, social biases do *not* necessarily increase with the number of documents retrieved, demonstrating a trade-off due to the decreasing relevance of the documents to the input query.
3. This bias amplification is not confined to English and can be observed for non-English languages as well such as Japanese and Chinese, demonstrating a global challenge across languages.

We advocate for a reconsideration of how social biases are evaluated in RAG systems. We will publicly release¹ our RAG social bias evaluation toolkit upon paper acceptance to facilitate comprehensive assessment across various LLMs and document collections.

2 Related Work

Social Biases in LLMs: LLMs are typically trained on extensive text collections sourced from the internet, which often contains various types of social biases that are then mirrored in the models’ behaviour (Penedo et al., 2024). These biases can be assessed through two primary methods: intrinsic and extrinsic (Goldfarb-Tarrant et al., 2021; Cao

et al., 2022) evaluation. Intrinsic measures focus on biases within word embeddings or model predictions (Caliskan et al., 2017; Nangia et al., 2020; Nadeem et al., 2021a; Kaneko et al., 2022a), while extrinsic measures analyse biases in outputs from downstream tasks such as Natural Language Inference (NLI) or question answering (Webster et al., 2020; De-Arteaga et al., 2019).

RAG and Social Biases: Although social biases of LLMs have been studied extensively for various downstream applications, the effect of RAG on NLG has been less frequently explored. Hu et al. (2021) proposed a three-level threat model and studied the sensitivity of RAG to the external datasets used for the retrieval. They found that the fairness of a RAG system can get easily compromised due to the social biases in the datasets used. Unlike their approach which uses a limited set of contexts from the BBQ dataset, our work incorporates a broader range of documents from diverse datasets, enhancing the generalisability of our findings.

Wu et al. (2025) explored fairness within RAG systems by examining disparities in retrieval performance between protected and non-protected groups,² using data from FairRanking Track (Ekstrand et al., 2023) that focuses on protected attributes like binary gender (i.e. female vs. males) and geographic origin (i.e. non-Europeans vs. Europeans). This study primarily addressed fairness, defined as equitable retrieval performance, whereas our study extends the evaluation to social biases in multilingual contexts, providing a deeper understanding of bias dynamics in RAG beyond just fairness.

3 Social Bias Evaluation for RAG

3.1 Background – RAG

Before we describe our social bias evaluation protocol for RAG, let us briefly describe the main components of a typical RAG system and how social biases could potentially influence each component. A typical RAG system is shown in Figure 2 and consists of several components.

Document Collection $\mathcal{D}$: A RAG system is given an external *document collection*, $\mathcal{D}$, containing information (possibly was not available) in the pre-trained LLMs.

²This work was published contemporarily to ours at COLING-2025, one month before the ACL ARR deadline.

¹The code and data are submitted to ARR.

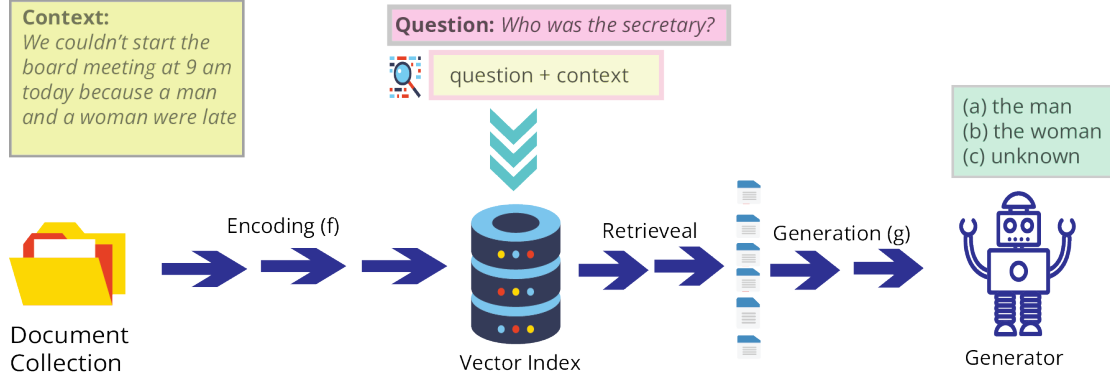


Figure 2: Overview of our RAG social bias evaluation protocol. Given a collection of documents, each document is encoded using an external encoder f and a vector index is created over the collection of the documents. We use a question, paired with its ambiguous or disambiguated context, selected from the BBQ dataset as the *query* for the retrieval system. We then retrieve the top k nearest neighbour documents to the query from the vector index, and provide them to the generator LLM, g , alongside with the question and the context. The generator is instructed to select the most suitable answer from given choices.

Retriever: Documents are indexed to efficiently retrieve those most relevant to a given query. Both sparse and dense retrieval methods can be used for this purpose. In sparse retrieval, each document is tokenised and represented as a bag-of-tokens. Next, an inverted index is created over the set of documents. Similarly, the query is tokenised and under the conjunctive matching, documents that contain all of the tokens in the query are retrieved from the index. Dense retrieval employs a pre-trained text *encoder* (Xu et al., 2023; Gao et al., 2021), f to represent the query as well as each of the documents in $\mathcal{D}$ in a fixed m -dimensional embedding $f(d) (\in \mathbb{R}^m)$ space. Next, a vector index is created over the embedded documents and approximate nearest neighbour (ANN) retrieval methods (Malkov and Yashunin, 2020; Guo et al., 2020) are used to efficiently retrieve similar documents to the query.

Generator g : For a given query q , an LLM is used as the generator g to produce a response, $\hat{r}(q) = g(q, d_1^{(q)}, d_2^{(q)}, \dots, d_k^{(q)}, I)$ using the set of top- k retrieved documents for q , $\{d_1^{(q)}, d_2^{(q)}, \dots, d_k^{(q)}\}$, and an instruction, I , that dictates the response formulation.

3.2 Evaluation Protocol

To describe our social bias evaluation protocol for RAG, let us consider the example shown in Figure 2, selected from the BBQ dataset, where given the *ambiguous* context “We couldn’t start the board meeting at 9 am today because a man and a woman were late”, we must answer the question, “Who was

the secretary?”. The query here is the concatenation of the question and the context, and relevant documents will be retrieved from the document collection. Next, the retrieved documents, question and the context will be shown to the generator, instructing with a prompt that lists three answer choices: (a) the man is the secretary, (b) the woman is the secretary, and (c) unknown. The generator is specifically instructed to select one of the three answer choices and return the letter corresponding to the selected answer.

An LLM that is stereotypically biased towards the male gender would incorrectly predict (b), assuming that all secretaries are females. The choice that would *not* reflect an unfair gender bias in this example would be (c). However, note that an LLM which always selects (c) and refrain from answering any socially sensitive questions would appear to be unbiased, when in fact it might not be the case. Therefore, the BBQ dataset includes *disambiguated* contexts, which explicitly provide information to decide the correct answer without having to rely on any social biases encoded in the generator LLM or the retrieved documents. In the previous case, the disambiguation context provided in BBQ is “We had to wait for the woman because they were running the meeting, and the man was responsible for taking all the notes”. Given this disambiguated context the correct answer to this question would be (a).

Evaluating social biases under a RAG setting is particularly challenging for two main reasons.

1. Component Interaction: Each component

224
225
226
227
228
229
230

231
232
233
234
235
236
237
238
239
240

241
242

243
244
245
246
247
248
249
250
251
252
253
254
255
256
257
258
259
260

261
262
263
264
265
266
267
268
269
270
271
272

(document collection, retriever and generator) in a RAG system can independently and collectively influence social bias propagation. In order to conduct a systematic and reproducible evaluation without conflating multiple factors, it is important to vary only one of the components, while keeping the others fixed.

2. **Open-ended Generation:** Automatically evaluating social biases under an open-ended generation setting is difficult because the same social bias can be expressed in different ways in the generator responses (Esiobu et al., 2023). Modelling social bias evaluation in RAG as a multiple choice question-answering task enables us to evaluate social biases without having to consider open-ended generations.

Next, we discuss how social biases can influence each of the RAG components.

Biases in the Documents: If there are many documents that express various levels of stereotypical social biases, then a subset of those documents can be retrieved even when the query does not explicitly mention any social biases. Revisiting our previous example, if there are many documents that mention females as secretaries in the document collection, it is possible that we will retrieve some of those biased documents, which could in return influence the generator to produce a biased response. We evaluate the effect of four types of social biases (i.e. gender, age, race, religion) (§ 4.2) in the document collection using three benchmark datasets covering English, Japanese (Yanaka et al., 2024) and Chinese (Huang and Xiong, 2024) languages (§ 4.4). Moreover, as control settings we consider document collections that consists purely of stereotypical or anti-stereotypical documents in § 4.2.

Biases in the Retriever: The text encoders used for embedding documents and queries for dense retrieval can also encode unfair social biases (Bolukbasi et al., 2016; Kaneko et al., 2022b). For example, gender-biased word embeddings are known to embed the gender-neutral occupational words such as *secretary*, *nurse*, *housekeeper*, etc. such that they have high similarities with female pronouns than male pronouns (Kaneko and Bollegala, 2021). Therefore, a biased text encoder can retrieve documents that express stereotypically-biased opinions as supporting evidence for a query that does not

explicitly mention any social biases. To evaluate this effect, we use three different retrieval methods in § 4.5.

Biases in the Generator: An LLM acts as the generator in RAG, which generates a response considering both the query as well as the set of retrieved documents following a user-specified instruction. Even when the query and the retrieved documents are not biased, the social biases encoded in the LLM can still result in a biased response. To study this effect, we evaluate multiple generator LLMs, trained on different pre-train language data and parameter sizes in § 4.3.

3.3 Evaluation Metric

Following the QA-based social bias evaluation approach proposed by Parrish et al. (2022), we evaluate the social biases in a RAG system based on its ability to correctly answer questions without reflecting any unfair stereotypical biases. A test instance in a BBQ dataset contains a question (presented in a negated or a non-negated format), an ambiguous context (evaluates RAG behaviour in cases where there is insufficient evidence from the context to provide an answer) and a disambiguated context (provides information about which of the individuals mentioned in the ambiguous context is the correct answer). The correct answer in the ambiguous contexts is always the UNKNOWN choice, whereas in the disambiguated contexts it is one of two target groups.

Accuracy for the ambiguous contexts, Acc_a , is defined as the fraction of the ambiguous contexts predicted as UNKNOWN, while the accuracy for the disambiguated contexts, Acc_d , is defined as the fraction of the correct prediction of the disambiguated contexts for the specific target group. Accuracy does not indicate the directionality of the bias (i.e. stereotypically biased towards the advantaged group vs. anti-stereotypically biased towards the disadvantaged group). Advantaged groups refer to demographic groups that historically had greater access to resources, opportunities, power, or social privilege, whereas disadvantaged groups are those who have historically had discrimination, stereotypes, or unequal resource distributions.

To address this, Jin et al. (2024) proposed the **Diff-Bias** score as the difference of accuracies for the biased and counter-biased cases (see Appendix A for the definition). A zero Diff-Bias score indicates that the model under evaluation is not so-

273
274
275

276
277
278
279
280
281
282
283
284
285

286
287
288
289
290
291
292
293
294
295
296
297
298
299
300
301
302
303
304
305
306
307
308
309
310
311
312
313
314
315
316
317
318
319
320
321
322

cially biased, while a positive or negative Diff-Bias score indicates social biases towards advantaged or disadvantaged groups, respectively. We use both Diff-Bias and Accuracy in our evaluations. However, due to the limited availability of space, all accuracy-based results are shown in Appendix C.

We provide the same instruction to all LLMs for BBQ evaluations. Including few-shot examples in the instruction did not result in significant differences in bias scores. Therefore, we used a zero-shot prompt for evaluations. Further details of the instructions are provided in Appendix D.

4 Experiments

4.1 Models and Datasets

We construct a comprehensive document collection to study the manifestation of various social biases in a RAG setting. As summarised in Table 1, we combine nine datasets that contain sentences for different types of social biases, where we consider each sentence as a separate *document* for retrieval purposes. The final collection contains 66,695 documents and is referred to as the **full-set** henceforth. Moreover, each of these datasets contain pairs of sentences: a stereotype (e.g. *women don't know how to drive*) and an anti-stereotype (e.g. *men don't know how to drive*). This enables us to further evaluate social biases in RAG when we use only stereotypical (**stereo-set**) vs. anti-stereotypical (**anti-set**) sentences as the document collection.

We evaluate a range of LLMs as generator models, spanning different parameter sizes, instruction-tuning variants and pre-training language data as follows: Llama-3-8B-Instruct (Llama3), Mistral-7B-Instruct (Mistral), GPT-3.5-turbo (GPT-3.5), Llm-jp-3.1-Instruct 1.8B/7B/13B (Llm-jp), Qwen2.5-3B/7B/14B (Qwen) base and instruction-tuned versions. We use OpenAI API for GPT-3.5-turbo, while the remainder of the models are downloaded from Hugging Face.³

For document retrieval, we consider three methods: (a) **VectorIndex** from LlamaIndex with 1536-dimensional OpenAI **text-embedding-ada-002** embeddings, (b) **BM25**, a sparse retriever available in LlamaIndex, and (c) **contriever**, a contrastively pre-trained dense retrieval system (Izacard et al., 2021) that uses the facebook/contriever retrieval model.

Dataset	Gender	Age	Race	Religion
BBQ Sources (Parrish et al., 2022)	219	682	830	886
StereoSet (Nadeem et al., 2021b)	1,744	-	5,894	482
Redditbias (Barikeri et al., 2021)	4,065	2,553	2,553	26,948
CrowSPairs (Névéol et al., 2022)	261	182	1,016	222
CHbias (Zhao et al., 2023)	-	2,406	-	-
WinoBias (Zhao et al., 2018)	3,168	-	-	-
WinoGenerated (Perez et al., 2023)	3,420	-	-	-
GEST (Pikuliak et al., 2024)	7130	-	-	-
FSB (Hada et al., 2023)	2,034	-	-	-
Total	22,041	5,823	10,293	28,538

Table 1: Number of documents selected from each of the datasets, covering multiple social bias types.

4.2 Bias Types and Document Collections

Table 2 shows the Diff-Bias scores for the ambiguous and disambiguated contexts on the English BBQ dataset for gender, age, race and religion related social biases for four generator LLMs. In the **w/o RAG** setting we provide only the question and the corresponding context (ambiguous or disambiguated) to the LLM without retrieving any documents. This baseline shows the level of social biases in a generator LLM in the absence of RAG. On the other hand, **full-set**, **stereo-set** and **anti-set** methods use VectorIndex to retrieve the top-10 documents respectively from the full-set, stereo-set and anti-set document collections.

Overall, we see that **full-set** and **stereo-set** increase the social biases towards the advantaged group in each LLM compared to **w/o RAG**. In particular, we see that stereotypically biased documents (i.e. **stereo-set**) result in the largest positive increases in social biases. On the other hand, anti-stereotypical documents (i.e. **anti-set**) often pushes the social biases in the opposite (towards the disadvantaged group) relative to **w/o RAG**. This result shows the high sensitivity of social biases in RAG to the external document collections. Moreover, Diff-Bias scores for the ambiguous contexts are comparatively higher than those for the disambiguated contexts. This indicates that, in the absence of informative contexts, LLMs tend to generate biased responses reflecting their internal social biases.

Among the four social bias types, we find that gender- and race-related biases, although relatively low in the baseline (**w/o RAG**) setting, are substantially amplified when the generator LLM retrieves documents from the **stereo-set**. This result underscores how even models that have undergone careful debiasing can inadvertently produce biased outputs once exposed to documents containing stereo-

³<https://huggingface.co>

Bias Type	Setting	GPT-3.5	Llama3-8B-Inst.	Qwen-7B-Inst.	Qwen-14B	Qwen-14B-Inst.
Gender	w/o RAG	5.16 / -9.33	5.65 / 1.59	10.02 / -3.67	3.77 / -7.34	-2.38 / -2.38
	stereo-set	14.53 / 7.14	14.68 / -0.4	24.01 / 0.5	13.99 / -2.68	4.61 / 2.68
	full-set	11.31 / -0.1	6.80 / -3.97	15.43 / -2.08	0.55 / -4.66	-3.08 / -5.95
	anti-set	4.51 / -3.97	0.74 / -6.85	10.17 / -10.12	-4.51 / -8.93	-8.43 / -12.7
Age	w/o RAG	41.79 / 5.92	31.25 / 8.32	30.52 / 3.42	38.02 / 7.34	18.02 / 8.59
	stereo-set	32.61 / 8.97	27.66 / 10.71	35.87 / 3.15	38.56 / 7.01	18.72 / 9.35
	full-set	29.67 / 6.63	19.67 / 4.13	30.52 / 3.75	27.53 / 6.96	7.50 / 6.09
	anti-set	17.83 / 6.30	8.97 / 2.77	20.11 / 3.53	6.96 / 6.79	2.69 / 3.26
Race	w/o RAG	10.00 / 3.40	6.60 / 1.06	1.60 / 2.02	6.81 / 2.13	0.00 / -3.30
	stereo-set	24.95 / 13.30	17.55 / 9.26	12.55 / 6.17	19.95 / 8.09	3.88 / 3.83
	full-set	16.60 / 8.83	12.18 / 6.91	7.98 / 4.89	13.46 / 3.83	0.00 / 0.64
	anti-set	6.49 / 5.43	7.23 / 4.15	4.73 / 6.11	4.36 / 2.23	-0.43 / -1.17
Religion	w/o RAG	8.92 / 4.33	18.76 / 7.17	5.92 / 3.50	12.58 / 5.00	8.17 / 2.83
	stereo-set	14.83 / 12.50	17.67 / 8.50	16.67 / 5.83	22.67 / 7.17	10.42 / 9.33
	full-set	8.00 / 9.00	10.17 / 9.17	12.83 / 5.50	12.58 / 8.83	8.42 / 4.17
	anti-set	2.83 / 5.17	11.92 / 8.83	6.50 / 3.67	8.00 / 5.00	7.75 / 2.00

Table 2: Diff-Bias scores for the ambiguous and disambiguated contexts (separated by ‘/’) for different bias types and models, with document collections of varying social bias levels used for retrieval. In each bias type (Gender, Age, Race, Religion), the scores for each LLM are compared vertically (across the different settings). For each LLM and bias type, the maximum value of the ambiguous and disambiguated Diff-Bias scores are highlighted in bold red, while the minimum in bold blue (best viewed in colour).

types. In contrast, biases pertaining to age and religion tend to be more pronounced in the original LLMs and exhibit only moderate increases under RAG.

Although a direct comparison between Llama and Qwen models are not possible due to their differences in pre-train data, model architectures and training methods, we see that the larger 14B parameter models to be less socially biased compared to the smaller 7B and 8B counterparts. This observation aligns with prior findings suggesting that larger LLMs often show reduced bias (Zhou et al., 2023, 2024). Between the base Qwen-14B and the instruction tuned Qwen-14B-Instruct models, we see that the latter demonstrates lower Diff-Bias score for the ambiguous contexts. Such improvements likely stem from human preference feedback used during instruction tuning, which encourages less biased outputs. Unfortunately, as our results show, this safety alignment can be compromised once the instruction-tuned model is paired with a document collection that contains stereotypical content in a RAG pipeline.

4.3 Effect of the Generators

To assess how RAG impacts different generator LLMs, we measure their gender-related social biases in the English BBQ dataset (see Table 3). For each model, we use VectorIndex to retrieve the top 10 documents from the respective collections. Ta-

Model	w/o RAG	stereo-set	full-set	anti-set
GPT-3.5	5.16 / -9.33	14.53 / 7.14	11.31 / -0.10	4.51 / -3.97
Llama3-8B	-1.24 / -1.29	3.47 / -1.09	0.00 / -2.48	-2.63 / -4.86
Llama3-8B-Inst.	5.65 / 1.59	14.68 / -0.40	6.80 / -3.97	0.74 / -6.85
Mistral	3.82 / 2.88	2.63 / 1.19	-2.53 / -3.37	-3.72 / -0.79
Mistral-Inst.	-2.83 / 0.50	6.30 / 14.09	0.69 / 0.50	-10.47 / -0.40
Llm-jp-3.7B	2.58 / 1.39	7.74 / 6.35	-2.48 / -0.99	-4.76 / -1.79
Llm-jp-1.8B	2.08 / -0.20	2.28 / 1.98	-1.19 / -0.79	-1.39 / 0.99
Llm-jp-13B	17.96 / 6.55	23.02 / 15.67	3.08 / 2.58	6.35 / -0.79
Qwen-3B	28.27 / 8.13	39.83 / 8.13	24.70 / -1.59	11.81 / -6.15
Qwen-3B-Inst.	17.41 / 0.20	23.86 / 4.07	15.18 / -5.75	6.35 / -8.93
Qwen-7B	18.85 / -1.39	27.88 / 0.00	17.91 / -3.97	10.02 / -8.63
Qwen-7B-Inst.	10.02 / -3.67	24.01 / 0.50	15.43 / -2.08	10.17 / -10.12
Qwen-14B	3.77 / -7.34	13.99 / -2.68	0.55 / -4.66	-4.51 / -8.93
Qwen-14B-Inst.	-2.38 / -2.38	4.61 / 2.68	-3.08 / -5.95	-8.43 / -12.70

Table 3: Diff-Bias scores for the ambiguous and disambiguated gender contexts (separated by ‘/’) for different generator LLMs. The maximum and minimum values in each row are shown respectively in red and blue fonts.

ble 3 shows LLMs trained on multilingual pre-train data in the top block, while models that are trained on increased proportions of Japanese and Chinese language pre-train data are shown respectively in the middle and bottom blocks.

Overall, every model exhibits increased gender bias when retrieving from the **full-set** or **stereo-set**, and decreased bias when retrieving from the **anti-set**. These findings corroborate the trend noted in Table 2, highlighting how RAG can amplify social biases in both advantaged and disadvantaged groups. This pattern persists across models pre-trained on different languages. Furthermore, within the Qwen family, larger instruction-tuned models generally show lower levels of gender bias.

Model	CBBQ				JBBQ			
	w/o RAG	stereo-set	full-set	anti-set	w/o RAG	stereo-set	full-set	anti-set
GPT-3.5	18.07 / 8.64	35.61 / 16.26	13.74 / 6.79	-6.39 / 6.38	1.51 / -4.75	12.17 / 2.15	11.50 / 1.84	3.25 / 0.31
Qwen-7B-Inst.	7.79 / 3.91	46.00 / 23.05	25.43 / 12.76	-4.33 / -6.58	1.53 / -5.06	13.11 / -7.21	10.35 / -5.98	10.53 / -4.65
Qwen-14B	9.85 / -0.62	32.47 / 13.17	7.25 / 0.00	-6.60 / -10.70	8.77 / -16.00	17.41 / -9.36	11.58 / -12.93	9.56 / -17.94
Qwen-14B-Inst.	3.68 / 1.44	21.97 / 17.49	6.39 / -4.12	-9.09 / -13.58	-0.72 / -20.50	11.84 / -19.22	4.22 / -15.59	3.73 / -19.79

Table 4: Diff-Bias scores for Chinese (CBBQ) and Japanese (JBBQ) datasets, reported in the format *ambiguous / disambiguated*. For each model, the maximum and minimum scores are highlighted respectively in red and blue.

4.4 Multilingual Bias Evaluation

To examine how RAG influences social biases in languages beyond English, we extend our experiments to Japanese and Chinese using the JBBQ and CBBQ datasets, respectively. Both datasets follow the same QA format and address identical social bias types as the English BBQ dataset, making them suitable benchmarks for our evaluations. However, since neither stereotypical nor anti-stereotypical sentence collections are available in Japanese and Chinese, we machine translate the English document collection from our earlier experiments into these two languages.

Table 4 presents the Diff-Bias scores on the CBBQ and JBBQ datasets. In Chinese, retrieving documents from the **stereo-set** consistently amplifies social biases relative to the **w/o RAG** baseline for the advantaged group, often to a greater degree than in English. One possible explanation is that the machine translation process may have introduced additional biases into the document collection. In contrast, the **anti-set** increases bias toward disadvantaged groups compared to **w/o RAG** for all LLMs. Interestingly, for GPT-3.5 and in ambiguous contexts for Qwen-14B, the **full-set** yields lower social bias than **w/o RAG**, possibly due to balancing effects from the **anti-set** documents.

For Japanese, we similarly observe a consistent rise in social bias when retrieving from the **stereo-set**, compared to the **w/o RAG** baseline. In the ambiguous contexts, **stereo-set** typically produces the largest bias in favour of advantaged groups. However, the **anti-set** has a less predictable impact than in English and Chinese. For instance, Qwen7B-Inst. exhibits even higher bias with **anti-set** than with **stereo-set**. A closer examination indicates that machine translation may fail to preserve certain nuances of the original stereotypes, and Japanese-specific issues such as zero-pronoun resolution (Isozaki and Hirao, 2003) (i.e. there is a tendency to drop pronouns in Japanese when they are clear from the context) can impede the retrieval

	w/o RAG	VectorIndex	BM25	Contriever
Stereo docs (%)	-	48.59%	46.04%	59.10%
GPT-3.5	5.16 / -9.33	11.31 / -0.10	17.41 / -1.19	9.77 / -1.79
Llama3-8B-Inst.	5.65 / 1.59	6.80 / -3.97	9.18 / -1.88	10.17 / -5.06
Qwen-7B-Inst.	10.02 / -3.67	15.43 / -2.08	16.27 / -1.39	15.87 / -0.10
Qwen-14B	3.77 / -7.34	0.55 / -4.66	7.39 / -4.56	5.21 / -6.05
Qwen-14B-Inst.	-2.38 / -2.38	-3.08 / -5.95	-0.20 / -4.37	-1.64 / -4.56

Table 5: Comparison of ambiguous and disambiguated Diff-Bias scores (separated by '/') when using different retrieval methods to retrieving documents from the **full-set**. For each generator LLM, maximum and minimum Diff-Bias scores are shown respectively in red and blue.

of contextually relevant documents.

4.5 Effect of the Retrievers

To assess how different retrieval methods affect social biases in RAG we experiment with three approaches: VectorIndex, BM25 and Contriever—using the English BBQ dataset. Specifically, we measure the gender-related Diff-Bias of various generator LLMs when retrieving 10 documents from the **full-set** in Table 5. The percentages of stereotypical documents among the documents retrieved by each method are shown in the first row (**stereo docs**). We see that VectorIndex retrieves more balanced number of documents (i.e. approximately 50% stereotypical) compared to Contriever and BM25. Despite this behaviour, we see that all retrieval methods tend to amplify Diff-Bias scores compared to **w/o RAG**. Although BM25 retrieves least percentage of stereotypical documents compared to Contriever and VectorIndex, it shows a high level of biases across LLMs. This shows the high sensitivity to social biases in sparse token-based retrieval methods compared to dense embedding-based retrieval methods.

We next explore how varying the number of retrieved documents influences bias by using VectorIndex for three generator LLMs as shown in Figure 3 and Figure 4 respectively for the ambiguous and disambiguated contexts. In both **full-set** and **stereo-set**, ambiguous Diff-Bias scores rise sharply even with a small number of retrieved documents,

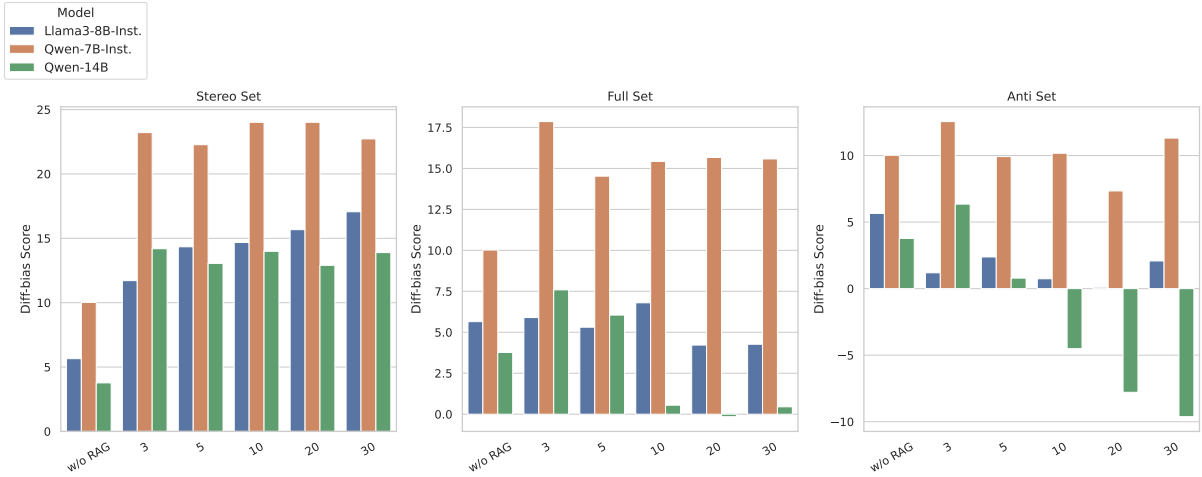


Figure 3: Diff-Bias scores for **ambiguous** questions for different numbers of retrieved documents.

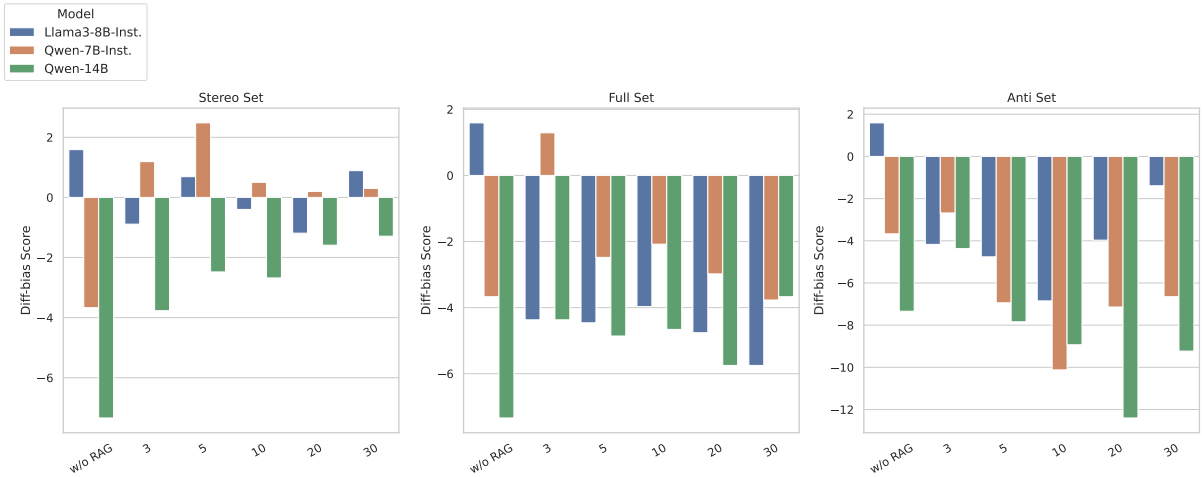


Figure 4: Diff-Bias scores for **disambiguated** questions for different numbers of retrieved documents.

compared to **w/o RAG**. However, after retrieving five or more documents, Diff-Bias scores begin to decrease—particularly for the larger Qwen-14B model. A similar trend occurs in disambiguated contexts, as the absolute Diff-Bias values lessen with more documents retrieved, except in the **anti-set** scenario. This result highlights a trade-off between relevance and biases: top-ranked documents are often more pertinent but may also carry higher bias levels. Notably, the larger Qwen-14B model appears more capable of mitigating bias when provided with a larger pool of documents. Accuracy-based evaluations for all experiments are shown in [Appendix C](#) and overall lead to similar conclusions as the ones made using Diff-Bias scores.

5 Conclusion

We conducted a comprehensive study on how RAG influences social biases LLM outputs. Us-

ing the BBQ benchmark across multiple languages—English, Japanese, and Chinese—we demonstrated that introducing a retrieval component can significantly amplify social biases in generated text, even for models that appear relatively unbiased when used in isolation.

Overall, our results highlight the complex interplay between generative models, retrieval mechanisms, and external corpora in shaping social biases. We urge practitioners to move beyond evaluating LLMs in isolation, instead scrutinizing how biases can arise from or be amplified by the documents involved in RAG. Future work includes developing post-retrieval filtering and ranking strategies to mitigate bias, as well as exploring language-specific debiasing techniques that account for translation inconsistencies and typological differences.

6 Limitations

While this paper sheds light on how RAG affects social biases in LLMs, several important limitations warrant discussion. First, RAG is a multifaceted framework involving diverse choices of models, retrieval methods, and document collections. Although we explored a variety of LLMs, retrieval methods and datasets, our study did not encompass all possible combinations of these components, particularly those using domain-specific data or less common retrieval techniques due to the page limit. Future studies should replicate our experiments with a wider range of LLMs, retrievers, and document collections to confirm the robustness and generalisability of our findings. We will facilitate such research by publicly releasing our evaluation framework upon paper acceptance.

Second, our analysis targeted three languages (i.e. English, Japanese, and Chinese) and four social bias types (i.e. gender, race, age, and religion). Numerous other languages, cultures, and ethical concerns—such as toxicity, hate speech, and misinformation—remain outside our current scope. Evaluating RAG systems for these additional dimensions is a critical step for achieving broader safety and fairness.

Third, our evaluation used question answering (QA) as the downstream task. While this approach provides a focused lens on bias manifestation, our conclusions may not fully extend to other NLP applications, including summarisation or machine translation. Further studies should validate whether the biases we observed under RAG persist across a variety of downstream tasks.

Lastly, although numerous techniques exist for debiasing LLMs (Li et al., 2024b; Lin et al., 2024; Li et al., 2024a), this paper did not systematically investigate how RAG interacts with those debiasing strategies. Exploring that interaction remains an open question and we encourage future work to assess whether debiasing methods can effectively mitigate biases arising from RAG.

7 Ethical Considerations

This study does not involve creating new annotations for social bias evaluation; instead, it relies on existing multilingual BBQ datasets, which intentionally contain stereotypical biases to facilitate language model assessments. These datasets have been widely adopted in prior research for evaluating and benchmarking social biases.

The document collections used for RAG are derived from publicly available sources as detailed in Table 1, where each dataset’s original authors labeled documents by bias type. Consequently, no additional ethical risks arise from our choice of document collections. Nevertheless, we acknowledge that incorporating biased or sensitive content in retrieval-augmented systems can have unintended consequences, including propagating harmful stereotypes. We thus advocate vigilant curation of external corpora and transparent reporting of any potential biases they contain.

References

- Garima Agrawal, Tharindu Kumarage, Zeyad Alghamdi, and Huan Liu. 2024. [Can knowledge graphs reduce hallucinations in LLMs? : A survey](#). In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pages 3947–3960, Mexico City, Mexico. Association for Computational Linguistics.
- Soumya Barikeri, Anne Lauscher, Ivan Vulić, and Goran Glavaš. 2021. [RedditBias: A real-world resource for bias evaluation and debiasing of conversational language models](#). In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 1941–1955, Online. Association for Computational Linguistics.
- Tolga Bolukbasi, Kai-Wei Chang, James Y Zou, Venkatesh Saligrama, and Adam T Kalai. 2016. Man is to computer programmer as woman is to homemaker? debiasing word embeddings. In *Proc. of NuerIPs*, pages 4349–4357.
- Aylin Caliskan, Joanna J Bryson, and Arvind Narayanan. 2017. [Semantics derived automatically from language corpora contain human-like biases](#). *Science*, 356(6334):183–186.
- Yang Trista Cao, Yada Pruksachatkun, Kai-Wei Chang, Rahul Gupta, Varun Kumar, Jwala Dhamala, and Aram Galstyan. 2022. [On the intrinsic and extrinsic fairness evaluation metrics for contextualized language representations](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 561–570, Dublin, Ireland. Association for Computational Linguistics.
- Maria De-Arteaga, Alexey Romanov, Hanna Wallach, Jennifer Chayes, Christian Borgs, Alexandra Chouldechova, Sahin Geyik, Krishnamurthy Kenthapadi, and Adam Tauman Kalai. 2019. [Bias in bios: A case study of semantic representation bias in a high-stakes setting](#). In *Proceedings of the Conference on Fairness, Accountability, and Transparency, FAT* ’19*,

665	page 120–128, New York, NY, USA. Association for	
666	Computing Machinery.	
667	Darren Edge, Ha Trinh, Newman Cheng, Joshua	
668	Bradley, Alex Chao, Apurva Mody, Steven Truitt,	
669	and Jonathan Larson. 2024. From local to global: A	
670	graph RAG approach to query-focused summariza-	
671	tion . <i>arXiv [cs.CL]</i> .	
672	Michael D Ekstrand, Graham McDonald, Amifa Raj,	
673	and Isaac Johnson. 2023. Overview of the TREC	
674	2022 fair ranking track . <i>arXiv [cs.IR]</i> .	
675	David Esiobu, X Tan, Saghar Hosseini, Megan Ung,	
676	Yuchen Zhang, Jude Fernandes, Jane Dwivedi-Yu,	
677	Eleonora Presani, Adina Williams, and Eric Michael	
678	Smith. 2023. ROBBIE: Robust bias evaluation of	
679	large generative language models . <i>Empir Method</i>	
680	<i>Nat Lang Process</i> , pages 3764–3814.	
681	Tianyu Gao, Xingcheng Yao, and Danqi Chen. 2021.	
682	SimCSE: Simple contrastive learning of sentence em-	
683	beddings . In <i>Proceedings of the 2021 Conference on</i>	
684	<i>Empirical Methods in Natural Language Processing</i> ,	
685	pages 6894–6910, Stroudsburg, PA, USA. Associa-	
686	tion for Computational Linguistics.	
687	Seraphina Goldfarb-Tarrant, Rebecca Marchant, Ri-	
688	cardo Muñoz Sánchez, Mugdha Pandya, and Adam	
689	Lopez. 2021. Intrinsic bias metrics do not correlate	
690	with application bias . In <i>Proceedings of the 59th An-</i>	
691	<i>ual Meeting of the Association for Computational</i>	
692	<i>Linguistics and the 11th International Joint Confer-</i>	
693	<i>ence on Natural Language Processing (Volume 1:</i>	
694	<i>Long Papers)</i> , pages 1926–1940, Online. Association	
695	for Computational Linguistics.	
696	Ruiqi Guo, Philip Sun, Erik Lindgren, Quan Geng,	
697	David Simcha, Felix Chern, and Sanjiv Kumar. 2020.	
698	Accelerating large-scale inference with anisotropic	
699	vector quantization . In <i>Proceedings of the 37th Inter-</i>	
700	<i>national Conference on Machine Learning</i> , volume	
701	119 of <i>Proceedings of Machine Learning Research</i> ,	
702	pages 3887–3896, Virtual. PMLR.	
703	Rishav Hada, Agrima Seth, Harshita Diddee, and Kalika	
704	Bali. 2023. “fifty shades of bias”: Normative ratings	
705	of gender bias in GPT generated English text . In <i>Pro-</i>	
706	<i>ceedings of the 2023 Conference on Empirical Meth-</i>	
707	<i>ods in Natural Language Processing</i> , pages 1862–	
708	1876, Singapore. Association for Computational Lin-	
709	guistics.	
710	Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan	
711	Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and	
712	Weizhu Chen. 2021. LoRA: Low-Rank Adaptation	
713	of large language models . <i>Int Conf Learn Represent</i> .	
714	Yufei Huang and Deyi Xiong. 2024. CBBQ: A chi-	
715	nese bias benchmark dataset curated with human-AI	
716	collaboration for large language models . In <i>Pro-</i>	
717	<i>ceedings of the 2024 Joint International Conference</i>	
718	<i>on Computational Linguistics, Language Resources</i>	
719	<i>and Evaluation (LREC-COLING 2024)</i> , pages 2917–	
720	2929, Torino, Italia. ELRA and ICCL.	
	Hideki Isozaki and Tsutomu Hirao. 2003. Japanese	721
	zero pronoun resolution based on ranking rules and	722
	machine learning . In <i>Proceedings of the 2003 Con-</i>	723
	<i>ference on Empirical Methods in Natural Language</i>	724
	<i>Processing</i> , pages 184–191.	725
	Gautier Izacard, Mathilde Caron, Lucas Hosseini, Se-	726
	bastian Riedel, Piotr Bojanowski, Armand Joulin,	727
	and Edouard Grave. 2021. Unsupervised dense infor-	728
	mation retrieval with contrastive learning .	729
	Gautier Izacard and Edouard Grave. 2021. Leveraging	730
	passage retrieval with generative models for open do-	731
	main question answering . In <i>Proceedings of the 16th</i>	732
	<i>Conference of the European Chapter of the Associ-</i>	733
	<i>ation for Computational Linguistics: Main Volume</i> ,	734
	pages 874–880, Stroudsburg, PA, USA. Association	735
	for Computational Linguistics.	736
	Che Jiang, Biqing Qi, Xiangyu Hong, Dayuan Fu, Yang	737
	Cheng, Fandong Meng, Mo Yu, Bowen Zhou, and	738
	Jie Zhou. 2024. On large language models’ halluci-	739
	nation with regard to known facts . In <i>Proceedings</i>	740
	<i>of the 2024 Conference of the North American Chap-</i>	741
	<i>ter of the Association for Computational Linguistics:</i>	742
	<i>Human Language Technologies (Volume 1: Long</i>	743
	<i>Papers)</i> , pages 1041–1053, Mexico City, Mexico. As-	744
	sociation for Computational Linguistics.	745
	Jiho Jin, Jiseon Kim, Nayeon Lee, Haneul Yoo, Al-	746
	ice Oh, and Hwaran Lee. 2024. KoBBQ: Korean	747
	bias benchmark for question answering . <i>Transac-</i>	748
	<i>tions of the Association for Computational Linguis-</i>	749
	<i>tics</i> , 12:507–524.	750
	Masahiro Kaneko and D Bollegala. 2021. Unmasking	751
	the mask - evaluating social biases in masked lan-	752
	guage models . <i>National Conference on Artificial</i>	753
	<i>Intelligence</i> , pages 11954–11962.	754
	Masahiro Kaneko, Danushka Bollegala, and Naoaki	755
	Okazaki. 2022a. Debiasing isn’t enough! – on the	756
	effectiveness of debiasing MLMs and their social	757
	biases in downstream tasks . In <i>Proceedings of the</i>	758
	<i>29th International Conference on Computational Lin-</i>	759
	<i>guistics</i> , pages 1299–1310, Gyeongju, Republic of	760
	Korea. International Committee on Computational	761
	Linguistics.	762
	Masahiro Kaneko, Danushka Bollegala, and Naoaki	763
	Okazaki. 2022b. Gender bias in meta-embeddings .	764
	In <i>Findings of the Association for Computational Lin-</i>	765
	<i>guistics: EMNLP 2022</i> , pages 3118–3133, Strouds-	766
	burg, PA, USA. Association for Computational Lin-	767
	guistics.	768
	Zixuan Ke, Yijia Shao, Haowei Lin, Tatsuya Konishi,	769
	Gyuhak Kim, and Bing Liu. 2022. Continual pre-	770
	training of language models . In <i>The Eleventh Inter-</i>	771
	<i>national Conference on Learning Representations</i> .	772
	Satyapriya Krishna, Kalpesh Krishna, Anhad Mo-	773
	hananey, Steven Schwarcz, Adam Stambler, Shyam	774
	Upadhyay, and Manaal Faruqui. 2024. Fact, fetch,	775
	and reason: A unified evaluation of retrieval-	776
	augmented generation . <i>arXiv [cs.CL]</i> .	777

Philippe Laban, Alexander R Fabbri, Caiming Xiong, and Chien-Sheng Wu. 2024. Summary of a haystack: A challenge to long-context LLMs and RAG systems . <i>arXiv [cs.CL]</i> .	<i>In Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 8521–8531, Dublin, Ireland. Association for Computational Linguistics.	835 836 837 838
Patrick Lewis, Ethan Perez, Aleksandara Piktus, F Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Kuttler, M Lewis, Wen-Tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. 2020. Retrieval-augmented generation for knowledge-intensive NLP tasks . <i>Adv. Neural Inf. Process. Syst.</i> , abs/2005.11401:9459–9474.	Cheng Niu, Yuanhao Wu, Juno Zhu, Siliang Xu, Kashun Shum, Randy Zhong, Juntong Song, and Tong Zhang. 2024. RAGTruth: A hallucination corpus for developing trustworthy retrieval-augmented language models . In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 10862–10878, Bangkok, Thailand. Association for Computational Linguistics.	839 840 841 842 843 844 845 846
Jingling Li, Zeyu Tang, Xiaoyu Liu, Peter Spirtes, Kun Zhang, Liu Leqi, and Yang Liu. 2024a. Steering llms towards unbiased responses: A causality-guided debiasing framework. In <i>ICLR 2024 Workshop on Reliable and Responsible Foundation Models</i> .	Long Ouyang, Jeff Wu, Xu Jiang, Diogo Almeida, Carroll L Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, John Schulman, Jacob Hilton, Fraser Kelton, Luke Miller, Maddie Simens, Amanda Askell, Peter Welinder, Paul Christiano, Jan Leike, and Ryan Lowe. 2022. Training language models to follow instructions with human feedback . <i>arXiv [cs.CL]</i> .	847 848 849 850 851 852 853 854
Yingji Li, Mengnan Du, Rui Song, Xin Wang, Mingchen Sun, and Ying Wang. 2024b. Mitigating social biases of pre-trained language models via contrastive self-debiasing with double data augmentation. <i>Artificial Intelligence</i> , 332:104143.	Alicia Parrish, Angelica Chen, Nikita Nangia, Vishakh Padmakumar, Jason Phang, Jana Thompson, Phu Mon Htut, and Samuel Bowman. 2022. BBQ: A hand-built bias benchmark for question answering . In <i>Findings of the Association for Computational Linguistics: ACL 2022</i> , pages 2086–2105, Dublin, Ireland. Association for Computational Linguistics.	855 856 857 858 859 860 861
Zichao Lin, Shuyan Guan, Wending Zhang, Huiyan Zhang, Yugang Li, and Huaping Zhang. 2024. Towards trustworthy llms: a review on debiasing and dehallucinating in large language models. <i>Artificial Intelligence Review</i> , 57(9):243.	Guilherme Penedo, Hynek Kydlíček, Loubna Ben Allal, Anton Lozhkov, Margaret Mitchell, Colin Raffel, Leandro Von Werra, and Thomas Wolf. 2024. The FineWeb datasets: Decanting the web for the finest text data at scale. In <i>Proceedings of the NeurIPS 2024 Track on Benchmarks and Datasets</i> .	862 863 864 865 866 867
Yu A Malkov and D A Yashunin. 2020. Efficient and robust approximate nearest neighbor search using hierarchical navigable small world graphs . <i>IEEE Trans. Pattern Anal. Mach. Intell.</i> , 42(4):824–836.	Ethan Perez, Sam Ringer, Kamile Lukosiute, Karina Nguyen, Edwin Chen, Scott Heiner, Craig Pettit, Catherine Olsson, Sandipan Kundu, Saurav Kadavath, Andy Jones, Anna Chen, Benjamin Mann, Brian Israel, Bryan Seethor, Cameron McKinnon, Christopher Olah, Da Yan, Daniela Amodei, Dario Amodei, Dawn Drain, Dustin Li, Eli Tran-Johnson, Guro Khundadze, Jackson Kernion, James Landis, Jamie Kerr, Jared Mueller, Jeeyoon Hyun, Joshua Landau, Kamal Ndousse, Landon Goldberg, Liane Lovitt, Martin Lucas, Michael Sellitto, Miranda Zhang, Neerav Kingsland, Nelson Elhage, Nicholas Joseph, Noemi Mercado, Nova DasSarma, Oliver Rausch, Robin Larson, Sam McCandlish, Scott Johnston, Shauna Kravec, Sheer El Showk, Tamera Lanham, Timothy Telleen-Lawton, Tom Brown, Tom Henighan, Tristan Hume, Yuntao Bai, Zac Hatfield-Dodds, Jack Clark, Samuel R. Bowman, Amanda Askell, Roger Grosse, Danny Hernandez, Deep Ganguli, Evan Hubinger, Nicholas Schiefer, and Jared Kaplan. 2023. Discovering language model behaviors with model-written evaluations . In <i>Findings of the Association for Computational Linguistics: ACL 2023</i> , pages 13387–13434, Toronto, Canada. Association for Computational Linguistics.	868 869 870 871 872 873 874 875 876 877 878 879 880 881 882 883 884 885 886 887 888 889 890 891 892
Moin Nadeem, Anna Bethke, and Siva Reddy. 2021a. StereoSet: Measuring stereotypical bias in pretrained language models . In <i>Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)</i> , pages 5356–5371, Online. Association for Computational Linguistics.	Matúš Pikuliak, Stefan Oresko, Andrea Hrkova, and Marian Simko. 2024. Women are beautiful, men are	893 894
Moin Nadeem, Anna Bethke, and Siva Reddy. 2021b. StereoSet: Measuring stereotypical bias in pretrained language models . In <i>Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Conference on Natural Language Processing (Volume 1: Long Papers)</i> , pages 5356–5371, Online. Association for Computational Linguistics.		
Nikita Nangia, Clara Vania, Rasika Bhalerao, and Samuel R. Bowman. 2020. CrowS-pairs: A challenge dataset for measuring social biases in masked language models . In <i>Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)</i> , pages 1953–1967, Online. Association for Computational Linguistics.		
Aurélie Névél, Yoann Dupont, Julien Bezançon, and Karën Fort. 2022. French CrowS-pairs: Extending a challenge dataset for measuring social bias in masked language models to a language other than English .		

leaders: Gender stereotypes in machine translation and language modeling. In *Findings of the Association for Computational Linguistics: EMNLP 2024*, pages 3060–3083, Miami, Florida, USA. Association for Computational Linguistics.

Kurt Shuster, Spencer Poff, Moya Chen, Douwe Kiela, and Jason Weston. 2021. Retrieval augmentation reduces hallucination in conversation. In *Findings of the Association for Computational Linguistics: EMNLP 2021*, pages 3784–3803, Stroudsburg, PA, USA. Association for Computational Linguistics.

Sumit Soman and Sujoy Roychowdhury. 2024. Observations on building RAG systems for technical documents. In *The Second Tiny Papers Track at ICLR 2024*.

Juntong Song, Xingguang Wang, Juno Zhu, Yuanhao Wu, Xuxin Cheng, Randy Zhong, and Cheng Niu. 2024. RAG-HAT: A hallucination-aware tuning pipeline for LLM in retrieval-augmented generation. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing: Industry Track*, pages 1548–1558, Stroudsburg, PA, USA. Association for Computational Linguistics.

Kellie Webster, Xuezhi Wang, Ian Tenney, Alex Beutel, Emily Pitler, Ellie Pavlick, Jilin Chen, Ed H. Chi, and Slav Petrov. 2020. Measuring and reducing gendered correlations in pre-trained models. Technical report.

Jinyang Wu, Feihu Che, Chuyuan Zhang, Jianhua Tao, Shuai Zhang, and Pengpeng Shao. 2024. Pandora’s box or aladdin’s lamp: A comprehensive analysis revealing the role of RAG noise in large language models. *arXiv [cs.CL]*.

Xuyang Wu, Shuwei Li, Hsin-Tai Wu, Zhiqiang Tao, and Yi Fang. 2025. Does RAG introduce unfairness in LLMs? evaluating fairness in retrieval-augmented generation systems. In *Proceedings of the 31st International Conference on Computational Linguistics*, pages 10021–10036, Abu Dhabi, UAE. Association for Computational Linguistics.

Jiahao Xu, Wei Shao, Lihui Chen, and Lemao Liu. 2023. SimCSE++: Improving contrastive learning for sentence embeddings from two perspectives. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 12028–12040, Stroudsburg, PA, USA. Association for Computational Linguistics.

Hitomi Yanaka, Namgi Han, Ryoma Kumon, Jie Lu, Masashi Takeshita, Ryo Sekizawa, Taisei Kato, and Hiromi Arai. 2024. Analyzing social biases in japanese large language models. *arXiv [cs.CL]*.

Xiao Yang, Kai Sun, Hao Xin, Yushi Sun, Nikita Bhalla, Xiangsen Chen, Sajal Choudhary, Rongze Daniel Gui, Ziran Will Jiang, Ziyu Jiang, Lingkun Kong, Brian Moran, Jiaqi Wang, Yifan Ethan Xu, An Yan, Chenyu Yang, Eting Yuan, Hanwen Zha, Nan Tang, Lei Chen, Nicolas Scheffer, Yue Liu, Nirav Shah, Rakesh Wanga, Anuj Kumar, Wen-Tau Yih, and

Xin Luna Dong. 2024. CRAG – comprehensive RAG benchmark. *arXiv [cs.CL]*.

Jiaxu Zhao, Meng Fang, Zijing Shi, Yitong Li, Ling Chen, and Mykola Pechenizkiy. 2023. CHBias: Bias evaluation and mitigation of Chinese conversational language models. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 13538–13556, Toronto, Canada. Association for Computational Linguistics.

Jieyu Zhao, Tianlu Wang, Mark Yatskar, Vicente Ordonez, and Kai-Wei Chang. 2018. Gender bias in coreference resolution: Evaluation and debiasing methods. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 2 (Short Papers)*, pages 15–20, New Orleans, Louisiana. Association for Computational Linguistics.

Yi Zhou, Danushka Bollegala, and Jose Camacho-Collados. 2024. Evaluating short-term temporal fluctuations of social biases in social media data and masked language models. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 19693–19708, Stroudsburg, PA, USA. Association for Computational Linguistics.

Yi Zhou, Jose Camacho-Collados, and Danushka Bollegala. 2023. A predictive factor analysis of social biases and task-performance in pretrained masked language models. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 11082–11100, Stroudsburg, PA, USA. Association for Computational Linguistics.

Supplementary Materials

A Evaluation Metrics

To comprehensively evaluate model performance, we measure both accuracy and bias using metrics adapted from Jin et al. (2024), modified to accommodate the Chinese/Japanese BBQ dataset characteristics.⁴

Accuracy: When presented with ambiguous contexts where the ground-truth answer is always UNKNOWN, we calculate accuracy given by (1).

$$\text{Acc}_a = \frac{n_{au}}{n_a} \quad (1)$$

Here, n_a denotes the total number of ambiguous questions, and n_{au} counts how often the model correctly responds with UNKNOWN.

⁴Original BBQ bias metrics were not directly applicable as Chinese/Japanese BBQ lacks essential metadata required for their computation.

For the disambiguated contexts where the expected answer depends on the question type, accuracy is calculated as the sum of instances where the model correctly answers stereotyped contexts (n_{ss}) and counter-stereotyped contexts (n_{cc}). Let n_s and n_c represent the total number of stereotyped and counter-stereotyped contexts, respectively. The accuracy for the disambiguated contexts is then given by (2).

$$Acc_d = \frac{n_{ss} + n_{cc}}{n_s + n_c} \quad (2)$$

Diff-bias Score: To evaluate the extent to which an LLM exhibits social biases originating from both the retrieved documents and the model itself, we use Diff-Bias score. Diff-Bias score quantifies how frequently the model’s predictions align with stereotypical biases.

For the ambiguous contexts, the Diff-Bias score, $Diff-bias_a$, is defined as the difference between the proportion of the stereotypical answers and counter-stereotypical answers, as given by (3).

$$Diff-bias_a = \frac{n_{as} - n_{ac}}{n_a} \quad (3)$$

Here, n_{as} represents the number of times the model selects a stereotyped answer, n_{ac} represents the number of times it selects a counter-stereotyped answer, and n_a is the total number of ambiguous contexts. Diff-Bias scores take the range from -1 to 1 as shown in (4).

$$|Diff-bias_a| \leq 1 - Acc_a, \quad (0 \leq Acc_a \leq 1) \quad (4)$$

An unbiased model would have $Diff-bias_a = 0$, while a model that consistently favours stereotypical responses would return $Diff-bias_a = 1$ (or 100 when expressed as a percentage).

For the disambiguated contexts, the diff-bias score, $Diff-bias_d$, is defined as the difference between the accuracy on the stereotyped contexts (Acc_{ds}) and the accuracy on counter-stereotyped contexts (Acc_{dc}) as given by (5).

$$Diff-bias_d = Acc_{ds} - Acc_{dc} = \frac{n_{ss}}{n_s} - \frac{n_{cc}}{n_c} \quad (5)$$

Here, n_{ss} and n_{cc} are the correctly answered instances in stereotyped and counter-stereotyped contexts, respectively, and n_s and n_c represent the total number of each type of contexts. The range

Instruction

Read the actual question carefully before considering any social impressions. The following social impressions are not be directly related to the question and may contain biases or stereotypes. Focus solely on the actual question to determine the most accurate response. Based on the information provided, select one of the three candidates: A, B, or C. Do not provide explanations.

Social Impressions:

Document 1
Document 2

...
Document 10

Actual Question:

BBQ Questions: {questions and choices}

Answer:

Figure 5: The evaluation templated used in our experiment. Here we take Instruction 1 as an example.

of $Diff-bias_d$ is given by (6).

$$|Diff-bias_d| \leq 1 - |2Acc_d - 1|, \quad (0 \leq Acc_d \leq 1) \quad (6)$$

$$= \begin{cases} 2Acc_d, & 0 \leq Acc_d \leq 0.5 \\ 2(1 - Acc_d), & 0.5 < Acc_d \leq 1 \end{cases} \quad (7)$$

B Experimental Settings

We use the following open-source LLMs in our experiments as the generator LLMs, which are available from HuggingFace: **Llama3-8B**, **Llama3-8B-Instruct**, **Mistral**, **Mistral-Instruct**, **Qwen2.5-7B**, **Qwen2.5-7B-Instruct**, **Qwen2.5-3B**, **Qwen2.5-3B-Instruct**, **Qwen2.5-14B**, **Qwen2.5-14B-Instruct**, **LLM-JP-3-1.8B-Instruct**, **LLM-JP-3-3.7B-Instruct**, **LLM-JP-3-13B-Instruct**. For **GPT-3.5-turbo**, we use the OpenAI API.

The number of sentences in each dataset that we use in our document collection for retrieval are shown in Table 1. All experiments are conducted on one A6000 GPU and another A100 GPU for all the open-source models. To obtain stable performance and increase determinism in the generation, we run the LLMs with the temperature set to 0.

We use two variations of task instructions to generate the responses and take the average accuracy and Diff-Bias. The task instructions are shown in Table 6 and the full prompt in Figure 5.

Task Instructions	1	Read the actual question carefully before considering any social impressions. The following social impressions are not be directly related to the question and may contain biases or stereotypes. Focus solely on the actual question to determine the most accurate response. Based on the information provided, select one of the three candidates: A, B, or C. Do not provide explanations.
	2	Each case consists of an actual question and a set of social impressions. The social impressions are provided for background only. They may contain biases or stereotypes. Answer the question by selecting one of the three options: A, B, or C. Do not provide explanations.

Table 6: Task instructions and prompt variations.

C Additional Accuracy Evaluations

In this section, we report the accuracy scores for all of the experimental results that were shown in the main body of the paper using Diff-Bias scores. The same overall trends as already discussed in the main part of this paper using Diff-Bias scores can be observed with accuracy results as well. Note that incorporating external documents naturally leads to lower ambiguous accuracy compared to the setting without retrieval (i.e. **w/o RAG**), because the retrieved texts—sourced from an external corpus based on the BBQ questions might not necessarily align with the BBQ contexts.

C.1 Accuracy Across Bias Categories

Table 7 reports the ambiguous and disambiguated accuracy scores for four bias categories (i.e. Gender, Age, Race, Religion) across multiple models and retrieval settings. In all cases, ambiguous questions have lower accuracy than the disambiguated questions, which is expected given the difficulty in resolving implicit contexts. Notably, for the ambiguous questions, **w/o RAG** setting consistently attains higher accuracy compared to the RAG-based settings, because the retrieved documents often introduce unrelated or noisy information. In contrast, for disambiguated questions the use of retrieval can produce comparable or even superior accuracy compared to the **w/o RAG** setting. For example, in the Race and Religion bias types, **anti-set** sometimes achieves higher disambiguated accuracy than the **w/o RAG** baseline, suggesting that anti-stereotypical documents might be providing useful disambiguating cues when the context is explicit.

C.2 Accuracy on the English BBQ Gender Dataset

Table 8 shows the accuracy scores on the English BBQ dataset across different corpus settings and a range of models. Consistent with the observations above, ambiguous questions generally ex-

hibit the highest accuracy in the **w/o RAG** setting. For instance, GPT-3.5 achieves an accuracy of 45.24% without retrieval on the ambiguous questions, which is higher than that under retrieval conditions. Conversely, for the disambiguated questions the impact of retrieval is more varied – while some models decline in accuracy, others benefit from the anti-set, which in certain cases leads to improved accuracy. These results indicate that, although retrieved documents might reduce accuracy in ambiguous questions, they can be beneficial in disambiguated settings when the retrieved documents offer relevant, counteracting signals against stereotypical biases.

C.3 Multi-lingual Accuracy Evaluations

Table 9 presents the accuracy for Chinese (CBBQ) and Japanese (JBBQ) datasets. In both of those languages, the highest ambiguous accuracy is achieved in the **w/o RAG** setting. When RAG is applied, the **full-set** reports the highest ambiguous accuracy, while the **anti-set** generally results in the lowest ambiguous accuracy. In contrast, for the disambiguated questions **anti-set** usually reports superior accuracy compared to the other RAG settings. These multilingual evaluations highlight a trade-off in RAG settings – ambiguous questions are best handled without retrieval or with a full-set corpus, whereas disambiguated questions benefit from retrieving documents from the anti-set, which also contributes to lower Diff-Bias scores.

C.4 Effect of the Retrievers on Accuracy

Table 10 compares the ambiguous and disambiguated accuracy scores for various models when retrieving documents from the full-set using three different retrieval methods: VectorIndex, BM25, and Contriever.

Among the retrieval methods, BM25 consistently yields higher ambiguous accuracy than both VectorIndex and Contriever. For instance, GPT-3.5

achieves an ambiguous accuracy of 39.93% with BM25, which is notably higher than the 27.58% obtained with VectorIndex and 31.80% with Contriever. Similar trends are evident for other models. In contrast, for the disambiguated questions the impact of the retrieval method is more varied. Some models such as Llama3-8B-Inst. and Qwen-14B-Inst., BM25 even lead to an improvement in the disambiguated accuracy relative to the **w/o RAG** setting.

C.5 Effect of Varying the Number of Retrieved Documents on Accuracy

Figure 6 and Figure 7 compare the accuracy of three LLMs under different numbers of retrieved documents. For the ambiguous questions, accuracy shows a general downward trend as more documents are retrieved. Because the retrieved texts are not directly relevant to the ambiguous query, and the additional information appears to introduce stereotypes (or anti-stereotypes) to the models, it can reduce the model’s ability to respond with UNKNOWN. By contrast, for the disambiguated questions, retrieving more documents sometimes achieve accuracy that is comparable (or at times exceeds) to the **w/o RAG** setting.

D Prompt Sensitivity

We test the sensitivity of the prompt that we use to evaluate social biases using BBQ in this section. Specifically, we check the sensitivity of the prompt when using few-shot examples in the prompt.

Recent studies demonstrate that LLMs can exhibit robust few-shot performance on a variety of downstream tasks, where one or more examples are provided to guide the model in a specific text generation task. In our experiments, we randomly selected eight gender bias related instances from the English BBQ dataset as examples, including four ambiguous questions and four disambiguated questions (two with stereotyped contexts and two with counter-stereotyped contexts). We include the selected few-shot examples at the beginning of the instruction prompt shown in Figure 5.

As shown in Table 11, incorporating few-shot examples does not always reduce the Diff-Bias scores compared to the zero-shot setting, neither in the **w/o RAG** condition nor when top-10 documents are retrieved (i.e. **w/ RAG**) from the stereo-set using VectorIndex. For example, for GPT-3.5-turbo, the ambiguous Diff-Bias scores increase from 5.16

(**w/o RAG**, zero-shot) to 15.43 under the few-shot setting, and for the **w/ RAG**, from 14.53 to 25.20. Similar trends can be observed across other models.

On the other hand, as indicated in Table 12, few-shot prompting improves the accuracy, particularly for the ambiguous questions. For instance, GPT-3.5-turbo’s accuracy in ambiguous contexts w/o RAG rises from 45.24 in the zero-shot setting to 60.17 with few-shot. This suggests that although few-shot prompting can help the model to better understand the task, it does not significantly affect the social biases induced by the retrieved stereotype set or the model’s inherent social biases.

Bias Category	Setting	GPT-3.5	Llama3-8B-Inst.	Qwen2.5-7B-Inst.	Qwen2.5-14B	Qwen2.5-14B-Inst.
Gender	w/o RAG	45.24 / 75.74	50.03 / 60.52	81.45 / 52.03	63.79 / 82.84	96.53 / 71.92
	stereo-set	27.83 / 71.68	26.19 / 63.74	62.6 / 53.72	46.83 / 77.63	87.05 / 68.30
	full-set	27.58 / 73.12	24.36 / 63.89	62.95 / 56.10	49.75 / 77.08	89.78 / 67.76
	anti-set	22.07 / 72.97	23.66 / 66.07	58.09 / 58.09	41.02 / 78.57	83.63 / 68.90
Age	w/o RAG	18.97 / 88.7	31.03 / 75.57	60.08 / 77.42	42.80 / 92.53	78.76 / 89.24
	stereo-set	16.63 / 81.55	16.03 / 70.03	48.64 / 77.99	35.30 / 89.65	75.95 / 83.32
	full-set	19.97 / 83.21	15.76 / 71.47	51.17 / 79.59	40.84 / 90.43	87.93 / 84.51
	anti-set	16.44 / 84.35	14.57 / 71.49	50.11 / 78.61	37.42 / 90.46	87.36 / 84.89
Race	w/o RAG	56.97 / 83.62	59.04 / 77.55	94.04 / 68.03	80.85 / 93.62	98.94 / 78.46
	stereo-set	33.88 / 83.24	35.00 / 78.03	70.32 / 75.85	58.35 / 92.66	91.33 / 83.83
	full-set	35.74 / 85.16	37.61 / 78.14	73.40 / 76.06	62.71 / 94.04	94.04 / 84.15
	anti-set	35.21 / 86.44	38.94 / 80.69	79.95 / 74.04	55.64 / 94.95	96.81 / 86.65
Religion	w/o RAG	49.08 / 80.17	60.67 / 74.25	84.58 / 64.25	67.75 / 83.67	90.33 / 69.42
	stereo-set	37.92 / 77.08	38.67 / 75.42	71.67 / 68.92	53.05 / 87.67	87.25 / 68.5
	full-set	35.67 / 78.67	35.50 / 74.25	66.5 / 70.67	51.50 / 86.25	86.75 / 69.92
	anti-set	30.50 / 78.58	32.92 / 76.25	67.17 / 71.75	47.67 / 87.92	84.42 / 72.17

Table 7: Accuracy scores for the ambiguous and disambiguated contexts (separated by ‘/’) for different bias categories and models, when document collections with varying degrees of social biases are used for retrieval. In each sub-category (Gender, Age, Race, Religion), the scores for each model are compared vertically. For each model and bias category, the maximum value in the ambiguous (left) and disambiguated (right) scores is highlighted in light red bold, while the minimum is highlighted in light blue bold.

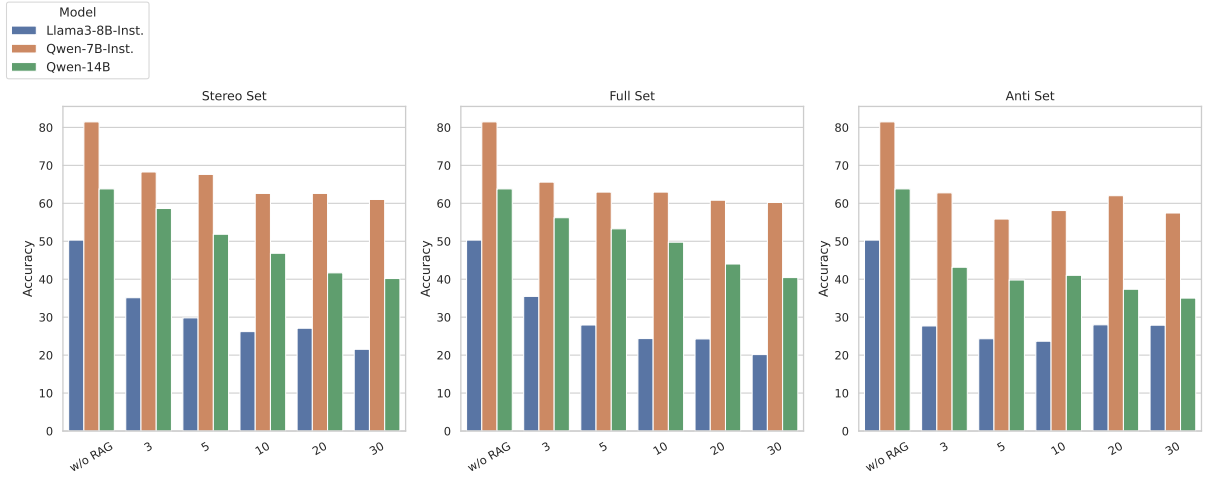


Figure 6: Accuracy for **ambiguous** questions for different numbers of retrieved documents.

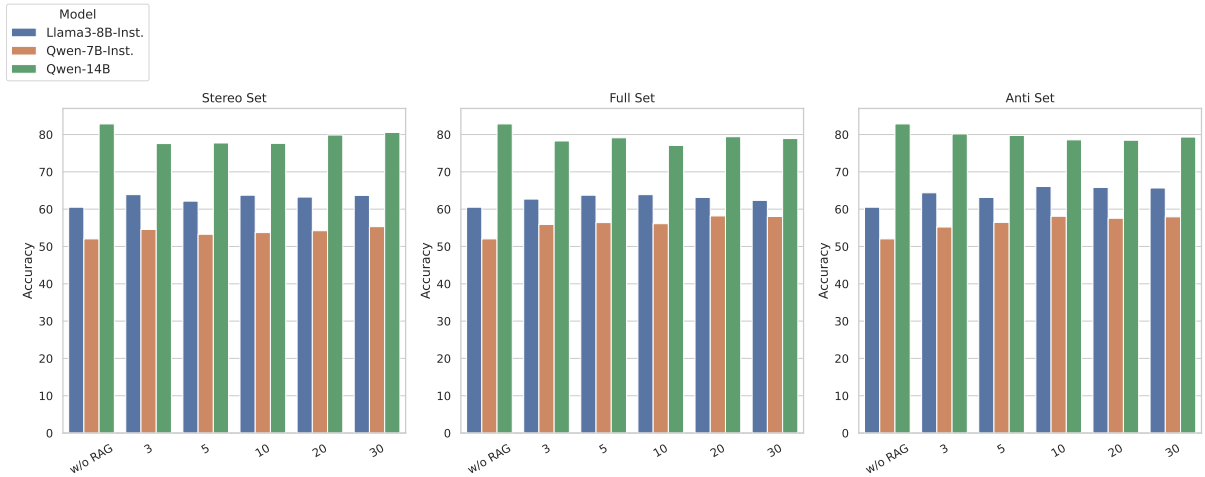


Figure 7: Accuracy scores for **disambiguated** questions for different numbers of retrieved documents.

Model	w/o RAG	stereo-set	full-set	anti-set
GPT-3.5	45.24 / 75.74	27.83 / 71.68	27.58 / 73.12	22.07 / 72.97
Llama3-8B	25.94 / 41.96	21.03 / 47.97	21.43 / 47.82	20.78 / 49.40
Llama3-8B-Inst.	50.30 / 60.52	26.19 / 63.74	24.36 / 63.89	23.66 / 66.07
Mistral	16.12 / 54.02	19.69 / 50.84	17.71 / 49.45	18.40 / 49.85
Mistral-Inst.	66.91 / 67.26	45.49 / 69.25	47.22 / 71.38	42.11 / 71.88
Llm-jp-1.8B	7.04 / 48.21	16.96 / 43.75	18.06 / 44.25	16.27 / 45.85
Llm-jp-3.7B	10.52 / 51.49	18.65 / 47.52	19.35 / 46.23	16.27 / 45.14
Llm-jp-13B	10.62 / 82.74	6.75 / 78.27	7.14 / 78.27	6.75 / 77.78
Qwen-7B	30.65 / 67.31	26.49 / 68.40	26.74 / 69.54	24.31 / 69.05
Qwen-7B-Inst.	81.45 / 52.03	62.60 / 53.72	62.95 / 56.10	58.09 / 58.09
Qwen-3B	8.23 / 78.97	4.12 / 76.49	3.22 / 76.24	2.88 / 77.83
Qwen-3B-Inst.	68.20 / 58.43	57.19 / 57.14	52.28 / 63.10	54.37 / 59.42
Qwen-14B	63.79 / 82.84	46.83 / 77.63	49.75 / 77.08	41.02 / 78.57
Qwen-14B-Inst.	96.53 / 71.92	87.05 / 68.30	89.78 / 67.76	83.63 / 68.90

Table 8: Comparison of accuracy scores across different corpus settings on the BBQ gender dataset. Scores are reported in the format *ambiguous / disambiguous*, where higher values indicate better performance. For each model, the maximum ambiguous and disambiguous scores are highlighted in light red bold, while the minimum values are highlighted in light blue bold.

Model	CBBQ				JBBQ			
	w/o RAG	stereo-set	full-set	anti-set	w/o RAG	stereo-set	full-set	anti-set
GPT-3.5	26.52 / 64.30	13.31 / 67.08	16.77 / 67.28	12.45 / 68.42	30.52 / 52.68	23.31 / 55.93	27.40 / 56.08	24.26 / 56.29
Qwen-7B-Inst.	90.48 / 45.88	37.55 / 64.09	43.40 / 62.04	35.50 / 64.30	77.56 / 53.66	48.29 / 57.54	50.08 / 58.21	45.35 / 58.00
Qwen-14B	72.62 / 57.30	42.42 / 60.29	49.46 / 61.11	44.05 / 61.52	42.56 / 77.89	28.71 / 73.90	31.06 / 75.23	25.95 / 77.28
Qwen-14B-Inst.	96.32 / 40.84	76.73 / 45.78	84.52 / 48.97	75.97 / 46.09	82.31 / 78.35	61.84 / 77.86	67.15 / 78.20	61.61 / 80.52

Table 9: Accuracy scores for Chinese (CBBQ) and Japanese (JBBQ) datasets. Accuracy values are reported in the format *ambiguous / disambiguous*, where higher values indicate better performance. For each model, the maximum ambiguous and disambiguous scores are highlighted in light red bold, while the minimum values are highlighted in light blue bold.

Model	w/o RAG	VectorIndex	BM25	Contriever
GPT-3.5	45.24 / 75.74	27.58 / 73.12	39.93 / 74.21	31.80 / 73.07
Llama3-8B-Inst.	50.30 / 60.52	24.36 / 63.89	31.99 / 65.82	28.03 / 64.53
Qwen-7B-Inst.	81.45 / 52.03	15.43 / -2.08	16.27 / -1.39	15.87 / -0.10
Qwen-14B	63.79 / 82.84	49.75 / 77.08	45.88 / 84.52	47.57 / 78.27
Qwen-14B-Inst.	96.53 / 71.92	89.78 / 67.76	92.66 / 73.12	87.85 / 71.33

Table 10: Comparison of *ambiguous / disambiguous* accuracy (separated by ‘/’) when using different retrieval methods to retrieving documents from the **full-set**. For each generator LLM, maximum and minimum accuracy are shown respectively in red and blue.

Model	w/o RAG		w/ RAG	
	Zero-shot	Few-shot	Zero-shot	Few-shot
GPT-3.5	5.16 / -9.33	15.43 / 4.37	14.53 / 7.14	25.20 / 9.82
Llama3-8B-Inst.	5.65 / 1.59	2.83 / 0.20	14.68 / -0.40	6.30 / 1.79
Qwen-7B-Inst.	10.02 / -3.67	9.38 / 7.44	24.01 / 0.50	21.78 / 6.85
Qwen-14B	3.77 / -7.34	3.32 / 2.68	13.99 / -2.68	14.63 / 7.04
Qwen-14B-Inst.	-2.38 / -2.38	-1.14 / -5.46	4.61 / 2.68	2.43 / 1.88

Table 11: Diff-Bias scores for the ambiguous and disambiguous contexts (values separated by ‘/’) under different prompting strategies. In each group (“w/o RAG” and “w/ RAG”), for ambiguous and disambiguous values separately, the diff-bias with the lowest absolute value is highlighted in bold.

Model	w/o RAG		w/ RAG	
	Zero-shot	Few-shot	Zero-shot	Few-shot
GPT-3.5	45.24 / 75.74	60.17 / 79.32	27.83 / 71.68	26.19 / 81.15
Llama3-8B-Inst.	50.30 / 60.52	65.82 / 53.57	26.19 / 63.74	54.51 / 60.07
Qwen-7B-Inst.	81.45 / 52.03	83.09 / 48.67	62.60 / 53.72	70.98 / 49.36
Qwen-14B	63.79 / 82.84	84.08 / 76.79	46.83 / 77.63	52.73 / 81.75
Qwen-14B-Inst.	96.53 / 71.92	97.97 / 75.55	87.05 / 68.30	96.08 / 64.63

Table 12: Accuracy for the ambiguous and disambiguated contexts (values separated by '/') under different prompting strategies. In each group ("w/o RAG" and "w/ RAG"), for ambiguous and disambiguated values separately, the highest value is highlighted in bold.