

Multilingual Information Retrieval with a Monolingual Knowledge Base

Yingying Zhuang
yyzhuang@amazon.com
Amazon.com, Inc.
San Francisco, CA, USA

Aman Gupta
amangta@amazon.com
Amazon.com, Inc.
San Francisco, CA, USA

Anurag Beniwal
beanurag@amazon.com
Amazon.com, Inc.
San Francisco, CA, USA

Abstract

Multilingual information retrieval has emerged as powerful tools for expanding knowledge sharing across languages. On the other hand, resources on high quality knowledge base are often scarce and in limited languages, therefore an effective embedding model to transform sentences from different languages into a feature vector space same as the knowledge base language becomes the key ingredient for cross language knowledge sharing, especially to transfer knowledge available in high-resource languages to low-resource ones. In this paper we propose a novel strategy to fine-tune multilingual embedding models with weighted sampling for contrastive learning, enabling multilingual information retrieval with a monolingual knowledge base. We demonstrate that the weighted sampling strategy produces performance gains compared to standard ones by up to 31.03% in MRR and up to 33.98% in Recall@3. Additionally, our proposed methodology is language agnostic and applicable for both multilingual and code switching use cases.

CCS Concepts

• Information systems → Query representation.

Keywords

Information Retrieval, Text Embedding, Contrastive Learning

ACM Reference Format:

Yingying Zhuang, Aman Gupta, and Anurag Beniwal. 2025. Multilingual Information Retrieval with a Monolingual Knowledge Base. In *Proceedings of July 17, 2025 (SIGIR 2025)*. ACM, New York, NY, USA, 6 pages.

1 Introduction

The task of Information Retrieval is to find relevant information from a large collection and a knowledge base often complements this retrieval task and facilitates various downstream tasks as well. With the recent emergence of Large Language Models (LLMs) as a powerful tool to build dialogue systems [21, 25, 26] and conversation AIs, the retrieved information can play a critical role in increasing practicality and reliability of these systems. Application of LLMs has also achieved remarkable advancements in multilingual scenarios and many of the current frontier LLMs, such as Anthropic’s Claude 3 [1] and Cohere’s Command A [3] are focused to excel

in the multilingual settings to support global businesses. On the other hand, because knowledge base construction is labor-intensive and expensive [14], imbalanced data, quality and scarcity issues become the bottleneck for rapid knowledge base development, especially for low-resource languages. Additionally there has been a surge of interest in code-switching recently where communicators switch between two or more languages during linguistic interactions [7, 20]. This communication style is common in bilingual and multilingual communities [9, 23]. For example, in US mixing Spanish and English is a common phenomenon while in Canada mixing French and English has been often observed. However, it remains a challenging task in Information Retrieval as the language IDs are not specified and knowledge base with code-switching context is scarce. Existing work has been mainly focused on automatic multilingual knowledge base construction [10, 24] where knowledge bases of low-resource languages are automatically propagated based on knowledge bases in well-populated high resource languages.

Contributions. Different from previous approaches, we propose a novel embedding model development pipeline to align multilingual information retrieval with a **monolingual** knowledge base, therefore removing the bottleneck of multilingual knowledge base construction in information retrieval systems. Specifically, we design a weighted sampling strategy to select positive and negative pairs for contrastive learning, which guides the model to embed queries with similar semantic meanings into close embedding vector space across different languages. Our approach aims at a more data-efficient multilingual information retrieval system which only requires a high resource language (e.g., English or French) for the knowledge base. Experimental results demonstrate that our approach enables an embedding model to properly handle multilingual queries. The ability to share knowledge in multiple languages is a fundamental component to ensure the development of Conversation AIs systems and to enable both businesses and individuals to reach a broader range of communities, fostering inclusivity, accessibility and collaboration among international teams.

2 Methodology

Embedding models such as E5 [17], GTE [11], Nomic [13], and Arctic-Embed [12] are typically trained with three stages: pretraining via masked language modeling, large scale contrastive pretraining with in-batch negatives, and finally quality focused contrastive fine-tuning [2, 19, 22]. Our work focuses on the contrastive fine-tuning stage and the proposed methodology can be applied on any open-source text embedding models. One key ingredient for effective embedding is supervised fine-tuning with a high quality dataset which contains explicitly labeled positive and negative pairs.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

SIGIR 2025,

© 2025 Copyright held by the owner/author(s). Publication rights licensed to ACM.

Specifically, we focus on the query-to-query multilingual retrieval problem with a monolingual knowledge base.

Given a user query q in a target language T , the goal is to retrieve a similar example query q_{EX} from the knowledge base in language E where E is different from T . The knowledge base contains a set of example queries and their corresponding labels

$$KB_E = (q_{EX_1}, l_1), (q_{EX_2}, l_2), \dots, (q_{EX_N}, l_N).$$

For example, a key component of task-oriented dialogue systems is to identify the underlying purpose or intent of a user in a conversation [6, 21]. To facilitate information retrieval business can build a monolingual intent knowledge base in English which contains a set of customer query and customer intent pairs based on high quality data source or human annotation. As business expands to other languages, it becomes infeasible to build a separate knowledge base in each language therefore the goal is to share the English knowledge base across all languages. If a user input query is in Spanish (the target language T), the goal is to retrieve a semantically similar English query in the knowledge base, q_{EX_j} , and the corresponding intent label l_j . In order to build a multilingual Information Retrieval system with an English-only (or any other monolingual) knowledge base, it is crucial for the embedding model to map English and non-English queries into a shared representation space to enable non-English retrieval.

Contrastive Training Data. In contrastive learning, the model learns by distinguishing between similar and dissimilar queries from positive and negative pairs. An effective embedding model will project similar (positive) pairs closer together and dissimilar (negative) pairs further apart in a embedding vector space. For positive pairs, existing work often uses a noisy data source to identify relevant queries, then applies some heuristic and consistency quality filters to improve data quality [4, 8, 13, 18]. For negative pairs, [15] demonstrated the importance of training on hard negatives, where “hard” refers to the fact that it is not trivial to determine their lower relevance relative to the positive examples, therefore, many existing work has been following this paradigm to identify the hardest negatives for each training example. For example, [12] leveraged a pre-existing text embedding model to identify and score the hardest negatives for each training example while NV Retriever [5] also added a filtering step where any negative with a relevance score exceeding a specified percentage of the known-positive’s score is discarded as a potential false negative.

Contrastive Training Data Generation. In our methodology, we propose a new strategy to leverage the available labels, $l_1, l_2, \dots, l_n$, in the monolingual knowledge base to construct positive and negative multilingual pairs. Our approach demonstrates the benefits of a weighted sampling strategy of negative mining based on relevance, which in Section 3 we show that it is more effective than focusing purely on selecting the “hardest” negatives possible, or a rank for relevance scores. Algorithm 1 shows the details of our proposed approach step-by-step and Figure 1 is an illustration with examples.

3 Experiments and Results

In this paper, we conduct experiments to evaluate the plausibility of our proposed methods for the business use case of performing

Algorithm 1 Contrastive Training Data Generation

• **Require:**

A knowledge base in language E with N example query and label pairs

$$KB_E = (q_{EX_1}, l_1), (q_{EX_2}, l_2), \dots, (q_{EX_N}, l_N).$$

A set of M unlabelled queries q ’s in a target language T ($q_1, q_2, \dots, q_M$).

• **Data Generation**

◦ **Step 1: Initialize** an empty paired dataset $P = \{\}$

◦ **Step 2: Split** KB_E into index and training set. We split the N example query and label pairs into two subsets N_1 and N_2 , reserve the N_1 pairs for retrieval index, while using the remaining N_2 for synthetic data generation.

◦ **Step 3: Generate Positive Pairs:** For each query q_i in N_2 which is in language E , use an LLM to translate it to the target language T into q_i^T . Randomly sample one query from all the queries sharing the same label in the Index set N_1 , q_i' . Update P with the positive pair (q_i^T, q_i') , i.e. $P = P \cup (q_i^T, q_i')$.

◦ **Step 4: Generate Negative Pairs based on weighted sampling:** For each query q_i in N_2 , use an LLM to translate it to the target language T into q_i^T (or reuse q_i^T from Step 3). For each query q_j in the index set N_1 , calculate a similarity score s_{ij} between q_i ’s label l_i and q_j ’s label l_j . We generate two types of negative pairs:

• **Random negative:** Randomly sample one query q_i'' from all the queries in set N_1 with a different label than l_i .

• **Hard negatives:** Randomly sample k queries from the index set N_1 based on sampling weight s_{ij} . The higher the similarity score s_{ij} , the more likely a query is to be sampled. The intuition is that queries with similar labels are harder to distinguish, making them better candidates for hard negatives.

With $k = 2$, we yield three negative pairs in total: one random negative pair (q_i^T, q_i'') and two hard negative pairs (q_i^T, q_i''') and (q_i^T, q_i''''') . This maintains our desired positive to negative ratio of 1:3 when using contrastive loss. Update P with these negative pairs, i.e., $P = P \cup \{(q_i^T, q_i''), (q_i^T, q_i'''), (q_i^T, q_i''''')\}$.

◦ **Step 5: Synthetic Data Augmentation:** Compared to the abundance of unlabeled data, high-quality labeled examples in KB_E are more scarce. To address this data scarcity issue, we also use synthetic data creation to construct additional positive and negative pairs. For any query in the target language T , use a LLM to generate a semantically similar query to form a positive pair, and 3 semantically different queries in language E to form 3 negative pairs. Update P with the these pairs.

Information Retrieval for Hinglish queries using an English knowledge base. Hinglish refers to the conversation style where speakers are mixing Hindi and English in one conversation interaction and often within one single sentence, which is a very common way of communication for the India marketplace.

Our proposed methodology is language agnostic and can be applied to any language with low-resource constraints. Therefore, we simply treat code-switching as a new language. In our experiments we fine-tune an embedding model optimized to embed English and

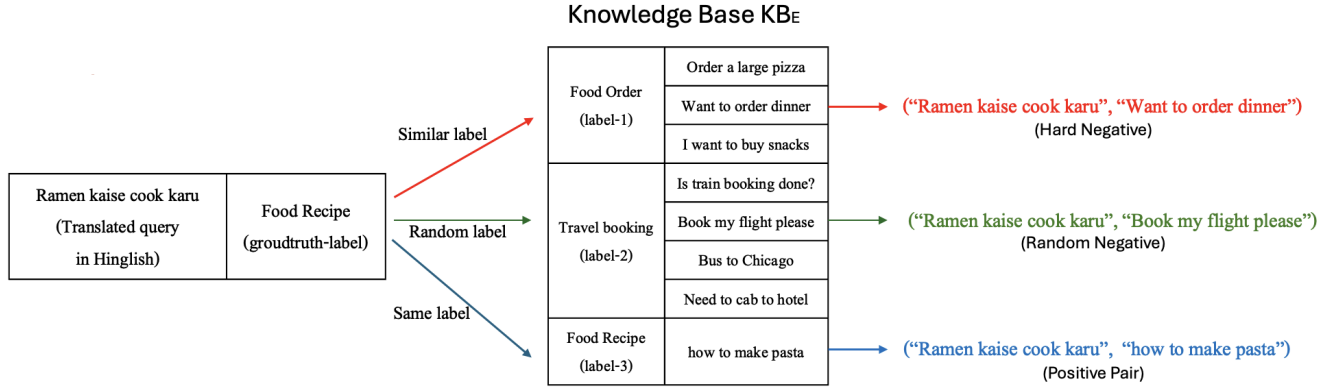


Figure 1: Illustration of our contrastive training data generation algorithm. For each translated query (left), we create a positive pair (blue) with queries sharing the same label, k hard negative pairs (red) from queries with similar labels selected through weighted sampling, and 1 random negative pair (green) from queries with different labels selected randomly from the knowledge base.

Hinglish queries into the same vector space for Information Retrieval.

3.1 Datasets

Data Anonymization. Due to business considerations, we are not permitted to share the results using the original customer queries. As a result, we manually anonymized both the labels and transcripts to ensure no personal information is included. Additionally, specific product and service names were denonymized to prevent the identification of the company from the transcript or label descriptions. Despite these modifications, the conclusions drawn from our experiments remain valid.

We collect about 35,000 dialogue transcripts with customer intent labels in English to construct our KB_E . Following the steps in Algorithm 1 we split these transcripts into two subsets: 3000 to construct the retrieval index and the remaining 32,000 to generate contrastive training data. For Step 5, we collect 10,000 unlabelled dialogue transcripts in Hinglish for data augmentation.

3.2 Experiments

As discussed earlier, we proposed Algorithm 1, which employs a hybrid approach combining both hard and random negative mining, along with synthetic data generation. This mixed strategy helps reduce excessive reliance on translated queries alone, with synthetic data generation performed directly on Hinglish data available natively.

In our experimental setup, we first evaluated several open-source multilingual retrieval models as the foundation for our fine-tuning process:

- paraphrase-multilingual-mpnet-base-v2
- multilingual-e5-base
- paraphrase-multilingual-MiniLM-L12-v2
- stsb-xlm-r-multilingual

To validate the effectiveness of our proposed approach and to better understand the contribution of different components, we

conducted ablation studies comparing Algorithm 1 against several alternative strategies:

Negative Sampling Variations.

- **Random Negative Mining:** This represents the simplest approach where we set $k = 0$ in Step 4 of Algorithm 1, selecting all negatives randomly with equal probability regardless of semantic similarity.
- **Hard Negative Mining:** In this variation, we set $k = 3$ in Step 4, selecting all negatives using weighted sampling based on similarity scores. This targets examples that are more challenging to distinguish.
- **Hardest Negative Mining:** Taking the hard negative approach further, we not only set $k = 3$ but also restrict sampling to only the top-3 most similar labels (excluding exact matches). This focuses exclusively on the most challenging negative examples.

Data Source Variations. To evaluate the impact of synthetic data augmentation described in Step 5, we also tested:

- **Labeled Data Only:** This variant trains solely on the high-quality translated examples from the original English dataset, omitting the synthetic data generation Step 5.
- **Synthetic Data Only:** Conversely, this approach relies exclusively on synthetically generated data in Step 5, representing the scenario where no labeled data is available in any language.

These methodical comparisons allowed us to isolate the contribution of each component in our algorithmic design and confirm the superiority of our hybrid approach.

3.3 Implementation Details

We fine-tuned various pretrained multilingual embedding models using the InfoNCE contrastive loss [16] with a maximum sequence length of 512 tokens. Training was performed on 4 NVIDIA A10G GPUs using PyTorch’s DistributedDataParallel framework with a

batch size of 32. Each model was trained for 15 epochs, with early convergence typically observed around epochs 8-9. For optimization, we used AdamW with a learning rate of $2e-5$ and a linear warmup period. Claude Sonnet 3.5¹ was employed for both query translation and generation of synthetic positive and negative pairs for training.

3.4 Evaluation Metrics

For evaluation, we compute Recall@1, Recall@3, Recall@5, and Recall@10 metrics, along with Mean Reciprocal Rank (MRR). These metrics are particularly appropriate for our task since we have a single ground truth label for each query, while our retrieval system returns ranked predictions. Recall@k measures whether the correct label appears within the top-k retrieved results, providing a clear assessment of our model’s retrieval performance at various thresholds of interest.

3.5 Results and Discussion

Figure 2 shows the performance comparison of different multilingual embedding models during training. The *multilingual-e5-base* model consistently outperforms other models across all epochs, reaching the highest Recall@1 of approximately 0.54 by the final epoch. The *paraphrase-multilingual-mpnet-base-v2* model performs comparably well, while *stsb-xtlm-r-multilingual* and *paraphrase-multilingual-MiniLM-L12-v2* show significantly lower performance. Based on these results, we selected *multilingual-e5-base* as our foundation model for all subsequent experiments. Appendix Section A presents the comparison plots between these models on Recall@3, Recall@5, Recall@10, and MRR, which show a similar pattern as Recall@1.

The evaluation results for different negative sampling and data strategies are presented in Table 1. Our ablation study reveals that combining random and hard negatives yields better performance than approaches using only one type of negative examples. While the Labeled Data Only approach performs comparably to our proposed method, we attribute this to the composition of our test set, which primarily contains samples similar to the labeled data. Nevertheless, Algorithm 1’s hybrid approach provides better protection against performance degradation on original queries in the low-resource language. The significantly poorer performance of the purely synthetic data approach indicates the crucial importance of high-quality labeled examples in the training process.

Why does mixed-mining strategy work better? The superior performance of Algorithm 1 can be attributed to its hybrid approach. By combining both random and hard negatives, the model optimizes the embedding space globally while simultaneously emphasizing differentiation in local neighborhoods. Pure hard negative mining (0.4613) or hardest negative mining (0.4012) approaches underperform because they overly focus on disambiguating local neighborhoods without maintaining global embedding structure. Additionally, the incorporation of synthetic data in Algorithm 1 helps improve performance on low-resource language queries, reducing overfitting to the translated query distribution while maintaining strong performance on the original dataset.

¹<https://www.anthropic.com/news/claude-3-5-sonnet>

Table 1: Performance comparison of different negative sampling and data strategies

Method	Top-1	Top-3	Top-10	MRR
Random Negative Mining	0.5308	0.7271	0.8794	0.6520
Hard Negative Mining	0.4613	0.6829	0.8535	0.5982
Hardest Negative Mining	0.4012	0.6678	0.8499	0.5610
Labeled Data Only	0.5474	0.7356	0.8803	0.6639
Synthetic Data Only	0.2191	0.4012	0.6444	0.3550
Using Algorithm 1	0.5450	0.7410	0.8842	0.6653

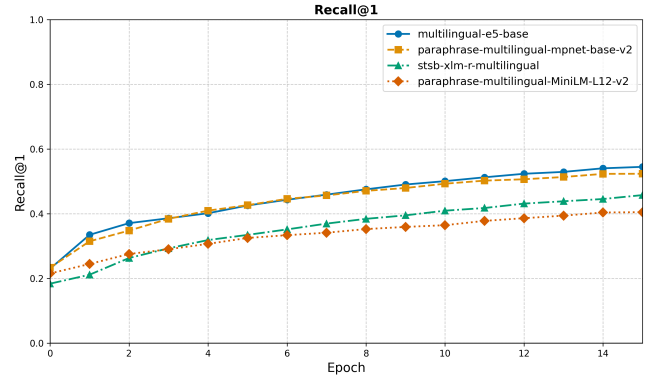


Figure 2: Performance of Recall@1 at different points during contrastive fine-tuning.

4 Conclusion and Future Work

Our proposed embedding model development pipeline to align multilingual information retrieval with a monolingual knowledge base removes the bottleneck of multilingual knowledge base construction in information retrieval systems. The ability to share knowledge across languages facilitates faster and easier global expansion of conversation AI systems, specifically to lower resource languages or code-switching use cases. Our comprehensive experiment results in a code-switching use case demonstrate the effectiveness and robustness of our proposed weighted sampling strategy for contrastive learning. Our framework is language agnostic therefore applicable for any target language and we hope our paper can help push forward the research in the multilingual information retrieval communities and facilitate faster global expansion for businesses.

In the future, we seek to continue our experimentation to leverage monolingual knowledge base for multilingual dialogue systems. Additionally, we hope to explore more robust embedding models with compression approaches such as binarization or quantization to further improve performance.

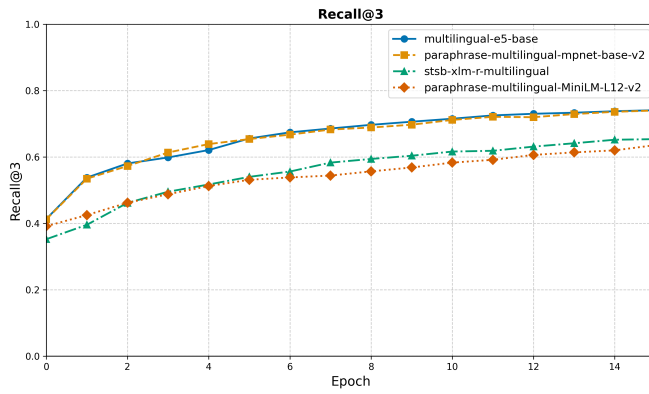
References

- [1] [n. d.]. Introducing the next generation of Claude. <https://www.anthropic.com/news/claude-3-family>. Accessed: 2025-04-20.
- [2] Jianlv Chen, Shitao Xiao, Peitian Zhang, Kun Luo, Defu Lian, and Zheng Liu. 2024. BGE M3-Embedding: Multi-Lingual, Multi-Functionality, Multi-Granularity Text Embeddings Through Self-Knowledge Distillation. arXiv:2402.03216 [cs.CL] <https://arxiv.org/abs/2402.03216>

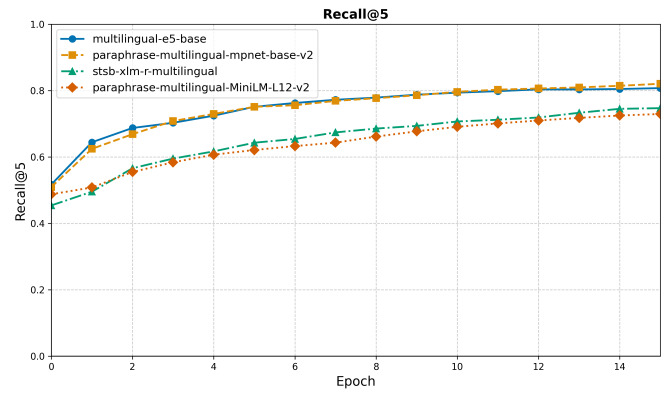
- [3] Team Cohere, :, Aakanksha, Arash Ahmadian, Marwan Ahmed, Jay Alammam, Milad Alizadeh, Yazeed Alnumay, Sophia Althammer, Arkady Arkhangorodsky, Viraat Aryabumi, Dennis Aumiller, Raphaël Avalos, Zahara Aviv, Sammie Bae, Saurabh Bajj, Alexandre Barbet, Max Bartolo, Björn Bebensee, Neeral Beladia, Walter Beller-Morales, Alexandre Bérard, Andrew Berneshaw, Anna Bialas, Phil Blunsom, Matt Bobkin, Adi Bongale, Sam Braun, Maxime Brunet, Samuel Cahyawijaya, David Cairuz, Jon Ander Campos, Cassie Cao, Kris Cao, Roman Castagné, Julián Cendrero, Leila Chan Currie, Yash Chandak, Diane Chang, Giannis Chatziveroglou, Hongyu Chen, Claire Cheng, Alexis Chevalier, Justin T. Chiu, Eugene Cho, Eugene Choi, Eujeong Choi, Tim Chung, Volkan Cirik, Ana Cismaru, Pierre Clavier, Henry Conklin, Lucas Crawhall-Stein, Devon Crouse, Andres Felipe Cruz-Salinas, Ben Cyrus, Daniel D'souza, Hugo Dalla-Torre, John Dang, William Darling, Omar Darwiche Domingues, Saurabh Dash, Antoine Debugne, Théo Dehaze, Shaan Desai, Joan Devassy, Rishit Dholakia, Kyle Duffy, Ali Edalati, Ace Eldeib, Abdullah Elkady, Sarah Elsharkawy, Irem Ergün, Beyza Ermiş, Marzieh Fadaee, Boyu Fan, Lucas Fayoux, Yannis Flet-Berliac, Nick Frost, Matthias Gallé, Wojciech Galuba, Utsav Garg, Matthieu Geist, Mohammad Gheshlaghi Azar, Ellen Gilsenan-McMahon, Seraphina Goldfarb-Tarrant, Tomas Goldsack, Aidan Gomez, Victor Machado Gonzaga, Nithya Govindarajan, Manoj Govindassamy, Nathan Grinsztajn, Nikolas Gritsch, Patrick Gu, Shangmin Guo, Kilian Haefeli, Rod Hajjar, Tim Hawes, Jingyi He, Sebastian Hofstätter, Sungjin Hong, Sara Hooker, Tom Hosking, Stephanie Howe, Eric Hu, Renjie Huang, Hemant Jain, Ritika Jain, Nick Jakobi, Madeline Jenkins, JJ Jordan, Dhruvi Joshi, Jason Jung, Trushant Kalyanpur, Siddhartha Rao Kamalakara, Julia Kedrzycki, Gokce Keskin, Edward Kim, Joon Kim, Wei-Yin Ko, Tom Kocmi, Michael Kozakov, Wojciech Kryściński, Arnav Kumar Jain, Komal Kumar Teru, Sander Land, Michael Lasby, Olivia Lasche, Justin Lee, Patrick Lewis, Jeffrey Li, Jonathan Li, Hangyu Lin, Acyr Locatelli, Kevin Luong, Raymond Ma, Lukáš Mach, Marina Machado, Joanne Magbitang, Brenda Malacara Lopez, Aryan Mann, Kelly Marchisio, Olivia Markham, Alexandre Matton, Alex McKinney, Dominic McLoughlin, Jozef Mokry, Adrien Morisot, Autumn Moulder, Harry Moynihan, Maximilian Mozes, Vivek Muppalla, Lidiya Murakhovska, Hemangani Nagarajan, Alekhyia Nandula, Hisham Nasir, Shauna Nehra, Josh Ndetto-Rosen, Daniel Ohashi, James Owers-Bardsley, Jason Ozuzu, Dennis Padilla, Gloria Park, Sam Passaglia, Jeremy Pekmez, Laura Penstone, Aleksandra Piktus, Case Ploeg, Andrew Poulton, Youran Qi, Shubha Raghvendra, Miguel Ramos, Ekagra Ranjan, Pierre Richemond, Cécile Robert-Michon, Aurélien Rodriguez, Sudip Roy, Sebastian Ruder, Laura Ruis, Louise Rust, Anubhav Sachan, Alejandro Salamanca, Kailash Karthik Saravanakumar, Isha Satyakam, Alice Schoenauer Sebag, Priyanka Sen, Sholeh Sepehri, Preethi Seshadri, Ye Shen, Tom Sherborne, Sylvie Shang Shi, Sanal Shivaprasad, Vladyslav Shmyhlo, Anirudh Shrinivasan, Inna Shteinbuk, Amir Shukayev, Mathieu Simard, Ella Snyder, Ava Spataru, Victoria Spooner, Trisha Starostina, Florian Strub, Yixuan Su, Jimin Sun, Dwarak Talupuru, Eugene Tarasov, Elena Tommasone, Jennifer Tracey, Billy Trend, Evren Tümer, Ahmet Üstün, Bharat Venkitesh, David Venuto, Pat Verga, Maxime Voisin, Alex Wang, Donglu Wang, Shijian Wang, Edmond Wen, Naomi White, Jesse Willman, Marysia Winkels, Chen Xia, Jessica Xi, Minjie Xu, Bowen Yang, Tan Yi-Chern, Ivan Zhang, Zhenyu Zhao, and Zhoujie Zhao. 2025. Command A: An Enterprise-Ready Large Language Model. arXiv:2504.00698 [cs.CL] <https://arxiv.org/abs/2504.00698>
- [4] Zhuyn Dai, Vincent Y. Zhao, Ji Ma, Yi Luan, Jianmo Ni, Jing Lu, Anton Bakalov, Kelvin Guu, Keith B. Hall, and Ming-Wei Chang. 2022. Promptagator: Few-shot Dense Retrieval From 8 Examples. arXiv:2209.11755 [cs.CL] <https://arxiv.org/abs/2209.11755>
- [5] Gabriel de Souza P. Moreira, Radek Osmulski, Mengyao Xu, Ronay Ak, Benedikt Schifferer, and Even Oldridge. 2025. NV-Retriever: Improving text embedding models with effective hard-negative mining. arXiv:2407.15831 [cs.IR] <https://arxiv.org/abs/2407.15831>
- [6] Liesbeth Degand and Philippe Muller. 2020. Introduction to the Special Issue on Dialogue and Dialogue Systems. *Traitement Automatique des Langues* 61, 3 (2020), 7–15. <https://aclanthology.org/2020.tal-3.1/>
- [7] A. Seza Doğruöz, Sunayana Sitaram, Barbara E. Bullock, and Almeida Jacqueline Toribio. 2021. A Survey of Code-switching: Linguistic and Social Perspectives for Language Technologies. 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)*, Chengqing Zong, Fei Xia, Wenjie Li, and Roberto Navigli (Eds.). Association for Computational Linguistics, Online, 1654–1666. doi:10.18653/v1/2021.acl-long.131
- [8] Michael Günther, Louis Milliken, Jonathan Geuter, Georgios Mastrapas, Bo Wang, and Han Xiao. 2023. Jina Embeddings: A Novel Set of High-Performance Sentence Embedding Models. arXiv:2307.11224 [cs.CL] <https://arxiv.org/abs/2307.11224>
- [9] Kuan-Po Huang, Chih-Kai Yang, Yu-Kuan Fu, Ewan Dunbar, and Hung yi Lee. 2024. Zero Resource Code-switched Speech Benchmark Using Speech Utterance Pairs For Multiple Spoken Languages. arXiv:2310.03018 [eess.AS] <https://arxiv.org/abs/2310.03018>
- [10] Shaoxiong Ji, Shirui Pan, Erik Cambria, Pekka Marttinen, and Philip S. Yu. 2022. A Survey on Knowledge Graphs: Representation, Acquisition, and Applications. *IEEE Transactions on Neural Networks and Learning Systems* 33, 2 (Feb. 2022), 494–514. doi:10.1109/tnnls.2021.3070843
- [11] Zehan Li, Xin Zhang, Yanzhao Zhang, Dingkun Long, Pengjun Xie, and Meishan Zhang. 2023. Towards General Text Embeddings with Multi-stage Contrastive Learning. *ArXiv abs/2308.03281* (2023). <https://api.semanticscholar.org/CorpusID:260682258>
- [12] Luke Merrick, Danmei Xu, Gaurav Nuti, and Daniel Campos. 2024. Arctic-Embed: Scalable, Efficient, and Accurate Text Embedding Models. arXiv:2405.05374 [cs.CL] <https://arxiv.org/abs/2405.05374>
- [13] Zach Nussbaum, John X. Morris, Brandon Duderstadt, and Andriy Mulyar. 2025. Nomic Embed: Training a Reproducible Long Context Text Embedder. arXiv:2402.01613 [cs.CL] <https://arxiv.org/abs/2402.01613>
- [14] Heiko Paulheim. 2018. How much is a triple? estimating the cost of knowledge graph creation. (2018).
- [15] Yingqi Qu, Yuchen Ding, Jing Liu, Kai Liu, Ruiyang Ren, Wayne Xin Zhao, Daxiang Dong, Hua Wu, and Haifeng Wang. 2021. RocketQA: An Optimized Training Approach to Dense Passage Retrieval for Open-Domain Question Answering. arXiv:2010.08191 [cs.CL] <https://arxiv.org/abs/2010.08191>
- [16] Aaron van den Oord, Yazhe Li, and Oriol Vinyals. 2018. Representation Learning with Contrastive Predictive Coding. *ArXiv abs/1807.03748* (2018). <https://api.semanticscholar.org/CorpusID:49670925>
- [17] Liang Wang, Nan Yang, Xiaolong Huang, Bingxing Jiao, Linjun Yang, Daxin Jiang, Rangan Majumder, and Furu Wei. 2022. Text Embeddings by Weakly-Supervised Contrastive Pre-training. *ArXiv abs/2212.03533* (2022). <https://api.semanticscholar.org/CorpusID:254366618>
- [18] Liang Wang, Nan Yang, Xiaolong Huang, Linjun Yang, Rangan Majumder, and Furu Wei. 2024. Improving Text Embeddings with Large Language Models. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Lun-Wei Ku, Andre Martins, and Vivek Srikumar (Eds.). Association for Computational Linguistics, Bangkok, Thailand, 11897–11916. doi:10.18653/v1/2024.acl-long.642
- [19] Liang Wang, Nan Yang, Xiaolong Huang, Linjun Yang, Rangan Majumder, and Furu Wei. 2024. Multilingual E5 Text Embeddings: A Technical Report. arXiv:2402.05672 [cs.CL] <https://arxiv.org/abs/2402.05672>
- [20] Genta Indra Winata, Samuel Cahyawijaya, Zihan Liu, Zhaojiang Lin, Andrea Madotto, and Pascale Fung. 2021. Are Multilingual Models Effective in Code-Switching?. In *Proceedings of the Fifth Workshop on Computational Approaches to Linguistic Code-Switching*, Thamar Solorio, Shuguang Chen, Alan W. Black, Mona Diab, Sunayana Sitaram, Victor Soto, Emre Yilmaz, and Anirudh Srinivasan (Eds.). Association for Computational Linguistics, Online, 142–153. doi:10.18653/v1/2021.calcs-1.20
- [21] Michael; Chen Zhiyu; Zhuang Yingying; Pi Shu-Ting; Liu Qun; Maragoud Rajashekar; Nguyen Vy; Beniwal Anurag Zhang; Ziji; Yang. 2025. REIC: RAG-Enhanced Intent Classification at Scale. *arXiv preprint arXiv:6498308* (2025).
- [22] Xin Zhang, Yanzhao Zhang, Dingkun Long, Wen Xie, Ziqi Dai, Jialong Tang, Huan Lin, Baosong Yang, Pengjun Xie, Fei Huang, Meishan Zhang, Wenjie Li, and Min Zhang. 2024. mGTE: Generalized Long-Context Text Representation and Reranking Models for Multilingual Text Retrieval. arXiv:2407.19669 [cs.CL] <https://arxiv.org/abs/2407.19669>
- [23] Ziqiang Zhang, Long Zhou, Chengyi Wang, Sanyuan Chen, Yu Wu, Shujie Liu, Zhuo Chen, Yanqing Liu, Huaming Wang, Jinyu Li, Lei He, Sheng Zhao, and Furu Wei. 2023. Speak Foreign Languages with Your Own Voice: Cross-Lingual Neural Codec Language Modeling. arXiv:2303.03926 [cs.CL] <https://arxiv.org/abs/2303.03926>
- [24] Wenxuan Zhou, Fangyu Liu, Ivan Vulić, Nigel Collier, and Muhao Chen. 2021. PRIX-1m: Pretraining for multilingual knowledge base construction. *arXiv preprint arXiv:2110.08443* (2021).
- [25] Yingying Zhuang, Yichao Lu, and Simi Wang. 2021. Weakly Supervised Extractive Summarization with Attention. In *Proceedings of the 22nd Annual Meeting of the Special Interest Group on Discourse and Dialogue*, Haizhou Li, Gina-Anne Levow, Zhou Yu, Chitralekha Gupta, Berrak Sisman, Siqi Cai, David Vandyke, Nina Dethlefs, Yan Wu, and Junyi Jessy Li (Eds.). Association for Computational Linguistics, Singapore and Online, 520–529. doi:10.18653/v1/2021.sigdia-1.54
- [26] Yingying Zhuang, Jiecheng Song, Narayanan Sadagopan, and Anurag Beniwal. 2023. Self-supervised Pre-training and Semi-supervised Learning for Extractive Dialog Summarization. In *Companion Proceedings of the ACM Web Conference 2023* (Austin, TX, USA) (WWW '23 Companion). Association for Computing Machinery, New York, NY, USA, 1069–1076. doi:10.1145/3543873.3587680

A Plots for Model Comparison

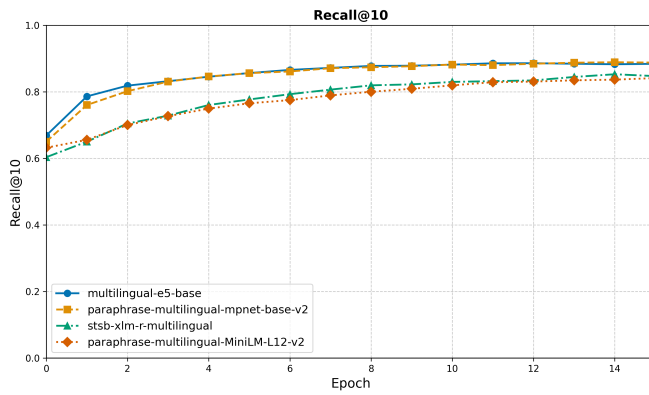
The following plots represent some other comparisons between the models we compared between.



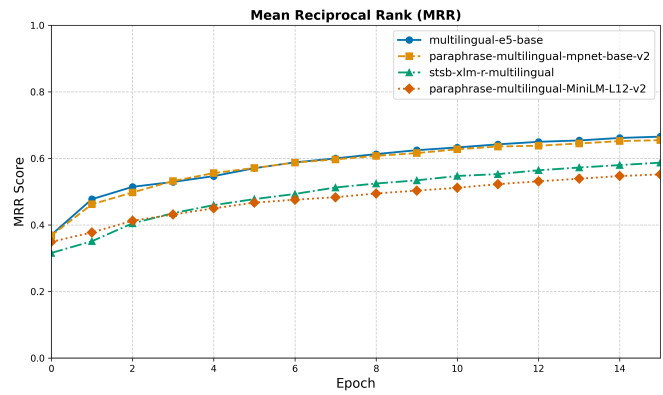
(a) Performance of Recall@3



(b) Performance of Recall@5



(c) Performance of Recall@10



(d) Mean Reciprocal Rank (MRR)

Figure 3: Performance metrics at different points during contrastive fine-tuning for various multilingual embedding models