
Enhancing Multilingual LLM Pretraining with Model-Based Data Selection

Bettina Messmer*

EPFL

`bettina.messmer@epfl.ch`

Vinko Sabolčec*

EPFL

`vinko.sabolcec@epfl.ch`

Martin Jaggi

EPFL

`martin.jaggi@epfl.ch`

Abstract

Dataset curation has become a basis for strong large language model (LLM) performance. While various rule-based filtering heuristics exist for English and multilingual datasets, model-based filtering techniques have primarily focused on English. To address the disparity stemming from limited research on non-English languages, we develop a model-based filtering framework for multilingual datasets that aims to identify a diverse set of structured and knowledge-rich samples. Our approach emphasizes transparency, simplicity, and efficiency, leveraging Transformer- and FastText-based classifiers to ensure the broad accessibility of our technique and data. We conduct comprehensive ablation studies on the FineWeb-2 web crawl dataset across diverse language families, scripts, and resource availability to demonstrate the effectiveness of our method. Training a 1B-parameter Llama model for 70B and 119B tokens, our approach can match the baseline MMLU score with as little as 15% of the training tokens, while also improving across other benchmarks and mitigating the curse of multilinguality. These findings provide strong evidence for the generalizability of our approach to other languages. As a result, we extend our framework to 20 languages for which we release the refined pretraining datasets.

1 Introduction

Large Language Models (LLMs) have demonstrated impressive performance improvements when trained on increasingly larger datasets and model sizes [Brown et al., 2020]. While Brown et al. [2020] already observed the importance of using a cleaned version of Common Crawl for improved performance, the high cost of LLM training has further motivated research into better pretraining quality filters.

Deduplication and heuristic-based dataset cleaning have become standard practices in data curation [Rae et al., 2021, Raffel et al., 2020, De Gibert et al., 2024]. These quality filters are often complemented by additional filters, such as the removal of personally identifiable information (PII) [Penedo et al., 2024a] or model-based toxicity filtering [Soldaini et al., 2024]. Recently, model-based filtering has also emerged as a promising method for quality filtering. The release of FineWeb-Edu [Penedo et al., 2024a] demonstrated that pretraining on just 10% of the tokens (38B) from an English dataset filtered using a model-based approach can achieve performance comparable to models trained on 350B tokens of unfiltered data. Moreover, when trained on equivalent amounts of data, this model largely outperforms the baseline. Concurrently, the release of DataComp-LM (DCLM) [Li et al., 2024b] showed that competitive performance can be achieved using a simple and efficient model-based approach, namely a FastText [Joulin et al., 2017] classifier trained on a carefully selected training dataset.

*Equal contribution

However, these recent advances have primarily focused on English data. This emphasis risks further widening the disparity in LLM performance between languages, as less than half of internet content is written in English². To address this concern, we aim to extend model-based filtering frameworks to multilingual datasets. While model perplexity-based filtering is commonly applied to multilingual datasets [Wenzek et al., 2019, Laurençon et al., 2022, Nguyen et al., 2023], the current state-of-the-art, FineWeb-2 [Penedo et al., 2024c], primarily relies on heuristic-based filters. In this work, we focus on model-based filtering with a quality definition that emphasizes: 1) structured data and 2) knowledge-rich data samples, to enhance multilingual pretraining datasets.

To achieve this, we leverage embedding-based classification models. Firstly, we adopt the FastText quality filtering approach from DCLM to develop a unified framework for multilingual datasets that span diverse language families, scripts, and resource availability, focusing on Chinese, German, French, Arabic, and Danish as representative languages for our experiments. Additionally, we extend this embedding-based approach by incorporating Transformer [Vaswani et al., 2023] embeddings, specifically XLM-RoBERTa [Conneau et al., 2020], for filtering. Figure 1 shows the clear performance gains of our best FastText and Transformer embedding-based approaches over the state-of-the-art baseline FineWeb-2 data.

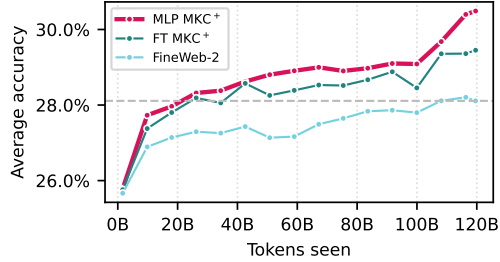


Figure 1: Average accuracy on Chinese (CMMLU), German (MMLU), and French (MMLU) benchmarks during training: FineWeb-2 baseline compared to our methods (10% data retention).

In summary, our contributions are as follows:

- We develop a transparent, simple, and unified framework for multilingual model-based filtering at web scale, enabling data curation across diverse language families, scripts and resource availability.
- We present comprehensive per-language ablation studies of embedding-based multilingual quality filtering on top of the FineWeb-2 dataset [Penedo et al., 2024c], achieving performance comparable to the baseline while using as little as 15% of the tokens. We additionally analyze the impact of dataset contamination. Lastly, our experiments show that our dataset doesn’t suffer from the *curse of multilinguality* [Chang et al., 2023].
- We evaluate the impact of different data selection classifiers, in particular their training datasets, on the downstream performance of LLMs.
- We release the refined pretraining dataset³ covering 20 languages⁴, filtered using our proposed framework, along with the codebase⁵, to advance multilingual language modeling.

2 Related Work

Data Curation. In order to pretrain LLMs on a large amount of diverse texts, Common Crawl⁶ is often used as the base dataset. However, early works already observed that performing data curation on Common Crawl is crucial for model performance [Brown et al., 2020]. In fact, data curation is important across NLP tasks [Peter et al., 2023, Finkelstein et al., 2024]. Specifically for pretraining data, there exist various data curation approaches, such as deduplication [Lee et al., 2022], PII removal [Subramani et al., 2023], or toxicity filtering [Arnett et al., 2024]. Another important aspect is quality filtering of the documents. For this, the definition of quality is an important aspect. A common approach is to use heuristics to remove documents outside of the target distribution, such as filtering based on average word length, existence of punctuation, or document length [Rae et al., 2021, Raffel et al., 2020]. Another approach is to define model-based filters, where research has focused on perplexity measure of the text [Wenzek et al., 2019, Marion et al.,

²w3techs.com/technologies/overview/content_language

³huggingface.co/datasets/epfml/FineWeb2-HQ

⁴Russian, Chinese, German, Japanese, Spanish, French, Italian, Portuguese, Polish, Dutch, Indonesian, Turkish, Czech, Vietnamese, Swedish, Persian, Arabic, Greek, Danish, Hungarian (dataset details in Appendix A)

⁵github.com/epfml/fineweb2-hq

⁶commoncrawl.org

2023, Ankner et al., 2024], distributional similarity measures [Li et al., 2024b] and LLM-based quality assessment [Gunasekar et al., 2023, Wettig et al., 2024, Sachdeva et al., 2024, Penedo et al., 2024a]. In this work, we build upon previous curated datasets based on heuristic filtering, namely the state-of-the-art dataset FineWeb-2 [Penedo et al., 2024c], and focus on model-based filtering for structured and knowledge-rich documents relying on textual embeddings.

Curated English datasets. One of the early curated datasets was C4 [Raffel et al., 2020], followed by MassiveText [Rae et al., 2021]. RefinedWeb [Penedo et al., 2023] was an important step forward, demonstrating that filtered web data can outperform selected high-quality data sources. Although these datasets have not been made fully publicly available, their filtering techniques have been expanded upon in recent fully public datasets, such as Dolma [Soldaini et al., 2024], FineWeb, FineWeb-Edu [Penedo et al., 2024a] and DCLM [Li et al., 2024b]. While FineWeb primarily relies on filter heuristics for data quality, Dolma adopts model perplexity filtering. FineWeb-Edu takes model-based filtering a step further and relies on LLM-based quality assessment. DCLM, a concurrent work, has achieved competitive performance using a FastText [Joulin et al., 2017] classifier trained on a carefully selected training dataset. In this work we adapt and extend this approach to the multilingual context.

Curated Multilingual Datasets. Analogously to English datasets, significant work has been done in the multilingual space. For example, CCNet [Wenzek et al., 2019] has been influential, with its language identification and model perplexity filtering for data quality being adopted in subsequent datasets. Similar to earlier English datasets, CCNet was not published directly, but rather provided tools for data cleaning. RedPajama [Together Computer, 2023] is a prominent multilingual dataset relying on these filtering techniques, offering data in 5 European languages. Other datasets, such as OSCAR [Ortiz Suárez et al., 2019, Abadji et al., 2021, Abadji et al., 2022], mC4 [Xue et al., 2021], ROOTS [Laurençon et al., 2022], MADLAD-400 [Kudugunta et al., 2023], CulturaX [Nguyen et al., 2023], and HPLT [de Gibert et al., 2024], expanded coverage across a variety of language families and scripts. These datasets offer refined content for hundreds of languages, while FineWeb-2 [Penedo et al., 2024c] pushes the limit to thousands of languages and further improves performance. A concurrent work by Martins et al. [2025] uses translation to train a multilingual quality filter based on the English FineWeb-Edu [Penedo et al., 2024a] scores. Our work similarly focuses on filtering high-quality samples across various language families and scripts. However, we take a different approach: we train a separate classifier for each language from scratch using structured and knowledge-rich representative samples. We limit our scope to 20 languages as the number of documents drops quickly and there is trade-off between retaining a sufficient number of pretraining tokens and ensuring data quality [Muennighoff et al., 2023, Held et al., 2025].

Multilingual Embedding Models. Early word embedding models like Word2Vec [Mikolov et al., 2013] and GloVe [Pennington et al., 2014] lacked contextual understanding. FastText [Bojanowski et al., 2017] built upon them and improved performance by incorporating subword information. Transformer [Vaswani et al., 2023] models like BERT [Devlin et al., 2019] and GPT [Radford et al., 2018] then revolutionized the field with context-aware embeddings. Multilingual models like mBERT, XLM [Lample and Conneau, 2019], and XLM-RoBERTa [Conneau et al., 2020] further advanced cross-lingual understanding, with recent open-source LLMs pushing performance even higher [Llama Team, 2024, Mistral AI, 2025]. Using such Transformer models, documents and representative samples can be mapped into a shared embedding space to estimate their similarity. Focusing on transparency, simplicity and efficiency in our work, we use FastText and XLM-RoBERTa for our model-based filtering.

Multilingual Evaluation. Evaluating LLMs requires diverse benchmarks testing linguistic and cognitive abilities like reading comprehension, reasoning, and knowledge. While established benchmarks such as MMLU [Hendrycks et al., 2020] and ARC [Clark et al., 2018] exist for English evaluation, assessments in other languages often rely on translations from English sources, as seen in XNLI [Conneau et al., 2018] and the machine-translated version of MMLU [Lai et al., 2023]. However, translations can be problematic, failing to capture cultural nuances or introducing "translationese" [Romanou et al., 2024]. Recent work by Romanou et al. [2024] and Singh et al. [2024a] emphasizes the importance of culturally sensitive, natively collected benchmarks. Task difficulty and formulation also impact model performance when trained for shorter durations [Kydliček et al., 2024]. In our work, we follow FineTasks, a recent evaluation tasks suite by Kydliček et al. [2024] to assess our model-based filtering approaches across multiple languages.

3 Methods

In this work, we present our model-based filtering approaches. Our methodology is structured into two key components: 1) we select suitable training datasets, aiming to identifying a diverse set of structured and knowledge-rich samples and 2) we describe the different models, namely FastText and Transformer embedding-based filters, used to capture and leverage these characteristics.

3.1 Classifier Training Dataset

Representative Sample Selection. Our goal is to identify a diverse set of structured and knowledge-rich samples, especially within a multilingual context. We define two criteria for our training datasets: 1) the samples must be informative and well-structured and 2) the datasets must be available in multiple languages. While some multilingual benchmark datasets meet these criteria precisely, it is important to note that we do not train the LLM directly on this data. Instead, we train a proxy model to assess pretraining data quality. Nevertheless, we must remain cautious about potentially increased pretraining data contamination stemming from this approach, as discussed in Section 4.2.5.

Based on our criteria, we selected the following datasets as representative examples.

- **Aya Collection.** A prompt completion dataset comprising ~ 514 M samples covering a variety of tasks, generated using instruction-style templates in 101 languages [Singh et al., 2024b].
- **Aya Dataset.** Human-annotated instruction fine-tuning dataset consisting of ~ 202 K prompt-completion pairs in 65 languages [Singh et al., 2024b].
- **MMLU.** Dataset contains ~ 14 K multiple-choice knowledge questions on various topics in English [Hendrycks et al., 2020]. Multilingual version was translated into 14 languages by professional translators [OpenAI, 2024].
- **OpenAssistant-2.** The dataset contains ~ 14 K user-assistant conversations with multiple messages in 28 languages [Fischer et al., 2024].
- **Include-Base-44.** Multiple-choice questions focused on general and regional knowledge, extracted from academic and professional exams. Spanning 44 languages, it includes a total of ~ 23 K samples [Romanou et al., 2024].

Representative Sample Collection. *MMLU* and *Include-Base-44* are highly curated benchmark datasets, containing structured, knowledge-rich samples. The *Aya Dataset* is human-curated, while *OpenAssistant-2* is partially human-curated and partially generated by large language models (LLMs). In contrast, the *Aya Collection* consists of various AI-generated samples without quality guarantee, though it represents the largest and most multilingual of the five.

To address the quality difference, we create two *Multilingual Knowledge Collection (MKC)* configurations which allow us to evaluate the trade-off between data quality and scale:

- **MKC:** Includes *Include-Base-44*, *OpenAssistant-2*, *MMLU*, and the *Aya Dataset*
- **MKC⁺:** Includes *MKC* and the *Aya Collection*

Dataset Creation. For our model-based filtering approaches, our goal is to identify documents from the pretraining dataset that are most similar to our representative samples, with the notion of similarity determined by the specific classifier used. We can directly measure similarity to our training data, for example, by calculating cosine similarity with training samples in the embedding space. Alternatively, following the approach of Li et al. [2024b], the task can be framed as a binary classification problem, with the representative samples as the positive class. For the negative class, we can subsample documents from our pretraining dataset, under the assumption that the majority of these documents are not well-structured or knowledge-rich. We use both approaches for our classifiers.

To create the binary classification training dataset, we selected up to 80K positive samples by using all examples from the smaller source datasets (e.g., *Include-Base-44*) and randomly subsampling from the *Aya Collection* for *MKC⁺*. The positive samples were combined with the same number of randomly sampled negative examples from FineWeb-2. The same training dataset was utilized across all model-based filtering approaches, disregarding negative samples when unnecessary. Additionally, we created a training dataset for each language individually to avoid leaking language-specific biases to data of other languages.

Sample Pre-processing. We applied no pre-processing to the FineWeb-2 (negative) samples but performed minimal pre-processing on the representative (positive) samples. For instance, in datasets like *MMLU* or *OpenAssistant-2*, we concatenated various sample components. For the *Aya Collection*, we resolved encoding issues in non-Latin languages and removed samples containing `<unk>` tokens, which were particularly prevalent in Arabic data (37.1%).

3.2 FastText-based Filtering (FT)

To efficiently process datasets with over 100 million documents [Penedo et al., 2024c], similar to DCLM [Li et al., 2024b], we used a binary FastText classifier [Joulin et al., 2017]. FastText runs on CPU and can be deployed across multiple cores, for example using DataTrove [Penedo et al., 2024b].

We trained our FastText classifier on the processed training set using 2-gram features (4-gram for Chinese). These classifiers were then used to assign scores to all documents in the pretraining dataset. To filter the dataset, we applied a score threshold based on the desired retention percentage of documents. This approach balances dataset size and the predicted quality of the samples.

3.3 Transformer Embedding-based Filtering

To leverage rich semantic information based on contextual relationships, we utilized Transformer model embeddings. Specifically, we selected a pretrained XLM-RoBERTa base model [Conneau et al., 2020] due to its support of 100 languages, a relatively small size of 279M parameters, and its transparent training procedure. This choice enabled us to process web-scale data efficiently without being restricted to a single language and aligned with our commitment to open science.

To retain general embeddings that can be reused across methods, we opted against fine-tuning the model. For each document from our datasets, we computed the 768-dimensional embedding by mean pooling the embeddings of the output sequence. Since the model has a fixed maximum sequence length of 512 tokens, we considered only the first 512 tokens of each document, assuming they are representative of the entire document.

After computing the embeddings of our corpora, we experimented with two methods: 1) classification of embeddings using a multi-layer perceptron and 2) cosine similarity between the embeddings. As in the FastText approach, we scored each document and applied a threshold to retain the desired percentage of the highest-scoring documents.

Multi-Layer Perceptron (MLP). We trained a single-hidden-layer neural network with a dimension of 256, the ReLU activation function, a 20% dropout, and the sigmoid function on the output. The network was trained for 6 epochs using the AdamW optimizer [Loshchilov and Hutter, 2019] with a constant learning rate 0.0003 and binary cross-entropy loss. We computed document scores using the output layer of the MLP model, which used XLM-RoBERTa document embeddings as input.

Cosine Similarity (CS). We computed the document scores as the maximum cosine similarity between its embeddings and a set of K randomly sampled positive sample embeddings. We experimented with varying values of K , including 1024, 2048, 4096, 8192, and 16384. However, we did not observe significant differences in the documents with high scores across these variations when manually inspecting the data. To strike a balance between the diversity of the positive samples and computational efficiency, we chose $K = 8192$ for our experiments.

4 Experiments

4.1 Experimental Setup

Technical Details. We evaluate 1B-parameter Llama models [Llama Team, 2024] to demonstrate the effectiveness of our model-based filtering approaches. The models are trained on either 70B or 119B tokens, balancing token quality and diversity. The smaller dataset (70B tokens) exposes the model to each token at most once (with a few exceptions where some tokens appear twice). The larger dataset (119B tokens) simulates longer training, resulting in increased token repetition. Training utilizes the HuggingFace Nanotron library [Hugging Face, 2024a] with the AdamW optimizer [Loshchilov and Hutter, 2019] and a WSD learning rate schedule [Hägele et al., 2024].

To minimize the need for costly hyperparameter tuning, we maintain a consistent setup across all experiments. Specifically, we adopt the DeepSeek scaling law [DeepSeek-AI et al., 2024] with a batch size of 1.6M tokens, learning rate of 0.0008, and 2000 warmup steps.

As the base dataset, we use FineWeb-2 [Penedo et al., 2024c], which has been shown to provide a strong baseline across a variety of languages. Since FineWeb-2 is globally deduplicated, we rehydrate both filtered and unfiltered data using the hyperparameters recommended by Penedo et al. [2024c].

To validate our method on English, we use three datasets: FineWeb [Penedo et al., 2024a] as the baseline, along with FineWeb-Edu [Penedo et al., 2024a] and DCLM [Li et al., 2024b], both of which represent the current state-of-the-art. Tokenization is performed using the multilingual Mistral v3 (Tekken) tokenizer [Mistral AI, 2024].

Evaluation. Our evaluation prioritizes a diverse range of tasks to ensure the models retain well-rounded capabilities, rather than focusing exclusively on knowledge-based tasks. Specifically, we include tasks covering reading comprehension, general knowledge, natural language understanding, common-sense reasoning, and generative tasks in the target language. To evaluate our approach, we use the HuggingFace LightEval library [Fourrier et al., 2023].

For French, Chinese, and Arabic, we utilize the FineTasks [Kydliček et al., 2024] multilingual evaluation suite, which is designed to provide meaningful signals even for models trained in the order of 100B tokens. We select analogous tasks for German and Danish. For English, we rely on the SmolLM tasks suite [Hugging Face, 2024b]. A complete list of tasks and their evaluation metrics for each language is provided in Appendix D.

Model Selection. We follow the approach used in FineTasks [Kydliček et al., 2024] for filter selection, computing a global rank score across individual metrics and languages to determine the optimal approach. For a detailed description of the average rank computation, please refer to Appendix E.

Computational Cost. We run our experiments on a compute cluster containing four GH200 chips per node, with each GH200 chip containing 72 CPU cores and one H100 GPU. Model training on 119B tokens costs approximately 1.1K H100 compute hours, while the embedding computation of all data for 20 languages costs approximately 4K H100 compute hours. We release the embeddings publicly so this computation does not have to be repeated⁷. In total, we use approximately 152K H100 compute hours for running our experiments and generating the document embeddings. Data selection classifier training (i.e. FastText and MLP) takes a few minutes on a CPU. Filtering using any of the approaches (i.e., FastText, MLP, or CS) is computationally inexpensive, parallelizable, and is run on CPU. Overall, filtering German, Chinese, French, Arabic, and Danish data for one filtering approach in the ablations costs approximately 60 CPU hours, which is distributed over multiple CPU cores depending on the dataset size in each language to have filtering finish in approximately 30 minutes.

4.2 Experimental Results & Discussion

4.2.1 Model Selection

In Section 3, we introduced several model-based filtering approaches. *But which of these performs the best?* We evaluate which combination of our defined classifier training datasets (*MKC* or *MKC*⁺) and filtering methods (*FT*, *MLP* or *CS*) achieve the highest performance. Table 1 presents the overall ranking across our representative language selection (Chinese, German, French, Arabic, Danish) and training runs of 70B and 119B tokens. Analogous to the DCLM filtering recipe [Li et al., 2024b], the results are based on a dataset that retains 10% of the documents for the high-resource datasets (Chinese, German, French) and keeps 56% and 65% of the documents for the lower-resource languages

Table 1: Benchmark performance comparison: Average rank between FineWeb-2 baseline and our proposed filtering methods (*FT*, *MLP*, and *CS*) trained on *MKC*⁺ or *MKC*, retaining top 10% for Chinese, German, and French, 56% for Arabic, and 65% for Danish. The average rank is computed across FineTasks for 1B-parameter models evaluated after 70B and 119B tokens.

Approach	Average Rank
<i>MLP MKC</i> ⁺	4.35
<i>MLP MKC</i>	6.11
<i>FT MKC</i> ⁺	7.17
<i>FT MKC</i>	8.04
<i>CS MKC</i>	8.10
Baseline	8.72
<i>CS MKC</i> ⁺	8.79

⁷huggingface.co/datasets/epfml/FineWeb2-embedded

(Arabic and Danish, respectively). These percentages maintain approximately 70B tokens, under the assumption of uniform token distribution across documents. We also exclude approaches that use *MKC* for training on Danish, as it lacks sufficient training data. For detailed, per-language results, please refer to Appendix C.1.

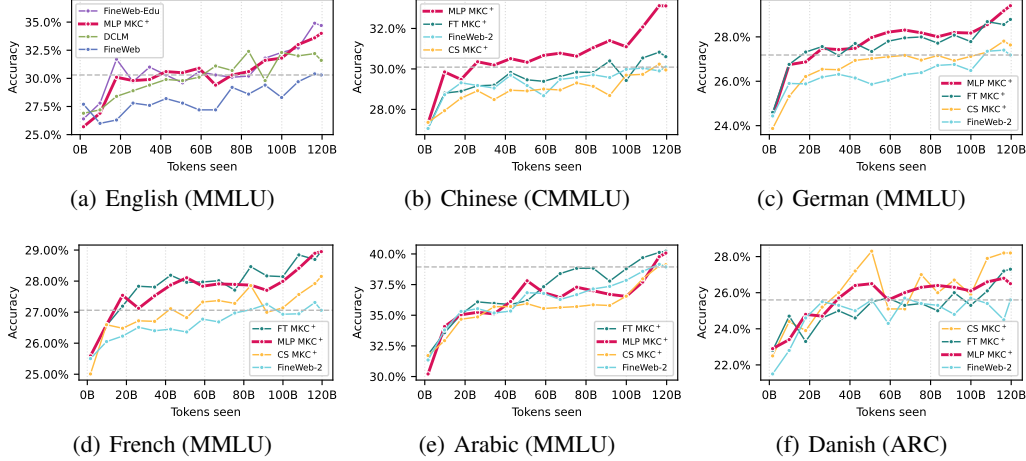


Figure 2: Benchmark performance comparison: Accuracy during 119B token training between baseline methods (FineWeb, DCLM, FineWeb-Edu, FineWeb-2) and our proposed filtering approaches (*FT*, *MLP*, and *CS*), trained on *MKC*⁺. Our approaches use 10% data retention for English, Chinese, German, and French, 56% for Arabic, and 65% for Danish. For English, Chinese, German, and French, baseline-level performance is reached at approximately 20B tokens (16.7% of total).

Table 1 demonstrates that *MLP MKC*⁺ approach outperforms all other approaches. Interestingly, the high- and low-scored samples presented in Appendix F align with the observed rankings. Figure 2 further highlights the strong performance of *MLP MKC*⁺, particularly for high-resource languages, where it largely outperforms the baseline. For lower-resource languages—where less data was filtered—the performance gains are less pronounced. Notably, *FT* filtering is also competitive. Given the computational expense of XLM-RoBERTa embeddings, FastText can be a promising alternative in resource-constrained setups.

4.2.2 Threshold Selection

In Section 4.2.1, we base our model selection on experiments that retain top 10% of the data for high-resource languages. *But is this the optimal threshold?* Following the methodology of Li et al. [2024b], we analyze the impact of varying filter strengths on performance for Chinese, German, and French, using our *MLP* and *FT* filtering methods. The results are summarized in Table 2, with a comprehensive analysis, including results for *CS*, provided in Appendix C.2 (Table 14). Consistent with their findings, we observe that retaining top 10% of the data is a competitive threshold, particularly for approaches using the *MKC*⁺ dataset. Interestingly, approaches using *MKC* perform better with higher retention. In Appendix C.2, we investigate how some filters’ bias toward shorter documents affects threshold selection, though our analysis indicates multiple factors contribute to optimal threshold determination.

Table 2: Benchmark performance comparison: Average rank between FineWeb-2 baseline and our proposed filtering methods (*FT*, *MLP*) trained on *MKC*⁺ or *MKC*, retaining top 10%, 15% or 20% of documents. The average rank is computed across FineTasks for 1B-parameter models evaluated on Chinese, German and French after 70B and 119B tokens.

Approach	Threshold	Average Rank
<i>MLP MKC</i> ⁺	10%	8.85
<i>MLP MKC</i> ⁺	15%	9.44
<i>MLP MKC</i>	20%	11.37
<i>MLP MKC</i>	15%	11.70
<i>MLP MKC</i>	10%	11.95
<i>MLP MKC</i> ⁺	20%	11.97
<i>FT MKC</i> ⁺	10%	13.92
<i>FT MKC</i>	15%	14.62
<i>FT MKC</i>	10%	14.74
<i>FT MKC</i>	20%	15.62
<i>FT MKC</i> ⁺	15%	16.27
<i>FT MKC</i> ⁺	20%	16.51
Baseline	—	18.55

4.2.3 Training Data Analysis

The experiments in Sections 4.2.1 and 4.2.2 are based on the training datasets *MKC* and *MKC*⁺. *But is the diversity introduced by combining various base datasets truly necessary?* We evaluate the impact of each base dataset individually and compare it to the combined *MKC*⁺ dataset. For this ablation study, we use our best filtering method (*MLP* with a top 10% retention) and train the models on 30B tokens. This token count is chosen to match the size of the smallest filtered dataset, ensuring consistency across comparisons.

The results, presented in Table 3, show that despite the absence of a quality guarantee for all samples in the *Aya Collection*, this dataset yields strong performance, making our approach applicable for various languages. Overall, we observe that the diversity resulting from combining all individual training datasets gives the best results.

Interestingly, models trained exclusively on *Include-Base-44* and *OpenAssistant-2* perform worse overall than the baseline. This may reflect dataset characteristics—*Include-Base-44* is small and domain-specific, containing mostly driving license exam questions in its German subset. *OpenAssistant-2* includes a limited number of samples, with fewer than 2K positive samples per training set, which likely negatively impacts classifier performance. In Appendix C.3, we reexamine how document length bias relates to model performance, confirming our Section 4.2.2 finding that performance depends on factors beyond document length. In Appendix C.4, we further verify our filtering approach preserves sufficient dataset diversity.

Table 3: Benchmark performance comparison: Average rank between FineWeb-2 baseline and *MLP* filtering trained on either full *MKC*⁺ or its individual components, retaining top 10% for Chinese, German, and French, 56% for Arabic, and 65% for Danish. The average rank is computed across FineTasks for 1B-parameter models trained on 30B tokens per language.

Dataset	Average Rank
<i>MKC</i> ⁺	2.52
<i>Aya Collection</i>	2.91
<i>Aya Dataset</i>	3.17
<i>MMLU</i>	3.57
Baseline	4.09
<i>OpenAssistant-2</i>	4.53
<i>Include-Base-44</i>	5.42

4.2.4 Approach Validation on English

Previous experiments have shown strong performance of our *MLP MKC*⁺ approach. *But do these results translate to English?* Table 4 presents the performance of *MLP MKC*⁺ with 10% retention applied to the English FineWeb dataset Penedo et al. [2024a]. Our method is compared against FineWeb and baselines using model-based filtered datasets, including DCLM Li et al. [2024b] and FineWeb-Edu [Penedo et al., 2024a]. To save computational resources, we use the 6 most recent FineWeb and FineWeb-Edu dumps and the first partition of DCLM⁸, which we denote with *. Each of these subsets contains more than 119B tokens, with FineWeb retaining this size even after applying our filtering retaining top 10% of the documents.

Table 4: English benchmark performance: Our *MLP MKC*⁺ approach (top 10% documents) compared to FineWeb, DCLM, and FineWeb-Edu baselines. The average rank is computed across SmolLM tasks using 1B-parameter models trained on 119B tokens.

Dataset	Ours	DCLM*	FW-Edu*	FW*
Average Rank	1.8333	2.3889	2.4444	3.3333
ARC (Challenge)	0.3550	0.3530	0.3850	0.3010
ARC (Easy)	0.6670	0.6470	0.6970	0.5880
CommonsenseQA	0.3870	0.4100	0.3770	0.3850
HellaSwag	0.6040	0.5960	0.5700	0.5930
MMLU	0.3400	0.3160	0.3470	0.3030
OpenBookQA	0.3860	0.3840	0.4180	0.3560
PIQA	0.7510	0.7510	0.7410	0.7620
WinoGrande	0.5720	0.5610	0.5660	0.5550
TriviaQA	0.0820	0.1240	0.0320	0.0370

While each approach demonstrates strengths in different benchmarks, as seen from Table 4 and Figure 2, the overall average rank results indicate that our method outperforms all other baselines.

4.2.5 Data Contamination Analysis

Our LLMs are never trained on benchmark datasets. *But is the strong performance observed in the previous sections primarily due to an increased ratio of data contamination?* To ensure the validity of our approach, we conduct decontamination experiments, as web crawl data may include evaluation benchmark tasks. While Li et al. [2024b] addressed similar concerns, our approach follows the

⁸huggingface.co/datasets/mlfoundations/dclm-baseline-1.0-parquet

methodology of Brown et al. [2020]. Specifically, we perform 13-gram decontamination of the LLM training data separately for English and French evaluation benchmarks. However, unlike the original approach, we remove the entire document if it is flagged as contaminated, using the implementation provided in DataTrove [Penedo et al., 2024b].

Tables 5 and 6 present the results of decontamination experiments for English and French, respectively. We used the following experimental setup (removed document contamination rates): baseline FineWeb English (0.16%), *MLP MKC*⁺ English with 10% retention (0.19%), baseline FineWeb-2 French (0.14%), and *MLP MKC*⁺ French with 10% retention (0.14%). All models were trained on 119B tokens. Additionally, we compare the results against equivalent training runs without decontamination to further analyze its impact. For an example of a contaminated sample, see Appendix G.

For English models, decontamination slightly reduces performance both for our approach and baseline FineWeb data. Even after decontamination, our approach still outperforms the baseline trained on non-decontaminated data. For French models, our approach performs similarly on decontaminated and non-decontaminated data, both outperforming baseline FineWeb-2. Interestingly, decontaminated baseline data yields better results than its non-decontaminated counterpart.

Table 5: English benchmark performance: Our *MLP MKC*⁺ approach (retaining top 10% documents) in both decontaminated (*D*) and non-decontaminated versions, compared to baseline FineWeb datasets with the same variants. The average rank is computed across SmoLLM tasks for 1B-parameter models trained on 119B tokens.

Dataset	Ours	Ours _D	FW*	FW _D *
Average Rank	1.5000	2.1111	3.0556	3.3333
ARC (Challenge)	0.3550	0.3440	0.3010	0.2880
ARC (Easy)	0.6670	0.6520	0.5880	0.5700
CommonsenseQA	0.3870	0.4000	0.3850	0.3820
HellaSwag	0.6040	0.6040	0.5930	0.5890
MMLU	0.3400	0.3220	0.3030	0.3050
OpenBookQA	0.3860	0.3840	0.3560	0.3740
PIQA	0.7510	0.7590	0.7620	0.7600
WinoGrande	0.5720	0.5550	0.5550	0.5570
TriviaQA	0.0820	0.0380	0.0370	0.0250

Table 6: French Benchmark performance: Our *MLP MKC*⁺ approach (retaining top 10% of the documents) in both decontaminated (*D*) and non-decontaminated versions, compared to baseline FineWeb-2 datasets with the same variants. The average rank is computed across FineTasks for 1B-parameter models trained on 119B tokens.

Dataset	Ours	Ours _D	FW-2 _D	FW-2
Average Rank	2.0556	2.0556	2.7222	3.1667
Belebele	0.3533	0.3400	0.3778	0.3444
HellaSwag	0.5380	0.5350	0.5180	0.5180
X-CSQA	0.2740	0.2810	0.2730	0.2870
XNLI 2.0	0.7400	0.7400	0.7070	0.7180
FQuAD	0.2803	0.2620	0.2890	0.2401
MMLU	0.2895	0.2875	0.2711	0.2706
Mintaka	0.0438	0.0797	0.0658	0.0712
X-CODAH	0.2667	0.2900	0.2800	0.2633
ARC (Challenge)	0.3180	0.3110	0.2880	0.2850

4.2.6 Impact on Multilingual Training - Mitigating the Curse of Multilinguality

Although not our main focus, we found that our refined datasets boost the performance of multilingual models. We trained a multilingual 1B-parameter model on 595B tokens (119B per language), covering all five languages: Chinese, German, French, Arabic and Danish. We compared each language’s results to its monolingual counterpart trained on 119B tokens. Training is performed once for our filtered data and once for original (unfiltered) FineWeb-2.

The results for French are presented in Table 7. Surprisingly, the *curse* of multilinguality [Chang et al., 2023] turns into a *benefit* for our quality filtered datasets: The multilingual model outperforms its monolingual counterpart, when both models have seen an equal amount of tokens of the language of interest. Meanwhile, for unfiltered training data, the multilingual LLM suffers from the *curse* as expected. The disappearance of the *curse* is consistent across all languages except Chinese. Detailed results for the other languages are provided in Appendix C.5.

Table 7: French benchmark performance: Multilingual LLMs (*M*) trained on FineWeb-2 or our *MLP MKC*⁺ refined dataset (retaining top 10% for Chinese, German and French, 56% for Arabic, 65% for Danish) with 595B tokens, compared to monolingual models trained on 119B tokens. The average rank is computed across FineTasks for 1B-parameter models.

Dataset	Ours _M	Ours	FW-2	FW-2 _M
Average Rank	1.8333	2.0556	3.0000	3.1111
Belebele	0.3667	0.3533	0.3444	0.3511
HellaSwag	0.5270	0.5380	0.5180	0.4970
X-CSQA	0.2740	0.2740	0.2870	0.2750
XNLI 2.0	0.7660	0.7400	0.7180	0.7330
FQuAD	0.3212	0.2803	0.2401	0.2459
MMLU	0.2841	0.2895	0.2706	0.2735
Mintaka	0.0456	0.0438	0.0712	0.0579
X-CODAH	0.2900	0.2667	0.2633	0.2567
ARC (Challenge)	0.2970	0.3180	0.2850	0.2670

5 Conclusion

In this work, we developed a framework for model-based filtering of web-scale multilingual pretraining datasets, demonstrating consistent improvements on LLM benchmarks across a wide range of languages. Our Transformer embedding-based classifier, *MLP MKC*⁺, outperforms state-of-the-art methods on both English and multilingual datasets, even when decontaminating the datasets or using them for training multilingual LLMs. While our FastText-based filtering approach performed well and shows promise in resource-constrained setups, *MLP MKC*⁺ consistently outperformed all other methods and can be easily scaled to other languages. These results provide strong empirical evidence supporting our expansion of the framework to 20 languages. We release the corresponding refined pretraining datasets and code, contributing to the advancement of multilingual language modeling.

Acknowledgments and Disclosure of Funding

We thank Guilherme Penedo, Hynek Kydlíček, and Leandro von Werra for their help with FineWeb-2 data, and to Alex Hägele for providing feedback on the paper draft.

This work was supported as part of the Swiss AI Initiative by a grant from the Swiss National Supercomputing Centre (CSCS) under project ID a06 on Alps.

References

- J. Abadji, P. J. O. Suárez, L. Romary, and B. Sagot. Ungoliant: An optimized pipeline for the generation of a very large-scale multilingual web corpus. *Proceedings of the Workshop on Challenges in the Management of Large Corpora (CMLC-9) 2021*. Limerick, 12 July 2021 (Online-Event), pages 1 – 9, Mannheim, 2021. Leibniz-Institut für Deutsche Sprache. doi: 10.14618/ids-pub-10468. URL <https://nbn-resolving.org/urn:nbn:de:bsz:mh39-104688>.
- J. Abadji, P. Ortiz Suarez, L. Romary, and B. Sagot. Towards a Cleaner Document-Oriented Multilingual Crawled Corpus. *arXiv e-prints*, art. arXiv:2201.06642, Jan. 2022.
- E. Almazrouei, R. Cojocaru, M. Baldo, Q. Malartic, H. Alobeidli, D. Mazzotta, G. Penedo, G. Campesan, M. Farooq, M. Alhammedi, J. Launay, and B. Noune. AlGhafa evaluation benchmark for Arabic language models. In H. Sawaf, S. El-Beltagy, W. Zaghoulani, W. Magdy, A. Abdelali, N. Tomeh, I. Abu Farha, N. Habash, S. Khalifa, A. Keleg, H. Haddad, I. Zitouni, K. Mrini, and R. Almatham, editors, *Proceedings of ArabicNLP 2023*, pages 244–275, Singapore (Hybrid), Dec. 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.arabicnlp-1.21. URL <https://aclanthology.org/2023.arabicnlp-1.21/>.
- Z. Ankner, C. Blakeney, K. Sreenivasan, M. Marion, M. L. Leavitt, and M. Paul. Perplexed by perplexity: Perplexity-based data pruning with small reference models, 2024. URL <https://arxiv.org/abs/2405.20541>.
- C. Arnett, E. Jones, I. P. Yamshchikov, and P.-C. Langlais. Toxicity of the Commons: Curating Open-Source Pre-Training Data. *arXiv preprint arXiv:2410.22587*, 2024. URL <https://arxiv.org/pdf/2410.22587>.
- L. Bandarkar, D. Liang, B. Muller, M. Artetxe, S. N. Shukla, D. Husa, N. Goyal, A. Krishnan, L. Zettlemoyer, and M. Khabsa. The Belebele Benchmark: a Parallel Reading Comprehension Dataset in 122 Language Variants. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, page 749–775. Association for Computational Linguistics, 2024. doi: 10.18653/v1/2024.acl-long.44. URL <http://dx.doi.org/10.18653/v1/2024.acl-long.44>.
- L. Bethune, D. Grangier, D. Busbridge, E. Gualdoni, M. Cuturi, and P. Ablin. Scaling laws for forgetting during finetuning with pretraining data injection, 2025. URL <https://arxiv.org/abs/2502.06042>.
- Y. Bisk, R. Zellers, R. L. Bras, J. Gao, and Y. Choi. Piqa: Reasoning about physical commonsense in natural language, 2019. URL <https://arxiv.org/abs/1911.11641>.
- P. Bojanowski, E. Grave, A. Joulin, and T. Mikolov. Enriching word vectors with subword information, 2017. URL <https://arxiv.org/abs/1607.04606>.
- T. Brown, B. Mann, N. Ryder, M. Subbiah, J. D. Kaplan, P. Dhariwal, A. Neelakantan, P. Shyam, G. Sastry, A. Askell, et al. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901, 2020.
- T. A. Chang, C. Arnett, Z. Tu, and B. K. Bergen. When is multilinguality a curse? language modeling for 250 high- and low-resource languages, 2023. URL <https://arxiv.org/abs/2311.09205>.
- Z. Chen, A. H. Cano, A. Romanou, A. Bonnet, K. Matoba, F. Salvi, M. Pagliardini, S. Fan, A. Köpf, A. Mohtashami, A. Sallinen, A. Sakhaeirad, V. Swamy, I. Krawczuk, D. Bayazit, A. Marmet, S. Montariol, M.-A. Hartley, M. Jaggi, and A. Bosselut. Meditron-70b: Scaling medical pretraining for large language models, 2023. URL <https://arxiv.org/abs/2311.16079>.
- J. H. Clark, E. Choi, M. Collins, D. Garrette, T. Kwiatkowski, V. Nikolaev, and J. Palomaki. TyDi QA: A Benchmark for Information-Seeking Question Answering in Typologically Diverse Languages, 2020. URL <https://arxiv.org/abs/2003.05002>.
- P. Clark, I. Cowhey, O. Etzioni, T. Khot, A. Sabharwal, C. Schoenick, and O. Tafjord. Think you have Solved Question Answering? Try ARC, the AI2 Reasoning Challenge, 2018. URL <https://arxiv.org/abs/1803.05457>.

- A. Conneau, R. Rinott, G. Lample, A. Williams, S. Bowman, H. Schwenk, and V. Stoyanov. XNLI: Evaluating cross-lingual sentence representations. In E. Riloff, D. Chiang, J. Hockenmaier, and J. Tsujii, editors, *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, pages 2475–2485, Brussels, Belgium, Oct.-Nov. 2018. Association for Computational Linguistics. doi: 10.18653/v1/D18-1269. URL <https://aclanthology.org/D18-1269/>.
- A. Conneau, K. Khandelwal, N. Goyal, V. Chaudhary, G. Wenzek, F. Guzmán, E. Grave, M. Ott, L. Zettlemoyer, and V. Stoyanov. Unsupervised cross-lingual representation learning at scale, 2020. URL <https://arxiv.org/abs/1911.02116>.
- Y. Cui, T. Liu, W. Che, L. Xiao, Z. Chen, W. Ma, S. Wang, and G. Hu. A span-extraction dataset for chinese machine reading comprehension. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*. Association for Computational Linguistics, 2019. doi: 10.18653/v1/d19-1600. URL <http://dx.doi.org/10.18653/v1/D19-1600>.
- O. de Gibert, G. Nail, N. Arefyev, M. Bañón, J. van der Linde, S. Ji, J. Zaragoza-Bernabeu, M. Aulamo, G. Ramírez-Sánchez, A. Kutuzov, S. Pyysalo, S. Oepen, and J. Tiedemann. A new massive multilingual dataset for high-performance language technologies. In N. Calzolari, M.-Y. Kan, V. Hoste, A. Lenci, S. Sakti, and N. Xue, editors, *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 1116–1128, Torino, Italia, May 2024. ELRA and ICCL. URL <https://aclanthology.org/2024.lrec-main.100>.
- O. De Gibert, G. Nail, N. Arefyev, M. Bañón, J. Van Der Linde, S. Ji, J. Zaragoza-Bernabeu, M. Aulamo, G. Ramírez-Sánchez, A. Kutuzov, et al. A new massive multilingual dataset for high-performance language technologies. *arXiv preprint arXiv:2403.14009*, 2024.
- DeepSeek-AI, :, X. Bi, D. Chen, G. Chen, S. Chen, D. Dai, C. Deng, H. Ding, K. Dong, Q. Du, Z. Fu, H. Gao, K. Gao, W. Gao, R. Ge, K. Guan, D. Guo, J. Guo, G. Hao, Z. Hao, Y. He, W. Hu, P. Huang, E. Li, G. Li, J. Li, Y. Li, Y. K. Li, W. Liang, F. Lin, A. X. Liu, B. Liu, W. Liu, X. Liu, X. Liu, Y. Liu, H. Lu, S. Lu, F. Luo, S. Ma, X. Nie, T. Pei, Y. Piao, J. Qiu, H. Qu, T. Ren, Z. Ren, C. Ruan, Z. Sha, Z. Shao, J. Song, X. Su, J. Sun, Y. Sun, M. Tang, B. Wang, P. Wang, S. Wang, Y. Wang, Y. Wang, T. Wu, Y. Wu, X. Xie, Z. Xie, Z. Xie, Y. Xiong, H. Xu, R. X. Xu, Y. Xu, D. Yang, Y. You, S. Yu, X. Yu, B. Zhang, H. Zhang, L. Zhang, L. Zhang, M. Zhang, M. Zhang, W. Zhang, Y. Zhang, C. Zhao, Y. Zhao, S. Zhou, S. Zhou, Q. Zhu, and Y. Zou. DeepSeek LLM: Scaling Open-Source Language Models with Longtermism, 2024. URL <https://arxiv.org/abs/2401.02954>.
- M. A. Desai, I. V. Pasquetto, A. Z. Jacobs, and D. Card. An archival perspective on pretraining data. *Patterns*, 5(4):100966, 2024. ISSN 2666-3899. doi: <https://doi.org/10.1016/j.patter.2024.100966>. URL <https://www.sciencedirect.com/science/article/pii/S2666389924000746>.
- J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding, 2019. URL <https://arxiv.org/abs/1810.04805>.
- M. d’Hoffschmidt, W. Belblidia, T. Brendlé, Q. Heinrich, and M. Vidal. FQuAD: French Question Answering Dataset, 2020. URL <https://arxiv.org/abs/2002.06071>.
- M. Finkelstein, D. Vilar, and M. Freitag. Introducing the newspalm mbr and qe dataset: Llm-generated high-quality parallel data outperforms traditional web-crawled data. In *Proceedings of the Ninth Conference on Machine Translation*, pages 1355–1372, 2024.
- S. Fischer, F. Rossetto, C. Gemmell, A. Ramsay, I. Mackie, P. Zubel, N. Tecklenburg, and J. Dalton. Open assistant toolkit–version 2. *arXiv preprint arXiv:2403.00586*, 2024.
- C. Fourier, N. Habib, T. Wolf, and L. Tunstall. LightEval: A lightweight framework for LLM evaluation, 2023. URL <https://github.com/huggingface/lighteval>.
- S. Gunasekar, Y. Zhang, J. Aneja, C. C. T. Mendes, A. D. Giorno, S. Gopi, M. Javaheripi, P. Kauffmann, G. de Rosa, O. Saarikivi, A. Salim, S. Shah, H. S. Behl, X. Wang, S. Bubeck, R. Eldan, A. T. Kalai, Y. T. Lee, and Y. Li. Textbooks are all you need, 2023. URL <https://arxiv.org/abs/2306.11644>.

- A. Hägele, E. Bakouch, A. Kosson, L. B. Allal, L. Von Werra, and M. Jaggi. Scaling laws and compute-optimal training beyond fixed training durations. *arXiv preprint arXiv:2405.18392*, 2024.
- M. Hardalov, T. Mihaylov, D. Zlatkova, Y. Dinkov, I. Koychev, and P. Nakov. Exams: A multi-subject high school examinations dataset for cross-lingual and multilingual question answering, 2020. URL <https://arxiv.org/abs/2011.03080>.
- W. Held, B. Paranjape, P. S. Koura, M. Lewis, F. Zhang, and T. Mihaylov. Optimizing pretraining data mixtures with llm-estimated utility. *arXiv preprint arXiv:2501.11747*, 2025.
- D. Hendrycks, C. Burns, S. Basart, A. Zou, M. Mazeika, D. Song, and J. Steinhardt. Measuring massive multitask language understanding. *arXiv preprint arXiv:2009.03300*, 2020.
- H. Hu, K. Richardson, L. Xu, L. Li, S. Kübler, and L. Moss. OCNLI: Original Chinese Natural Language Inference. In T. Cohn, Y. He, and Y. Liu, editors, *Findings of the Association for Computational Linguistics: EMNLP 2020*, pages 3512–3526, Online, Nov. 2020. Association for Computational Linguistics. doi: 10.18653/v1/2020.findings-emnlp.314. URL <https://aclanthology.org/2020.findings-emnlp.314/>.
- Y. Huang, Y. Bai, Z. Zhu, J. Zhang, J. Zhang, T. Su, J. Liu, C. Lv, Y. Zhang, J. Lei, Y. Fu, M. Sun, and J. He. C-eval: A multi-level multi-discipline chinese evaluation suite for foundation models, 2023. URL <https://arxiv.org/abs/2305.08322>.
- Hugging Face. Nanotron, 2024a. URL <https://github.com/huggingface/nanotron>. Accessed 30 Jan. 2025.
- Hugging Face. SmolLM - blazingly fast and remarkably powerful, 2024b. URL <https://huggingface.co/blog/smollm>. Accessed 30 Jan. 2025.
- M. Joshi, E. Choi, D. Weld, and L. Zettlemoyer. TriviaQA: A large scale distantly supervised challenge dataset for reading comprehension. In R. Barzilay and M.-Y. Kan, editors, *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1601–1611, Vancouver, Canada, July 2017. Association for Computational Linguistics. doi: 10.18653/v1/P17-1147. URL <https://aclanthology.org/P17-1147/>.
- A. Joulin, E. Grave, P. Bojanowski, and T. Mikolov. Bag of tricks for efficient text classification. In *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics: Volume 2, Short Papers*, pages 427–431. Association for Computational Linguistics, April 2017.
- F. Koto, H. Li, S. Shatnawi, J. Doughman, A. B. Sadallah, A. Alraeesi, K. Almubarak, Z. Alyafeai, N. Sengupta, S. Shehata, N. Habash, P. Nakov, and T. Baldwin. Arabicmmlu: Assessing massive multitask language understanding in arabic, 2024. URL <https://arxiv.org/abs/2402.12840>.
- S. Kudugunta, I. Caswell, B. Zhang, X. Garcia, C. A. Choquette-Choo, K. Lee, D. Xin, A. Kusupati, R. Stella, A. Bapna, and O. Firat. MADLAD-400: A Multilingual And Document-Level Large Audited Dataset, 2023. URL <https://arxiv.org/abs/2309.04662>.
- H. Kydlíček, G. Penedo, C. Fourier, N. Habib, and T. Wolf. FineTasks: Finding signal in a haystack of 200+ multilingual tasks, 2024. URL <https://huggingface.co/spaces/HuggingFaceFW/blogpost-fine-tasks>. Accessed 30 Jan. 2025.
- V. Lai, C. Nguyen, N. Ngo, T. Nguyen, F. Dernoncourt, R. Rossi, and T. Nguyen. Okapi: Instruction-tuned large language models in multiple languages with reinforcement learning from human feedback. In Y. Feng and E. Lefever, editors, *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 318–327, Singapore, Dec. 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.emnlp-demo.28. URL <https://aclanthology.org/2023.emnlp-demo.28/>.
- G. Lample and A. Conneau. Cross-lingual language model pretraining, 2019. URL <https://arxiv.org/abs/1901.07291>.

- H. Laurençon, L. Saulnier, T. Wang, C. Akiki, A. Villanova del Moral, T. Le Scao, L. Von Werra, C. Mou, E. González Ponferrada, H. Nguyen, et al. The bigscience roots corpus: A 1.6 tb composite multilingual dataset. *Advances in Neural Information Processing Systems*, 35:31809–31826, 2022.
- K. Lee, D. Ippolito, A. Nystrom, C. Zhang, D. Eck, C. Callison-Burch, and N. Carlini. Deduplicating training data makes language models better. In S. Muresan, P. Nakov, and A. Villavicencio, editors, *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 8424–8445, Dublin, Ireland, May 2022. Association for Computational Linguistics. doi: 10.18653/v1/2022.acl-long.577. URL <https://aclanthology.org/2022.acl-long.577/>.
- P. Lewis, B. Oğuz, R. Rinott, S. Riedel, and H. Schwenk. Mlqa: Evaluating cross-lingual extractive question answering, 2020. URL <https://arxiv.org/abs/1910.07475>.
- H. Li, Y. Zhang, F. Koto, Y. Yang, H. Zhao, Y. Gong, N. Duan, and T. Baldwin. Cmmlu: Measuring massive multitask language understanding in chinese, 2024a. URL <https://arxiv.org/abs/2306.09212>.
- J. Li, A. Fang, G. Smyrnis, M. Ivgi, M. Jordan, S. Gadre, H. Bansal, E. Guha, S. Keh, K. Arora, et al. DataComp-LM: In search of the next generation of training sets for language models. *arXiv preprint arXiv:2406.11794*, 2024b.
- B. Y. Lin, S. Lee, X. Qiao, and X. Ren. Common sense beyond English: Evaluating and improving multilingual language models for commonsense reasoning. In C. Zong, F. Xia, W. Li, and R. Navigli, editors, *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 1274–1287, Online, Aug. 2021a. Association for Computational Linguistics. doi: 10.18653/v1/2021.acl-long.102. URL <https://aclanthology.org/2021.acl-long.102/>.
- X. V. Lin, T. Mihaylov, M. Artetxe, T. Wang, S. Chen, D. Simig, M. Ott, N. Goyal, S. Bhosale, J. Du, R. Pasunuru, S. Shleifer, P. S. Koura, V. Chaudhary, B. O’Horo, J. Wang, L. Zettlemoyer, Z. Kozareva, M. T. Diab, V. Stoyanov, and X. Li. Few-shot learning with multilingual language models. *CoRR*, abs/2112.10668, 2021b. URL <https://arxiv.org/abs/2112.10668>.
- Llama Team. The Llama 3 Herd of Models, 2024. URL <https://arxiv.org/abs/2407.21783>.
- I. Loshchilov and F. Hutter. Decoupled weight decay regularization. 2019. URL <https://arxiv.org/abs/1711.05101>.
- M. Marion, A. Üstün, L. Pozzobon, A. Wang, M. Fadaee, and S. Hooker. When less is more: Investigating data pruning for pretraining llms at scale, 2023. URL <https://arxiv.org/abs/2309.04564>.
- P. H. Martins, J. Alves, P. Fernandes, N. M. Guerreiro, R. Rei, A. Farajian, M. Klimaszewski, D. M. Alves, J. Pombal, N. Boizard, et al. Eurollm-9b: Technical report. *arXiv preprint arXiv:2506.04079*, 2025.
- T. Mihaylov, P. Clark, T. Khot, and A. Sabharwal. Can a suit of armor conduct electricity? a new dataset for open book question answering, 2018. URL <https://arxiv.org/abs/1809.02789>.
- T. Mikolov, K. Chen, G. Corrado, and J. Dean. Efficient estimation of word representations in vector space, 2013. URL <https://arxiv.org/abs/1301.3781>.
- Mistral AI. v3 (tekken) tokenizer, 2024. URL <https://docs.mistral.ai/guides/tokenization/>. Accessed 30 Jan. 2025.
- Mistral AI. Mistral small 3, 2025. URL <https://mistral.ai/news/mistral-small-3/>. Accessed 30 Jan. 2025.
- H. Mozannar, K. E. Hajal, E. Maamary, and H. Hajj. Neural arabic question answering, 2019. URL <https://arxiv.org/abs/1906.05394>.

- N. Muennighoff, A. M. Rush, B. Barak, T. L. Scao, A. Piktus, N. Tazi, S. Pyysalo, T. Wolf, and C. Raffel. Scaling data-constrained language models, 2023. URL <https://arxiv.org/abs/2305.16264>.
- T. Nguyen, C. Van Nguyen, V. D. Lai, H. Man, N. T. Ngo, F. Dernoncourt, R. A. Rossi, and T. H. Nguyen. Culturax: A cleaned, enormous, and multilingual dataset for large language models in 167 languages. *arXiv preprint arXiv:2309.09400*, 2023.
- OpenAI. MMMLU, 2024. URL <https://huggingface.co/datasets/openai/MMMLU>. Accessed 30 Jan. 2025.
- P. J. Ortiz Suárez, B. Sagot, and L. Romary. Asynchronous pipelines for processing huge corpora on medium to low resource infrastructures. *Proceedings of the Workshop on Challenges in the Management of Large Corpora (CMLC-7) 2019*, Cardiff, 22nd July 2019, pages 9 – 16, Mannheim, 2019. Leibniz-Institut für Deutsche Sprache. doi: 10.14618/ids-pub-9021. URL <http://nbn-resolving.de/urn:nbn:de:bsz:mh39-90215>.
- G. Penedo, Q. Malartic, D. Hesslow, R. Cojocar, H. Alobeidli, A. Cappelli, B. Pannier, E. Almazrouei, and J. Launay. The RefinedWeb dataset for Falcon LLM: Outperforming curated corpora with web data only. *Advances in Neural Information Processing Systems*, 36:79155–79172, 2023.
- G. Penedo, H. Kydlíček, A. Lozhkov, M. Mitchell, C. Raffel, L. Von Werra, T. Wolf, et al. The FineWeb Datasets: Decanting the Web for the Finest Text Data at Scale. *arXiv preprint arXiv:2406.17557*, 2024a.
- G. Penedo, H. Kydlíček, A. Cappelli, M. Sasko, and T. Wolf. DataTrove: large scale data processing, 2024b. URL <https://github.com/huggingface/datatrove>. Accessed 30 Jan. 2025.
- G. Penedo, H. Kydlíček, V. Sabolčec, B. Messmer, N. Foroutan, M. Jaggi, L. von Werra, and T. Wolf. FineWeb2: A sparkling update with 1000s of languages, Dec. 2024c. URL <https://huggingface.co/datasets/HuggingFaceFW/fineweb-2>. Accessed 30 Jan. 2025.
- J. Pennington, R. Socher, and C. Manning. GloVe: Global vectors for word representation. In A. Moschitti, B. Pang, and W. Daelemans, editors, *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 1532–1543, Doha, Qatar, Oct. 2014. Association for Computational Linguistics. doi: 10.3115/v1/D14-1162. URL <https://aclanthology.org/D14-1162/>.
- J.-T. Peter, D. Vilar, D. Deutsch, M. Finkelstein, J. Juraska, and M. Freitag. There’s no data like better data: Using qe metrics for mt data filtering. In *Proceedings of the Eighth Conference on Machine Translation*, pages 561–577, 2023.
- Pluto-Junzeng. pluto-junzeng/chinesesquad, 2019. URL <https://github.com/pluto-junzeng/ChineseSquad>. Accessed 30 Jan. 2025.
- E. M. Ponti, G. Glavaš, O. Majewska, Q. Liu, I. Vulić, and A. Korhonen. XCOPA: A multilingual dataset for causal commonsense reasoning. In B. Webber, T. Cohn, Y. He, and Y. Liu, editors, *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 2362–2376, Online, Nov. 2020. Association for Computational Linguistics. doi: 10.18653/v1/2020.emnlp-main.185. URL <https://aclanthology.org/2020.emnlp-main.185/>.
- A. Radford, K. Narasimhan, T. Salimans, and I. Sutskever. Improving language understanding by generative pre-training. 2018.
- J. W. Rae, S. Borgeaud, T. Cai, K. Millican, J. Hoffmann, F. Song, J. Aslanides, S. Henderson, R. Ring, S. Young, et al. Scaling language models: Methods, analysis & insights from training gopher. *arXiv preprint arXiv:2112.11446*, 2021.
- C. Raffel, N. Shazeer, A. Roberts, K. Lee, S. Narang, M. Matena, Y. Zhou, W. Li, and P. J. Liu. Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of machine learning research*, 21(140):1–67, 2020.

- A. Romanou, N. Foroutan, A. Sotnikova, Z. Chen, S. H. Nelaturu, S. Singh, R. Maheshwary, M. Altomare, M. A. Haggag, A. Amayuelas, et al. Include: Evaluating multilingual language understanding with regional knowledge. *arXiv preprint arXiv:2411.19799*, 2024.
- N. Sachdeva, B. Coleman, W.-C. Kang, J. Ni, L. Hong, E. H. Chi, J. Caverlee, J. McAuley, and D. Z. Cheng. How to train data-efficient llms, 2024. URL <https://arxiv.org/abs/2402.09668>.
- K. Sakaguchi, R. L. Bras, C. Bhagavatula, and Y. Choi. WinoGrande: An Adversarial Winograd Schema Challenge at Scale. *arXiv preprint arXiv:1907.10641*, 2019.
- P. Sen, A. F. Aji, and A. Saffari. Mintaka: A Complex, Natural, and Multilingual Dataset for End-to-End Question Answering, 2022. URL <https://arxiv.org/abs/2210.01613>.
- S. Singh, A. Romanou, C. Fourrier, D. I. Adelani, J. G. Ngui, D. Vila-Suero, P. Limkonchotiwat, K. Marchisio, W. Q. Leong, Y. Susanto, R. Ng, S. Longpre, W.-Y. Ko, M. Smith, A. Bosselut, A. Oh, A. F. T. Martins, L. Choshen, D. Ippolito, E. Ferrante, M. Fadaee, B. Ermiş, and S. Hooker. Global mmlu: Understanding and addressing cultural and linguistic biases in multilingual evaluation, 2024a. URL <https://arxiv.org/abs/2412.03304>.
- S. Singh, F. Vargus, D. Dsouza, B. F. Karlsson, A. Mahendiran, W.-Y. Ko, H. Shandilya, J. Patel, D. Mataciunas, L. OMahony, M. Zhang, R. Hettiarachchi, J. Wilson, M. Machado, L. S. Moura, D. Krzemiński, H. Fadaei, I. Ergün, I. Okoh, A. Alaagib, O. Mudannayake, Z. Alyafeai, V. M. Chien, S. Ruder, S. Guthikonda, E. A. Alghamdi, S. Gehrmann, N. Muennighoff, M. Bartolo, J. Kreutzer, A. Üstün, M. Fadaee, and S. Hooker. Aya dataset: An open-access collection for multilingual instruction tuning, 2024b.
- L. Soldaini, R. Kinney, A. Bhagia, D. Schwenk, D. Atkinson, R. Authur, B. Bogin, K. Chandu, J. Dumas, Y. Elazar, et al. Dolma: An open corpus of three trillion tokens for language model pretraining research. *arXiv preprint arXiv:2402.00159*, 2024.
- N. Subramani, S. Luccioni, J. Dodge, and M. Mitchell. Detecting personal information in training corpora: an analysis. In A. Ovalle, K.-W. Chang, N. Mehrabi, Y. Pruksachatkun, A. Galystan, J. Dhamala, A. Verma, T. Cao, A. Kumar, and R. Gupta, editors, *Proceedings of the 3rd Workshop on Trustworthy Natural Language Processing (TrustNLP 2023)*, pages 208–220, Toronto, Canada, July 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.trustnlp-1.18. URL <https://aclanthology.org/2023.trustnlp-1.18/>.
- K. Sun, D. Yu, D. Yu, and C. Cardie. Investigating prior knowledge for challenging Chinese machine reading comprehension. *Transactions of the Association for Computational Linguistics*, 8:141–155, 2020. doi: 10.1162/tac1_a_00305. URL <https://aclanthology.org/2020.tac1-1.10/>.
- A. Talmor, J. Herzig, N. Lourie, and J. Berant. CommonsenseQA: A question answering challenge targeting commonsense knowledge. In J. Burstein, C. Doran, and T. Solorio, editors, *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4149–4158, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics. doi: 10.18653/v1/N19-1421. URL <https://aclanthology.org/N19-1421/>.
- A. Tikhonov and M. Ryabinin. It’s all in the heads: Using attention heads as a baseline for cross-lingual transfer in commonsense reasoning, 2021. URL <https://arxiv.org/abs/2106.12066>.
- Together Computer. Redpajama: An open source recipe to reproduce llama training dataset, 2023. URL <https://github.com/togethercomputer/RedPajama-Data>. Accessed 30 Jan. 2025.
- A. K. Upadhyay and H. K. Upadhyay. Xnli 2.0: Improving xnli dataset and performance on cross lingual understanding (xlu), 2023. URL <https://arxiv.org/abs/2301.06527>.
- A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin. Attention is all you need, 2023. URL <https://arxiv.org/abs/1706.03762>.
- G. Wenzek, M.-A. Lachaux, A. Conneau, V. Chaudhary, F. Guzmán, A. Joulin, and E. Grave. Ccnet: Extracting high quality monolingual datasets from web crawl data. *arXiv preprint arXiv:1911.00359*, 2019.

- A. Wettig, A. Gupta, S. Malik, and D. Chen. Qurating: Selecting high-quality data for training language models, 2024. URL <https://arxiv.org/abs/2402.09739>.
- L. Xue, N. Constant, A. Roberts, M. Kale, R. Al-Rfou, A. Siddhant, A. Barua, and C. Raffel. mT5: A massively multilingual pre-trained text-to-text transformer. In K. Toutanova, A. Rumshisky, L. Zettlemoyer, D. Hakkani-Tur, I. Beltagy, S. Bethard, R. Cotterell, T. Chakraborty, and Y. Zhou, editors, *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 483–498, Online, June 2021. Association for Computational Linguistics. doi: 10.18653/v1/2021.naacl-main.41. URL <https://aclanthology.org/2021.naacl-main.41/>.
- R. Zellers, A. Holtzman, Y. Bisk, A. Farhadi, and Y. Choi. Hellaswag: Can a machine really finish your sentence?, 2019. URL <https://arxiv.org/abs/1905.07830>.
- W. Zhang, S. M. Aljunied, C. Gao, Y. K. Chia, and L. Bing. M3exam: A multilingual, multimodal, multilevel benchmark for examining large language models, 2023. URL <https://arxiv.org/abs/2306.05179>.
- W. Zhong, R. Cui, Y. Guo, Y. Liang, S. Lu, Y. Wang, A. Saied, W. Chen, and N. Duan. Agieval: A human-centric benchmark for evaluating foundation models, 2023. URL <https://arxiv.org/abs/2304.06364>.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: Claims are backed by experimental results in Section 4.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We discuss the limitations in Appendix B.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[NA\]](#)

Justification: The paper does not include theoretical results.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We provide a detailed description of the methods in Section 3, experimental setup in Section 4, evaluation benchmarks selection in Appendix D, and the codebase.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: We provide the dataset and the codebase with sufficient information to reproduce the main experimental results.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: We provide all details in Section 3 and Section 4.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [No]

Justification: Penedo et al. [2024c] show that standard deviation when pretraining large language models in different languages on multiple seeds is small. We do not provide error bars because training large language models is expensive and we use our compute allocation budget for the main experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).

- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We provide the type of compute workers and the number of GPU hours used in our experiments in Section 4, and the storage usage of the data in Appendix A.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We conduct our research in accordance with the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: We do not introduce societal impacts beyond those already associated with large language models.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.

- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: **[No]**

Justification: We base our dataset on the FineWeb-2 dataset which conforms to Common Crawl robots.txt opt-outs (at crawl time), removes personally identifiable content, and offers a form for requesting data removal.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: **[Yes]**

Justification: We credit all assets with their licenses in Appendix H and respect their licenses.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.

- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [Yes]

Justification: We document the experiments as a part of the dataset creation process in Section 4 and include dataset information in Appendix A.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: We do not perform crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: We do not perform research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.

- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: Beyond pretraining and evaluating large language models, the core methodology does not involve the use of large language models.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Dataset Information

Based on the results of our experiments, we create the dataset, named ***FineWeb2-HQ***, by filtering all available FineWeb-2 data (version 2.0.1) in 20 languages using the *MLP MKC⁺* approach with 10% retention rate. The statistics of the resulting dataset are presented in Table 8. We release the dataset under the *Open Data Commons Attribution License (ODC-By) v1.0* license at huggingface.co/datasets/epfml/FineWeb2-HQ.

The main use case of our dataset is LLM pretraining, however, the dataset may also be used for other natural language processing tasks.

Table 8: Statistics (number of documents and disk size) of the dataset resulting from filtering FineWeb-2 using the *MLP MKC⁺* approach with 10% retention rate in 20 languages.

Language	Number of documents	Disk size
Russian	55,220,956	1.2TB
Chinese	54,211,986	784GB
German	43,095,728	618GB
Spanish	40,057,637	515GB
Japanese	34,185,427	393GB
French	32,248,772	483GB
Italian	21,180,304	269GB
Portuguese	18,135,468	222GB
Polish	13,384,885	168GB
Dutch	12,920,963	160GB
Indonesian	8,911,149	125GB
Turkish	8,578,808	100GB
Czech	5,995,459	104GB
Arabic	5,560,599	94GB
Persian	5,107,187	69GB
Hungarian	4,527,332	79GB
Swedish	4,382,454	61GB
Greek	4,346,440	84GB
Danish	4,082,751	61GB
Vietnamese	4,003,956	59GB

B Limitations

A limitation of our work is that we perform experiments on relatively small 1B models with one seed per experiment. We use 1B models to balance the trade-off between the cost of pretraining and the measured signal from the experiments, as found in prior work [Penedo et al., 2024a,c, Li et al., 2024b]. Additionally, we compare our method to one multilingual baseline, FineWeb-2. However, since FineWeb-2 is the current state-of-the-art and due to our limited computational budget, we decided to allocate more compute towards understanding the mechanics of the data selection process, rather than confirming our results across previous datasets. Nevertheless, computational constraints prevented us from ablating every decision—such as our choice to use only the first 512 tokens for classification. We assume that if the first 512 tokens demonstrate good quality, the remainder of the document likely does as well. Given the strong performance achieved using the first 512 tokens, we prioritized this methodological simplicity. To facilitate further exploration of alternative selection strategies, we have made FineWeb2-embedded⁹ available to the community, which contains embeddings for all 512-token chunks.

Although we develop our framework on languages from diverse language families, with different writing systems and with varying resource availability to find an approach that best generalizes for general web crawl text data across languages, classifier training datasets have no quality guarantees for other languages and may result in performance differences that are not visible in our experiments.

Since we focus on simple methods with broad availability and low computational cost, we discuss the computational cost difference between FastText and Transformer embeddings-based methods.

⁹huggingface.co/datasets/epfml/FineWeb2-embedded

While FastText classifiers are cheap to train and inference and can be efficiently run on CPU, Transformer-based methods require an initial computation of embeddings. To mitigate the higher cost of Transformer embeddings, we use a relatively small XLM-RoBERTa model and additionally release the dataset with precomputed embeddings⁹. The total cost for computing the embeddings is approximately 4K compute hours for the 20 languages.

We base our dataset on the FineWeb-2 dataset which conforms to Common Crawl robots.txt opt-outs (at crawl time), removes personally identifiable content, and offers a form for requesting data removal. Since ensuring privacy and fairness of our dataset is beyond the scope of this work, we make the dataset publicly available. This allows other researchers and the public to analyze potential biases, a critical task given that data curation is a political process that can introduce cultural and political impacts [Desai et al., 2024].

C Additional Results

C.1 Model Selection - Per Language Results

For clarity, we present the individual benchmark results of the 1B-parameter model trained on 119B tokens for each language in the following tables: Table 9 for Chinese, Table 10 for French, Table 11 for German, Table 12 for Arabic, and Table 13 for Danish.

Table 9: Chinese Benchmark performance comparison: Average rank between FineWeb-2 baseline and our proposed filtering methods (*FT*, *MLP*, and *CS*) trained on *MKC*⁺ or *MKC*, retaining top 10% of documents. The average rank is computed across FineTasks for 1B-parameter models evaluated after 119B tokens.

Approach	<i>MLP MKC</i> ⁺	<i>MLP MKC</i>	<i>CS MKC</i>	<i>FT MKC</i>	<i>FT MKC</i> ⁺	Baseline	<i>CS MKC</i> ⁺
Average Rank	1.7333	2.4333	4.0667	4.0667	4.4667	5.2333	6.0000
AGIEval	0.2995	0.2948	0.2897	0.2919	0.2817	0.2853	0.2773
Belebele	0.3300	0.3233	0.3178	0.3133	0.3133	0.3056	0.3022
C ³	0.4550	0.4480	0.4400	0.4500	0.4400	0.4400	0.4370
C-Eval	0.3095	0.3060	0.2760	0.2903	0.2906	0.2878	0.2805
CMMLU	0.3312	0.3259	0.3041	0.3043	0.3060	0.3009	0.2995
CMRC 2018	0.2224	0.2125	0.1614	0.2251	0.2164	0.1949	0.1866
HellaSwag	0.3790	0.3800	0.3530	0.3680	0.3660	0.3510	0.3370
M3Exam	0.3319	0.3245	0.3084	0.3201	0.3245	0.3216	0.3245
X-CODAH	0.3033	0.3000	0.3233	0.3100	0.2900	0.2967	0.3067
X-CSQA	0.2740	0.2680	0.2690	0.2610	0.2520	0.2510	0.2650
XCOPA	0.6200	0.6400	0.6180	0.5740	0.5740	0.6000	0.5620
OCNLI	0.5470	0.5470	0.5340	0.5250	0.5600	0.5420	0.5060
Chinese-SQuAD	0.0929	0.1097	0.0865	0.0889	0.0850	0.0777	0.0585
XStoryCloze	0.5800	0.5630	0.5710	0.5560	0.5610	0.5580	0.5570
XWINO	0.6429	0.6528	0.6587	0.6131	0.5992	0.6429	0.6111

C.2 Threshold Selection

Complete Result. To confirm that the *CS* filtering method is not competitive with *MLP* and *FT*, even when a higher percentage of documents is retained, we present the complete threshold selection results, including the *CS* method, in Table 14 in addition to the results shown in Section 4.2.2 (Table 2).

Document Length Bias. Motivated by the observed bias in certain approaches favoring the selection of shorter documents, as seen in Figure 3, Figure 4 and Table 15, we examine how this bias interacts with performance when retaining more documents. As demonstrated in Table 15, the *MLP MKC* approach shows a tendency to retain shorter documents, while achieving higher performance with an increased number of retained documents. In contrast, the *CS* and *FT* filtering methods present mixed results, suggesting that the optimal threshold selection may be influenced by additional factors.

Table 10: French Benchmark performance comparison: Average rank between FineWeb-2 baseline and our proposed filtering methods (*FT*, *MLP*, and *CS*) trained on *MKC*⁺ or *MKC*, retaining top 10% of documents. The average rank is computed across FineTasks for 1B-parameter models evaluated after 119B tokens.

Approach	<i>FT MKC</i> ⁺	<i>MLP MKC</i> ⁺	<i>MLP MKC</i>	<i>FT MKC</i>	<i>CS MKC</i>	<i>CS MKC</i> ⁺	Baseline
Average Rank	3.2222	3.5000	3.5556	3.7778	4.0000	4.6667	5.2778
Belebele	0.3378	0.3533	0.3678	0.3489	0.3444	0.3344	0.3444
HellaSwag	0.5380	0.5380	0.4990	0.5150	0.5280	0.5070	0.5180
X-CSQA	0.2820	0.2740	0.2730	0.2990	0.2850	0.2900	0.2870
XNLI 2.0	0.7340	0.7400	0.7430	0.7230	0.7450	0.7330	0.7180
FQuAD	0.2597	0.2803	0.3032	0.2981	0.2411	0.2476	0.2401
MMLU	0.2896	0.2895	0.2925	0.2886	0.2806	0.2815	0.2706
Mintaka	0.0710	0.0438	0.0334	0.0670	0.0610	0.0976	0.0712
X-CODAH	0.3000	0.2667	0.2867	0.2767	0.3000	0.2800	0.2633
ARC (Challenge)	0.3120	0.3180	0.3090	0.3060	0.2950	0.2830	0.2850

Table 11: German Benchmark performance comparison: Average rank between FineWeb-2 baseline and our proposed filtering methods (*FT*, *MLP*, and *CS*) trained on *MKC*⁺ or *MKC*, retaining top 10% of documents. The average rank is computed across FineTasks for 1B-parameter models evaluated after 119B tokens.

Approach	<i>MLP MKC</i> ⁺	<i>FT MKC</i> ⁺	<i>FT MKC</i>	<i>CS MKC</i>	<i>MLP MKC</i>	<i>CS MKC</i> ⁺	Baseline
Average Rank	3.1250	3.1250	3.5000	3.7500	4.5000	4.7500	5.2500
MMLU	0.2940	0.2879	0.2926	0.2770	0.2905	0.2764	0.2718
ARC (Challenge)	0.2760	0.2850	0.2820	0.2880	0.2830	0.2640	0.2680
Mintaka	0.0580	0.0548	0.0735	0.0576	0.0494	0.0766	0.0498
Belebele	0.3611	0.3578	0.3544	0.3544	0.3567	0.3422	0.3544
X-CODAH	0.3367	0.3500	0.3300	0.3567	0.3400	0.3600	0.3467
X-CSQA	0.2978	0.3008	0.2877	0.2887	0.2857	0.2918	0.2787
HellaSwag	0.4640	0.4710	0.4870	0.4820	0.4540	0.4390	0.4470
XNLI 2.0	0.6620	0.6530	0.6740	0.6440	0.6610	0.6520	0.6890

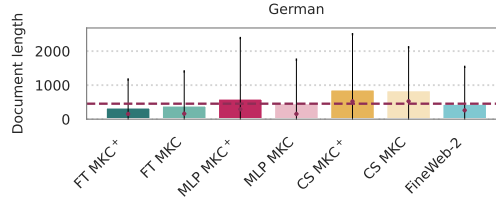


Figure 3: Document length comparison: Average length and standard deviation in FineWeb-2 before and after 10% retention filtering. Red horizontal line shows average document length, red dots indicate medians. Length measured by space-separated tokens.

C.3 Training Data Analysis

We give details on the variation in the average length of documents retained by our model-based filtering method *MLP* for Chinese, French, Arabic, and Danish with different training datasets. The results are shown for German in Figure 5 and for all other languages in Figure 6.

C.4 Replay of Original Data

We explore whether incorporating a small percentage of original raw data (replay) can help improve performance. We do this for our best FastText (*FT MKC*⁺) and Transformer approaches (*MLP MKC*⁺). Table 16 presents the results of experiments where 5% and 10% unfiltered data were mixed

Table 12: Arabic Benchmark performance comparison: Average rank between FineWeb-2 baseline and our proposed filtering methods (*FT*, *MLP*, and *CS*) trained on *MKC*⁺ or *MKC*, retaining top 56% of documents. The average rank is computed across FineTasks for 1B-parameter models evaluated after 119B tokens.

Approach	<i>MLP MKC</i> ⁺	<i>MLP MKC</i>	<i>FT MKC</i> ⁺	Baseline	<i>CS MKC</i> ⁺	<i>CS MKC</i>	<i>FT MKC</i>
Average Rank	2.7812	3.2500	3.6875	3.9688	3.9688	5.0312	5.3125
EXAMS	0.3537	0.3656	0.3552	0.3582	0.3443	0.3262	0.3346
MMLU	0.4007	0.3909	0.4023	0.3894	0.3912	0.3781	0.3885
ARC (Easy)	0.4330	0.4230	0.4210	0.4120	0.4020	0.3940	0.4080
AlGhafa SciQ	0.6915	0.7005	0.6965	0.6854	0.6724	0.6683	0.6804
Belebele	0.3456	0.3356	0.3322	0.3311	0.3356	0.3567	0.3233
SOQAL	0.7333	0.6867	0.7000	0.7200	0.7267	0.6867	0.7133
MLQA	0.2386	0.2402	0.1928	0.1901	0.2189	0.2154	0.1793
TyDi QA	0.1547	0.1476	0.1230	0.1441	0.1223	0.1097	0.1182
AlGhafa RACE	0.3720	0.3740	0.3640	0.3710	0.3590	0.3660	0.3730
ARCD	0.3638	0.3505	0.3235	0.3354	0.3358	0.3432	0.3043
X-CODAH	0.2600	0.2533	0.2567	0.2633	0.2633	0.2500	0.2600
AlGhafa PIQA	0.6360	0.6320	0.6400	0.6240	0.6320	0.6320	0.6370
X-CSQA	0.2740	0.2810	0.2770	0.2900	0.2880	0.2720	0.2770
XNLI 2.0	0.6570	0.6910	0.6990	0.7010	0.6910	0.6900	0.6770
HellaSwag	0.4270	0.4220	0.4280	0.4250	0.4260	0.4320	0.4150
XStoryCloze	0.6150	0.6100	0.6100	0.6070	0.6130	0.6180	0.5930

Table 13: Danish Benchmark performance comparison: Average rank between FineWeb-2 baseline and our proposed filtering methods (*FT*, *MLP*, and *CS*) trained on *MKC*⁺ or *MKC*, retaining top 65% of documents. The average rank is computed across FineTasks for 1B-parameter models evaluated after 119B tokens.

Approach	<i>CS MKC</i> ⁺	<i>MLP MKC</i> ⁺	<i>FT MKC</i> ⁺	Baseline
Average Rank	1.0000	2.5000	3.1667	3.3333
ARC (Challenge)	0.2820	0.2650	0.2730	0.2560
HellaSwag	0.4950	0.4850	0.4750	0.4750
Belebele	0.3333	0.3289	0.3189	0.3289

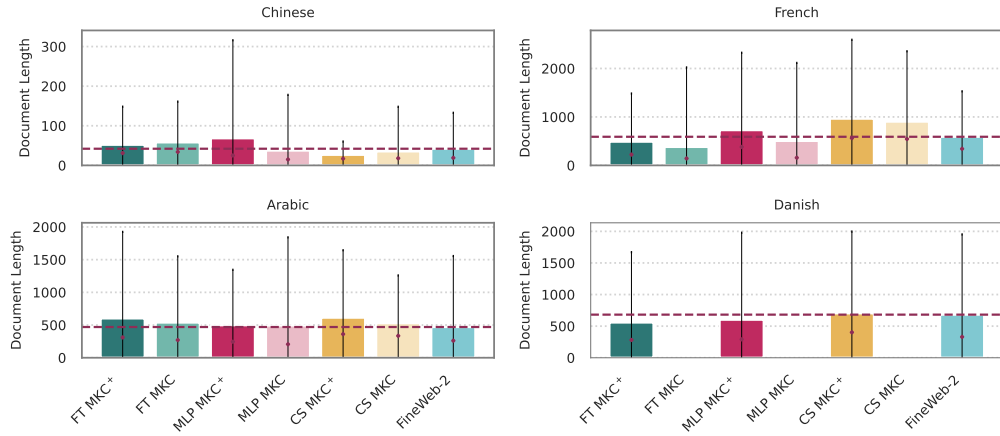


Figure 4: Document length comparison: Average length and standard deviation in FineWeb-2 before and after 10% retention filtering. Red horizontal line shows average document length, red dots indicate medians. Length measured by space-separated tokens.

into the training dataset, alongside results from training without any replay. Although, the *FT MKC*⁺ filters shows mixed signal, our *MLP MKC*⁺ approach clearly demonstrates that replay does not

Table 14: Benchmark performance comparison: Average rank between FineWeb-2 baseline and our proposed filtering methods (*FT*, *MLP*) trained on *MKC*⁺ or *MKC*, retaining top 10%, 15% or 20% of documents. The average rank is computed across FineTasks for 1B-parameter models evaluated on Chinese, German and French after 70B and 119B tokens.

Approach	Threshold	Average Rank
<i>MLP MKC</i> ⁺	10%	11.73
<i>MLP MKC</i> ⁺	15%	12.13
<i>MLP MKC</i>	20%	15.07
<i>MLP MKC</i>	15%	15.09
<i>MLP MKC</i> ⁺	20%	15.40
<i>MLP MKC</i>	10%	16.09
<i>FT MKC</i> ⁺	10%	18.61
<i>CS MKC</i>	15%	19.02
<i>CS MKC</i>	20%	19.24
<i>FT MKC</i>	15%	19.84
<i>FT MKC</i>	10%	20.02
<i>CS MKC</i>	10%	20.67
<i>FT MKC</i>	20%	20.80
<i>FT MKC</i> ⁺	15%	22.05
<i>FT MKC</i> ⁺	20%	22.52
<i>CS MKC</i> ⁺	15%	24.66
<i>CS MKC</i> ⁺	20%	25.08
Baseline	–	25.54
<i>CS MKC</i> ⁺	10%	26.94

Table 15: Token retention comparison: Counts in FineWeb-2 before and after filtering using our approach with 10% document retention for Chinese, French and German, 56% for Arabic, and 65% for Danish. Token counts represent tokenized dataset sizes using the multilingual Mistral v3 (Tekken) tokenizer [Mistral AI, 2024].

Approach	Chinese	French	German	Arabic	Danish
<i>MLP MKC</i> ⁺	150B (9%)	89B (12%)	119B (12%)	78B (61%)	71B (66%)
<i>MLP MKC</i>	105B (7%)	72B (10%)	87B (9%)	75B (59%)	–
<i>FT MKC</i> ⁺	221B (14%)	70B (10%)	63B (6%)	77B (61%)	70B (65%)
<i>FT MKC</i>	190B (12%)	43B (6%)	65B (7%)	80B (63%)	–
<i>CS MKC</i> ⁺	170B (11%)	126B (17%)	166B (17%)	82B (65%)	77B (71%)
<i>CS MKC</i>	161B (10%)	132B (18%)	172B (18%)	83B (65%)	–
Baseline	1597B	730B	973B	127B	108B

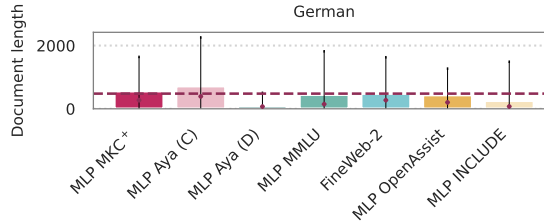


Figure 5: Document length comparison: Average length and standard deviation in FineWeb-2 before and after filtering using *MLP* method with 10% retention on different training datasets. Red horizontal line shows average document length, red dots indicate medians. Length measured by space-separated tokens.

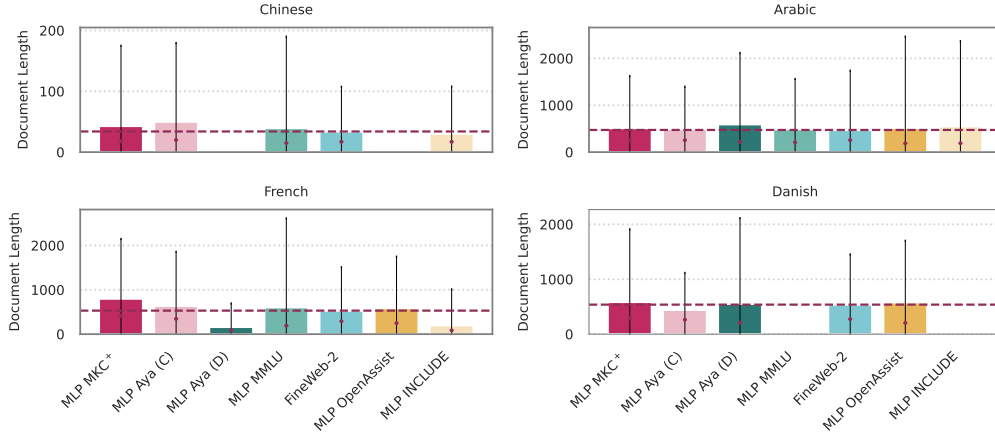


Figure 6: Document length comparison: Average length and standard deviation in FineWeb-2 before and after filtering using *MLP* method with 10% retention for Chinese and French, 56% for Arabic and 65% for Danish on different training datasets. Red horizontal line shows average document length, red dots indicate medians. Length measured by space-separated tokens.

improve performance, indicating the data selection already retains enough diversity. In cases of less diverse datasets, replay was shown to offer benefits [Bethune et al., 2025, Chen et al., 2023].

Table 16: Benchmark performance comparison: Average rank of our *MLP MKC*⁺ and *FT MKC*⁺ approaches with 10% document retention, mixed with 0%, 5%, or 10% of original FineWeb-2 dataset. The average rank is computed across FineTasks for 1B-parameter models evaluated on Chinese, German and French after 70B and 119B tokens.

Approach	Mixture Rate	Average Rank
<i>MLP MKC</i> ⁺	5%	5.09
<i>MLP MKC</i> ⁺	0%	5.16
<i>MLP MKC</i> ⁺	10%	5.40
<i>FT MKC</i> ⁺	10%	7.17
<i>FT MKC</i> ⁺	0%	7.51
<i>FT MKC</i> ⁺	5%	8.66

C.5 Impact on multilingual model training

This section presents the results of our *MLP MKC*⁺ approach on multilingual model training for Chinese (Table 17), Arabic (Table 18), German (Table 19), and Danish (Table 20), in addition to the results for French discussed in Section 4.2.6.

Table 17: Chinese benchmark performance: Multilingual LLMs (M) trained on FineWeb-2 or our $MLP MKC^+$ refined dataset (retaining top 10% for Chinese, German and French, 56% for Arabic, 65% for Danish) with 595B tokens, compared to monolingual models trained on 119B tokens. The average rank is computed across FineTasks for 1B-parameter models.

Dataset	Ours	Ours $_M$	FW-2 $_M$	FW-2
Average Rank	1.5667	2.1667	2.9000	3.3667
AGIEval	0.2995	0.2863	0.2894	0.2853
Belebele	0.3300	0.3456	0.3189	0.3056
C ³	0.4550	0.4520	0.4480	0.4400
C-Eval	0.3095	0.2848	0.2683	0.2878
CMMLU	0.3312	0.3064	0.2967	0.3009
CMRC 2018	0.2224	0.2689	0.2090	0.1949
HellaSwag	0.3790	0.3740	0.3740	0.3510
M3Exam	0.3319	0.3040	0.3304	0.3216
X-CODAH	0.3033	0.3067	0.2800	0.2967
X-CSQA	0.2740	0.2810	0.2780	0.2510
XCOPA	0.6200	0.6020	0.5860	0.6000
OCNLI	0.5470	0.5320	0.4910	0.5420
Chinese-SQuAD	0.0929	0.1304	0.1017	0.0777
XStoryCloze	0.5800	0.5760	0.5650	0.5580
XWINO	0.6429	0.6409	0.6468	0.6429

Table 18: Arabic benchmark performance: Multilingual LLMs (M) trained on FineWeb-2 or our $MLP MKC^+$ refined dataset (retaining top 10% for Chinese, German and French, 56% for Arabic, 65% for Danish) with 595B tokens, compared to monolingual models trained on 119B tokens. The average rank is computed across FineTasks for 1B-parameter models.

Dataset	Ours $_M$	Ours	FW-2	FW-2 $_M$
Average Rank	1.9688	2.0000	2.7500	3.2812
EXAMS	0.3336	0.3537	0.3582	0.3076
MMLU	0.3828	0.4007	0.3894	0.3599
ARC (Easy)	0.4190	0.4330	0.4120	0.3760
AlGhafa SciQ	0.6764	0.6915	0.6854	0.6563
Belebele	0.3511	0.3456	0.3311	0.3344
SOQAL	0.7000	0.7333	0.7200	0.6533
MLQA	0.2208	0.2386	0.1901	0.2085
TyDi QA	0.1634	0.1547	0.1441	0.1429
AlGhafa RACE	0.3830	0.3720	0.3710	0.3770
ARCD	0.3377	0.3638	0.3354	0.2970
X-CODAH	0.2767	0.2600	0.2633	0.2767
AlGhafa PIQA	0.6170	0.6360	0.6240	0.6160
X-CSQA	0.2860	0.2740	0.2900	0.2660
XNLI 2.0	0.7080	0.6570	0.7010	0.7340
HellaSwag	0.4390	0.4270	0.4250	0.4240
XStoryCloze	0.6370	0.6150	0.6070	0.6160

Table 19: German benchmark performance: Multilingual LLMs (M) trained on FineWeb-2 or our *MLP MKC⁺* refined dataset (retaining top 10% for Chinese, German and French, 56% for Arabic, 65% for Danish) with 595B tokens, compared to monolingual models trained on 119B tokens. The average rank is computed across FineTasks for 1B-parameter models.

Dataset	Ours _{M}	Ours	FW-2	FW-2 _{M}
Average Rank	1.5000	2.1250	2.9375	3.4375
MMLU	0.2918	0.2940	0.2718	0.2691
ARC (Challenge)	0.2740	0.2760	0.2680	0.2640
Mintaka	0.0821	0.0580	0.0498	0.0660
Belebele	0.3956	0.3611	0.3544	0.3633
X-CODAH	0.3500	0.3367	0.3467	0.3167
X-CSQA	0.3048	0.2978	0.2787	0.2787
HellaSwag	0.4690	0.4640	0.4470	0.4430
XNLI 2.0	0.6420	0.6620	0.6890	0.6340

Table 20: Danish benchmark performance: Multilingual LLMs (M) trained on FineWeb-2 or our *MLP MKC⁺* refined dataset (retaining top 10% for Chinese, German and French, 56% for Arabic, 65% for Danish) with 595B tokens, compared to monolingual models trained on 119B tokens. The average rank is computed across FineTasks for 1B-parameter models.

Dataset	Ours _{M}	Ours	FW-2 _{M}	FW-2
Average Rank	1.6667	2.1667	3.0000	3.1667
ARC (Challenge)	0.2920	0.2650	0.2600	0.2560
HellaSwag	0.4710	0.4850	0.4560	0.4750
Belebele	0.3700	0.3289	0.3311	0.3289

D List of evaluation benchmarks and metrics

We provide a detailed overview of the evaluation benchmarks used to assess our models’ performance, along with their respective evaluation metrics in Table 21. For non-English tasks and English MMLU, we use the *cloze* multiple-choice prompt, which allows the model to directly predict each option instead of using the standard prompt format with A/B/C/D letter prefixes as targets. This approach was chosen because it has been shown to serve as a more reliable performance indicator earlier in training [Kydliček et al., 2024]. We evaluate the models in a 0-shot setting.

Table 21: List of Evaluation Benchmarks and Metrics used in our setup for Chinese, French, German, Arabic, Danish, and English.

Benchmark	Chinese	French	German	Arabic	Danish	English	Evaluation metric
AGIEval [Zhong et al., 2023]	✓						Normalized accuracy
AlGhafa ARC [Almazrouei et al., 2023]				✓			Normalized accuracy
AlGhafa PIQA [Almazrouei et al., 2023]				✓			Normalized accuracy
AlGhafa RACE [Almazrouei et al., 2023]				✓			Normalized accuracy
AlGhafa SciQ [Almazrouei et al., 2023]				✓			Normalized accuracy
ArabicMMLU [Koto et al., 2024]				✓			Normalized accuracy
ARC [Clark et al., 2018]						✓	Normalized accuracy
ARCD [Mozannar et al., 2019]				✓			F1 score
Belebele [Bandarkar et al., 2024]	✓	✓	✓	✓	✓		Normalized accuracy
C ³ [Sun et al., 2020]	✓						Normalized accuracy
C-Eval [Huang et al., 2023]	✓						Normalized accuracy
Chinese-SQuAD [Pluto-Junzeng, 2019]	✓						F1 score
CMMLU [Li et al., 2024a]	✓						Normalized accuracy
CMRC 2018 [Cui et al., 2019]	✓						F1 score
CommonsenseQA [Talmor et al., 2019]						✓	Normalized accuracy
EXAMS [Herdalov et al., 2020]				✓			Normalized accuracy
FQuAD [d’Hoffschmidt et al., 2020]		✓					F1 score
HellaSwag [Zellers et al., 2019]						✓	Normalized accuracy
M3Exam [Zhang et al., 2023]	✓						Normalized accuracy
Meta MMLU [Llama Team, 2024]		✓	✓				Normalized accuracy
Mintaka [Sen et al., 2022]		✓	✓				F1 score
MLMM ARC [Lai et al., 2023]		✓	✓		✓		Normalized accuracy
MLMM HellaSwag [Lai et al., 2023]	✓	✓	✓	✓	✓		Normalized accuracy
MLQA [Lewis et al., 2020]				✓			F1 score
MMLU [Hendrycks et al., 2020]						✓	Normalized accuracy
OCNLI [Hu et al., 2020]	✓						Normalized accuracy
OpenBookQA [Mihaylov et al., 2018]						✓	Normalized accuracy
PIQA [Bisk et al., 2019]						✓	Normalized accuracy
SOQAL [Mozannar et al., 2019]				✓			Normalized accuracy
TriviaQA [Joshi et al., 2017]						✓	Quasi-exact match
TyDi QA [Clark et al., 2020]				✓			F1 score
WinoGrande [Sakaguchi et al., 2019]						✓	Normalized accuracy
X-CODAH [Lin et al., 2021a]	✓	✓	✓	✓			Normalized accuracy
XCOPA [Ponti et al., 2020]	✓						Normalized accuracy
X-CSQA [Lin et al., 2021a]	✓	✓	✓	✓			Normalized accuracy
XNLI 2.0 [Upadhyay and Upadhyay, 2023]		✓	✓	✓			Normalized accuracy
XStoryCloze [Lin et al., 2021b]	✓			✓			Normalized accuracy
XWINO [Tikhonov and Ryabinin, 2021]	✓						Normalized accuracy

E Average Rank Computation

Analogous to the method in FineTasks [Kydliček et al., 2024], we compute the average rank for our ablations as follows:

1. We train a model for each parameter configuration we want to ablate on.
2. We evaluate each model on all the selected benchmarks.
3. We compute the rank of each model (individual experiment) with regard to each benchmark and language.
4. We compute the average rank for each model across all benchmarks and languages.

F FineWeb documents in different scoring approaches

To illustrate the types of documents each classifier scores highly or poorly, we present the highest- and lowest-scoring FineWeb examples for each of our classifier approaches (*FT MKC*⁺, *MLP MKC*⁺, *CS MKC*⁺). These examples were selected from the randomly chosen FineWeb test dataset (10K samples) used to validate the training of our model-based classifiers.

F.1 FastText Classifier (FT)

Highest score:

hi. i couldn't solve my problem because it has two conditional logical propositions. the problem is:can anyone help me about this, thanks =)we're expected to know that: . is equivalent to find a logically equivalent proposition for:by first writing its contrapositive, and then applying demorgan's lawand the equality forthey were trying to be helpful by outlining the steps we should follow,. . but i think they made it more confusing.i don't see the purpose of using the contrapositive here.. . i wouldn't have done it that way.besides, the statement is a tautology . . .which gives us: .and this is a tautology: "a thing implies itself" ... which is always true.i don't know of any "logically equivalent proposition" we can write . . .

Lowest score:

lstarts||23 sep 2016 (fri) (one day only)||want to travel soon but don't wish to fork out a fortune for flights? check out today's promotion from jetstar featuring promo fares fr \$35 all-in valid for travel period commencing 12 october 2016don't miss out! all-in frenzy fares to hong kong, penang and more from \$35.sale ends 23 sep, 11pm! travelling||price||travel period||find flight||penang||\$35^|| [...]

F.2 Multi-Layer Perceptron (MLP)

Highest score:

Naqhadeh County is a county in West Azerbaijan Province in Iran. The capital of the county is Naqhadeh. At the 2006 census, the county's population was 117,831, in 27,937 families. The county is subdivided into two districts: the Central District and Mohammadyar District. The county has two cities: Naqhadeh and Mohammadyar.

Lowest score:

Custom Wedding Gifts
Personalized photo frames, albums & keepsakes. Heirloom quality!
Custom Engraved Journals
Handmade in Florence Italy. Dozens of sizes and paper styles!
Awesome Leather Journals
Personalized, Customizable, Artisan made in Santa Fe, NM.
Ink Rendering from Photos
100% Hand painted with unique style by pro artists. From \$49.

F.3 Cosine Similarity (CS)

Highest score:

When you are renting a 5, 10, 15, 20, 30 or 40 yard dumpster, you want a company you can trust with prices that make you smile. Give us a call today and see the difference we can make in your next construction or clean out project.
Simply give us a call and we will help you figure out your dumpster rental needs.
Our dumpsters usually go out same-day or next-day depending on when you call.
We provide top-notch service, while going easy on your bottom line. What more could you ask for?
Our trained operators are here to give you a fast and hassle-free experience from start to finish.[...]

Lowest score:

Cooperative flat 206/J
- Cooperative flat 201/J - Sold
2(1)+kitchenette, 50,1 m2Cooperative flat 202/J - Sold
2(1)+kitchenette, 44,9 m2Cooperative flat 203/J - Sold
2(1)+kitchenette, 50,6 m2Cooperative flat 204/J - Sold
1+kitchenette, 27,1 m2Cooperative flat 205/J - Sold
2(1)+kitchenette, 50,1 m2Cooperative flat 206/J - On sale
3+kitchenette 86,7 m2[...]

G Example of a contaminated document

We present an example of a FineWeb document that was removed during our decontamination pipeline.

MMLU contaminated document (matched 13-gram in bold):

Here is our diagram of the Preamble to the Constitution of the United States. It is based on our understanding of the use of "in order to" as a subordinating conjunction that introduces a series of infinitival clauses (without subjects) that, in turn, modify the compound verbs "do ordain" and "establish."
See A Grammar of Contemporary English by Randolph Quirk, Sidney Greenbaum, Geoffrey Leech, and Jan Svartvik. Longman Group: London. 1978. p. 753.
We the People of the United States, in Order to form a more perfect Union, establish Justice, insure domestic Tranquility, **provide for the common defence, promote the general Welfare, and secure the Blessings** of Liberty to ourselves and our Posterity, do ordain and establish this Constitution for the United States of America.
If you have alternative rendering for this sentence, we would be happy to hear of it. Use the e-mail icon to the left.

H License Information

H.1 Dataset Licenses

We use the following pretraining datasets:

- FineWeb-2 (ODC-By license)
- FineWeb (ODC-By license)
- FineWeb-Edu (ODC-By license)
- DCLM (CC-BY 4.0)

We use the following classifier training datasets:

- Aya Collection (Apache 2.0 license)
- Aya Dataset (Apache 2.0 license)
- Translated multilingual MMLU [OpenAI, 2024] (MIT license)
- OpenAssistant-2 (Apache 2.0 license)
- Include-Base-44 (Apache 2.0 license)

H.2 Code Licenses

We use the following open source code:

- Nanotron (Apache 2.0 license)
- Datatrove (Apache 2.0 license)
- Lighteval (MIT license)
- FastText (MIT license)

H.3 Model Licenses

We use the following models:

- Mistral v3 (Tekken) (Apache 2.0 license)
- XLM-RoBERTa (MIT license)