

# Parallel Tokenizers: Rethinking Vocabulary Design for Cross-Lingual Transfer

Muhammad Dehan Al Kautsar Fajri Koto

Mohamed bin Zayed University of Artificial Intelligence

{muhammad.dehan,fajri.koto}@mbzuai.ac.ae

## Abstract

Tokenization defines the foundation of multilingual language models by determining how words are represented and shared across languages. However, existing methods often fail to support effective cross-lingual transfer because semantically equivalent words are assigned distinct embeddings. For example, “*I eat rice*” in English and “*Ina cin shinkafa*” in Hausa are typically mapped to different vocabulary indices, preventing shared representations and limiting cross-lingual generalization. We introduce **parallel tokenizers**. This new framework trains tokenizers monolingually and then aligns their vocabularies exhaustively using bilingual dictionaries or word-to-word translation, ensuring consistent indices for semantically equivalent words. This alignment enforces a shared semantic space across languages while naturally improving fertility balance. To assess their effectiveness, we pretrain a transformer encoder from scratch on thirteen low-resource languages and evaluate it on sentiment analysis, hate speech detection, emotion classification, and sentence embedding similarity. Across all tasks, models trained with parallel tokenizers outperform conventional multilingual baselines, confirming that rethinking tokenization is essential for advancing multilingual representation learning—especially in low-resource settings.<sup>1</sup>

## 1 Introduction

Tokenization is a fundamental component of language models, converting text into discrete units such as words, subwords, or bytes (Sennrich et al., 2016). In multilingual settings (Devlin et al., 2019; Conneau et al., 2020a; Xue et al., 2021), a single tokenizer is typically applied across languages, which disproportionately disadvantages underrepresented languages: the same semantic content is often segmented into longer token sequences compared to

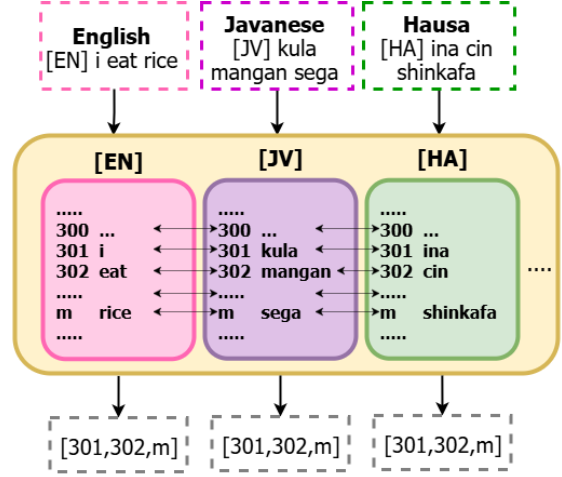


Figure 1: The overview of the parallel tokenizer. Tokens with equivalent meanings across languages are mapped to the same index and thus share the same embedding representation in the model.

high-resource languages (Rust et al., 2021). This imbalance is captured by the fertility score (Rust et al., 2021), which quantifies the number of tokens required to represent a given unit of meaning and exposes systematic inefficiencies in multilingual tokenization.

Beyond fertility, multilingual tokenizers also fail to align semantics across languages. Semantically equivalent words (e.g., “*eat*” in English, “*ci*” in Hausa, and “*食べる*” in Japanese) are mapped to entirely distinct vocabulary indices (Devlin et al., 2019; Conneau et al., 2020a; Xue et al., 2021; Achiam et al., 2023; Bai et al., 2023; Dubey et al., 2024), erecting artificial barriers to cross-lingual sharing. This fragmentation undermines the promise of transfer learning in multilingual models and introduces redundancy and inefficiency at the very foundation of representation learning.

Several strategies have been proposed to address these challenges, especially in the pretraining of medium- and low-resource language models. For

<sup>1</sup>Code can be found at anonymous.com

instance, the Sherkala model (Koto et al., 2025)—a multilingual model covering Kazakh, Russian, Turkic, and English—employs vocabulary balancing and data sampling techniques to oversample low-resource languages during tokenizer training, thereby improving fertility scores. However, semantic redundancy remains unresolved. Alternative approaches, such as byte- or character-level tokenizers (Xue et al., 2022; Clark et al., 2022; Cao, 2023), alleviate token length disparities but at the cost of efficiency and semantic interpretability.

We propose a **parallel tokenizer** to address these challenges (Figure 1). Instead of relying on a single multilingual tokenizer, we first train a monolingual tokenizer as a pivot and then align its vocabulary with target languages via word-level mapping, ensuring that token indices correspond to semantically equivalent words across languages. While achieving perfect alignment is inherently difficult, we restrict alignment to exact word forms. To reduce cross-lingual interference, we further incorporate language identity embeddings. (Figure 2), ensuring that token representations remain unique to each language while still enabling cross-lingual semantic sharing. This design yields fairer representations, facilitates more effective transfer, and provides a scalable framework for adding new languages without retraining the entire vocabulary.

Our contributions can be summarized as follows:

- We propose a parallel tokenizer for the multilingual language model that enhances the balance of fertility across individual languages and strengthens cross-lingual transfer. We benchmark pretraining from scratch and continual pretraining strategies, each under three setups: our parallel tokenizer, the traditional multilingual tokenizer, and a pretrained tokenizer from mBERT.
- We demonstrate the effectiveness of parallel tokenizers across diverse downstream tasks, including sentiment analysis, emotion classification, hate speech detection, and sentence embedding, all in low-resource contexts where tokenization quality has a disproportionate impact on performance.
- We present a detailed analysis of training under varying levels of data availability (e.g., 1%, 10%, 50%, 100%), auxiliary languages, and conduct an extensive evaluation of the fertility score to quantify tokenization efficiency.

## 2 Related Works

The standard practice in multilingual language modeling is to train a single tokenizer on a combined multilingual corpus (Achiam et al., 2023; Devlin et al., 2019; Xue et al., 2021; Dubey et al., 2024; Bai et al., 2023), typically using algorithms such as WordPiece (Devlin et al., 2019), Byte Pair Encoding (BPE) (Sennrich et al., 2016), the Unigram Language Model (Kudo, 2018), or SentencePiece (Kudo and Richardson, 2018). Despite their simplicity and scalability, these approaches face the well-documented *curse of multilinguality* (Conneau et al., 2020a): monolingual tokenizers consistently outperform shared multilingual ones in fertility and segmentation quality (Rust et al., 2021; Petrov et al., 2023), revealing a persistent trade-off between coverage and linguistic specialization. To mitigate these issues, prior studies have proposed balancing vocabulary allocation across languages (Zheng et al., 2021; Chung et al., 2020; Liang et al., 2023; Koto et al., 2025). However, while such methods improve fertility for low-resource languages, they often lead to redundant semantic representations, where tokens with equivalent meanings occupy distinct indices.

Beyond these engineering efforts, several studies have examined the conceptual role of tokenization in multilingual transfer. Conneau et al. (2020b) argue that parameter sharing is the primary factor enabling cross-lingual generalization. However, this line of research has not explored multilingual token embeddings in a parallel setting, where alignment across languages could better capture shared semantics. Limisiewicz et al. (2023) focus on lexical overlap, showing that surface-form similarity alone does not ensure semantic equivalence, leaving open the challenge of designing vocabularies that promote semantic alignment across languages.

Recent theoretical work further highlights the enduring importance of tokenization. Rajaraman et al. (2024) demonstrate that transformers trained on structured distributions struggle to learn effectively without discrete token boundaries, calling into question the feasibility of tokenization-free paradigms (Deiseroth et al., 2024). In this context, our work revisits multilingual tokenization from a new perspective, emphasizing that tokenization remains a foundational component for building semantically aligned and efficient multilingual language models.

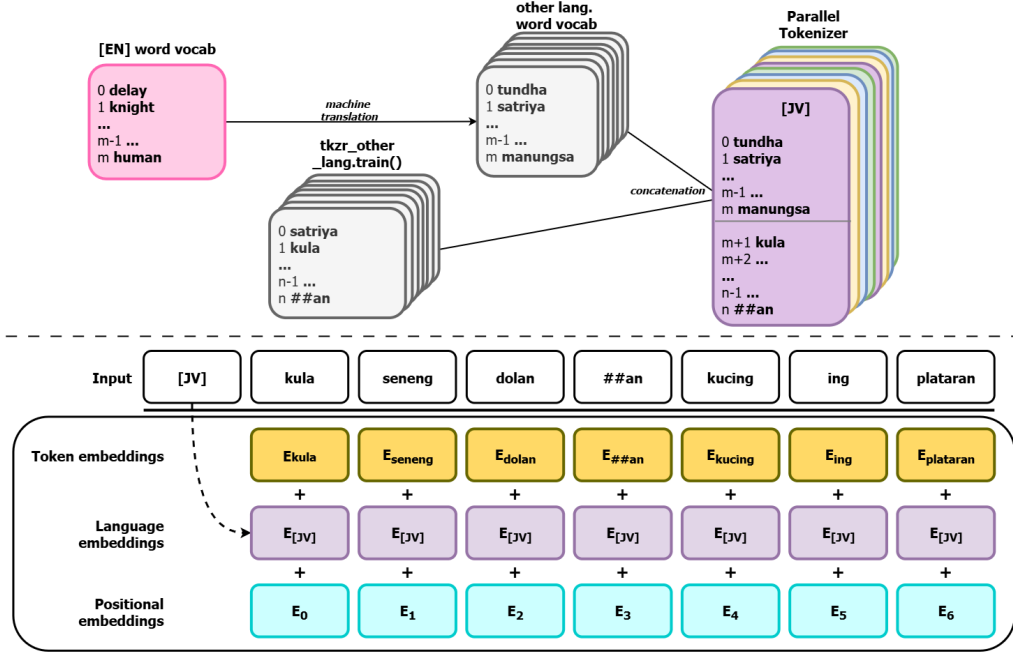


Figure 2: Design of the parallel vocabularies (top) and the model input representation (bottom). The language token (e.g., [JV] in the example) is not used as an explicit input token; instead, it functions as a signal to select the corresponding language identity embedding during input representation.

### 3 Parallel Tokenizer

To address fertility imbalance and the separation of semantically equivalent words into different token indices (Rust et al., 2021; Bai et al., 2023), we introduce the Parallel Tokenizer. This section outlines (1) the construction of the parallel vocabulary and (2) the design of the model input representation.

#### 3.1 Parallel Vocabulary Design

To construct parallel vocabularies across languages, we first train an English monolingual tokenizer using WordPiece (Devlin et al., 2019) on the English Wikipedia dump as the corpus. We set the vocabulary size to 30,522, which matches the monolingual BERT model (Devlin et al., 2019). The resulting vocabulary entries can be grouped into four categories:

1. subword, e.g., ‘##ing’, ‘##able’ (28.75% of vocabulary);
2. short word, e.g., ‘ple’, ‘szo’, ‘rae’; tokens with a length shorter than four characters (3.28% of vocabulary);
3. number, numerical tokens such as ‘192’, ‘392’, etc. (2.10% of vocabulary); and
4. word, which encompasses all other tokens, e.g., ‘care’, ‘drink’, ‘administration’ (65.87% of vocabulary).

We retain only the English word-type vocabulary and expand it into other languages through token-level machine translation (see Figure 2). This filtering is necessary because other token types (subword, short word, and number) often yield unreliable or meaningless translations that rarely correspond to valid equivalents in the target languages. Using this filtered vocabulary as the base, we construct the parallel tokenizer in three steps: (1) translating the word-type vocabulary, (2) training a monolingual tokenizer for each target language, and (3) concatenating the resulting vocabularies.

**Translating the word-type vocabulary.** We translate each English word-type token into its counterpart in the target language using Google Translate.<sup>2</sup> We initially considered bilingual dictionaries such as PanLex,<sup>3</sup> but found substantial limitations in both coverage and data quality, discussed in Section 6. Machine translation offers broader and more up-to-date lexical coverage. To further ensure quality, we apply back-translation to the invalid outputs, such as multiword or malformed phrases. This process yields semantically aligned vocabularies that form the word-level back-

<sup>2</sup><https://translate.google.com/>

<sup>3</sup><https://old.panlex.org/>

bone of our parallel tokenizer.

**Training a monolingual tokenizer on the target language.** Since the English tokenizer contains only 65.87% word-type tokens, we expand each parallel tokenizer to the full vocabulary size of 30,052 tokens by training a monolingual tokenizer on its corresponding Wikipedia corpus. The resulting vocabularies serve as sources for additional token types. We also introduce *language identity tokens* (e.g., [JV] in Figure 2) to explicitly mark the input language and ensure consistent handling across all languages (see Section 3.2).

**Concatenating the vocabulary.** For each target language, we concatenate two components: (1) the translated word-type tokens and (2) the corresponding monolingual vocabulary. We prioritize the translated word-type tokens, special tokens (e.g., [SEP], [UNK], [CLS]), language identity tokens,<sup>4</sup> and a 1,000-character-level subset from the English tokenizer. We then append the monolingual vocabulary, remove duplicates, and cap each final vocabulary at 30,522 entries for consistency across languages. On average, across all languages defined in Section 4, we align 82% of word-type tokens (slightly below 100% due to back-translation filtering), and 61% of all tokens are aligned in the final parallel tokenizer.

### 3.2 Model Input Representation

Let  $x = (l, x_1, x_2, \dots, x_i)$  denote an input sequence of length  $i$ , where  $l \in \mathcal{L}$  represents a language identity token and  $\mathcal{L}$  is the set of all such tokens (e.g., [AM], [HA], etc.; see Table 6 in Appendix A), with  $|\mathcal{L}| = k$ . We define a collection of  $k$  parallel tokenizers, each corresponding to a specific language, as

$$\mathcal{T} = \{T_1, T_2, \dots, T_k\}.$$

Based on Figure 2, the language identity token  $l$  determines which tokenizer  $T_j \in \mathcal{T}$  is used for tokenization. We then tokenize the remaining sequence  $x' = (x_1, x_2, \dots, x_i)$ , excluding  $l$ , using  $T_j$  to obtain the token IDs  $\mathcal{I} = (I_1, I_2, \dots, I_m)$ , where  $i \leq m$ , along with any auxiliary representations such as attention masks and token type IDs. The final input representation is constructed by summing the token embeddings, segment embeddings, positional embeddings, and a language identity embedding. The language identity embedding

<sup>4</sup>The complete list of language tokens is provided in Appendix A, Table 6, column ‘code’.

is broadcast across all  $m$  tokens to provide a unique representation for each language. This additional embedding serves both as a syntactic cue and as a disambiguation signal for unaligned tokens (recall from Section 3.1 that approximately 61% of tokens are aligned across languages).

## 4 Experimental Setup

### 4.1 Language Setup

We first define the set of languages considered in this work. English is included as the base language, given its large corpus size and its role in training the base tokenizer. We then select a diverse set of low-resource languages, some represented in mBERT (Devlin et al., 2019) and others absent from it, to evaluate the generalization of our approach in settings where linguistic resources are limited. We use mBERT as a lower-bound baseline, particularly for languages it has not been trained on. The selected languages are as follows:

- Seen by mBERT: Javanese (jav), Minangkabau (min), Sundanese (sun), and Swahili (swa).
- Unseen by mBERT: Acehnese (ace), Amharic (amh), Balinese (ban), Hausa (hau), Igbo (ibo), Kinyarwanda (kin), Oromo (orm), Tigrinya (tir), and Twi (twi).

Most of these languages use the Latin script, while only Amharic and Tigrinya use the Amharic script. The selection is guided by the availability of essential resources, including Wikipedia dumps<sup>5</sup> for tokenizer and transformer pretraining, Google Translate<sup>6</sup> for constructing parallel tokenizers, and FLORES+ (Costa-Jussà et al., 2022) for evaluating tokenization metrics.

### 4.2 Tokenizers & Models

We adopt the encoder-based transformer mBERT (Multilingual BERT, uncased) (Devlin et al., 2019) as the backbone of our experiments. Focusing on an encoder-only architecture allows us to examine how well cross-lingual transfer generalizes in low-resource settings. mBERT was pretrained on 102 languages using a shared 110K WordPiece vocabulary. As our baseline, we use the original mBERT model and refer to its tokenizer as *Single-102L* (i.e., a single tokenizer vocabulary covering 102 languages), which serves as the reference multilin-

<sup>5</sup><https://dumps.wikimedia.org/>. Resources were downloaded in November 2024.

<sup>6</sup><https://translate.google.com/>

gual tokenizer for evaluating our proposed parallel tokenization approach.

To enable a fair comparison, we also train a single tokenizer on the Wikipedia data of the 13 languages listed in Section 4.1 (hereafter referred to as *Single-13L*). Using the same data, we train our proposed parallel tokenizer, denoted as *Parallel-13L*. For both tokenizers, we adopt a pretraining-from-scratch setup, where all model weights are randomly initialized.

### 4.3 Evaluation Setup

To evaluate the effectiveness of our proposed parallel tokenizer, we conduct three main experiments in low-resource languages: (i) tokenization analysis, (ii) sequence classification, and (iii) cross-lingual representation similarity.

**Tokenization Analysis.** We use the FLORES+ dataset (Costa-Jussà et al., 2022; Abdulmumin et al., 2024)<sup>7</sup>, a parallel corpus covering 200 languages, to evaluate the effectiveness of our parallel tokenizer against both multilingual and monolingual tokenizers in 1,012 parallel texts. All 13 languages studied in our experiments are represented in FLORES+. We report two metrics: (i) Fertility score: the average number of tokens per word, indicating how compact or fragmented the tokenization is; and (ii) Parity score (Petrov et al., 2023): a cross-lingual metric that compares token counts across two languages,<sup>8</sup> assigning the best score when the counts are identical.

**Sequence Classification.** We conduct experiments on sentiment analysis, hate speech detection, and emotion classification in low-resource languages by fine-tuning the language model jointly on data from all languages. To analyze cross-lingual transfer, we vary the training size at 1%, 10%, 50%, and 100% of the available data. For sentiment analysis and hate speech detection, we use the NusaX-Senti (Winata et al., 2023) and AfriHate (Muhammad et al., 2025a) datasets, respectively, and evaluate both using macro-F1 over three multi-classes. For emotion classification, we use EthioEmo (Belay et al., 2025) and BRIGHTER (Muhammad et al., 2025b), evaluated with weighted F1 for consistency across six multi-label tasks. A detailed overview of the dataset

<sup>7</sup>[https://huggingface.co/datasets/openlanguage/flores\\_plus](https://huggingface.co/datasets/openlanguage/flores_plus)

<sup>8</sup>We select Hausa as the reference language for comparison, as it provides the largest available resource during pretraining.

languages is provided in Table 7 (Appendix B).

**Cross-Lingual Representation Similarity.** We perform bitext mining (Artetxe and Schwenk, 2019) on FLORES+ using the contextualized representations from the final layer of the language models. Following the setup of Huang et al. (2025), we evaluate all possible language pairs covered by our models. Performance is measured using the error rate of the xsim score (Artetxe and Schwenk, 2019). Additionally, we apply Principal Component Analysis (Jolliffe, 2002) to visualize cross-lingual clustering, comparing our parallel tokenizer with the multilingual model that uses a single tokenizer. Representations are considered semantically aligned when sentences from different languages cluster together, rather than being separated by language identity.

### 4.4 Hyperparameters and Resources

We use an NVIDIA RTX A6000 GPU with 48GB of VRAM for both pretraining and fine-tuning. Pretraining follows the masked language modeling (MLM) setup of Devlin et al. (2019), using a learning rate of 5e-5 for the *Single-13L* model and 1e-4 for our proposed *Parallel-13L* model. Training is performed on 394M tokens per epoch (6M tokens for the development set) for a total of 122,850 steps or 50 epochs. The value is determined empirically based on the convergence of the pretraining loss on the development set.

For evaluation, we set the batch size to 20, the maximum sequence length to 128, and the maximum number of epochs to 100, with early stopping triggered after 5, 2, 3, and 3 epochs for NusaX-senti, AfriHate, EthioEmo, and BRIGHTER, respectively. Learning rates were adopted from the original benchmark recommendations: 1e-5 for NusaX-senti and BRIGHTER, and 5e-5 for AfriHate and EthioEmo. To ensure robustness and fairness, each experiment was repeated three times with different seeds (1, 12, and 123), and we report the mean performance together with the corresponding standard deviations.

## 5 Results and Analysis

### 5.1 Tokenization Qualities

Table 1 compares our *Parallel-13L* tokenizer with *Single-102L* (mBERT), *Single-13L*, and the monolingual tokenizers. In terms of fertility score, our method substantially outperforms the multilingual baselines, exceeding the closest one by an average

| lang.                  | tokens/word (fertility score) ↓ |            |              |              | parity score ↓         |             |              |              |
|------------------------|---------------------------------|------------|--------------|--------------|------------------------|-------------|--------------|--------------|
|                        | Multilingual Tokenizer          |            |              | Mono-lingual | Multilingual Tokenizer |             |              | Mono-lingual |
|                        | Single-102L                     | Single-13L | Parallel-13L |              | Single-102L            | Single-13L  | Parallel-13L |              |
| <i>Seen by mBERT</i>   |                                 |            |              |              |                        |             |              |              |
| jav                    | 1.94                            | 1.65       | <b>1.48</b>  | 1.43         | 1.24                   | <b>1.06</b> | 1.13         | 1.69         |
| min                    | 1.96                            | 1.79       | <b>1.50</b>  | 1.46         | 1.20                   | <b>1.05</b> | 1.09         | 1.68         |
| sun                    | 2.01                            | 1.74       | <b>1.50</b>  | 1.44         | 1.21                   | <b>1.01</b> | 1.13         | 1.74         |
| swa                    | 2.10                            | 1.65       | <b>1.42</b>  | 1.38         | 1.07                   | <b>1.01</b> | 1.10         | 1.63         |
| <i>Unseen by mBERT</i> |                                 |            |              |              |                        |             |              |              |
| ace                    | 2.27                            | 1.95       | <b>1.54</b>  | 1.61         | <b>1.01</b>            | 1.20        | 1.02         | 1.51         |
| amh                    | -                               | 2.43       | <b>1.82</b>  | 1.66         | 2.19                   | 1.21        | <b>1.06</b>  | 2.33         |
| ban                    | 2.13                            | 1.77       | <b>1.53</b>  | 1.48         | 1.11                   | <b>1.03</b> | 1.07         | 1.63         |
| hau                    | 1.91                            | 1.38       | <b>1.33</b>  | 1.30         | 1.00                   | 1.00        | 1.00         | 1.00         |
| ibo                    | 2.16                            | 1.54       | <b>1.50</b>  | 1.48         | 1.12                   | <b>1.10</b> | 1.12         | 1.32         |
| kin                    | 2.78                            | 2.10       | <b>1.70</b>  | 1.63         | 1.16                   | 1.21        | <b>1.03</b>  | 1.53         |
| orm                    | 3.01                            | 2.40       | <b>1.82</b>  | 1.72         | 1.23                   | 1.35        | <b>1.07</b>  | 1.55         |
| tir                    | -                               | 2.61       | <b>1.83</b>  | 1.85         | 1.98                   | 1.45        | <b>1.06</b>  | 2.14         |
| twi                    | 2.17                            | 1.55       | <b>1.38</b>  | 1.35         | 1.17                   | 1.16        | <b>1.07</b>  | 1.42         |
| avg                    | 2.22                            | 1.89       | <b>1.57</b>  | 1.52         | 1.28                   | 1.14        | <b>1.07</b>  | 1.63         |

Table 1: Fertility and parity scores for each tokenizer across all languages. The **bold** scores indicate the best performance for each language. Monolingual tokenizers are also included to provide a comparative perspective between monolingual and multilingual tokenizers. Fertility scores for amh and tir are marked as N/A for *Single-102L* because mBERT consistently produces [UNK] tokens for these languages.

of 0.32 points, indicating more compact segmentation (i.e., fewer tokens per word). The only slight gap (0.05) appears against the monolingual tokenizer, which is expected since monolingual models are optimized for their respective languages.

For the parity score, which quantifies consistency in tokenization across parallel texts and reflects cross-lingual alignment, *Parallel-13L* achieves the best overall performance, outperforming the next-best baseline (*Single-13L*) by 0.07 points on average. The multilingual baselines (*Single-102L* and *Single-13L*) perform notably worse on Amharic-script languages (Amharic and Tigrinya) due to script-specific segmentation differences, while monolingual tokenizers score lowest by design, as parity inherently measures multilingual generalization.

We also analyze the number of [UNK] tokens generated when tokenizing FLORES+ inputs, with results shown in Figure 3. Among the languages covered by mBERT, our method produces the fewest [UNK] tokens. In contrast, *Parallel-13L* yields slightly more unknown tokens, about eight more on average, than *Single-13L*. The mBERT tokenizer (*Single-102L*) performs worst overall, primarily due to its limited handling of Amharic-script languages.

## 5.2 Sequence Classification

Across three random seeds, Table 2 shows that our parallel tokenizer (*Parallel-13L*) consistently outperforms all baselines across training sizes and evaluation metrics. On average, *Parallel-13L* ex-

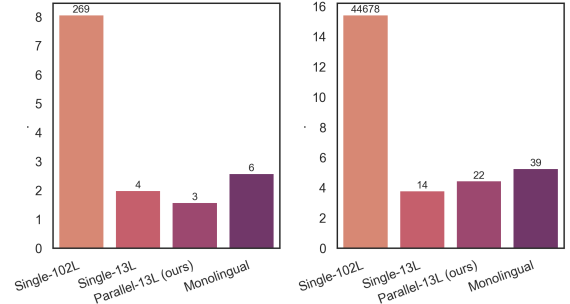


Figure 3: Total number of [UNK] tokens occurs when tokenizing the FLORES+ dataset for each language. Left: languages seen by mBERT (jav, min, sun, swa); Right: languages unseen by mBERT (ace, amh, ban, hau, ibo, kin, orm, tir, twi). Values are plotted on a  $\log_2$  scale.

ceeds the closest baseline by 0.92%, 1.28%, 0.72%, and 1.22% F1 at the 100%, 50%, 10%, and 1% training data levels, respectively. These results demonstrate the effectiveness of a parallelized tokenizer design in enhancing cross-lingual transfer, particularly in low-resource scenarios.

The improvements in our approach are consistent across the NusaX-senti, AfriHate, and BRIGHTER datasets, where *Parallel-13L* outperforms all other baselines regardless of training size. The only exception occurs with EthioEmo, where *Single-13L* achieves higher scores than *Parallel-13L* only at the 10% and 1% settings. Moreover, *Single-102L* performs notably worse on EthioEmo due to its limited coverage of Amharic, Oromo, and Tigrinya, which are languages that rely pre-

| #training data | Multilingual Tokenizer | Sentiment                   | Hate Speech                 | Emotion Classification      |                             | Avg $\uparrow$ |
|----------------|------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|----------------|
|                |                        | NusaX-senti $\uparrow$      | AfriHate $\uparrow$         | EthioEmo $\uparrow$         | BRIGHTER $\uparrow$         |                |
| 100%           | Single-102L            | 73.11 ( $\pm 1.70$ )        | 62.83 ( $\pm 2.38$ )        | 30.10 ( $\pm 6.97$ )        | 48.47 ( $\pm 1.32$ )        | 53.63          |
|                | Single-13L             | 76.09 ( $\pm 1.27$ )        | 69.61 ( $\pm 1.19$ )        | 54.98 ( $\pm 0.98$ )        | 48.26 ( $\pm 0.94$ )        | 62.24          |
|                | Parallel-13L (ours)    | <b>76.16</b> ( $\pm 1.05$ ) | <b>69.80</b> ( $\pm 1.20$ ) | <b>57.01</b> ( $\pm 0.84$ ) | <b>49.68</b> ( $\pm 1.89$ ) | <b>63.16</b>   |
| 50%            | Single-102L            | 69.81 ( $\pm 1.09$ )        | 61.52 ( $\pm 1.47$ )        | 31.96 ( $\pm 2.20$ )        | 44.68 ( $\pm 1.13$ )        | 51.99          |
|                | Single-13L             | 72.47 ( $\pm 1.36$ )        | 67.26 ( $\pm 1.36$ )        | 51.33 ( $\pm 1.48$ )        | 45.77 ( $\pm 1.46$ )        | 59.21          |
|                | Parallel-13L (ours)    | <b>73.52</b> ( $\pm 1.17$ ) | <b>67.40</b> ( $\pm 0.94$ ) | <b>54.27</b> ( $\pm 1.78$ ) | <b>46.76</b> ( $\pm 1.51$ ) | <b>60.49</b>   |
| 10%            | Single-102L            | 59.08 ( $\pm 2.63$ )        | 54.69 ( $\pm 2.98$ )        | 22.07 ( $\pm 7.47$ )        | 36.71 ( $\pm 1.87$ )        | 43.14          |
|                | Single-13L             | 64.76 ( $\pm 1.84$ )        | 59.80 ( $\pm 2.09$ )        | <b>42.10</b> ( $\pm 2.19$ ) | 36.64 ( $\pm 2.23$ )        | 50.83          |
|                | Parallel-13L (ours)    | <b>66.16</b> ( $\pm 2.00$ ) | <b>60.17</b> ( $\pm 1.50$ ) | 41.35 ( $\pm 2.21$ )        | <b>38.51</b> ( $\pm 1.27$ ) | <b>51.55</b>   |
| 1%             | Single-102L            | 27.79 ( $\pm 7.41$ )        | 40.82 ( $\pm 4.15$ )        | 10.73 ( $\pm 11.17$ )       | 0.00 ( $\pm 0.00$ )         | 19.84          |
|                | Single-13L             | 31.54 ( $\pm 3.87$ )        | 44.70 ( $\pm 4.45$ )        | <b>27.37</b> ( $\pm 1.88$ ) | 17.54 ( $\pm 2.57$ )        | 30.29          |
|                | Parallel-13L (ours)    | <b>33.83</b> ( $\pm 3.90$ ) | <b>46.73</b> ( $\pm 2.52$ ) | 26.58 ( $\pm 1.37$ )        | <b>18.88</b> ( $\pm 3.96$ ) | <b>31.51</b>   |

Table 2: Performance comparison of tokenizer setups on sentiment analysis, hate speech detection, and emotion classification tasks. Best scores per benchmark setups are highlighted in **bold**.

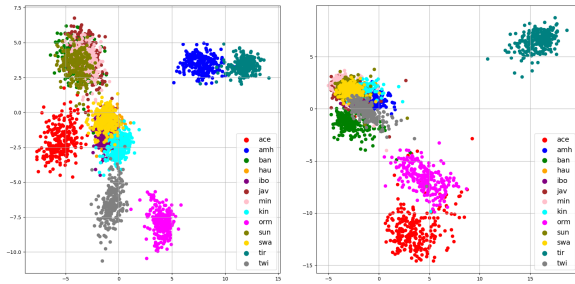


Figure 4: PCA visualization of the last hidden states from *Single-13L* (left) and *Parallel-13L* (right) models on the FLORES+ dataset.

dominantly on the Amharic script. Detailed per-language benchmark results are provided in Appendix C (Tables 8, 9, 10, and 11 for NusaX-senti, AfriHate, EthioEmo, and BRIGHTER, respectively).

### 5.3 Cross-Lingual Representation Similarity

Using the parallel FLORES+ corpus, we visualize cross-lingual representation similarity for each model by extracting the final hidden states of 250 sampled parallel sentences. If a model learns shared representations across languages, PCA should reveal clusters based on semantic similarity rather than language identity. Figure 4 compares *Single-13L* and *Parallel-13L* in 2D space. *Single-13L* tends to cluster representations by language family, for example, Indonesian languages (Sundanese, Minangkabau, Balinese, and Javanese) group together, while African languages form a separate cluster. In contrast, *Parallel-13L* yields more compact cross-lingual clusters, indicating stronger semantic alignment, with only Acehnese, Oromo, and Tigrinya appearing as outliers for both due to limited data.

We further evaluate cross-lingual similarity

|                             | Single-102L | Single-13L | Parallel-13L |
|-----------------------------|-------------|------------|--------------|
| Avg. err. xsim $\downarrow$ | 91.33       | 83.56      | <b>74.08</b> |
| # of best xsim $\uparrow$   | 3           | 12         | <b>63</b>    |

Table 3: Bitext mining metric meta-scores on pretrained models under different setups.

| #data in tgt language | S-102L      | S-13L | P-13L (ours) |
|-----------------------|-------------|-------|--------------|
| 50%                   | NusaX-senti | 71.76 | 75.22        |
|                       | AfriHate    | 61.45 | <b>67.06</b> |
|                       | EthioEmo    | 29.75 | 51.73        |
|                       | BRIGHTER    | 44.11 | 45.00        |
| 0%                    | NusaX-senti | 68.80 | 71.83        |
|                       | AfriHate    | 33.72 | 37.27        |
|                       | EthioEmo    | 9.73  | <b>23.75</b> |
|                       | BRIGHTER    | 15.66 | 15.04        |

Table 4: Benchmarking results under limited target-language data. S-102L, S-13L, and P-13L denote *Single-102L*, *Single-13L*, and *Parallel-13L*, respectively.

through bitext mining. As shown in Table 3, *Parallel-13L* achieves the lowest xsim error rate and the highest number of best scores across language pairs (details in Appendix D). These results confirm that *Parallel-13L* learns more semantically consistent representations across languages than *Single-13L*, reinforcing its advantage for cross-lingual tasks.

### 5.4 Cross-Lingual Transfer under Limited Target-Language Data

While our main experiments use data from all languages during training, we further analyze the model’s behavior when the target-language data is scarce or unavailable—a common scenario in low-resource settings. To simulate this condition, we exclude the target language from the training corpus and evaluate two setups: (i) 0% target-language data (zero-shot) and (ii) 50% target-language data (partial-shot). Table 4 shows that our method consistently achieves the highest F1 scores in both

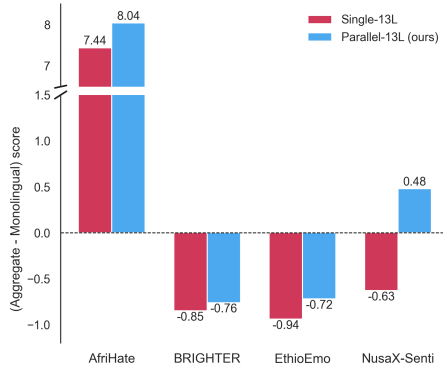


Figure 5: Performance differences across benchmarks on both cross-lingual and monolingual finetuning data for models pretrained with different tokenizer setups.

settings. The only exceptions are a slight decrease of 0.06% on AfriHate with 50% target data and 0.55% on EthioEmo with 0% target train data. Overall, *Parallel-13L* outperforms the baselines, demonstrating stronger cross-lingual transfer than the traditional single tokenizer. The detailed results are presented in Appendix E in Table 15, 16, 17, and 18.

### 5.5 How Much Does Multilingual Training Improve Over Single-Language Training?

This analysis compares fine-tuning on concatenated multilingual datasets with training solely on a target language’s data. As shown in Figure 5, both baselines and our method benefit from multilingual training in the AfriHate benchmark, with average gains of 7-8%. In contrast, improvements are smaller for BRIGHTER and EthioEmo, while in NusaX-senti, our method still achieves a modest 0.48% increase. Between *Single-13L* and *Parallel-13L*, the latter benefits more consistently from multilingual concatenation, suggesting that it leverages shared cross-lingual information more effectively. Detailed results for monolingual fine-tuning are provided in Appendix F, showing that overall performance is comparable, though *Parallel-13L* maintains slightly higher scores on average.

### 5.6 Continual Pre-Training Experiments

We perform continual pretraining (CPT) from mBERT by adapting its vocabulary and tokenizer to our parallel setup. Newly introduced token embeddings are initialized by averaging the mBERT embeddings of their constituent subwords, as determined by the original mBERT tokenizer (Koto et al., 2021). Unlike pretraining from scratch, this

| #data       | S-102L*            | S-13L*                    | P-13L* (ours)             |
|-------------|--------------------|---------------------------|---------------------------|
| NusaX-senti | 73.1 ( $\pm 1.7$ ) | 78.3 ( $\pm 1.0$ )        | <b>78.6</b> ( $\pm 1.2$ ) |
| Afrihate    | 62.8 ( $\pm 2.4$ ) | <b>71.3</b> ( $\pm 1.9$ ) | 70.1 ( $\pm 0.9$ )        |
| EthioEmo    | 30.1 ( $\pm 7.0$ ) | 56.7 ( $\pm 0.9$ )        | <b>57.4</b> ( $\pm 0.8$ ) |
| BRIGHTER    | 48.5 ( $\pm 1.3$ ) | <b>50.5</b> ( $\pm 0.9$ ) | 50.3 ( $\pm 1.2$ )        |

Table 5: Benchmarking results under the continual pre-training setting. S-102L\*, S-13L\*, and P-13L\* denote *Single-102L*, *Single-13L*, and *Parallel-13L*, respectively, which have been continually pretrained from the mBERT model.

approach retains all mBERT parameters while updating the token embedding layer to accommodate the new vocabularies. In *Single-13L*, each token in the new vocabulary is mapped to its corresponding English subwords in mBERT, and its embedding is initialized by averaging the associated subword embeddings. In *Parallel-13L*, each token is tokenized into subwords across all 13 languages, and its embedding is initialized by averaging the pooled subword embeddings across languages.

Table 5 shows that *Single-13L* and *Parallel-13L* achieve similar performance under continual pretraining (CPT). However, as detailed in Appendix G, our method yields stronger cross-lingual alignment, languages cluster by meaning rather than family, and bitext mining confirms more consistent representations. This indicates that even when task scores converge, *Parallel-13L* maintains a structural advantage in preserving cross-lingual coherence.

## 6 Conclusion

We introduce parallel tokenizers, constructed by first training a monolingual tokenizer as a pivot and then aligning its word vocabulary with other languages using machine translation, ensuring that semantically equivalent words across languages share the same representation. We evaluate this approach along three dimensions. In terms of tokenization quality, it yields the lowest fertility and parity scores, indicating more efficient tokenization while reducing disparities across languages. In sequence classification, our method shows improvements over the baselines, yielding higher average F1 scores across benchmarks. Finally, in cross-lingual sentence representation, our approach produces stronger cross-lingual alignment representation than ordinary tokenization methods, making it especially beneficial for cross-lingual training and for incorporating auxiliary languages into multilingual models.

## Limitations

To develop a robust parallel tokenizer, each stage of designing the parallel vocabulary and model input representation required careful consideration. For mapping the English vocabulary in the tokenizer, we employed machine translation (MT). Although bilingual dictionaries can potentially provide more precise word-level mappings due to their predefined vocabularies, our attempt to utilize the PanLex bilingual dictionary (as discussed in Section 3.1) revealed substantial limitations in both coverage and data quality, particularly for lower-resource languages. While MT enabled us to align approximately 61% of tokens, the bilingual dictionary achieved only about 21.5%. Therefore, we adopted MT for our final approach.

This choice, however, introduces a minor limitation related to the quality of the machine translation itself, as several mapping errors were observed, including invalid outputs such as multiword or malformed phrases. Consequently, there remains potential to improve token alignment through more refined resources or hybrid methods. Furthermore, there is room to enhance the model’s input representation, as detailed in Appendix H, which we plan to investigate further in future work.

Our current study focuses on sequence classification, aligning with our focus on low-resource languages that have limited data. In future research, as more diverse datasets become available, we aim to expand the benchmarking to additional tasks and languages, ensuring that the parallel tokenizer can effectively scale and maintain performance across a broader linguistic spectrum.

## Ethical Consideration

All datasets used in this work are publicly available and comply with their respective licenses, including Wikipedia, FLORES+, and established benchmarks for sentiment, hate speech, and emotion classification. No personally identifiable or sensitive information was collected or used. Our study aims to improve multilingual representation quality, particularly for low-resource languages. However, since the underlying data may contain cultural or societal biases, we recommend that future work include bias analysis and community engagement when extending models to additional linguistic or cultural contexts.

## References

- Idris Abdulmumin, Sthembiso Mkhwanazi, Mahlatse S. Mbooi, Shamsuddeen Hassan Muhammad, Ibrahim Said Ahmad, Neo N. Putini, Miehleketo Mathebula, Matimba Shingange, Tajuddeen Gwadabe, and Vukosi Marivate. 2024. Correcting FLORES evaluation dataset for four African languages. In *Proceedings of the Ninth Conference on Machine Translation*, Miami, USA. Association for Computational Linguistics.
- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, and 1 others. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Mikel Artetxe and Holger Schwenk. 2019. [Margin-based parallel corpus mining with multilingual sentence embeddings](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 3197–3203, Florence, Italy. Association for Computational Linguistics.
- Jinze Bai, Shuai Bai, Yunfei Chu, Zeyu Cui, Kai Dang, Xiaodong Deng, Yang Fan, Wenbin Ge, Yu Han, Fei Huang, and 1 others. 2023. Qwen technical report. *arXiv preprint arXiv:2309.16609*.
- Tadesse Destaw Belay, Israel Abebe Azime, Abinew Ali Ayele, Grigori Sidorov, Dietrich Klakow, Philip Slusallek, Olga Kolesnikova, and Seid Muhie Yimam. 2025. [Evaluating the capabilities of large language models for multi-label emotion understanding](#). In *Proceedings of the 31st International Conference on Computational Linguistics*, pages 3523–3540, Abu Dhabi, UAE. Association for Computational Linguistics.
- Kris Cao. 2023. [What is the best recipe for character-level encoder-only modelling?](#) In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 5924–5938, Toronto, Canada. Association for Computational Linguistics.
- Hyung Won Chung, Dan Garrette, Kiat Chuan Tan, and Jason Riesa. 2020. [Improving multilingual models with language-clustered vocabularies](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 4536–4546, Online. Association for Computational Linguistics.
- Jonathan H. Clark, Dan Garrette, Iulia Turc, and John Wieting. 2022. [Canine: Pre-training an efficient tokenization-free encoder for language representation](#). *Transactions of the Association for Computational Linguistics*, 10:73–91.
- Alexis Conneau, Kartikay Khandelwal, Naman Goyal, Vishrav Chaudhary, Guillaume Wenzek, Francisco Guzmán, Edouard Grave, Myle Ott, Luke Zettlemoyer, and Veselin Stoyanov. 2020a. [Unsupervised](#)

- cross-lingual representation learning at scale. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 8440–8451, Online. Association for Computational Linguistics.
- Alexis Conneau, Shijie Wu, Haoran Li, Luke Zettlemoyer, and Veselin Stoyanov. 2020b. [Emerging cross-lingual structure in pretrained language models](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 6022–6034, Online. Association for Computational Linguistics.
- Marta R Costa-Jussà, James Cross, Onur Çelebi, Maha Elbayad, Kenneth Heafield, Kevin Heffernan, Elahe Kalbassi, Janice Lam, Daniel Licht, Jean Maillard, and 1 others. 2022. No language left behind: Scaling human-centered machine translation. *arXiv preprint arXiv:2207.04672*.
- Björn Deiseroth, Manuel Brack, Patrick Schramowski, Kristian Kersting, and Samuel Weinbach. 2024. [T-FREE: Subword tokenizer-free generative LLMs via sparse representations for memory-efficient embeddings](#). In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 21829–21851, Miami, Florida, USA. Association for Computational Linguistics.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: Pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, and 1 others. 2024. The llama 3 herd of models. *arXiv e-prints*, pages arXiv–2407.
- Yongxin Huang, Kexin Wang, Goran Glavaš, and Iryna Gurevych. 2025. [Modular sentence encoders: Separating language specialization from cross-lingual alignment](#). In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2167–2187, Vienna, Austria. Association for Computational Linguistics.
- I. T. Jolliffe. 2002. *Principal Component Analysis*, 2nd edition. Springer Series in Statistics. Springer New York, NY, New York.
- Fajri Koto, Rituraj Joshi, Nurdaulet Mukhituly, Yuxia Wang, Zhuohan Xie, Rahul Pal, Daniil Orel, Parvez Mullah, Diana Turmakhan, Maiya Goloburda, and 1 others. 2025. Sherkala-chat: Building a state-of-the-art llm for kazakh in a moderately resourced setting. In *Second Conference on Language Modeling*.
- Fajri Koto, Jey Han Lau, and Timothy Baldwin. 2021. [IndoBERTweet: A pretrained language model for Indonesian Twitter with effective domain-specific vocabulary initialization](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 10660–10668, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Taku Kudo. 2018. [Subword regularization: Improving neural network translation models with multiple subword candidates](#). In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 66–75, Melbourne, Australia. Association for Computational Linguistics.
- Taku Kudo and John Richardson. 2018. [SentencePiece: A simple and language independent subword tokenizer and detokenizer for neural text processing](#). In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 66–71, Brussels, Belgium. Association for Computational Linguistics.
- Davis Liang, Hila Gonen, Yuning Mao, Rui Hou, Namann Goyal, Marjan Ghazvininejad, Luke Zettlemoyer, and Madian Khabza. 2023. [XLM-V: Overcoming the vocabulary bottleneck in multilingual masked language models](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 13142–13152, Singapore. Association for Computational Linguistics.
- Tomasz Limisiewicz, Jiří Balhar, and David Mareček. 2023. [Tokenization impacts multilingual language modeling: Assessing vocabulary allocation and overlap across languages](#). In *Findings of the Association for Computational Linguistics: ACL 2023*, pages 5661–5681, Toronto, Canada. Association for Computational Linguistics.
- Shamsuddeen Hassan Muhammad, Idris Abdulmumin, Abinew Ali Ayele, David Ifeoluwa Adelani, Ibrahim Said Ahmad, Saminu Mohammad Aliyu, Paul Röttger, Abigail Oppong, Andiswa Bukula, Chiamaka Ijeoma Chukwunke, Ebrahim Chekol Jibril, Elyas Abdi Ismail, Esubalew Alemneh, Hagos Tesfahun Gebremichael, Lukman Jibril Aliyu, Meriem Beloucif, Oumaima Hourrane, Rooweither Mabuya, Salomey Osei, and 8 others. 2025a. [AfriHate: A multilingual collection of hate speech and abusive language datasets for African languages](#). In *Proceedings of the 2025 Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pages 1854–1871, Albuquerque, New Mexico. Association for Computational Linguistics.
- Shamsuddeen Hassan Muhammad, Nedjma Ousidhoum, Idris Abdulmumin, Jan Philip Wahle, Terry Ruas, Meriem Beloucif, Christine de Kock, Nirmal Surange, Daniela Teodorescu, Ibrahim Said Ahmad, David Ifeoluwa Adelani, Alham Fikri Aji, Felermينو

- D. M. A. Ali, Ilseyar Alimova, Vladimir Araujo, Nikolay Babakov, Naomi Baes, Ana-Maria Bucur, Andiswa Bukula, and 29 others. 2025b. [BRIGHTER: BRIdging the gap in human-annotated textual emotion recognition datasets for 28 languages](#). In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 8895–8916, Vienna, Austria. Association for Computational Linguistics.
- Aleksandar Petrov, Emanuele La Malfa, Philip Torr, and Adel Bibi. 2023. [Language model tokenizers introduce unfairness between languages](#). In *Advances in Neural Information Processing Systems*, volume 36, pages 36963–36990. Curran Associates, Inc.
- Nived Rajaraman, Jiantao Jiao, and Kannan Ramchandran. 2024. [An analysis of tokenization: Transformers under markov data](#). In *Advances in Neural Information Processing Systems*, volume 37, pages 62503–62556. Curran Associates, Inc.
- Phillip Rust, Jonas Pfeiffer, Ivan Vulić, Sebastian Ruder, and Iryna Gurevych. 2021. [How good is your tokenizer? on the monolingual performance of multilingual language models](#). In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 3118–3135, Online. Association for Computational Linguistics.
- Rico Sennrich, Barry Haddow, and Alexandra Birch. 2016. [Neural machine translation of rare words with subword units](#). In *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1715–1725, Berlin, Germany. Association for Computational Linguistics.
- Genta Indra Winata, Alham Fikri Aji, Samuel Cahyawijaya, Rahmad Mahendra, Fajri Koto, Ade Romadhony, Kemal Kurniawan, David Moeljadi, Radityo Eko Prasajo, Pascale Fung, Timothy Baldwin, Jey Han Lau, Rico Sennrich, and Sebastian Ruder. 2023. [NusaX: Multilingual parallel sentiment dataset for 10 Indonesian local languages](#). In *Proceedings of the 17th Conference of the European Chapter of the Association for Computational Linguistics*, pages 815–834, Dubrovnik, Croatia. Association for Computational Linguistics.
- Linting Xue, Aditya Barua, Noah Constant, Rami Al-Rfou, Sharan Narang, Mihir Kale, Adam Roberts, and Colin Raffel. 2022. [ByT5: Towards a token-free future with pre-trained byte-to-byte models](#). *Transactions of the Association for Computational Linguistics*, 10:291–306.
- Linting Xue, Noah Constant, Adam Roberts, Mihir Kale, Rami Al-Rfou, Aditya Siddhant, Aditya Barua, and Colin Raffel. 2021. [mT5: A massively multilingual pre-trained text-to-text transformer](#). In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 483–498, Online. Association for Computational Linguistics.
- Bo Zheng, Li Dong, Shaohan Huang, Saksham Singhal, Wanxiang Che, Ting Liu, Xia Song, and Furu Wei. 2021. [Allocating large vocabulary capacity for cross-lingual language model pre-training](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 3203–3215, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.

## A Language Resource Availability Details

Table 6 presents the availability of language resources under study in the Wikipedia dumps. English is highlighted in lime, as it serves both as the reference for comparison to other lower-resource languages and the pivot language when constructing the parallel tokenizers, due to its large corpus size and richer vocabulary coverage.

| Lang.       | iso | code | Script  | Resource |
|-------------|-----|------|---------|----------|
| English     | eng | [EN] | Latin   | 10.7GB   |
| Acehnese    | ace | [AC] | Latin   | 2.52MB   |
| Amharic     | amh | [AM] | Amharic | 18MB     |
| Balinese    | ban | [BA] | Latin   | 14.24MB  |
| Hausa       | hau | [HA] | Latin   | 108MB    |
| Igbo        | ibo | [IG] | Latin   | 104MB    |
| Javanese    | jav | [JV] | Latin   | 50.7MB   |
| Kinyarwanda | kin | [RW] | Latin   | 12.25MB  |
| Minangkabau | min | [MI] | Latin   | 85.4MB   |
| Oromo       | orm | [OR] | Latin   | 3.71MB   |
| Sundanese   | sun | [SU] | Latin   | 34.1MB   |
| Swahili     | swa | [SW] | Latin   | 54.4MB   |
| Tigrinya    | tir | [TI] | Amharic | 1.1MB    |
| Twi         | twi | [TW] | Latin   | 8.6MB    |

Table 6: Language resource availability. The ‘iso’ column lists language codes following the ISO 639-3 convention, while the ‘code’ column specifies the language identifiers used to determine the language embeddings in the model’s input representation.

## B Benchmark Language Coverage

Table 7 reports the language coverage of each benchmark included in our study. For brevity, we use the abbreviations NusaX, Afri, Ethio, and BRIGHT to refer to the NusaX-senti, AfriHate, EthioEmo, and BRIGHTER datasets, respectively. The numbers next to the checkmarks indicate the number of instances used in the benchmarking process.

| lang. | NusaX | Afri     | Ethio    | BRIGHT   |
|-------|-------|----------|----------|----------|
| ace   | ✓(1K) |          |          |          |
| amh   |       | ✓(4.96K) | ✓(5.92K) |          |
| ban   | ✓(1K) |          |          |          |
| hau   |       | ✓(6.64K) |          | ✓(5.02K) |
| ibo   |       | ✓(5K)    |          | ✓(6.73K) |
| jav   | ✓(1K) |          |          |          |
| min   | ✓(1K) |          |          |          |
| kin   |       | ✓(4.72K) |          | ✓(5.73K) |
| orm   |       | ✓(5.03K) | ✓(5.74K) |          |
| sun   | ✓(1K) |          |          | ✓(3.17K) |
| swa   |       | ✓(21.1K) |          | ✓(7.72K) |
| tir   |       | ✓(5.07K) | ✓(6.14K) |          |
| twi   |       | ✓(3.9K)  |          |          |

Table 7: Benchmark language coverage.

## C Benchmarking Result Details

Tables 8, 9, 10, and 11 present the benchmarking results for NusaX-senti, AfriHate, EthioEmo, and BRIGHTER, respectively. Results are reported for each language in the benchmark under training data settings of 100%, 50%, 10%, and 1%. To ensure stability and reliability, each experiment was run three times with different random seeds, and the reported results include the corresponding standard deviations. Based on the average results, our method yields better results than the other baselines, highlighting the performance of our method in cross-lingual settings.

## D Bitext Mining Details

Tables 12, 13, and 14 present the bitext mining performance of each model using their corresponding tokenizers: *Single-102L*, *Single-13L*, and *Parallel-13L*. Performance is evaluated using the xsim error rate, which measures how accurately a model selects the equivalent sentence across languages based on the FLORES+ parallel data. A lower error rate indicates better cross-lingual alignment, making this metric an effective indicator of a model’s ability to capture cross-lingual representations.

## E Cross-Lingual Transfer under Limited Target-Language Data Details

Similar to Appendix C, we report benchmarking results under a limited target-language data setting, where the auxiliary languages are fully utilized (100% of the data) during training, while the target language is limited to 50% or 0% of its data. This setup highlights cross-lingual transfer from the auxiliary languages to the target language. Tables 15, 16, 17, and 18 present the results for NusaX-senti, AfriHate, EthioEmo, and BRIGHTER, respectively. For each language, we repeat the experiments three times with different seeds to ensure robustness.

## F Monolingual Benchmarking Details

The benchmarking results using monolingual training data are shown in Tables 19, 20, 21, and 22. Each experiment was repeated three times with pre-selected seeds, and we report the averages along with the standard deviations for each language. Overall, the results are comparable: *Single-13L* performs better on the NusaX-senti and AfriHate benchmarks, while *Parallel-13L* performs better on emotion classification tasks such as EthioEmo and

| #data | Tokenizer           | ace                         | ban                         | jav                         | min                         | sun                         | avg                         |
|-------|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| 100%  | Single-102L         | 72.02 ( $\pm 4.04$ )        | 71.02 ( $\pm 0.43$ )        | 76.20 ( $\pm 0.99$ )        | 74.96 ( $\pm 1.39$ )        | 71.32 ( $\pm 1.67$ )        | 73.11 ( $\pm 1.70$ )        |
|       | Single-13L          | <b>76.61</b> ( $\pm 1.31$ ) | <b>73.23</b> ( $\pm 1.00$ ) | 77.43 ( $\pm 1.45$ )        | 76.48 ( $\pm 1.37$ )        | 76.68 ( $\pm 1.19$ )        | 76.09 ( $\pm 1.27$ )        |
|       | Parallel-13L (ours) | 74.22 ( $\pm 1.91$ )        | 72.62 ( $\pm 1.22$ )        | <b>78.26</b> ( $\pm 0.50$ ) | <b>78.43</b> ( $\pm 0.42$ ) | <b>77.26</b> ( $\pm 1.22$ ) | <b>76.16</b> ( $\pm 1.05$ ) |
| 50%   | Single-102L         | 67.86 ( $\pm 0.57$ )        | 68.24 ( $\pm 2.08$ )        | 71.56 ( $\pm 1.11$ )        | 71.76 ( $\pm 0.62$ )        | 69.62 ( $\pm 1.09$ )        | 69.81 ( $\pm 1.09$ )        |
|       | Single-13L          | <b>73.44</b> ( $\pm 1.41$ ) | 69.68 ( $\pm 1.76$ )        | 71.86 ( $\pm 2.03$ )        | 73.14 ( $\pm 1.01$ )        | 74.22 ( $\pm 0.61$ )        | 72.47 ( $\pm 1.36$ )        |
|       | Parallel-13L (ours) | 71.53 ( $\pm 0.52$ )        | <b>72.83</b> ( $\pm 2.03$ ) | <b>74.34</b> ( $\pm 0.86$ ) | <b>74.38</b> ( $\pm 2.15$ ) | <b>74.52</b> ( $\pm 0.31$ ) | <b>73.52</b> ( $\pm 1.17$ ) |
| 10%   | Single-102L         | 56.95 ( $\pm 2.34$ )        | 55.81 ( $\pm 0.60$ )        | 57.59 ( $\pm 6.97$ )        | 64.25 ( $\pm 1.65$ )        | 60.80 ( $\pm 1.57$ )        | 59.08 ( $\pm 2.63$ )        |
|       | Single-13L          | 63.63 ( $\pm 2.35$ )        | 63.02 ( $\pm 1.92$ )        | <b>66.59</b> ( $\pm 1.86$ ) | 65.06 ( $\pm 1.60$ )        | 65.51 ( $\pm 1.46$ )        | 64.76 ( $\pm 1.84$ )        |
|       | Parallel-13L (ours) | <b>65.38</b> ( $\pm 0.47$ ) | <b>65.07</b> ( $\pm 1.38$ ) | 65.65 ( $\pm 5.20$ )        | <b>68.62</b> ( $\pm 1.19$ ) | <b>66.07</b> ( $\pm 1.77$ ) | <b>66.16</b> ( $\pm 2.00$ ) |
| 1%    | Single-102L         | 30.56 ( $\pm 6.21$ )        | 27.26 ( $\pm 8.12$ )        | 27.53 ( $\pm 8.04$ )        | 30.42 ( $\pm 7.12$ )        | 23.17 ( $\pm 7.53$ )        | 27.79 ( $\pm 7.41$ )        |
|       | Single-13L          | <b>32.98</b> ( $\pm 3.42$ ) | 30.46 ( $\pm 5.47$ )        | 28.12 ( $\pm 4.77$ )        | 34.90 ( $\pm 2.78$ )        | 31.26 ( $\pm 2.93$ )        | 31.54 ( $\pm 3.87$ )        |
|       | Parallel-13L (ours) | 24.78 ( $\pm 6.28$ )        | <b>33.14</b> ( $\pm 2.55$ ) | <b>39.18</b> ( $\pm 6.41$ ) | <b>37.75</b> ( $\pm 2.91$ ) | <b>34.29</b> ( $\pm 1.34$ ) | <b>33.83</b> ( $\pm 3.90$ ) |

Table 8: NusaX-senti benchmarking results. ‘ace’, ‘ban’, ‘jav’, ‘min’, and ‘sun’ denote Acehnese, Balinese, Javanese, Minangkabau, and Sundanese, respectively.

| #data | Tokenizer           | amh                         | hau                         | ibo                         | kin                         | orm                         | swa                         | tir                         | twi                         | avg                         |
|-------|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| 100%  | Single-102L         | 32.78 ( $\pm 5.06$ )        | 69.25 ( $\pm 0.61$ )        | 86.33 ( $\pm 1.59$ )        | 69.93 ( $\pm 1.95$ )        | 62.11 ( $\pm 3.26$ )        | 90.28 ( $\pm 0.31$ )        | 36.39 ( $\pm 4.05$ )        | 55.58 ( $\pm 2.27$ )        | 62.83 ( $\pm 2.38$ )        |
|       | Single-13L          | 59.72 ( $\pm 1.80$ )        | 69.51 ( $\pm 1.94$ )        | 85.36 ( $\pm 0.74$ )        | 72.96 ( $\pm 1.54$ )        | <b>62.21</b> ( $\pm 0.73$ ) | <b>90.29</b> ( $\pm 0.39$ ) | 60.91 ( $\pm 1.09$ )        | <b>55.91</b> ( $\pm 1.27$ ) | 69.61 ( $\pm 1.19$ )        |
|       | Parallel-13L (ours) | <b>59.91</b> ( $\pm 0.35$ ) | <b>70.93</b> ( $\pm 2.19$ ) | <b>87.62</b> ( $\pm 0.85$ ) | <b>73.44</b> ( $\pm 1.56$ ) | 61.45 ( $\pm 0.59$ )        | 89.94 ( $\pm 0.22$ )        | <b>61.61</b> ( $\pm 1.07$ ) | 53.47 ( $\pm 2.77$ )        | <b>69.80</b> ( $\pm 1.20$ ) |
| 50%   | Single-102L         | 34.47 ( $\pm 2.34$ )        | 67.86 ( $\pm 0.28$ )        | 83.70 ( $\pm 1.90$ )        | 67.68 ( $\pm 1.98$ )        | 59.34 ( $\pm 2.45$ )        | 88.50 ( $\pm 0.37$ )        | 37.96 ( $\pm 0.18$ )        | 52.66 ( $\pm 2.24$ )        | 61.52 ( $\pm 1.47$ )        |
|       | Single-13L          | 55.15 ( $\pm 1.98$ )        | 66.23 ( $\pm 1.76$ )        | 84.30 ( $\pm 0.20$ )        | <b>71.21</b> ( $\pm 1.57$ ) | 59.44 ( $\pm 0.77$ )        | 87.41 ( $\pm 0.46$ )        | <b>61.33</b> ( $\pm 1.21$ ) | <b>53.01</b> ( $\pm 2.95$ ) | 67.26 ( $\pm 1.36$ )        |
|       | Parallel-13L (ours) | <b>57.23</b> ( $\pm 0.65$ ) | <b>67.89</b> ( $\pm 2.15$ ) | <b>85.47</b> ( $\pm 0.67$ ) | 69.39 ( $\pm 1.45$ )        | <b>60.16</b> ( $\pm 0.86$ ) | <b>88.60</b> ( $\pm 0.08$ ) | 59.50 ( $\pm 0.76$ )        | 50.96 ( $\pm 0.92$ )        | <b>67.40</b> ( $\pm 0.94$ ) |
| 10%   | Single-102L         | 32.07 ( $\pm 6.33$ )        | 58.86 ( $\pm 1.08$ )        | 74.01 ( $\pm 1.08$ )        | 58.65 ( $\pm 4.15$ )        | 48.25 ( $\pm 5.85$ )        | <b>85.12</b> ( $\pm 0.72$ ) | 35.44 ( $\pm 3.44$ )        | 45.15 ( $\pm 1.17$ )        | 54.69 ( $\pm 2.98$ )        |
|       | Single-13L          | 50.83 ( $\pm 0.29$ )        | 61.34 ( $\pm 0.87$ )        | 77.39 ( $\pm 1.80$ )        | 56.89 ( $\pm 1.52$ )        | 50.49 ( $\pm 6.57$ )        | 82.78 ( $\pm 1.13$ )        | <b>51.52</b> ( $\pm 2.53$ ) | <b>47.13</b> ( $\pm 1.98$ ) | 59.80 ( $\pm 2.09$ )        |
|       | Parallel-13L (ours) | <b>51.27</b> ( $\pm 2.07$ ) | <b>62.50</b> ( $\pm 1.78$ ) | <b>79.93</b> ( $\pm 1.09$ ) | <b>63.93</b> ( $\pm 0.15$ ) | <b>51.08</b> ( $\pm 0.38$ ) | 84.24 ( $\pm 1.19$ )        | 49.10 ( $\pm 1.47$ )        | 39.29 ( $\pm 3.87$ )        | <b>60.17</b> ( $\pm 1.50$ ) |
| 1%    | Single-102L         | 38.30 ( $\pm 2.98$ )        | 40.12 ( $\pm 5.33$ )        | 33.43 ( $\pm 5.68$ )        | 37.84 ( $\pm 1.97$ )        | <b>42.42</b> ( $\pm 7.43$ ) | 70.75 ( $\pm 0.55$ )        | 35.88 ( $\pm 9.04$ )        | 27.81 ( $\pm 0.22$ )        | 40.82 ( $\pm 4.15$ )        |
|       | Single-13L          | 41.41 ( $\pm 2.79$ )        | <b>43.31</b> ( $\pm 7.82$ ) | 47.25 ( $\pm 6.60$ )        | 42.37 ( $\pm 3.92$ )        | 41.98 ( $\pm 4.11$ )        | 69.11 ( $\pm 2.53$ )        | <b>41.49</b> ( $\pm 4.06$ ) | 30.70 ( $\pm 3.79$ )        | 44.70 ( $\pm 4.45$ )        |
|       | Parallel-13L (ours) | <b>42.70</b> ( $\pm 1.03$ ) | 43.09 ( $\pm 3.84$ )        | <b>55.44</b> ( $\pm 3.25$ ) | <b>48.10</b> ( $\pm 1.81$ ) | 42.35 ( $\pm 3.00$ )        | <b>71.33</b> ( $\pm 2.22$ ) | 39.56 ( $\pm 2.28$ )        | <b>31.26</b> ( $\pm 2.74$ ) | <b>46.73</b> ( $\pm 2.52$ ) |

Table 9: AfriHate benchmarking results. ‘amh’, ‘hau’, ‘ibo’, ‘kin’, ‘orm’, ‘swa’, ‘tir’, and ‘twi’ denote Amharic, Hausa, Igbo, Kinyarwanda, Oromo, Swahili, Tigrinya, and Twi, respectively.

| #data | Tokenizer           | amh                         | orm                         | tir                         | avg                         |
|-------|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| 100%  | Single-102L         | 12.06 ( $\pm 10.55$ )       | <b>60.80</b> ( $\pm 0.73$ ) | 17.43 ( $\pm 9.62$ )        | 30.10 ( $\pm 6.97$ )        |
|       | Single-13L          | 55.07 ( $\pm 1.01$ )        | 58.96 ( $\pm 1.18$ )        | 50.92 ( $\pm 0.75$ )        | 54.98 ( $\pm 0.98$ )        |
|       | Parallel-13L (ours) | <b>60.22</b> ( $\pm 0.32$ ) | 59.59 ( $\pm 1.18$ )        | <b>51.21</b> ( $\pm 1.03$ ) | <b>57.01</b> ( $\pm 0.84$ ) |
| 50%   | Single-102L         | 18.16 ( $\pm 2.62$ )        | <b>58.22</b> ( $\pm 0.29$ ) | 19.49 ( $\pm 3.68$ )        | 31.96 ( $\pm 2.20$ )        |
|       | Single-13L          | 52.98 ( $\pm 2.09$ )        | 55.53 ( $\pm 1.06$ )        | 45.48 ( $\pm 1.31$ )        | 51.33 ( $\pm 1.48$ )        |
|       | Parallel-13L (ours) | <b>58.33</b> ( $\pm 1.12$ ) | 55.64 ( $\pm 1.78$ )        | <b>48.86</b> ( $\pm 2.44$ ) | <b>54.27</b> ( $\pm 1.78$ ) |
| 10%   | Single-102L         | 12.51 ( $\pm 5.24$ )        | 39.85 ( $\pm 5.11$ )        | 13.84 ( $\pm 12.06$ )       | 22.07 ( $\pm 7.47$ )        |
|       | Single-13L          | <b>44.32</b> ( $\pm 2.86$ ) | 43.18 ( $\pm 1.03$ )        | <b>38.79</b> ( $\pm 2.66$ ) | <b>42.10</b> ( $\pm 2.19$ ) |
|       | Parallel-13L (ours) | 44.00 ( $\pm 2.86$ )        | <b>43.36</b> ( $\pm 1.73$ ) | 36.68 ( $\pm 2.04$ )        | 41.35 ( $\pm 2.21$ )        |
| 1%    | Single-102L         | 10.97 ( $\pm 9.50$ )        | 6.45 ( $\pm 11.17$ )        | 14.78 ( $\pm 12.85$ )       | 10.73 ( $\pm 11.17$ )       |
|       | Single-13L          | <b>34.85</b> ( $\pm 0.34$ ) | 24.37 ( $\pm 3.50$ )        | <b>22.91</b> ( $\pm 1.81$ ) | <b>27.37</b> ( $\pm 1.88$ ) |
|       | Parallel-13L (ours) | 33.77 ( $\pm 1.83$ )        | <b>24.42</b> ( $\pm 1.67$ ) | 21.55 ( $\pm 0.62$ )        | 26.58 ( $\pm 1.37$ )        |

Table 10: EthioEmo benchmarking results. ‘amh’, ‘orm’, and ‘tir’ denote Amharic, Oromo, and Tigrinya, respectively.

| #data | Tokenizer           | hau                         | ibo                         | kin                         | sun                         | swa                         | avg                         |
|-------|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| 100%  | Single-102L         | 56.84 ( $\pm 1.34$ )        | 57.71 ( $\pm 1.16$ )        | 41.42 ( $\pm 1.79$ )        | <b>61.01</b> ( $\pm 1.39$ ) | <b>25.37</b> ( $\pm 0.94$ ) | 48.47 ( $\pm 1.32$ )        |
|       | Single-13L          | 58.66 ( $\pm 0.16$ )        | <b>58.72</b> ( $\pm 0.63$ ) | 41.65 ( $\pm 0.58$ )        | 60.11 ( $\pm 0.61$ )        | 22.16 ( $\pm 2.71$ )        | 48.26 ( $\pm 0.94$ )        |
|       | Parallel-13L (ours) | <b>60.20</b> ( $\pm 1.96$ ) | 58.37 ( $\pm 1.57$ )        | <b>45.41</b> ( $\pm 2.10$ ) | 60.49 ( $\pm 1.05$ )        | 23.95 ( $\pm 2.76$ )        | <b>49.68</b> ( $\pm 1.89$ ) |
| 50%   | Single-102L         | 51.34 ( $\pm 0.63$ )        | 54.09 ( $\pm 0.91$ )        | 37.82 ( $\pm 1.16$ )        | 58.25 ( $\pm 2.64$ )        | <b>21.89</b> ( $\pm 0.33$ ) | 44.68 ( $\pm 1.13$ )        |
|       | Single-13L          | 53.89 ( $\pm 1.29$ )        | <b>55.08</b> ( $\pm 0.58$ ) | 39.36 ( $\pm 1.83$ )        | <b>59.77</b> ( $\pm 1.05$ ) | 20.76 ( $\pm 2.52$ )        | 45.77 ( $\pm 1.46$ )        |
|       | Parallel-13L (ours) | <b>56.65</b> ( $\pm 0.81$ ) | 55.05 ( $\pm 1.65$ )        | <b>43.42</b> ( $\pm 0.81$ ) | 57.56 ( $\pm 2.02$ )        | 21.13 ( $\pm 2.26$ )        | <b>46.76</b> ( $\pm 1.51$ ) |
| 10%   | Single-102L         | 39.67 ( $\pm 1.50$ )        | 42.25 ( $\pm 1.83$ )        | 28.33 ( $\pm 1.57$ )        | <b>55.10</b> ( $\pm 3.32$ ) | <b>18.21</b> ( $\pm 1.13$ ) | 36.71 ( $\pm 1.87$ )        |
|       | Single-13L          | 43.03 ( $\pm 1.65$ )        | <b>45.12</b> ( $\pm 2.80$ ) | 27.62 ( $\pm 0.60$ )        | 50.63 ( $\pm 4.83$ )        | 16.78 ( $\pm 1.25$ )        | 36.64 ( $\pm 2.23$ )        |
|       | Parallel-13L (ours) | <b>45.73</b> ( $\pm 0.70$ ) | 45.00 ( $\pm 2.27$ )        | <b>33.12</b> ( $\pm 0.68$ ) | 51.95 ( $\pm 2.11$ )        | 16.74 ( $\pm 0.58$ )        | <b>38.51</b> ( $\pm 1.27$ ) |
| 1%    | Single-102L         | 0.00 ( $\pm 0.00$ )         | 0.00 ( $\pm 0.00$ )         | 0.00 ( $\pm 0.00$ )         | 0.00 ( $\pm 0.00$ )         | 0.00 ( $\pm 0.00$ )         | 0.00 ( $\pm 0.00$ )         |
|       | Single-13L          | 9.02 ( $\pm 8.00$ )         | <b>21.28</b> ( $\pm 3.23$ ) | 14.89 ( $\pm 0.76$ )        | 42.50 ( $\pm 0.84$ )        | 0.00 ( $\pm 0.00$ )         | 17.54 ( $\pm 2.57$ )        |
|       | Parallel-13L (ours) | <b>14.91</b> ( $\pm 3.24$ ) | 15.97 ( $\pm 13.91$ )       | <b>17.20</b> ( $\pm 1.93$ ) | <b>46.30</b> ( $\pm 0.72$ ) | 0.00 ( $\pm 0.00$ )         | <b>18.88</b> ( $\pm 3.96$ ) |

Table 11: BRIGHTER benchmarking results. ‘hau’, ‘ibo’, ‘kin’, ‘sun’, and ‘swa’ denote Hausa, Igbo, Kinyarwanda, Sundanese, and Swahili, respectively.

|     | ace | amh   | ban          | hau   | ibo   | jav   | min   | kin   | orm          | sun   | swa   | tir          | twi   |
|-----|-----|-------|--------------|-------|-------|-------|-------|-------|--------------|-------|-------|--------------|-------|
| ace |     | 98.81 | <b>79.74</b> | 93.58 | 93.87 | 79.35 | 83.70 | 95.06 | <b>96.15</b> | 82.61 | 90.22 | 98.91        | 91.70 |
| amh |     |       | 99.11        | 99.90 | 99.80 | 99.41 | 99.31 | 99.80 | 99.60        | 98.62 | 99.51 | <b>88.34</b> | 99.80 |
| ban |     |       |              | 93.28 | 93.68 | 56.52 | 72.33 | 95.16 | 96.15        | 61.96 | 80.24 | 98.81        | 90.71 |
| hau |     |       |              |       | 93.48 | 93.48 | 94.96 | 95.16 | 96.44        | 95.65 | 96.34 | 98.91        | 94.37 |
| ibo |     |       |              |       |       | 94.17 | 95.45 | 96.15 | 97.92        | 96.44 | 95.95 | 99.41        | 96.84 |
| jav |     |       |              |       |       |       | 59.88 | 94.47 | 96.25        | 36.56 | 57.11 | 99.31        | 91.21 |
| min |     |       |              |       |       |       |       | 96.15 | 96.74        | 55.93 | 75.89 | 99.21        | 92.98 |
| kin |     |       |              |       |       |       |       |       | 98.52        | 96.15 | 95.65 | 99.21        | 97.33 |
| orm |     |       |              |       |       |       |       |       |              | 98.02 | 98.42 | 99.11        | 97.83 |
| sun |     |       |              |       |       |       |       |       |              |       | 68.28 | 99.31        | 92.59 |
| swa |     |       |              |       |       |       |       |       |              |       |       | 99.41        | 95.65 |
| tir |     |       |              |       |       |       |       |       |              |       |       |              | 99.90 |
| twi |     |       |              |       |       |       |       |       |              |       |       |              |       |

Table 12: *Single-102L* bitext mining scores, calculated using xsim error rate score.

|     | ace | amh          | ban   | hau   | ibo   | jav          | min          | kin          | orm   | sun          | swa   | tir          | twi          |
|-----|-----|--------------|-------|-------|-------|--------------|--------------|--------------|-------|--------------|-------|--------------|--------------|
| ace |     | <b>97.83</b> | 82.21 | 85.67 | 88.34 | 73.12        | <b>72.83</b> | <b>91.01</b> | 97.04 | <b>77.57</b> | 86.26 | <b>98.52</b> | <b>89.62</b> |
| amh |     |              | 96.74 | 96.34 | 96.94 | 96.44        | 96.74        | 97.53        | 97.73 | 96.94        | 96.54 | 90.51        | 98.02        |
| ban |     |              |       | 72.83 | 77.87 | <b>35.97</b> | 48.02        | 84.19        | 96.34 | 47.43        | 71.54 | 98.22        | 85.28        |
| hau |     |              |       |       | 64.43 | 62.15        | 72.13        | 82.91        | 95.26 | 68.38        | 63.54 | 97.63        | 81.72        |
| ibo |     |              |       |       |       | 70.95        | 78.46        | 87.15        | 96.84 | 75.69        | 71.05 | 98.02        | 85.97        |
| jav |     |              |       |       |       |              | 43.58        | 85.38        | 96.44 | <b>34.88</b> | 63.64 | 98.42        | 83.70        |
| min |     |              |       |       |       |              |              | 84.19        | 96.34 | <b>42.89</b> | 69.76 | 98.32        | 84.09        |
| kin |     |              |       |       |       |              |              |              | 97.04 | <b>84.68</b> | 78.36 | 98.32        | 88.54        |
| orm |     |              |       |       |       |              |              |              |       | 97.23        | 96.64 | <b>98.12</b> | <b>96.84</b> |
| sun |     |              |       |       |       |              |              |              |       |              | 67.29 | 98.62        | 84.58        |
| swa |     |              |       |       |       |              |              |              |       |              |       | 98.52        | 82.11        |
| tir |     |              |       |       |       |              |              |              |       |              |       |              | 98.91        |
| twi |     |              |       |       |       |              |              |              |       |              |       |              |              |

Table 13: *Single-13L* bitext mining scores, calculated using xsim error rate score.

|     | ace | amh   | ban          | hau          | ibo          | jav          | min          | kin          | orm          | sun          | swa          | tir          | twi          |
|-----|-----|-------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|
| ace |     | 98.32 | 88.04        | <b>76.98</b> | <b>85.77</b> | <b>71.34</b> | 81.23        | 98.62        | 97.92        | 82.31        | <b>70.36</b> | 98.81        | 98.22        |
| amh |     |       | <b>77.87</b> | <b>73.62</b> | <b>85.08</b> | <b>69.96</b> | <b>72.43</b> | <b>83.30</b> | <b>95.45</b> | <b>81.92</b> | <b>65.61</b> | 97.33        | <b>86.96</b> |
| ban |     |       |              | <b>57.11</b> | <b>68.18</b> | 50.79        | <b>46.25</b> | <b>70.95</b> | <b>85.87</b> | <b>46.25</b> | <b>47.73</b> | <b>95.45</b> | <b>76.48</b> |
| hau |     |       |              |              | <b>52.87</b> | <b>35.57</b> | <b>41.40</b> | <b>60.18</b> | <b>82.91</b> | <b>57.11</b> | <b>30.93</b> | <b>89.43</b> | <b>70.55</b> |
| ibo |     |       |              |              |              | <b>58.10</b> | <b>68.28</b> | <b>78.36</b> | <b>87.35</b> | <b>72.53</b> | <b>54.25</b> | <b>97.23</b> | <b>77.27</b> |
| jav |     |       |              |              |              |              | <b>22.92</b> | <b>56.92</b> | <b>79.84</b> | 44.27        | <b>24.51</b> | <b>92.98</b> | <b>75.30</b> |
| min |     |       |              |              |              |              |              | <b>60.77</b> | <b>79.45</b> | 87.15        | <b>29.74</b> | <b>95.16</b> | <b>71.44</b> |
| kin |     |       |              |              |              |              |              |              | <b>92.29</b> | 89.03        | <b>59.19</b> | <b>96.15</b> | <b>81.92</b> |
| orm |     |       |              |              |              |              |              |              |              | <b>93.08</b> | <b>83.70</b> | 98.62        | 97.23        |
| sun |     |       |              |              |              |              |              |              |              |              | <b>47.33</b> | <b>92.89</b> | <b>72.92</b> |
| swa |     |       |              |              |              |              |              |              |              |              |              | <b>87.65</b> | <b>70.36</b> |
| tir |     |       |              |              |              |              |              |              |              |              |              |              | <b>98.62</b> |
| twi |     |       |              |              |              |              |              |              |              |              |              |              |              |

Table 14: *Parallel-13L* bitext mining scores, calculated using xsim error rate score.

| #data | Tokenizer           | ace                         | ban                         | jav                         | min                         | sun                         | avg                         |
|-------|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| 50%   | Single-102L         | 71.15 ( $\pm 1.11$ )        | 68.03 ( $\pm 1.24$ )        | 74.70 ( $\pm 1.05$ )        | 71.64 ( $\pm 1.93$ )        | 71.83 ( $\pm 0.98$ )        | 71.76 ( $\pm 1.26$ )        |
|       | Single-13L          | <b>75.99</b> ( $\pm 0.88$ ) | 70.69 ( $\pm 2.97$ )        | <b>78.00</b> ( $\pm 0.82$ ) | 74.92 ( $\pm 2.56$ )        | <b>76.50</b> ( $\pm 0.72$ ) | 75.22 ( $\pm 1.59$ )        |
|       | Parallel-13L (ours) | 73.75 ( $\pm 0.62$ )        | <b>71.94</b> ( $\pm 1.05$ ) | 77.01 ( $\pm 1.52$ )        | <b>78.15</b> ( $\pm 0.70$ ) | 76.22 ( $\pm 2.10$ )        | <b>75.42</b> ( $\pm 1.20$ ) |
| 0%    | Single-102L         | 63.10 ( $\pm 1.92$ )        | 67.80 ( $\pm 2.28$ )        | 72.52 ( $\pm 1.66$ )        | 68.41 ( $\pm 0.97$ )        | 70.44 ( $\pm 1.28$ )        | 68.80 ( $\pm 1.62$ )        |
|       | Single-13L          | 66.77 ( $\pm 2.59$ )        | <b>69.76</b> ( $\pm 2.20$ ) | 73.93 ( $\pm 2.37$ )        | 74.09 ( $\pm 1.85$ )        | <b>74.60</b> ( $\pm 0.95$ ) | 71.83 ( $\pm 1.99$ )        |
|       | Parallel-13L (ours) | <b>69.44</b> ( $\pm 1.94$ ) | 67.43 ( $\pm 1.27$ )        | <b>76.66</b> ( $\pm 0.71$ ) | <b>75.87</b> ( $\pm 2.05$ ) | 74.12 ( $\pm 1.50$ )        | <b>72.70</b> ( $\pm 1.49$ ) |

Table 15: NusaX-senti benchmarking results under a limited target-language data setting. ‘ace’, ‘ban’, ‘jav’, ‘min’, and ‘sun’ denote Acehnese, Balinese, Javanese, Minangkabau, and Sundanese, respectively.

| #data | Tokenizer           | amh                         | hau                         | ibo                         | kin                         | orm                         | swh                         | tir                         | twi                         | avg                         |
|-------|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| 50%   | Single-102L         | 34.81 ( $\pm 1.71$ )        | <b>67.85</b> ( $\pm 0.83$ ) | 83.68 ( $\pm 0.67$ )        | 58.28 ( $\pm 3.41$ )        | 69.06 ( $\pm 1.60$ )        | 88.58 ( $\pm 0.33$ )        | 34.69 ( $\pm 3.07$ )        | <b>54.65</b> ( $\pm 2.64$ ) | 61.45 ( $\pm 1.78$ )        |
|       | Single-13L          | <b>57.92</b> ( $\pm 0.69$ ) | 65.35 ( $\pm 1.71$ )        | 84.20 ( $\pm 0.43$ )        | 59.20 ( $\pm 1.52$ )        | 69.23 ( $\pm 0.80$ )        | 88.22 ( $\pm 0.55$ )        | 57.75 ( $\pm 3.59$ )        | 54.63 ( $\pm 2.30$ )        | <b>67.06</b> ( $\pm 1.45$ ) |
|       | Parallel-13L (ours) | 57.01 ( $\pm 0.13$ )        | 66.18 ( $\pm 1.73$ )        | <b>85.27</b> ( $\pm 1.52$ ) | <b>59.35</b> ( $\pm 2.04$ ) | <b>70.77</b> ( $\pm 0.62$ ) | <b>88.73</b> ( $\pm 0.11$ ) | <b>58.57</b> ( $\pm 1.41$ ) | 50.13 ( $\pm 3.01$ )        | 67.00 ( $\pm 1.32$ )        |
| 0%    | Single-102L         | 31.74 ( $\pm 3.35$ )        | 38.28 ( $\pm 0.81$ )        | 34.82 ( $\pm 1.71$ )        | 28.28 ( $\pm 1.67$ )        | 34.10 ( $\pm 1.50$ )        | 37.58 ( $\pm 2.75$ )        | <b>39.93</b> ( $\pm 2.01$ ) | 25.01 ( $\pm 1.15$ )        | 33.72 ( $\pm 1.87$ )        |
|       | Single-13L          | 42.32 ( $\pm 2.74$ )        | 39.32 ( $\pm 3.11$ )        | 38.91 ( $\pm 1.91$ )        | 36.86 ( $\pm 0.78$ )        | 40.06 ( $\pm 4.86$ )        | 42.06 ( $\pm 0.67$ )        | 35.27 ( $\pm 2.89$ )        | 23.37 ( $\pm 4.55$ )        | 37.27 ( $\pm 2.69$ )        |
|       | Parallel-13L (ours) | <b>46.38</b> ( $\pm 2.74$ ) | <b>45.13</b> ( $\pm 3.04$ ) | <b>45.13</b> ( $\pm 4.18$ ) | <b>41.32</b> ( $\pm 1.07$ ) | <b>44.78</b> ( $\pm 4.19$ ) | <b>44.11</b> ( $\pm 5.10$ ) | 32.45 ( $\pm 2.41$ )        | <b>26.08</b> ( $\pm 1.68$ ) | <b>40.67</b> ( $\pm 3.05$ ) |

Table 16: AfriHate benchmarking results under a limited target-language data setting. ‘amh’, ‘hau’, ‘ibo’, ‘kin’, ‘orm’, ‘swh’, ‘tir’, and ‘twi’ denote Amharic, Hausa, Igbo, Kinyarwanda, Oromo, Swahili, Tigrinya, and Twi, respectively.

| #data | Tokenizer           | amh                         | orm                         | tir                         | avg                         |
|-------|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| 50%   | Single-102L         | 14.89 ( $\pm 1.15$ )        | 53.83 ( $\pm 4.11$ )        | 20.53 ( $\pm 7.24$ )        | 29.75 ( $\pm 4.17$ )        |
|       | Single-13L          | 53.96 ( $\pm 1.26$ )        | 54.65 ( $\pm 1.48$ )        | 46.60 ( $\pm 0.91$ )        | 51.73 ( $\pm 1.22$ )        |
|       | Parallel-13L (ours) | <b>58.39</b> ( $\pm 0.51$ ) | <b>55.77</b> ( $\pm 0.82$ ) | <b>46.74</b> ( $\pm 3.27$ ) | <b>53.63</b> ( $\pm 1.53$ ) |
| 0%    | Single-102L         | 11.04 ( $\pm 6.25$ )        | 0.00 ( $\pm 0.00$ )         | 18.15 ( $\pm 5.83$ )        | 9.73 ( $\pm 4.03$ )         |
|       | Single-13L          | <b>32.58</b> ( $\pm 4.21$ ) | 12.10 ( $\pm 4.66$ )        | <b>26.55</b> ( $\pm 0.65$ ) | <b>23.75</b> ( $\pm 3.17$ ) |
|       | Parallel-13L (ours) | 27.46 ( $\pm 3.60$ )        | <b>18.61</b> ( $\pm 0.27$ ) | 23.53 ( $\pm 4.77$ )        | 23.20 ( $\pm 2.88$ )        |

Table 17: EthioEmo benchmarking results under a limited target-language data setting. ‘amh’, ‘orm’, and ‘tir’ denote Amharic, Oromo, and Tigrinya, respectively.

BRIGHTER. On average across all benchmarks, however, *Parallel-13L* outperforms *Single-13L* by 0.42%.

## G Continual Pre-Training: Representation Similarity

As discussed in Section 5.6, although the benchmarking performance of *Single-13L* and *Parallel-13L* after CPT is comparable, the cross-lingual representation similarity is stronger in *Parallel-13L*. This is illustrated in Figure 6, where we visualize the embedding spaces of FLORES+ (Costa-Jussà et al., 2022) sentences for each language in two dimensions using PCA. In *Single-13L*, inputs from the same language family remain clustered. For instance, Indonesian local languages (Acehnese, Minangkabau, Sundanese, etc.) cluster together, African languages cluster together, and Amharic-script languages (Amharic and Tigrinya) are separated into distinct clusters. In contrast, *Parallel-13L* shows tighter cross-lingual clustering, with languages grouped more uniformly, including Amharic. The main outliers are Tigrinya and Twi, which remain separated in both models, likely due to limited pretraining resources.

These observations are further supported by our bitext mining analysis on the models after CPT. As shown in Table 23, *Parallel-13L* consistently achieves the lowest error rates and the highest number of best xsim scores, outperforming the other models, although the gains are smaller than those observed when pretraining from scratch. Tables 24 and 25 provide the detailed xsim scores for each language pair in both *Single-13L* and *Parallel-13L*.

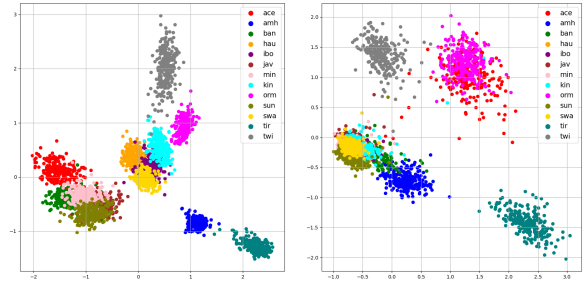


Figure 6: Principal Component Analysis (PCA) of continual pretrained model embeddings, comparing the *Single-13L* (left) and the *Parallel-13L* (right) on the parallel data FLORES+ dataset.

## H Model Input Representation Design

Since we require a ‘language token’ to select the appropriate tokenizer for a given input, our input representation differs slightly from the conventional approach. During the model design phase, as presented in Figure 7, we explored several alternatives and ultimately selected the most effective method. Specifically, we considered three designs: (1) placing the language token at the beginning of every input sequence, requiring the model to process it as the first token; (2) the *Parallel-13L* approach, where the language token serves as a signal that is added to each token embedding, effectively combining token and language embeddings; and (3) augmenting only the embeddings of unaligned vocabulary items (i.e., tokens not belonging to the word-type category<sup>9</sup>).

We pretrained models with each of these meth-

<sup>9</sup>We store unaligned vocabulary items at the end of the vocabulary indices

| #data | Tokenizer           | hau                         | ibo                         | kin                         | sun                         | swa                         | avg                         |
|-------|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| 50%   | Single-102L         | 50.75 ( $\pm 0.85$ )        | 54.30 ( $\pm 0.95$ )        | 35.29 ( $\pm 0.75$ )        | <b>58.34</b> ( $\pm 0.80$ ) | 21.89 ( $\pm 1.51$ )        | 44.11 ( $\pm 0.97$ )        |
|       | Single-13L          | 52.92 ( $\pm 1.09$ )        | 55.41 ( $\pm 0.32$ )        | 39.57 ( $\pm 2.08$ )        | 58.18 ( $\pm 1.78$ )        | 18.95 ( $\pm 1.33$ )        | 45.00 ( $\pm 1.32$ )        |
|       | Parallel-13L (ours) | <b>57.36</b> ( $\pm 0.62$ ) | <b>56.64</b> ( $\pm 0.95$ ) | <b>41.72</b> ( $\pm 0.83$ ) | 57.89 ( $\pm 1.38$ )        | <b>21.96</b> ( $\pm 1.39$ ) | <b>47.11</b> ( $\pm 1.04$ ) |
| 0%    | Single-102L         | 13.71 ( $\pm 2.70$ )        | 13.02 ( $\pm 2.54$ )        | 19.47 ( $\pm 2.37$ )        | 15.41 ( $\pm 1.41$ )        | <b>16.70</b> ( $\pm 2.54$ ) | 15.66 ( $\pm 2.31$ )        |
|       | Single-3L           | 15.13 ( $\pm 3.30$ )        | <b>16.31</b> ( $\pm 0.61$ ) | 9.72 ( $\pm 0.59$ )         | <b>18.93</b> ( $\pm 2.24$ ) | 15.13 ( $\pm 1.45$ )        | 15.04 ( $\pm 1.64$ )        |
|       | Parallel-13L (ours) | <b>26.92</b> ( $\pm 2.65$ ) | 16.08 ( $\pm 2.15$ )        | <b>23.37</b> ( $\pm 0.32$ ) | 15.90 ( $\pm 0.96$ )        | 16.38 ( $\pm 3.31$ )        | <b>19.73</b> ( $\pm 1.88$ ) |

Table 18: BRIGHTER benchmarking results under a limited target-language data setting. ‘hau’, ‘ibo’, ‘kin’, ‘sun’, and ‘swa’ denote Hausa, Igbo, Kinyarwanda, Sundanese, and Swahili, respectively.

| Tokenizer           | ace                         | ban                         | jav                         | min                         | sun                         | avg                         |
|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| Single-13L          | <b>74.86</b> ( $\pm 1.02$ ) | <b>77.01</b> ( $\pm 0.80$ ) | <b>77.92</b> ( $\pm 1.62$ ) | 77.14 ( $\pm 1.00$ )        | <b>76.65</b> ( $\pm 2.51$ ) | <b>76.72</b> ( $\pm 1.39$ ) |
| Parallel-13L (ours) | 73.27 ( $\pm 1.18$ )        | 73.76 ( $\pm 0.43$ )        | 77.75 ( $\pm 0.44$ )        | <b>77.85</b> ( $\pm 0.43$ ) | 75.76 ( $\pm 2.70$ )        | 75.68 ( $\pm 1.04$ )        |

Table 19: NusaX-senti benchmarking results under monolingual setting. ‘ace’, ‘ban’, ‘jav’, ‘min’, and ‘sun’ denote Acehnese, Balinese, Javanese, Minangkabau, and Sundanese, respectively.

| Tokenizer           | amh                         | hau                         | ibo                         | kin                         | orm                         | swh                         | tir                         | twi                         | avg                         |
|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| Single-13L          | <b>60.92</b> ( $\pm 2.41$ ) | 74.39 ( $\pm 0.65$ )        | 84.96 ( $\pm 0.83$ )        | <b>63.82</b> ( $\pm 1.20$ ) | <b>73.60</b> ( $\pm 1.29$ ) | 90.13 ( $\pm 0.18$ )        | 61.76 ( $\pm 1.66$ )        | <b>56.49</b> ( $\pm 3.00$ ) | <b>62.17</b> ( $\pm 1.75$ ) |
| Parallel-13L (ours) | 60.39 ( $\pm 0.89$ )        | <b>75.35</b> ( $\pm 1.52$ ) | <b>87.70</b> ( $\pm 0.26$ ) | 62.94 ( $\pm 1.54$ )        | 73.59 ( $\pm 0.41$ )        | <b>90.30</b> ( $\pm 0.42$ ) | <b>61.94</b> ( $\pm 2.77$ ) | 55.20 ( $\pm 1.15$ )        | 61.76 ( $\pm 1.73$ )        |

Table 20: AfriHate benchmarking results under monolingual setting. ‘amh’, ‘hau’, ‘ibo’, ‘kin’, ‘orm’, ‘swh’, ‘tir’, and ‘twi’ denote Amharic, Hausa, Igbo, Kinyarwanda, Oromo, Swahili, Tigrinya, and Twi, respectively.

| Tokenizer           | amh                         | orm                         | tir                         | avg                         |
|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| Single-13L          | 58.70 ( $\pm 1.32$ )        | <b>60.15</b> ( $\pm 0.75$ ) | 48.92 ( $\pm 1.24$ )        | 55.92 ( $\pm 1.10$ )        |
| Parallel-13L (ours) | <b>60.98</b> ( $\pm 1.09$ ) | 59.63 ( $\pm 1.63$ )        | <b>52.58</b> ( $\pm 0.20$ ) | <b>57.73</b> ( $\pm 0.97$ ) |

Table 21: EthioEmo benchmarking results under monolingual setting. ‘amh’, ‘orm’, and ‘tir’ denote Amharic, Oromo, and Tigrinya, respectively.

| Tokenizer           | hau                         | ibo                         | kin                         | sun                         | swa                         | avg                         |
|---------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|-----------------------------|
| Single-13L          | 57.90 ( $\pm 1.40$ )        | 58.89 ( $\pm 1.53$ )        | 43.89 ( $\pm 0.20$ )        | <b>61.05</b> ( $\pm 0.93$ ) | 23.81 ( $\pm 1.14$ )        | 49.11 ( $\pm 1.04$ )        |
| Parallel-13L (ours) | <b>60.23</b> ( $\pm 0.35$ ) | <b>59.23</b> ( $\pm 0.71$ ) | <b>47.74</b> ( $\pm 2.45$ ) | 60.97 ( $\pm 0.79$ )        | <b>24.04</b> ( $\pm 2.15$ ) | <b>50.44</b> ( $\pm 1.29$ ) |

Table 22: BRIGHTER benchmarking results under monolingual setting. ‘hau’, ‘ibo’, ‘kin’, ‘sun’, and ‘swa’, denote Hausa, Igbo, Kinyarwanda, Sundanese, and Swahili, respectively.

|                  | Single-102L | Single-13L | Parallel-13L |
|------------------|-------------|------------|--------------|
| Avg. xsim ↓      | 91.33       | 72.18      | <b>69.34</b> |
| # of best xsim ↑ | 1           | 29         | <b>48</b>    |

Table 23: Bitext mining metric meta-scores on CPT models.

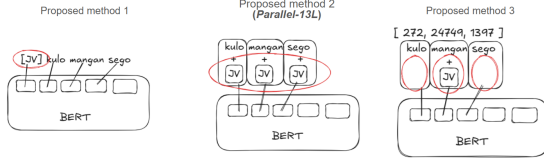


Figure 7: The design of the model input representation. We adopt the second method and refer to it as *Parallel-13L*.

ods from scratch using the same dataset (Wikipedia dumps for the 13 languages studied in this paper). The first and third approaches did not converge as quickly or effectively as *Parallel-13L*, which therefore became our chosen design for the experiments reported in Section 5 and beyond.

|     | ace | amh          | ban          | hau          | ibo          | jav          | min          | kin          | orm          | sun          | swa          | tir          | twi          |
|-----|-----|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|
| ace |     | <b>91.01</b> | <b>65.42</b> | 68.28        | <b>76.09</b> | <b>55.83</b> | <b>67.59</b> | <b>82.61</b> | <b>93.68</b> | 79.05        | 68.28        | <b>97.73</b> | <b>82.51</b> |
| amh |     |              | 93.48        | 90.42        | 89.72        | 88.24        | 90.51        | 91.90        | <b>94.96</b> | 89.92        | 90.81        | 92.29        | 94.76        |
| ban |     |              |              | <b>40.02</b> | <b>50.40</b> | <b>27.96</b> | <b>44.86</b> | 74.51        | 91.40        | 50.59        | 38.54        | 97.33        | 77.08        |
| hau |     |              |              |              | 36.26        | 29.25        | 41.80        | 67.00        | 90.61        | <b>40.32</b> | 24.90        | 95.55        | 69.76        |
| ibo |     |              |              |              |              | <b>43.68</b> | <b>55.14</b> | 74.11        | 92.59        | <b>53.56</b> | 40.51        | <b>96.44</b> | 76.38        |
| jav |     |              |              |              |              |              | 24.41        | 70.85        | 91.70        | <b>28.06</b> | 29.05        | 95.65        | 75.30        |
| min |     |              |              |              |              |              |              | <b>69.96</b> | 90.51        | <b>28.36</b> | 38.14        | 96.54        | 75.30        |
| kin |     |              |              |              |              |              |              |              | 94.37        | <b>73.62</b> | 66.21        | <b>96.84</b> | 85.57        |
| orm |     |              |              |              |              |              |              |              |              | <b>91.11</b> | 90.81        | <b>97.43</b> | <b>93.08</b> |
| sun |     |              |              |              |              |              |              |              |              |              | <b>33.10</b> | 96.05        | <b>75.79</b> |
| swa |     |              |              |              |              |              |              |              |              |              |              | 95.16        | 73.22        |
| tir |     |              |              |              |              |              |              |              |              |              |              |              | 98.02        |
| twi |     |              |              |              |              |              |              |              |              |              |              |              |              |

Table 24: Bitext mining metric scores on CPT models trained with *Single-13L* tokenizer.

|     | ace | amh   | ban          | hau          | ibo          | jav          | min          | kin          | orm          | sun          | swa          | tir          | twi          |
|-----|-----|-------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|--------------|
| ace |     | 98.02 | 77.37        | <b>62.55</b> | 80.53        | 57.41        | 90.32        | 98.62        | 97.83        | <b>68.58</b> | <b>59.98</b> | 98.52        | 89.43        |
| amh |     |       | <b>74.90</b> | <b>62.65</b> | <b>78.46</b> | <b>63.54</b> | <b>64.62</b> | <b>76.88</b> | 95.06        | <b>71.15</b> | <b>59.49</b> | 97.33        | <b>86.07</b> |
| ban |     |       |              | 41.01        | 61.07        | 30.43        | 56.82        | <b>63.34</b> | <b>85.18</b> | <b>37.06</b> | <b>34.19</b> | <b>92.98</b> | <b>76.68</b> |
| hau |     |       |              |              | <b>31.82</b> | <b>20.06</b> | <b>28.75</b> | <b>47.92</b> | <b>79.15</b> | 53.56        | <b>17.39</b> | <b>90.61</b> | <b>63.24</b> |
| ibo |     |       |              |              |              | 43.97        | 69.86        | <b>66.60</b> | <b>85.18</b> | 62.35        | <b>36.76</b> | 97.04        | <b>72.63</b> |
| jav |     |       |              |              |              |              | <b>22.43</b> | <b>50.79</b> | <b>76.98</b> | 52.87        | <b>22.13</b> | <b>92.69</b> | <b>69.57</b> |
| min |     |       |              |              |              |              |              | 76.98        | <b>89.53</b> | 57.21        | <b>25.10</b> | <b>92.79</b> | <b>72.23</b> |
| kin |     |       |              |              |              |              |              |              | <b>91.70</b> | 83.60        | <b>54.64</b> | 97.04        | <b>78.56</b> |
| orm |     |       |              |              |              |              |              |              |              | 96.34        | <b>80.53</b> | 98.81        | 98.22        |
| sun |     |       |              |              |              |              |              |              |              |              | 40.32        | <b>90.12</b> | 85.28        |
| swa |     |       |              |              |              |              |              |              |              |              |              | <b>91.21</b> | <b>68.08</b> |
| tir |     |       |              |              |              |              |              |              |              |              |              |              | <b>97.92</b> |
| twi |     |       |              |              |              |              |              |              |              |              |              |              |              |

Table 25: Bitext mining metric scores on CPT models trained with *Parallel-13L* tokenizer.