

# AFRIDOC-MT: Document-level MT Corpus for African Languages

Jesujoba O. Alabi<sup>1,†</sup>, Israel Abebe Azime<sup>1,†</sup>, Miaoran Zhang<sup>1</sup>, Cristina España-Bonet<sup>2,3</sup>, Rachel Bawden<sup>4</sup>, Dawei Zhu<sup>1</sup>, David Ifeoluwa Adelani<sup>5,†</sup>, Clement Oyeleke Odoje<sup>6</sup>, Idris Akinade<sup>6,†</sup>, Iffat Maab<sup>7</sup>, Davis David<sup>8,†</sup>, Shamsuddeen Hassan Muhammad<sup>9,†</sup>, Neo Putini<sup>10,†</sup>, David O. Ademuyiwa<sup>11</sup>, Andrew Caines<sup>12</sup>, Dietrich Klakow<sup>1</sup>

<sup>†</sup>Masakhane NLP, <sup>1</sup>Saarland University, Saarland Informatic Campus, <sup>2</sup>DFKI GmbH, <sup>3</sup>Barcelona Supercomputing Center,

<sup>4</sup>Inria, Paris, France, <sup>5</sup>Mila, McGill University & Canada CIFAR AI Chair, <sup>6</sup>University of Ibadan, Nigeria,

<sup>7</sup>National Institute of Informatics, Japan, <sup>8</sup>Selcom Tanzania, <sup>9</sup>Imperial College London, <sup>10</sup>University of KwaZulu-Natal,

<sup>11</sup>Loughborough University, U.K., <sup>12</sup>University of Cambridge, U.K.

## Abstract

This paper introduces AFRIDOC-MT, a document-level multi-parallel translation dataset covering English and five African languages: Amharic, Hausa, Swahili, Yorùbá, and Zulu. The dataset comprises 334 health and 271 information technology news documents, all human-translated from English to these languages. We conduct document-level translation benchmark experiments by evaluating the ability of neural machine translation (NMT) models and large language models (LLMs) to translate between English and these languages, at both the sentence and pseudo-document levels, the outputs being realigned to form complete documents for evaluation. Our results indicate that NLLB-200 achieves the best average performance among the standard NMT models, while GPT-4o outperforms general-purpose LLMs. Fine-tuning selected models leads to substantial performance gains, but models trained on sentences struggle to generalize effectively to longer documents. Furthermore, our analysis reveals that some LLMs exhibit issues such as under-generation, over-generation, repetition of words and phrases, and off-target translations, specifically for translation into African languages.<sup>1</sup>

## 1 Introduction

The field of machine translation (MT) has seen notable progress in the past years, with neural machine translation (NMT) models achieving close to human performance in many high-resource language directions (Vaswani et al., 2017; Akhbardeh et al., 2021; Mohammadshahi et al., 2022; Yuan et al., 2023; Kocmi et al., 2023; NLLB Team et al., 2024). However, efforts have primarily been concentrated on sentence-level translation, without the use of inter-sentential context.

In recent years, there has been interest in document-level translation (i.e. the holistic translation of multiple sentences), where sentences are translated with their context rather than in isolation. Document-level translation is important in order to capture discourse relations (Bawden et al., 2018; Voita et al., 2018; Maruf et al., 2021), maintain consistency and coherence across sentences (Herold and Ney, 2023), particularly for technical domains, but poses unique challenges, such as how to handle longer documents (Wang et al., 2024b) given the limited context size of translation models. Current efforts have primarily focused on high-resource language directions, where document-level datasets are readily available (Lopes et al., 2020; Feng et al., 2022; Wu et al., 2023; Wang et al., 2023; Wu et al., 2024), and so far there has been no work on low-resource African languages. Developing and evaluating document-level MT systems for low-resource languages is a useful and under-studied direction, which requires the creation of datasets.

To fill this gap, we present AFRIDOC-MT, a document-level translation dataset for English from and into five African languages: Amharic, Hausa, Swahili, Yorùbá, and Zulu, created through the manual translation of English documents. It consists of 334 *health* documents and 271 *tech* documents. In addition, AFRIDOC-MT supports multi-way translation, allowing translations not only between English and the African languages but also between any two of the languages covered.

We conduct a comprehensive set of document translation benchmark experiments on AFRIDOC-MT, using sentence-level and pseudo-document translation due to most models’ limited context length, and then realigning them to form complete documents. We evaluate performance using automatic metrics and compare the results of encoder-decoder models with decoder-only LLMs

<sup>1</sup>The data can be accessed on [HuggingFace](#) and [GitHub](#).

| Dataset                               | #Langs. | Multiway | Domain       | Type           | #Sents. | #Docs.  |
|---------------------------------------|---------|----------|--------------|----------------|---------|---------|
| TICO-19 (Anastasopoulos et al., 2020) | 12      | ✓        | health       | document-level | 4k      | 30      |
| MAFAND-MT (Adelani et al., 2022)      | 16      | ✗        | news         | sentence-level | 4k-35k  | -       |
| FLORES-200 (NLLB Team et al., 2022)   | 42      | ✓        | general      | sentence-level | 3k      | -       |
| NTREX-128 (Federmann et al., 2022)    | 24      | ✓        | news         | sentence-level | 1.9k    | -       |
| AFRIDOC-MT (Ours)                     | 5       | ✓        | tech, health | document-level | 10k     | 271-334 |

Table 1: Overview of highly related work, including for each dataset the number of African languages, the domain, the kind of MT task they can be used for and the range of the sentence numbers for each language direction.

across both domains. Our results demonstrate that NLLB-200, both before and after fine-tuning on AFRIDOC-MT, excels in sentence translation, surpassing all other models. GPT-4o performs equally well for sentences and pseudo-documents, while other decoder-only models lag behind. In addition to automatic metrics, we use GPT-4o as a judge, human evaluation, and qualitative assessment to compare documents translation carried out sentence-by-sentence and as pseudo-documents for selected models. The evaluation shows that GPT-4o is generally unreliable for assessing document translations into African languages. However, we observe agreement between other evaluation methods, all indicating that sentence-by-sentence translation results in better document-level translation into African languages. We conduct additional analyses of the models’ outputs to better understand their behavior and why they under-perform when translating pseudo-documents. They show that LLMs often under-generate, contain repetitions, and produce off-target translations, especially when translating into African languages.

## 2 Related Work

**MT Datasets for African Languages** Several MT datasets exist for African languages, including web-mined datasets such as WikiMatrix (Schwenk et al., 2021a) and CCMatrix (Schwenk et al., 2021b). However, they have been adjudged to be of poor quality for certain low-resource subsets, including African languages (Kreutzer et al., 2022). There are also well curated datasets for African languages including the Bible (McCarthy et al., 2020), JW300 (Agić and Vulić, 2019)<sup>2</sup> and MAFAND-MT (Adelani et al., 2022), which are from religious and news domains.

There exist several MT evaluation benchmark datasets for African languages. They can be categorized into two kinds. First, evaluation datasets specifically designed for translating into or from

African languages (Ezeani et al., 2020; Azunre et al., 2021; Adelani et al., 2021, 2022, *inter alia*). Second, benchmark datasets covering many languages, including African languages. For example, TICO-19 (Anastasopoulos et al., 2020), NTREX-128 (Federmann et al., 2022), FLORES-101 (Goyal et al., 2022) and FLORES-200 (NLLB Team et al., 2024) are a few such datasets. However, most of these datasets are designed for sentence-level MT, primarily drawn from religious or news domains, although some consist of translated sentences originating from the same document. To the best of our knowledge, only TICO-19, a health domain translation benchmark, has the potential to be used for document-level MT, while it is restricted to topics related to COVID-19. Table 1 gives a comparison of the most relevant existing benchmarks.

### Document-level Neural Machine Translation

Document-level NMT aims to overcome the limitations of sentence-level systems by translating an entire document as a whole. Similar to context-aware NMT, which involves translating segments with additional, localized context, it differs in that it involves in principle translating an entire document holistically. Both document-level and context-aware MT allow for the possibility of improving translation quality for context-dependent phenomena such as coreference resolution (Müller et al., 2018; Bawden et al., 2018; Voita et al., 2018; Herold and Ney, 2023), lexical disambiguation (Rios Gonzales et al., 2017; Martínez Garcia et al., 2019), and lexical cohesion (Wong and Kit, 2012; Garcia et al., 2014, 2017; Bawden et al., 2018; Voita et al., 2019). Various methods have been proposed to extend sentence-level models to capture document-level context (Tiedemann and Scherrer, 2017; Libovický and Helcl, 2017; Bawden et al., 2018; Miculicich et al., 2018; Sun et al., 2022). The emergence of LLMs, such as GPT-3 (Brown et al., 2020), Llama (Dubey et al., 2024) and Gemma (Gemma Team et al., 2024), has transformed NLP, including for MT (Zhu et al., 2024a,c;

<sup>2</sup>The dataset is no longer available for use.

| Language      | Classification          | Spkrs. (M) |
|---------------|-------------------------|------------|
| Amharic [amh] | Afro-Asiatic/Semitic    | 57.6       |
| Hausa [hau]   | Afro-Asiatic/Chadic     | 78.5       |
| Swahili [swa] | Niger-Congo/Bantu       | 71.6       |
| Yorùbá [yor]  | Niger-Congo/Volta-Niger | 45.9       |
| isiZulu [zul] | Niger-Congo/Bantu       | 27.8       |

Table 2: Languages in the AFRIDOC-MT corpus, their classification and number of speakers (in millions).

Lu et al., 2024). Pre-trained on vast amounts of text, LLMs can effectively manage long-range dependencies, making them in principle well-suited for document-level translation. While these models have shown promising results for high-resource languages (Wu et al., 2023; Wang et al., 2023; Wu et al., 2024), research remains limited for low-resource languages (Ul Haq et al., 2020).

### 3 AFRIDOC-MT Corpus

**Languages and their characteristics** We cover five languages from the two most common African language families: Afro-Asiatic and Niger-Congo. Three languages belong to the Niger-Congo family: Swahili (North-East Bantu), Yorùbá (Volta-Niger) and isiZulu (Southern Bantu). The other two languages belong to the Afro-Asiatic family: Amharic (Semitic) and Hausa (Chadic). The choice of languages was based on geographical representation, speaking population, and web coverage (which we consider as a proxy for the potential performance of existing models on these languages). Furthermore, each of these languages has over 20 million speakers. All of them use the Latin script except for Amharic, which uses the Ge’ez script. The Latin-script languages use the Latin alphabet with the omission of some letters and the addition of new ones, and the use of diacritics (e.g., Yorùbá). The languages are tonal, except for Amharic and Swahili. Just like English, all languages follow the subject-verb-object word order. Refer to Adelani (2022) for a comprehensive overview of the characteristics of these languages. Table 2 shows summarized details.

**Data Collection and Preprocessing** We scraped English articles from the websites of Techpoint Africa<sup>3</sup> and the World Health Organization (WHO).<sup>4,5</sup> The articles cover different topics of different lengths with an average length of 30 and 37

sentences for *health* and *tech* respectively. While our corpus is initially structured at the article level, we aim to make it suitable for sentence-level translation tasks as well. To achieve this, we segmented the raw articles into sentences using NLTK (Bird et al., 2009). To ensure high segmentation quality, we recruited a linguist and a professional translator to verify the correctness of the segmentation and make corrections as needed. Finally, we selected 334 and 271 English articles/documents from the *health* and *tech* domains respectively, which represents 10k sentences each per domain.

**Translation** We translated the extracted 10k English sentences to the 5 African languages through 4 expert translators per language.<sup>6</sup> The translators were recruited through a language coordinator who is also a native speaker of the language. The 10k sentences were distributed equally among the translators and the translations were done in-context (i.e. the translators translated on the sentence level but had access to the whole document). Due to the domain-specific nature of the task, before starting the translation process, we conducted a translation workshop, during which three translation experts shared their experiences in creating terminologies and they also shared existing resources with the translators including short translation guidelines (Appendix A.1).

**Quality Checks** Quality control was conducted using automated quality estimation, followed by manual inspections by our language coordinators. We also used Quality Estimation (QE), specifically AfriCOMET (Wang et al., 2024a),<sup>7</sup> Given a translated sentence in any African language and its corresponding source English sentence, AfriCOMET generates a score between 0 and 1, where 0 indicates poor quality and higher values signify better quality. The translators, in collaboration with the language coordinators, were tasked with reviewing instances that had quality estimation scores below 0.65. This step was essential to identify and correct low-quality translations.

Figure 3 shows the distribution of final quality scores for five languages across both domains. Manual inspection indicates that QE scores below 0.65 do not necessarily reflect poor translations, consistent with Adelani et al. (2024b), likely due to

<sup>3</sup><https://techpoint.africa/>

<sup>4</sup><https://www.who.int/health-topics>

<sup>5</sup><https://www.who.int/news-room/>

<sup>6</sup>Each translator was paid \$1,250 for 2,500 sentences.

<sup>7</sup><https://huggingface.co/masakhane/africomet-qe-stl>

| Domain                     | Train | Dev. | Test | Min/Max/Avg |
|----------------------------|-------|------|------|-------------|
| <b>Number of documents</b> |       |      |      |             |
| <i>health</i>              | 240   | 33   | 61   | 2/151/29.9  |
| <i>tech</i>                | 187   | 25   | 59   | 8/247/36.9  |
| <b>Number of sentences</b> |       |      |      |             |
| <i>health</i>              | 7041  | 977  | 1982 | -           |
| <i>tech</i>                | 7048  | 970  | 1982 | -           |

Table 3: The number of documents and sentences in AFRIDOC-MT, and (at the document level) minimum, maximum and average sentences per document.

domain shift, translation length, and other factors, which warrant further investigation.

**AFRIDOC-MT data split** We created train, development (dev), and test splits for each domain. To prevent data leakage, documents sharing sentences with the same English translation were assigned to the training set. The dev set contains 25–33 documents, and the test set 59–61 documents, drawn from the remaining data. Table 4 shows the average number of whitespace-separated tokens per sentence across domains and languages, including English. The *health* domain has more tokens than *tech*. Hausa and Yorùbá are longer than English, likely due to their descriptive nature, while Swahili is similar in length. Amharic and Zulu are relatively shorter, reflecting interesting linguistic properties. Table 3 provides additional data statistics. Table 3 shows some data statistics.

The *health* data is licensed under CC BY-NC-SA 3.0, while the *tech* data is licensed under CC BY-NC-SA 4.0.

## 4 Benchmark Experiments

Given the AFRIDOC-MT data, we conducted both sentence- and document-level translation, evaluating two types of models: encoder-decoder and decoder-only models. While the majority of these models are open-source, we also evaluated two proprietary models of the same type. Our evaluation primarily focuses on document-level translation, reflecting the availability of our document-level translation corpus. For completeness, we also conduct a series of sentence-level experiments, with the results presented in Appendix C.

### 4.1 Models

**Encoder-Decoder Models** We evaluate five kinds of open encoder-decoder model including Toucan (Elmadany et al., 2024; Adebara et al.,

| Domain          | eng   | amh   | hau   | swa   | yor   | zul   |
|-----------------|-------|-------|-------|-------|-------|-------|
| <b>Sentence</b> |       |       |       |       |       |       |
| <i>health</i>   | 21.6  | 19.3  | 28.1  | 23.2  | 27.9  | 16.7  |
| <i>tech</i>     | 17.8  | 15.6  | 22.2  | 18.0  | 23.7  | 13.4  |
| <b>Document</b> |       |       |       |       |       |       |
| <i>health</i>   | 647.3 | 576.7 | 841.7 | 695.4 | 834.8 | 500.1 |
| <i>tech</i>     | 658.2 | 575.0 | 821.6 | 665.4 | 873.4 | 495.9 |

Table 4: The average number of tokens in AFRIDOC-MT, both at sentence and document level.

2024), M2M-100 (Fan et al., 2020), NLLB-200 (NLLB Team et al., 2024), MADLAD-400 (Kudugunta et al., 2023), and Aya-101 (Üstün et al., 2024). Toucan is an Afro-centric multilingual MT model supporting 150 African language pairs. In comparison, M2M-100, NLLB-200, and MADLAD-400 cover between 100 and 450 language pairs. Aya-101, an instruction-tuned mT5 model (Xue et al., 2021), supports 100 languages and can translate between various languages, including those considered in AFRIDOC-MT.

**Decoder-only Models** We also evaluate open and closed decoder-only models. Open models include LLama3.1 (Dubey et al., 2024), Gemma2 (Gemma Team et al., 2024), their instruction-tuned variants, and LLaMAX3 (Lu et al., 2024)—a LLama3-based model further pre-trained on 100+ languages, including several African ones. The closed models include OpenAI GPT models (GPT-3.5 Turbo and GPT-4o) (OpenAI, 2024), which have been shown to have document-level translation ability (Wang et al., 2023). While their language coverage is not well documented, they show some ability to handle African languages (Adelani et al., 2024b; Bayes et al., 2024), though far below their performance in English, their primary training language.

We present the result of 12 models in total, including the 1.2B version of Toucan, 1.3B and 3.3B versions of NLLB-200, 3B and 13B versions of MADLAD-400 and Aya-101 respectively. We also have the 8B instruction tuned version of LLama3.1 (LLama3.1-IT), 9B version of Gemma-2 (Gemma2-IT), and LLaMAX3-Alpaca.<sup>8</sup> We provide more description of the models in Appendix B.1.

**Supervised fine-tuning of the models** For sentence-level evaluation, we jointly fine-tune NLLB-200 with 1.3B parameters on the 30 language directions and on the two domains to make the models more specialized. Similarly, we did supervised fine-tuning (SFT) on LLaMAX3 and

<sup>8</sup>We refer to it as LLaMAX3-Alp in the results tables.

LLama3.1 using the prompt augmentation approach from (Zhu et al., 2024b), as shown in Appendix B.4. We chose these two models because LLaMAX3 is already adapted to several languages including our languages of interest, and LLama3.1 because of its long context window. We perform SFT on LLaMAX3 and LLama3.1 for document-level translation, using pseudo-documents with  $k=10$ . We refer to each system as {model\_name}-SFT<sub>k</sub>.<sup>9</sup>

## 4.2 Experimental Setup

**Sentence-level Evaluation** Given that our created dataset can be used for sentence-level translation and as a baseline for document-level translation, we evaluate all models on the test splits for each domain. We evaluate the translation models (M2M-100, NLLB-200, and MADLAD-400) using the Fairseq (Ott et al., 2019) codebase for (M2M-100 and NLLB-200), and the Transformers (Wolf et al., 2020) codebase for MADLAD-400. However, for other models including Aya-101, we use the EleutherAI LM Evaluation Harness (lm-eval) tool (Biderman et al., 2024) using the three templates listed in Table 23 of Appendix B.4.

**Document-level Evaluation** We also perform document-level translation using a setup similar to the sentence-level experiment, but only with models that meet context length requirements. An initial analysis showed that some models were unable to process entire documents due to input length limits, which were exceeded by token counts in some languages (Amharic and Yorùbá). To address this, we adopted a similar approach to Lee et al. (2022), splitting documents into fixed-size chunks of  $k$  sentences to fit within token limits; the final chunk may contain fewer than  $k$  sentences. To select an appropriate chunk size, we conducted initial tests with  $k = 1$  (sentence-level), 5, 10, and 25, choosing  $k = 10$  for our experiments. We provide results from this analysis in Table 11.

## 4.3 Evaluation Metrics

Evaluating document-level translation remains challenging, as existing automatic metrics struggle to capture improvements and account for discourse phenomena (Jiang et al., 2022; Dahan et al., 2024), and embedding-based metrics have not been explored in this context for African languages due

to the lack of data. Hence, we realigned sentence-level or pseudo-translation outputs into full documents, then computed BLEU (Papineni et al., 2002) and chrF (Popović, 2015) to create document BLEU (d-BLEU) and document chrF (d-chrF). Metrics were computed using SacreBLEU<sup>10</sup> (Post, 2018) with bootstrap resampling ( $n = 1000$ ) to report 95% confidence intervals. We report d-chrF scores for the best prompt per model and language direction in the main text, as chrF better captures the morphological richness of African languages (Adelani et al., 2022), with full results provided in Appendix C.

We use GPT-4o as a judge to evaluate translation outputs, following recent work showing LLMs’ effectiveness in assessing translation quality and analyzing errors (Wu et al., 2024; Sun et al., 2025). Following Sun et al. (2025), we assess each translated document’s fluency, content errors (CE), and cohesion errors—specifically lexical (LE) and grammatical (GE) errors—using GPT-4o, with evaluation limited to a few model outputs due to cost constraints (Appendix B.6). We also complement this with human evaluation for direct assessment scores (Appendix B.7) and qualitative analysis through manual inspection (Appendix B.8).

## 5 Results

### 5.1 Sentence-level Evaluation

In Tables 5 and 6 we present d-chrF scores based on the realigned documents, created by merging the translated sentences into their corresponding documents. We highlight our main findings below, and sentence-level evaluation results using sentence-level metrics are reported in Appendix C.

#### NLLB-200 outperforms all other encoder-decoder models across languages and domains

On average the NLLB models obtain scores of 65.4/66.6 and 64.3/65.0 on *health* and *tech* domains respectively, with 3.3B outperforming 1.3B except when translating into Yorùbá. When translating to English, the worst performing model across the two domains is Toucan. However, it gives better results than MADLAD-400 and Aya-101 when translating to African languages. Furthermore, translating to African languages is significantly worse compared to translating to English for all the models.

#### GPT-4o outperforms other decoder-only counterparts

GPT-4o on average outperforms other

<sup>9</sup>We denote models finetuned on sentences as {model\_name}-SFT or {model\_name}-SFT<sub>1</sub>

<sup>10</sup>case:mixed|eff:no|tok:13a|smooth:exp|v:2.3.1

| Model             | Size | $eng \rightarrow X$ |                     |                     |                     |                     |      | $X \rightarrow eng$ |                     |                     |                     |                     |      | AVG  |
|-------------------|------|---------------------|---------------------|---------------------|---------------------|---------------------|------|---------------------|---------------------|---------------------|---------------------|---------------------|------|------|
|                   |      | amh                 | hau                 | swa                 | yor                 | zul                 | Avg. | amh                 | hau                 | swa                 | yor                 | zul                 | Avg. |      |
| Encoder-Decoder   |      |                     |                     |                     |                     |                     |      |                     |                     |                     |                     |                     |      |      |
| Toucan            | 1.2B | 33.8 <sub>1.2</sub> | 57.6 <sub>1.4</sub> | 70.3 <sub>0.8</sub> | 36.0 <sub>1.5</sub> | 58.0 <sub>1.0</sub> | 51.2 | 54.7 <sub>1.0</sub> | 57.7 <sub>1.3</sub> | 65.2 <sub>0.9</sub> | 54.0 <sub>1.2</sub> | 59.9 <sub>0.8</sub> | 58.3 | 54.7 |
| NLLB-200          | 1.3B | 49.8 <sub>1.5</sub> | 64.7 <sub>2.2</sub> | 75.5 <sub>0.8</sub> | 45.1 <sub>1.0</sub> | 69.0 <sub>1.3</sub> | 60.8 | 69.4 <sub>1.3</sub> | 65.3 <sub>1.7</sub> | 75.3 <sub>0.8</sub> | 66.3 <sub>1.1</sub> | 73.2 <sub>0.9</sub> | 69.9 | 65.4 |
| MADLAD-400        | 3B   | 36.5 <sub>0.9</sub> | 54.4 <sub>2.0</sub> | 74.2 <sub>0.9</sub> | 19.1 <sub>0.9</sub> | 57.1 <sub>1.4</sub> | 48.3 | 68.9 <sub>1.1</sub> | 63.8 <sub>1.6</sub> | 76.1 <sub>0.6</sub> | 51.4 <sub>1.8</sub> | 68.9 <sub>0.9</sub> | 65.8 | 57.0 |
| NLLB-200          | 3.3B | 53.0 <sub>1.9</sub> | 65.2 <sub>2.2</sub> | 76.7 <sub>0.7</sub> | 43.8 <sub>1.1</sub> | 70.7 <sub>1.3</sub> | 61.9 | 70.9 <sub>1.3</sub> | 66.5 <sub>1.7</sub> | 77.0 <sub>0.7</sub> | 67.6 <sub>1.1</sub> | 74.7 <sub>1.0</sub> | 71.3 | 66.6 |
| Aya-101           | 13B  | 36.6 <sub>0.9</sub> | 56.4 <sub>1.5</sub> | 44.7 <sub>2.4</sub> | 31.2 <sub>1.4</sub> | 58.6 <sub>0.8</sub> | 45.5 | 64.6 <sub>1.1</sub> | 61.5 <sub>1.4</sub> | 70.8 <sub>0.8</sub> | 57.9 <sub>1.3</sub> | 67.4 <sub>0.8</sub> | 64.4 | 55.0 |
| SFT on AFriDOC-MT |      |                     |                     |                     |                     |                     |      |                     |                     |                     |                     |                     |      |      |
| NLLB-SFT          | 1.3B | 55.9 <sub>1.6</sub> | 67.4 <sub>1.9</sub> | 81.3 <sub>0.7</sub> | 61.5 <sub>1.0</sub> | 73.7 <sub>1.6</sub> | 68.0 | 72.4 <sub>1.2</sub> | 67.5 <sub>1.6</sub> | 79.2 <sub>0.7</sub> | 71.8 <sub>1.1</sub> | 76.5 <sub>0.9</sub> | 73.5 | 70.7 |
| Decoder-only      |      |                     |                     |                     |                     |                     |      |                     |                     |                     |                     |                     |      |      |
| Gemma2-IT         | 9B   | 20.1 <sub>0.7</sub> | 56.4 <sub>1.4</sub> | 71.2 <sub>0.7</sub> | 21.0 <sub>0.6</sub> | 41.6 <sub>1.1</sub> | 42.1 | 61.6 <sub>0.9</sub> | 62.5 <sub>1.3</sub> | 74.2 <sub>0.7</sub> | 54.7 <sub>1.3</sub> | 63.9 <sub>0.9</sub> | 63.4 | 52.7 |
| LLama3.1-IT       | 8B   | 19.6 <sub>0.5</sub> | 45.9 <sub>1.4</sub> | 63.7 <sub>0.9</sub> | 19.7 <sub>0.6</sub> | 28.5 <sub>0.7</sub> | 35.5 | 53.9 <sub>0.9</sub> | 59.8 <sub>1.3</sub> | 69.1 <sub>0.9</sub> | 53.4 <sub>1.3</sub> | 54.0 <sub>1.1</sub> | 58.0 | 46.8 |
| LLaMAX3-Alp       | 8B   | 30.5 <sub>0.8</sub> | 56.3 <sub>1.5</sub> | 67.8 <sub>0.8</sub> | 19.3 <sub>0.8</sub> | 56.1 <sub>0.9</sub> | 46.0 | 63.3 <sub>1.0</sub> | 62.4 <sub>1.3</sub> | 71.7 <sub>0.8</sub> | 56.1 <sub>1.1</sub> | 65.3 <sub>0.9</sub> | 63.8 | 54.9 |
| GPT-3.5           | —    | 20.4 <sub>0.6</sub> | 44.3 <sub>0.9</sub> | 76.7 <sub>0.6</sub> | 21.3 <sub>0.9</sub> | 51.1 <sub>0.9</sub> | 42.8 | 48.3 <sub>0.9</sub> | 52.4 <sub>1.2</sub> | 75.0 <sub>0.6</sub> | 52.1 <sub>1.2</sub> | 59.5 <sub>0.9</sub> | 57.4 | 50.1 |
| GPT-4o            | —    | 36.7 <sub>0.8</sub> | 64.2 <sub>1.9</sub> | 79.8 <sub>0.6</sub> | 29.3 <sub>1.6</sub> | 69.0 <sub>1.3</sub> | 55.8 | 67.2 <sub>1.0</sub> | 66.5 <sub>1.5</sub> | 78.1 <sub>0.6</sub> | 69.1 <sub>1.1</sub> | 75.1 <sub>1.0</sub> | 71.2 | 63.5 |
| SFT on AFriDOC-MT |      |                     |                     |                     |                     |                     |      |                     |                     |                     |                     |                     |      |      |
| LLaMAX3-SFT       | 8B   | 46.8 <sub>1.2</sub> | 62.5 <sub>1.4</sub> | 73.1 <sub>0.9</sub> | 57.5 <sub>1.0</sub> | 67.5 <sub>1.0</sub> | 61.5 | 66.6 <sub>1.2</sub> | 58.9 <sub>1.6</sub> | 73.1 <sub>1.1</sub> | 64.7 <sub>1.5</sub> | 70.5 <sub>1.0</sub> | 66.8 | 64.1 |
| LLama3.1-SFT      | 8B   | 45.6 <sub>1.1</sub> | 61.8 <sub>1.5</sub> | 71.5 <sub>1.0</sub> | 57.0 <sub>1.1</sub> | 66.8 <sub>0.9</sub> | 60.6 | 64.3 <sub>1.2</sub> | 59.5 <sub>1.5</sub> | 72.1 <sub>0.8</sub> | 64.8 <sub>1.5</sub> | 69.0 <sub>1.0</sub> | 65.9 | 63.2 |

Table 5: Performance of the models in the Health domain, measured by d-chrF at the sentence-level, realigned to the document-level. For each model and language, the best result from three prompt variations is reported.

| Model              | Size | $eng \rightarrow X$ |                     |                     |                     |                     |      | $X \rightarrow eng$ |                     |                     |                     |                     |      | AVG  |
|--------------------|------|---------------------|---------------------|---------------------|---------------------|---------------------|------|---------------------|---------------------|---------------------|---------------------|---------------------|------|------|
|                    |      | amh                 | hau                 | swa                 | yor                 | zul                 | Avg. | amh                 | hau                 | swa                 | yor                 | zul                 | Avg. |      |
| Encoder-Decoder    |      |                     |                     |                     |                     |                     |      |                     |                     |                     |                     |                     |      |      |
| Toucan             | 1.2B | 32.0 <sub>1.6</sub> | 59.5 <sub>1.7</sub> | 66.1 <sub>1.7</sub> | 37.1 <sub>2.0</sub> | 58.5 <sub>1.4</sub> | 50.7 | 54.0 <sub>1.6</sub> | 59.9 <sub>1.5</sub> | 64.1 <sub>1.4</sub> | 54.3 <sub>1.3</sub> | 59.6 <sub>1.2</sub> | 58.4 | 54.5 |
| NLLB-200           | 1.3B | 49.3 <sub>2.0</sub> | 65.7 <sub>2.2</sub> | 72.3 <sub>1.6</sub> | 43.0 <sub>1.3</sub> | 70.3 <sub>1.3</sub> | 60.1 | 69.5 <sub>1.0</sub> | 66.8 <sub>1.5</sub> | 72.0 <sub>1.4</sub> | 63.0 <sub>1.2</sub> | 71.5 <sub>1.2</sub> | 68.5 | 64.3 |
| MADLAD-400         | 3B   | 37.3 <sub>1.3</sub> | 57.0 <sub>2.8</sub> | 62.1 <sub>2.9</sub> | 21.3 <sub>1.0</sub> | 58.5 <sub>1.8</sub> | 47.3 | 68.6 <sub>1.1</sub> | 66.0 <sub>1.4</sub> | 72.1 <sub>1.4</sub> | 53.1 <sub>1.4</sub> | 67.6 <sub>1.2</sub> | 65.5 | 56.4 |
| NLLB-200           | 3.3B | 52.2 <sub>2.4</sub> | 65.4 <sub>2.3</sub> | 72.8 <sub>1.5</sub> | 40.1 <sub>1.8</sub> | 71.6 <sub>1.3</sub> | 60.4 | 70.9 <sub>1.0</sub> | 67.7 <sub>1.5</sub> | 73.2 <sub>1.4</sub> | 63.9 <sub>1.1</sub> | 72.5 <sub>1.2</sub> | 69.6 | 65.0 |
| Aya-101            | 13B  | 37.3 <sub>1.1</sub> | 58.9 <sub>2.3</sub> | 42.4 <sub>2.6</sub> | 31.4 <sub>1.4</sub> | 58.9 <sub>1.5</sub> | 45.8 | 65.2 <sub>1.2</sub> | 64.8 <sub>1.2</sub> | 69.1 <sub>1.1</sub> | 58.5 <sub>1.3</sub> | 67.1 <sub>1.1</sub> | 64.9 | 55.4 |
| SFT on AFRI DOC-MT |      |                     |                     |                     |                     |                     |      |                     |                     |                     |                     |                     |      |      |
| NLLB-SFT           | 1.3B | 53.4 <sub>2.4</sub> | 67.9 <sub>2.2</sub> | 76.5 <sub>1.6</sub> | 59.5 <sub>1.3</sub> | 74.0 <sub>1.5</sub> | 66.2 | 72.1 <sub>1.0</sub> | 69.0 <sub>1.3</sub> | 74.1 <sub>1.4</sub> | 67.5 <sub>1.1</sub> | 74.3 <sub>1.1</sub> | 71.4 | 68.8 |
| Decoder-only       |      |                     |                     |                     |                     |                     |      |                     |                     |                     |                     |                     |      |      |
| Gemma2-IT          | 9B   | 20.6 <sub>0.6</sub> | 58.3 <sub>1.5</sub> | 68.7 <sub>1.6</sub> | 23.9 <sub>1.3</sub> | 46.5 <sub>1.8</sub> | 43.6 | 61.1 <sub>1.3</sub> | 65.4 <sub>1.4</sub> | 71.5 <sub>1.2</sub> | 56.7 <sub>1.3</sub> | 63.8 <sub>1.1</sub> | 63.7 | 53.7 |
| LLama3.1-IT        | 8B   | 19.5 <sub>0.9</sub> | 47.8 <sub>1.3</sub> | 63.4 <sub>1.5</sub> | 20.8 <sub>1.2</sub> | 30.4 <sub>1.3</sub> | 36.4 | 51.0 <sub>1.3</sub> | 61.0 <sub>1.4</sub> | 66.0 <sub>1.3</sub> | 53.5 <sub>1.2</sub> | 52.4 <sub>1.3</sub> | 56.8 | 46.6 |
| LLaMAX3-Alp        | 8B   | 30.3 <sub>1.1</sub> | 58.9 <sub>1.9</sub> | 64.9 <sub>1.7</sub> | 22.0 <sub>0.8</sub> | 58.6 <sub>1.7</sub> | 46.9 | 63.4 <sub>1.4</sub> | 64.9 <sub>1.5</sub> | 69.1 <sub>1.1</sub> | 56.5 <sub>1.3</sub> | 65.7 <sub>1.2</sub> | 63.9 | 55.4 |
| GPT-3.5            | –    | 22.6 <sub>0.8</sub> | 49.2 <sub>1.5</sub> | 72.6 <sub>1.6</sub> | 23.0 <sub>1.0</sub> | 53.6 <sub>1.5</sub> | 44.2 | 47.4 <sub>1.5</sub> | 56.5 <sub>1.3</sub> | 71.5 <sub>1.4</sub> | 54.0 <sub>1.3</sub> | 59.9 <sub>1.1</sub> | 57.9 | 51.0 |
| GPT-4o             | –    | 36.9 <sub>1.2</sub> | 65.2 <sub>2.3</sub> | 75.3 <sub>1.6</sub> | 29.4 <sub>1.5</sub> | 71.1 <sub>1.4</sub> | 55.6 | 67.2 <sub>1.0</sub> | 69.1 <sub>1.4</sub> | 74.4 <sub>1.4</sub> | 66.4 <sub>1.1</sub> | 73.4 <sub>1.1</sub> | 70.1 | 62.8 |
| SFT on AFRI DOC-MT |      |                     |                     |                     |                     |                     |      |                     |                     |                     |                     |                     |      |      |
| LLaMAX3-SFT        | 8B   | 42.8 <sub>1.5</sub> | 62.4 <sub>1.9</sub> | 67.6 <sub>1.4</sub> | 55.2 <sub>1.5</sub> | 66.0 <sub>1.2</sub> | 58.8 | 63.0 <sub>1.2</sub> | 53.5 <sub>1.9</sub> | 67.5 <sub>1.2</sub> | 57.3 <sub>1.3</sub> | 66.8 <sub>1.3</sub> | 61.6 | 60.2 |
| LLama3.1-SFT       | 8B   | 41.6 <sub>1.7</sub> | 61.8 <sub>2.0</sub> | 66.4 <sub>1.3</sub> | 54.9 <sub>1.4</sub> | 64.6 <sub>1.6</sub> | 57.9 | 62.0 <sub>1.2</sub> | 58.6 <sub>1.5</sub> | 67.1 <sub>1.2</sub> | 61.3 <sub>1.3</sub> | 65.6 <sub>1.3</sub> | 62.9 | 60.4 |

Table 6: Performance of the models in the Tech domain, measured by d-chrF at the sentence-level, realigned to the document-level. For each model and language, the best result from three prompt variations is reported.

decoder-only LMs, with average d-chrF scores of 63.5 and 62.8 for health and tech respectively. The next best performing decoder-only model is LLaMAX3-Alpaca, with d-chrF scores of 54.9 and 55.4. Unlike other open decoder-based LLMs, LLaMAX3-Alpaca was trained on African languages through continued pretraining and adapted via instruction tuning. It outperforms Gemma2-IT by +2.2 in the health domain and +1.7 in the *tech* domain, particularly when translating into African languages. In contrast, GPT-3.5 and LLaMa3.1-IT are the worst performing models.

**Fine-tuning models significantly improves translation quality** We obtain improved performance after fine-tuning NLLB-1.3B on AFRI DOC-MT, and the resulting model outperforms the 3.3B version without fine-tuning. Similarly, the SFT-based LLMs (LLaMAX3 and LLaMa3.1) become the best performing open LLMs and outperform their base-lines (LLaMAX3-Alpaca and LLaMa3.1-IT) but below GPT-4o. Overall, our fine-tuned NLLB-200

model is the state-of-the-art model, and our fine-tuned LLaMAX3 is competitive to GPT-4o.

## 5.2 Document-level Evaluation

In Tables 7 and 8 we present d-chrF scores based on the best prompt per language for the translation output of the models when evaluated on the realigned documents from pseudo-documents with  $k=10$  sentences per pseudo-document.

**Pseudo-document translation is worse than sentence-level translation when translating into African languages** Our results from pseudo-document translation show a performance drop across different models compared to sentence-level translation, especially when translating into African languages. However, GPT-4o demonstrates similar and consistent performance in both setups and domains. Additionally, we observe that GPT-3.5 is the next best performing decoder-only LLM, which contrasts with its performance in sentence-level translation. Gemma2-IT outper-

| Model                      | Size | $eng \rightarrow X$        |                            |                            |                            |                            |             | $X \rightarrow eng$        |                            |                            |                            |                            |             | AVG         |
|----------------------------|------|----------------------------|----------------------------|----------------------------|----------------------------|----------------------------|-------------|----------------------------|----------------------------|----------------------------|----------------------------|----------------------------|-------------|-------------|
|                            |      | amh                        | hau                        | swa                        | yor                        | zul                        | Avg.        | amh                        | hau                        | swa                        | yor                        | zul                        | Avg.        |             |
| Encoder-Decoder            |      |                            |                            |                            |                            |                            |             |                            |                            |                            |                            |                            |             |             |
| MADLAD-400                 | 3B   | 27.5 <sub>1.8</sub>        | 40.2 <sub>2.3</sub>        | 46.6 <sub>3.4</sub>        | 15.1 <sub>0.8</sub>        | 43.6 <sub>2.6</sub>        | 34.6        | 63.3 <sub>1.6</sub>        | 62.5 <sub>2.0</sub>        | 74.4 <sub>0.9</sub>        | 44.2 <sub>1.6</sub>        | 66.6 <sub>1.5</sub>        | 62.2        | 48.4        |
| Aya-101                    | 13B  | 28.7 <sub>1.6</sub>        | 48.5 <sub>2.3</sub>        | 34.7 <sub>3.4</sub>        | 18.7 <sub>1.3</sub>        | 54.9 <sub>1.4</sub>        | 37.1        | 61.6 <sub>1.7</sub>        | 62.3 <sub>1.8</sub>        | 71.2 <sub>0.9</sub>        | 56.1 <sub>2.1</sub>        | 69.0 <sub>1.0</sub>        | 64.0        | 50.6        |
| Decoder-only               |      |                            |                            |                            |                            |                            |             |                            |                            |                            |                            |                            |             |             |
| Gemma2-IT                  | 9B   | 6.5 <sub>0.6</sub>         | 37.0 <sub>3.4</sub>        | 52.9 <sub>3.6</sub>        | 6.4 <sub>0.5</sub>         | 12.0 <sub>1.0</sub>        | 23.0        | 36.5 <sub>3.0</sub>        | 51.8 <sub>3.4</sub>        | 65.0 <sub>3.0</sub>        | 44.8 <sub>2.9</sub>        | 56.1 <sub>3.3</sub>        | 50.8        | 36.9        |
| LLama3.1-IT                | 8B   | 7.5 <sub>0.5</sub>         | 14.0 <sub>1.2</sub>        | 43.2 <sub>3.9</sub>        | 6.4 <sub>0.7</sub>         | 8.7 <sub>0.6</sub>         | 16.0        | 23.8 <sub>2.3</sub>        | 49.3 <sub>4.1</sub>        | 62.8 <sub>3.3</sub>        | 31.7 <sub>3.9</sub>        | 34.0 <sub>3.7</sub>        | 40.3        | 28.1        |
| LLaMAX3-Alp                | 8B   | 11.4 <sub>0.9</sub>        | 28.9 <sub>2.9</sub>        | 40.4 <sub>3.2</sub>        | 9.2 <sub>0.8</sub>         | 23.6 <sub>1.8</sub>        | 22.7        | 29.2 <sub>2.1</sub>        | 41.7 <sub>3.8</sub>        | 55.4 <sub>4.9</sub>        | 23.5 <sub>3.0</sub>        | 40.5 <sub>4.7</sub>        | 38.1        | 30.4        |
| GPT-3.5                    | —    | 11.6 <sub>0.5</sub>        | 23.1 <sub>2.0</sub>        | 76.1 <sub>0.6</sub>        | 10.1 <sub>0.9</sub>        | 29.2 <sub>2.1</sub>        | 30.0        | 41.6 <sub>2.3</sub>        | 52.7 <sub>1.5</sub>        | 77.7 <sub>0.6</sub>        | 51.7 <sub>1.6</sub>        | 61.1 <sub>1.1</sub>        | 56.9        | 43.5        |
| GPT-4o                     | —    | 29.6 <sub>1.7</sub>        | <b>63.8</b> <sub>1.9</sub> | <b>80.2</b> <sub>0.6</sub> | 29.6 <sub>2.1</sub>        | <b>69.5</b> <sub>1.6</sub> | <b>54.5</b> | <b>69.5</b> <sub>1.1</sub> | <b>69.3</b> <sub>1.7</sub> | <b>81.0</b> <sub>0.6</sub> | <b>73.8</b> <sub>1.0</sub> | <b>78.2</b> <sub>1.1</sub> | <b>74.4</b> | <b>64.4</b> |
| SFT on AFriDOC-MT          |      |                            |                            |                            |                            |                            |             |                            |                            |                            |                            |                            |             |             |
| LLaMAX3-SFT                | 8B   | 24.1 <sub>1.6</sub>        | 29.0 <sub>3.2</sub>        | 42.2 <sub>4.2</sub>        | 33.8 <sub>2.8</sub>        | 33.7 <sub>3.1</sub>        | 32.6        | 22.6 <sub>1.8</sub>        | 22.9 <sub>2.6</sub>        | 33.1 <sub>4.4</sub>        | 27.2 <sub>3.6</sub>        | 31.5 <sub>6.7</sub>        | 27.5        | 30.0        |
| LLama3.1-SFT               | 8B   | 25.2 <sub>1.8</sub>        | 31.9 <sub>4.0</sub>        | 50.2 <sub>6.4</sub>        | 33.8 <sub>2.8</sub>        | 38.6 <sub>4.1</sub>        | 35.9        | 24.2 <sub>3.7</sub>        | 24.1 <sub>4.1</sub>        | 33.7 <sub>5.4</sub>        | 30.2 <sub>4.7</sub>        | 29.3 <sub>6.2</sub>        | 28.3        | 32.1        |
| LLaMAX3-SFT <sub>10</sub>  | 8B   | <b>37.8</b> <sub>2.2</sub> | 51.9 <sub>5.0</sub>        | 74.4 <sub>3.5</sub>        | <b>52.2</b> <sub>3.3</sub> | 55.0 <sub>5.5</sub>        | 54.2        | 64.0 <sub>3.4</sub>        | 66.7 <sub>2.8</sub>        | 77.8 <sub>0.7</sub>        | 71.8 <sub>1.0</sub>        | 74.1 <sub>0.9</sub>        | 70.9        | 62.6        |
| LLama3.1-SFT <sub>10</sub> | 8B   | 27.6 <sub>2.4</sub>        | 49.7 <sub>5.2</sub>        | 64.1 <sub>5.6</sub>        | 50.3 <sub>2.8</sub>        | 47.0 <sub>4.8</sub>        | 47.8        | 63.8 <sub>1.1</sub>        | 61.7 <sub>3.5</sub>        | 74.4 <sub>3.5</sub>        | 68.9 <sub>3.4</sub>        | 71.4 <sub>1.0</sub>        | 68.0        | 57.9        |

Table 7: Performance results of various models on the pseudo-documents ( $k=10$ ) translation task (Health domain), measured using d-chrF. The best prompt was selected for each language after evaluating three different prompts.

| Model                      | Size | $eng \rightarrow X$        |                            |                            |                            |                            |             | $X \rightarrow eng$        |                            |                            |                            |                            |             | AVG         |
|----------------------------|------|----------------------------|----------------------------|----------------------------|----------------------------|----------------------------|-------------|----------------------------|----------------------------|----------------------------|----------------------------|----------------------------|-------------|-------------|
|                            |      | amh                        | hau                        | swa                        | yor                        | zul                        | Avg.        | amh                        | hau                        | swa                        | yor                        | zul                        | Avg.        |             |
| Encoder-Decoder            |      |                            |                            |                            |                            |                            |             |                            |                            |                            |                            |                            |             |             |
| MADLAD-400                 | 3B   | 29.5 <sub>2.1</sub>        | 38.3 <sub>4.3</sub>        | 31.7 <sub>4.6</sub>        | 15.1 <sub>1.1</sub>        | 44.1 <sub>3.6</sub>        | 31.8        | 62.6 <sub>2.0</sub>        | 63.5 <sub>2.2</sub>        | 66.4 <sub>3.2</sub>        | 45.9 <sub>2.4</sub>        | 63.4 <sub>2.2</sub>        | 60.3        | 46.0        |
| Aya-101                    | 13B  | 30.1 <sub>1.5</sub>        | 55.0 <sub>3.2</sub>        | 51.7 <sub>3.5</sub>        | 22.3 <sub>1.7</sub>        | 55.0 <sub>1.9</sub>        | 42.8        | 62.5 <sub>1.4</sub>        | 65.5 <sub>1.3</sub>        | 68.8 <sub>1.8</sub>        | 55.7 <sub>2.4</sub>        | 68.4 <sub>1.0</sub>        | 64.2        | 53.5        |
| Decoder-only               |      |                            |                            |                            |                            |                            |             |                            |                            |                            |                            |                            |             |             |
| Gemma2-IT                  | 9B   | 6.2 <sub>0.7</sub>         | 42.1 <sub>3.9</sub>        | 51.0 <sub>5.3</sub>        | 6.6 <sub>0.8</sub>         | 15.4 <sub>1.7</sub>        | 24.3        | 35.9 <sub>4.8</sub>        | 50.1 <sub>4.6</sub>        | 57.7 <sub>3.7</sub>        | 48.2 <sub>3.4</sub>        | 51.7 <sub>3.7</sub>        | 48.7        | 36.5        |
| LLama3.1-IT                | 8B   | 7.4 <sub>0.9</sub>         | 15.3 <sub>1.9</sub>        | 43.3 <sub>4.4</sub>        | 6.2 <sub>1.1</sub>         | 8.8 <sub>0.7</sub>         | 16.2        | 26.1 <sub>2.0</sub>        | 48.7 <sub>3.4</sub>        | 59.0 <sub>2.7</sub>        | 34.4 <sub>3.2</sub>        | 34.7 <sub>3.1</sub>        | 40.6        | 28.4        |
| LLaMAX3-Alp                | 8B   | 11.4 <sub>1.2</sub>        | 32.5 <sub>4.4</sub>        | 38.1 <sub>4.1</sub>        | 12.0 <sub>1.4</sub>        | 26.1 <sub>2.2</sub>        | 24.0        | 29.4 <sub>2.9</sub>        | 51.4 <sub>4.3</sub>        | 62.4 <sub>2.5</sub>        | 24.7 <sub>3.6</sub>        | 48.8 <sub>3.3</sub>        | 43.3        | 33.7        |
| GPT-3.5                    | —    | 13.5 <sub>1.1</sub>        | 29.7 <sub>2.5</sub>        | 72.1 <sub>1.6</sub>        | 12.7 <sub>1.2</sub>        | 35.1 <sub>2.9</sub>        | 32.6        | 38.5 <sub>4.0</sub>        | 56.3 <sub>1.5</sub>        | 73.5 <sub>1.4</sub>        | 53.0 <sub>1.6</sub>        | 61.2 <sub>1.3</sub>        | 56.5        | 44.6        |
| GPT-4o                     | —    | 31.3 <sub>1.9</sub>        | <b>65.1</b> <sub>2.5</sub> | <b>75.1</b> <sub>1.6</sub> | 28.1 <sub>1.8</sub>        | <b>70.7</b> <sub>1.5</sub> | 54.0        | <b>68.6</b> <sub>1.1</sub> | <b>71.6</b> <sub>1.4</sub> | <b>76.5</b> <sub>1.6</sub> | <b>70.1</b> <sub>1.1</sub> | <b>76.5</b> <sub>1.1</sub> | <b>72.7</b> | <b>63.3</b> |
| SFT on AFriDOC-MT          |      |                            |                            |                            |                            |                            |             |                            |                            |                            |                            |                            |             |             |
| LLaMAX3-SFT                | 8B   | 21.7 <sub>2.0</sub>        | 29.9 <sub>3.2</sub>        | 37.0 <sub>3.4</sub>        | 30.5 <sub>2.7</sub>        | 31.7 <sub>3.5</sub>        | 30.2        | 24.2 <sub>2.6</sub>        | 27.6 <sub>4.2</sub>        | 32.3 <sub>4.5</sub>        | 28.5 <sub>3.3</sub>        | 29.8 <sub>5.4</sub>        | 28.5        | 29.3        |
| LLama3.1-SFT               | 8B   | 21.0 <sub>2.0</sub>        | 30.8 <sub>3.2</sub>        | 40.0 <sub>4.1</sub>        | 33.4 <sub>3.8</sub>        | 29.3 <sub>3.1</sub>        | 30.9        | 23.9 <sub>2.5</sub>        | 28.9 <sub>4.3</sub>        | 36.9 <sub>5.8</sub>        | 32.2 <sub>4.3</sub>        | 32.3 <sub>5.2</sub>        | 30.8        | 30.9        |
| LLaMAX3-SFT <sub>10</sub>  | 8B   | <b>37.7</b> <sub>2.1</sub> | 58.6 <sub>5.1</sub>        | 68.3 <sub>3.9</sub>        | 49.3 <sub>4.1</sub>        | 60.9 <sub>3.9</sub>        | <b>55.0</b> | 65.4 <sub>1.4</sub>        | 68.5 <sub>1.3</sub>        | 73.1 <sub>1.2</sub>        | 67.7 <sub>1.2</sub>        | 71.6 <sub>1.2</sub>        | 69.3        | 62.1        |
| LLama3.1-SFT <sub>10</sub> | 8B   | 23.7 <sub>1.9</sub>        | 47.0 <sub>5.2</sub>        | 58.6 <sub>5.6</sub>        | <b>49.7</b> <sub>3.8</sub> | 43.8 <sub>4.5</sub>        | 44.5        | 60.9 <sub>2.7</sub>        | 65.4 <sub>2.5</sub>        | 71.1 <sub>1.2</sub>        | 66.3 <sub>1.2</sub>        | 66.4 <sub>4.0</sub>        | 66.0        | 55.3        |

Table 8: Performance results of various models on the pseudo-documents ( $k=10$ ) translation task (Tech domain), measured using d-chrF. The best prompt was selected for each language after evaluating three different prompts.

forms LLaMAX3-Alpaca especially when translating into English, which also differs from the trends observed in the sentence-level setup.

**LLMs trained on longer documents are better for long document translation** Both LLaMa models trained via SFT on sentences (LLaMa3.1-SFT, and LLaMAX3-SFT) show a decline in performance in the pseudo-document setting compared to sentence-level translation. However, the same models trained via SFT on pseudo-documents with  $k=10$  demonstrate significant improvements on pseudo-documents. Interestingly, the LLaMAX3-SFT<sub>10</sub> model performs consistently well, achieving results comparable to its sentence-level counterpart on sentence-level tasks, and also outperforming LLaMa3.1-SFT<sub>10</sub>, particularly when translating into African languages.

### 5.3 GPT-4o based evaluation

Table 9 presents average GPT-4o evaluations for fluency and content errors (CE) of realigned outputs from sentence-level and pseudo-document-level tasks ( $k=10$ ) across four models in the *health* domain. When translating into English,

pseudo-document outputs are generally rated as more fluent and show fewer content errors, except for LLaMAX3-SFT<sub>1</sub>, which, when evaluated on pseudo-documents, shows lower fluency but still fewer content errors—an outcome that is counter-intuitive. However, when translating into African languages, the results are less consistent. Notably, GPT-3.5 achieves a fluency score of 4.9 for both fine-tuned versions of LLaMAX3-SFT—a score no model achieved when translating into English. Additionally, Yorùbá, a language that had some of the lowest d-chrF scores across models, achieved fluency scores of 4.0 and 4.2 with the two LLaMAX3-SFT versions. These inconsistencies raise concerns about GPT-4o’s reliability. Consequently, we focus on human evaluation going forward. Full GPT-4o results are provided in Appendix C.3.

### 5.4 Human evaluation

In Table 10 we report average direct assessment (DA) scores (on a scale from 0 to 100) from three annotators per language for the *health* domain, when translating into four African languages. For each language, we used 30 documents across models and both domains to compute inter-annotator

| Model                     | Setup | $eng \rightarrow X$ |      |      |      |      |      | $X \rightarrow eng$ |      |      |      |      |      |
|---------------------------|-------|---------------------|------|------|------|------|------|---------------------|------|------|------|------|------|
|                           |       | amh                 | hau  | swh  | yor  | zul  | avg  | amh                 | hau  | swh  | yor  | zul  | avg  |
| Fluency                   |       |                     |      |      |      |      |      |                     |      |      |      |      |      |
| Aya-101                   | Sent  | 2.8                 | 2.9  | 1.3  | 1.2  | 2.7  | 2.2  | 3.0                 | 2.9  | 3.3  | 2.3  | 2.9  | 2.9  |
|                           | Doc10 | 2.5                 | 2.9  | 1.8  | 1.4  | 3.1  | 2.3  | 3.6                 | 3.2  | 3.8  | 2.9  | 3.3  | 3.4  |
| GPT-3.5                   | Sent  | 1.1                 | 1.4  | 4.9  | 1.1  | 1.5  | 2.0  | 3.0                 | 2.5  | 3.6  | 2.6  | 2.6  | 2.9  |
|                           | Doc10 | 1.0                 | 1.1  | 4.9  | 1.0  | 1.3  | 1.9  | 4.2                 | 4.1  | 4.7  | 3.9  | 4.0  | 4.2  |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 3.5                 | 3.2  | 3.6  | 4.0  | 3.3  | 3.5  | 3.5                 | 2.9  | 3.9  | 3.2  | 3.6  | 3.4  |
|                           | Doc10 | 1.7                 | 2.3  | 3.2  | 2.7  | 2.4  | 2.5  | 3.0                 | 2.6  | 3.1  | 3.2  | 3.1  | 3.0  |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 2.6                 | 3.9  | 4.4  | 4.2  | 3.6  | 3.8  | 4.1                 | 4.2  | 4.6  | 4.4  | 4.4  | 4.3  |
| Content Error             |       |                     |      |      |      |      |      |                     |      |      |      |      |      |
| Aya-101                   | Sent  | 9.0                 | 9.1  | 9.5  | 8.2  | 13.7 | 9.9  | 15.8                | 17.1 | 17.5 | 23.8 | 19.1 | 18.7 |
|                           | Doc10 | 8.9                 | 8.3  | 8.2  | 6.7  | 16.1 | 9.6  | 12.6                | 14.6 | 14.1 | 18.7 | 15.7 | 15.1 |
| GPT-3.5                   | Sent  | 7.2                 | 11.3 | 4.2  | 7.8  | 20.9 | 10.3 | 9.5                 | 13.2 | 12.4 | 12.9 | 15.9 | 12.8 |
|                           | Doc10 | 3.3                 | 7.4  | 3.9  | 4.1  | 10.1 | 5.8  | 6.6                 | 9.2  | 7.1  | 9.0  | 10.8 | 8.5  |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 10.0                | 9.5  | 12.3 | 12.3 | 12.4 | 11.3 | 11.5                | 9.6  | 11.8 | 12.2 | 12.5 | 11.5 |
|                           | Doc10 | 10.4                | 8.8  | 9.5  | 8.1  | 8.4  | 9.0  | 9.0                 | 9.0  | 8.9  | 9.1  | 8.6  | 8.9  |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 7.3                 | 14.6 | 10.5 | 10.3 | 12.1 | 11.0 | 8.7                 | 9.6  | 8.9  | 10.2 | 9.4  | 9.4  |

Table 9: Document-level evaluation in the *health* domain, judged by GPT-4o. Compares sentence- vs. document-level outputs on Fluency (1–5 scale) and Content Errors (CE).

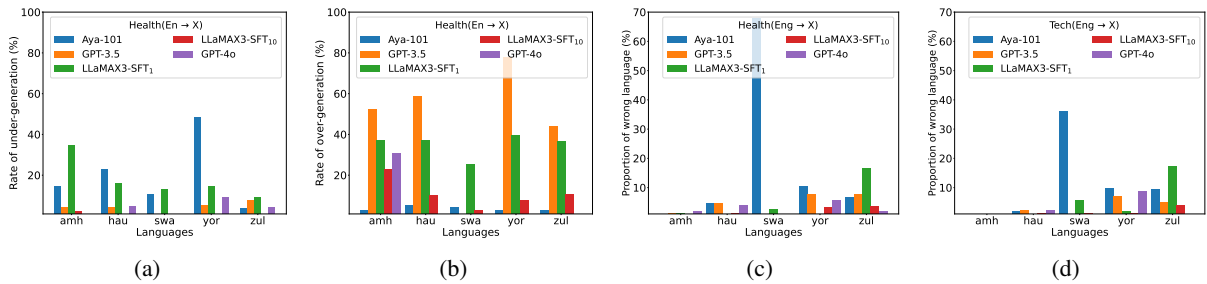


Figure 1: Rate of under-generation, over-generation, and off-target translation in pseudo-document translation ( $k = 10$ ).

agreement. We obtained Krippendorff’s alpha values of  $\geq 0.40$ , which are relatively low due to the fine granularity of the evaluation scale. **Human evaluation results align closely with d-chrF**, which favors sentence-level translations over pseudo-document translations when translating into African languages. Among the models, LLaMAX3-SFT<sub>1</sub> receives higher ratings at the sentence-level but is rated lower when translating pseudo-documents. In contrast, LLaMAX3-SFT<sub>10</sub> receives slightly lower ratings than LLaMAX3-SFT<sub>1</sub> at the sentence-level but is rated higher in the pseudo-document setting. GPT-3.5 is generally rated as the weakest model across languages, except for Swahili, where its performance is comparatively better (see Appendix C.4 for details).

## 5.5 Qualitative evaluation

Our qualitative analysis, based on feedback from native speakers who are also authors, indicates that GPT-3.5 frequently over-generates in the

| Model                     | Setup | amh         | hau         | swh         | yor         | zul         |
|---------------------------|-------|-------------|-------------|-------------|-------------|-------------|
| GPT-3.5                   | Sent  | 14.6        | 29.6        | 66.5        | 7.5         | 9.2         |
|                           | Doc10 | 1.7         | 16.4        | <b>68.3</b> | 4.2         | 3.1         |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | <b>64.5</b> | <b>81.5</b> | 57.9        | <b>65.1</b> | <b>48.1</b> |
|                           | Doc10 | 27.4        | 45.7        | 50.6        | 44.3        | 22.5        |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 38.5        | 76.7        | 62.4        | 64.9        | 46.7        |

Table 10: Average DA score (scale 0–100) from the human evaluators per language in the *health* domain.

pseudo-document setup by repeating words and phrases—except in Swahili, where it performs best. However, for Yorùbá, it often uses inconsistent or partial diacritics, resulting in inaccuracies. LLaMAX3-SFT<sub>1</sub> also exhibits repetition in pseudo-document translations, likely due to a length generalization problem (Anil et al., 2022), and does so more than LLaMAX3-SFT<sub>10</sub>. For the other four languages, LLaMAX3-SFT<sub>1</sub> with the sentence-level setup was rated higher than other models and configurations, owing to better context preservation and fewer repetitions. These obser-

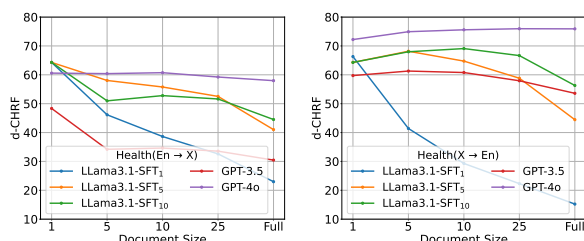


Figure 2: Comparison of Average d-chrF scores across models and pseudo-document lengths.

variations are consistent with both d-chrF and DA scores, although d-chrF scores tend to be inflated.

## 6 Discussion and Analysis

To better understand model behavior, we analyze their pseudo-document ( $k = 10$ ) translation outputs based on our findings and common issues in document-level MT with LLMs (Wu et al., 2024; Wang et al., 2024b). Additional results are provided in Appendix D.

**Are the outputs generated by translation models of appropriate length?** We compare model translations to the reference translations to identify empty outputs and cases of under- or over-generation. We found that all models rarely generate empty translations (refer to Appendix D), although GPT-3.5 and GPT-4o showed a slight tendency to do so for Yorùbá and Zulu, occurring in under 10% of cases. To evaluate this we compute the compute model output to reference translation length ratio across five models and for the languages, for the evaluated models and took the lowest 10th percentile, and this returns 73% and upper percentile of 550%. However, we defined under-generation as outputs <70% of reference length, and over-generation as >143%. As shown in Figures 1a and 1b and consistent with our qualitative findings, GPT-3.5 tends to over-generate more in African languages except for Swahili, while LLaMAX3-SFT<sub>1</sub> often under-generates as a sentence-level model. Moreover, all models over-generated by about 20% for Amharic, likely due to its unique script.

**Do LLMs generate translations in the correct target languages?** We evaluate whether these models understand the task by generating outputs in the target languages using the OpenLID (Burchell et al., 2023) language identification model. Our results show that these models rarely generate outputs in the wrong language when translating to English. However, when translating to African languages,

there is a higher likelihood of incorrect language translations, particularly with open models (see Figures 1c and 1d). These off-target languages often include English, and other languages including other African languages.

**What is the effect of document length on translation quality?** We compare average d-chrF scores of selected models, including GPT-3.5/4 and LLaMa3.1-SFT<sub>k</sub> ( $k = 1, 5, 10$ ), evaluated across pseudo-document lengths of 1, 5, 10, 25, and full length. As shown in Figure 2, d-chrF scores generally decline with increasing document length for African language translations. The reverse translation direction shows a similar trend, except for GPT-4o, which improves with length. Models trained on longer documents also generalize better to longer inputs than those trained on sentences.

## 7 Conclusion

We introduce AFRIDOC-MT, a document-level translation dataset in the *health* and *tech* domains for five African languages. We benchmarked various models, fine-tuning selected ones. Due to context length limits, documents were translated either sentence by sentence or as pseudo-documents. Outputs were evaluated using standard MT metrics, GPT-4o as a judge, and human direct assessment. NLLB-200 was the strongest built-in MT model, while GPT-4o outperformed general-purpose LLMs. However, our DA and qualitative analysis found GPT-4o’s judgments inconsistent for African languages, and sentence-by-sentence translation proved more effective for some languages.

## 8 Limitations

**Choice of LLMs and Prompts** We evaluated only a small subset of the numerous multilingual LLMs available. Our experiments were also limited by the context length of the LLMs, particularly for open LLMs. Except for LLaMa3.1, all other open LLMs have a context length of 8192 tokens, while encoder-decoder models were primarily based on T5. This makes it difficult to use the context length beyond a certain limit, making full document translation infeasible. Additionally, LLMs are prone to variance in performance based on the prompt. We therefore evaluated them for translation using three different prompts. However, it is possible that our prompts were not optimal.

**Language Coverage** Africa is home to thousands of indigenous languages, many of which exhibit unique linguistic properties. However, due to the high cost of translation using human translators and limited available funding, it is currently impossible to cover all languages. As a result, we focused on just five languages. We hope that future work will expand this dataset to include more languages and inspire the creation of additional datasets with broader coverage for document-level translation. Similarly, AFRIDOC-MT is a multi-way parallel dataset. However, due to the cost of running inference over three prompts and across all 30 translation directions for all the models evaluated, most of our analysis is limited to translation tasks between English and the five African languages. While we fine-tuned NLLB-200, LLama3.1 and LLaMAX3 on all 30 directions, we only provide results from NLLB-200 for all directions both before and after fine-tuning for sentence-level and pseudo-document tasks in the Appendix D.

**Evaluation Metrics** Quality evaluation in MT is an open and ongoing area of research, especially for document-level translation. Recent works have proposed embedding-based metrics for evaluation at both the sentence and document levels. While this has been well explored for high-resource language pairs, it remains underexplored for African languages, although there is a tool, AfriCOMET, that works for sentence-level evaluation in African languages. However, we evaluated the document-level translation outputs using *ModernBERT-base-long-context-qe-v1*<sup>11</sup>, trained on the WMT human evaluation dataset across 41 language pairs, including over 20 languages and three African languages (Hausa, Xhosa, and Zulu), two of which are included to our work. However, the scores were nearly identical across all models, offering no meaningful differentiation. Hence, for our document-level evaluation, in addition to lexical-based metrics, we incorporated three other evaluation approaches: using GPT-4o as a judge, human evaluation, and qualitative analysis. GPT-4o was employed to assess and rate the translation outputs of four models. While its ratings were consistent for translations into English, the same was not observed for translations into African languages, likely due to the model’s limited understanding of these languages. Therefore, we conducted a hu-

man evaluation for translations from English to African languages, comparing only three models due to cost constraints. However, we were unable to recruit annotators for Zulu.

**Model Coverage and Evaluation Focus** While we fine-tuned both NLLB-1.3B and LLaMAX3 models across all 30 language directions, due to computational constraints and the high cost of qualitative evaluation, our detailed analysis focuses only on translation between English and the 5 African languages. Nevertheless, we report quantitative results across all 30 directions for NLLB-1.3B. We will make all fine-tuned models publicly available to support future work, and we hope that further research will explore the remaining translation directions in greater depth.

**Translationese and English-Centric Bias** A potential limitation of our dataset is the influence of translationese (Koppel and Ordan, 2011). Since all source material translated originates in English, translated sentences in African languages may exhibit patterns such as unnatural syntax or overly literal phrasing. Although we have not conducted an analysis to quantify these effects, prior work suggests that they can affect MT model performance, generalization and evaluation including direct assessment (Freitag et al., 2019; Edunov et al., 2020). Furthermore, AFRIDOC-MT may reflect a bias toward English in terms of structure, semantics, and cultural framing. We leave a deeper investigation of these issues to future work.

## Ethics Statement

AFRIDOC-MT was created with the utmost consideration for ethical standards. The English texts translated were sourced from publicly available and ethically sourced materials. The data sources were selected to represent different cultural perspectives, with a focus on minimizing any potential bias. Efforts were made to ensure the dataset does not include harmful, biased, or offensive content via manual inspection.

## Acknowledgements

This work was carried out with support from Lacuna Fund, an initiative co-founded by The Rockefeller Foundation, Google.org, and Canada’s International Development Research Centre. This research project has benefited from the Microsoft Accelerate Foundation Models Research (AFMR)

<sup>11</sup><https://huggingface.co/yoslem/ModernBERT-base-long-context-qe-v1>

grant program. We acknowledge the support of OpenAI for providing API credits through their Researcher Access API programme. Jesujoba Alabi acknowledges the support of the BMBF's (German Federal Ministry of Education and Research) SLIK project under the grant 01IS22015C. The work has also been partially funded by BMBF through the TRAILS project (grant number 01IW24005). Miaoran Zhang received funding from the DFG (German Research Foundation) under project 232722074, SFB 1102. CEB acknowledges her AI4S fellowship within the "Generación D" initiative by Red.es, Ministerio para la Transformación Digital y de la Función Pública, for talent attraction (C005/24-ED CV1), funded by NextGenerationEU through PRTR. R. Bawden's participation was partly funded by her chair in the PRAIRIE institute, funded by the French national agency ANR, as part of the "Investissements d'avenir" programme under the reference ANR-19-P3IA-0001 and by the French *Agence Nationale de la Recherche* (ANR) under the projects TraLaLaM ("ANR-23-IAS1-0006") and MaTOS ("ANR-22-CE23-0033"). David Adelani acknowledges the funding of IVADO and the Canada First Research Excellence Fund.

Also, we also appreciate the translators whose names we have listed below:

1. **Amharic:** Bereket Tilahun, Hana M. Tamiru, Biniam Asmlash, Lidetewold Kebede
2. **Hausa:** Junaid Garba, Umma Abubakar, Jibrin Adamu, Ruqayya Nasir Iro
3. **Swahili:** Mohamed Mwinyimkuu, Laanyu koone, Baraka Karuli Mgasu, Said Athumani Said
4. **Yorùbá:** Ifeoluwa Akinsola, Simeon Onaolapo, Ganiyat Dasola Afolabi, Oluwatosin Koya
5. **Zulu:** Busisiwe Pakade, Rooweither Mabuya, Tholakele Zungu, Zanele Themban

Lastly, we thank the human annotators who were not part of the translation team.

## References

Ife Adebara, AbdelRahim Elmadany, and Muhammad Abdul-Mageed. 2024. [Cheetah: Natural language generation for 517 African languages](#). In *Proceedings of the 62nd Annual Meeting of the Association*

*for Computational Linguistics (Volume 1: Long Papers)*, pages 12798–12823, Bangkok, Thailand and virtual meeting. Association for Computational Linguistics.

David Adelani, Jesujoba Alabi, Angela Fan, Julia Kreutzer, Xiaoyu Shen, Machel Reid, Dana Ruiter, Dietrich Klakow, Peter Nabende, Ernie Chang, Tajudeen Gwadabe, Freshia Sackey, Bonaventure F. P. Dossou, Chris Emezue, Colin Leong, Michael Beukman, Shamsuddeen Muhammad, Guyo Jarso, Oreen Yousuf, Andre Niyongabo Rubungo, Gilles Hacheme, Eric Peter Wairagala, Muhammad Umair Nasir, Benjamin Ajibade, Tunde Ajayi, Yvonne Gitau, Jade Abbott, Mohamed Ahmed, Millicent Ochieng, Anuoluwapo Aremu, Perez Ogayo, Jonathan Mukiibi, Fatoumata Ouoba Kabore, Godson Kalipe, Derguene Mbaye, Allahsera Auguste Tapo, Victoire Memdjokam Koagne, Edwin Munkoh-Buabeng, Valencia Wagner, Idris Abdulmumin, Ayodele Awokoya, Happy Buzaaba, Blessing Sibanda, Andiswa Bukula, and Sam Manthulu. 2022. [A few thousand translations go a long way! leveraging pre-trained models for African news translation](#). In *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 3053–3070, Seattle, United States. Association for Computational Linguistics.

David Adelani, Hannah Liu, Xiaoyu Shen, Nikita Vassilyev, Jesujoba Alabi, Yanke Mao, Haonan Gao, and En-Shiun Lee. 2024a. [SIB-200: A simple, inclusive, and big evaluation dataset for topic classification in 200+ languages and dialects](#). In *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 226–245, St. Julian's, Malta. Association for Computational Linguistics.

David Adelani, Dana Ruiter, Jesujoba Alabi, Damilola Adebajo, Adesina Ayeni, Mofe Adeyemi, Ayodele Esther Awokoya, and Cristina España-Bonet. 2021. [The effect of domain and diacritics in Yoruba-English neural machine translation](#). In *Proceedings of Machine Translation Summit XVIII: Research Track*, pages 61–75, Virtual. Association for Machine Translation in the Americas.

David Ifeoluwa Adelani. 2022. [Natural language processing for african languages](#).

David Ifeoluwa Adelani, Jessica Ojo, Israel Abebe Azime, Jian Yun Zhuang, Jesujoba O. Alabi, Xuanli He, Millicent Ochieng, Sara Hooker, Andiswa Bukula, En-Shiun Annie Lee, Chiamaka Chukwunke, Happy Buzaaba, Blessing Sibanda, Godson Kalipe, Jonathan Mukiibi, Salomon Kabongo, Foutse Yuehgoh, Mmasibidi Setaka, Lolwethu Ndolela, Nkiruka Odu, Rooweither Mabuya, Shamsuddeen Hassan Muhammad, Salomey Osei, Sokhar Samb, Tadesse Kebede Guge, and Pontus Stenetorp. 2024b. [IrokoBench: A new benchmark for african languages in the age of large language models](#).

- Željko Agić and Ivan Vulić. 2019. [JW300: A wide-coverage parallel corpus for low-resource languages](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 3204–3210, Florence, Italy. Association for Computational Linguistics.
- Farhad Akhbardeh, Arkady Arkhangorodsky, Magdalena Biesialska, Ondřej Bojar, Rajen Chatterjee, Vishrav Chaudhary, Marta R. Costa-jussa, Cristina España-Bonet, Angela Fan, Christian Federmann, Markus Freitag, Yvette Graham, Roman Grundkiewicz, Barry Haddow, Leonie Harter, Kenneth Heafield, Christopher Homan, Matthias Huck, Kwabena Amponsah-Kaakyire, Jungo Kasai, Daniel Khashabi, Kevin Knight, Tom Kocmi, Philipp Koehn, Nicholas Lourie, Christof Monz, Makoto Morishita, Masaaki Nagata, Ajay Nagesh, Toshiaki Nakazawa, Matteo Negri, Santanu Pal, Allahsera Auguste Tapo, Marco Turchi, Valentin Vydrin, and Marcos Zampieri. 2021. [Findings of the 2021 conference on machine translation \(WMT21\)](#). In *Proceedings of the Sixth Conference on Machine Translation*, pages 1–88, Online. Association for Computational Linguistics.
- Jesujoba O. Alabi, David Ifeoluwa Adelani, Marius Mosbach, and Dietrich Klakow. 2022. [Adapting pre-trained language models to African languages via multilingual adaptive fine-tuning](#). In *Proceedings of the 29th International Conference on Computational Linguistics*, pages 4336–4349, Gyeongju, Republic of Korea. International Committee on Computational Linguistics.
- Antonios Anastasopoulos, Alessandro Cattelan, Zi-Yi Dou, Marcello Federico, Christian Federmann, Dmitriy Genzel, Francisco Guzmán, Junjie Hu, Macduff Hughes, Philipp Koehn, Rosie Lazar, Will Lewis, Graham Neubig, Mengmeng Niu, Alp Öktem, Eric Paquin, Grace Tang, and Sylwia Tur. 2020. [TICO-19: the translation initiative for COVID-19](#). In *Proceedings of the 1st Workshop on NLP for COVID-19 (Part 2) at EMNLP 2020*, Online. Association for Computational Linguistics.
- Cem Anil, Yuhuai Wu, Anders Johan Andreassen, Aitor Lewkowycz, Vedant Misra, Vinay Venkatesh Ramasesh, Ambrose Slone, Guy Gur-Ari, Ethan Dyer, and Behnam Neyshabur. 2022. [Exploring length generalization in large language models](#). In *Advances in Neural Information Processing Systems*.
- Paul Azunre, Salomey Osei, Salomey Addo, Lawrence Asamoah Adu-Gyamfi, Stephen Moore, Bernard Adabankah, Bernard Opoku, Clara Asare-Nyarko, Samuel Nyarko, Cynthia Amoaba, Esther Dansoa Appiah, Felix Akwerh, Richard Nii Lante Lawson, Joel Budu, Emmanuel Debrah, Nana Boateng, Wisdom Ofori, Edwin Buabeng-Munkoh, Franklin Adjei, Isaac Kojo Essel Ampomah, Joseph Otoo, Reindorf Borkor, Standyllove Birago Mensah, Lucien Mensah, Mark Amoako Marcel, Anokye Acheampong Amponsah, and James Ben Hayfron-Acquah. 2021. [English-twi parallel corpus for machine translation](#). In *2nd AfricaNLP Workshop Proceedings, AfricaNLP@EACL 2021, Virtual Event, April 19, 2021*.
- Rachel Bawden, Rico Sennrich, Alexandra Birch, and Barry Haddow. 2018. [Evaluating discourse phenomena in neural machine translation](#). In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*, pages 1304–1313, New Orleans, Louisiana. Association for Computational Linguistics.
- Edward Bayes, Israel Abebe Azime, Jesujoba O. Alabi, Jonas Kgomo, Tyna Eloundou, Elizabeth Proehl, Kai Chen, Imaan Khadir, Naome A. Etori, Shamsuddeen Hassan Muhammad, Choice Mpanza, Igneciah Pocia Thete, Dietrich Klakow, and David Ifeoluwa Adelani. 2024. [Uhura: A benchmark for evaluating scientific question answering and truthfulness in low-resource african languages](#).
- Stella Biderman, Hailey Schoelkopf, Lintang Sutawika, Leo Gao, Jonathan Tow, Baber Abbasi, Alham Fikri Aji, Pawan Sasanka Ammanamanchi, Sidney Black, Jordan Clive, Anthony DiPofi, Julien Etxaniz, Benjamin Fattori, Jessica Zosa Forde, Charles Foster, Jeffrey Hsu, Mimansa Jaiswal, Wilson Y. Lee, Haonan Li, Charles Lovering, Niklas Muennighoff, Ellie Pavlick, Jason Phang, Aviya Skowron, Samson Tan, Xiangru Tang, Kevin A. Wang, Genta Indra Winata, François Yvon, and Andy Zou. 2024. [Lessons from the trenches on reproducible evaluation of language models](#).
- Steven Bird, Ewan Klein, and Edward Loper. 2009. *Natural language processing with Python: analyzing text with the natural language toolkit*. "O'Reilly Media, Inc."
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel Ziegler, Jeffrey Wu, Clemens Winter, Chris Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. [Language models are few-shot learners](#). In *Advances in Neural Information Processing Systems*, volume 33, pages 1877–1901. Curran Associates, Inc.
- Laurie Burchell, Alexandra Birch, Nikolay Bogoychev, and Kenneth Heafield. 2023. [An open dataset and model for language identification](#). In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 865–879, Toronto, Canada. Association for Computational Linguistics.
- Nicolas Dahan, Rachel Bawden, and François Yvon. 2024. *Survey of Automatic Metrics for Evaluating*

*Machine Translation at the Document Level*. Ph.D. thesis, Inria Paris, Sorbonne Université; Sorbonne Université; Inria Paris.

Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurelien Rodriguez, Austen Gregerson, Ava Spataru, Baptiste Roziere, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux, Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret, Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Al-lonsius, Daniel Song, Danielle Pintz, Danny Livshits, David Esiobu, Dhruv Choudhary, Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Graeme Nail, Gregoire Milon, Guan Pang, Guillem Cucurell, Hailey Nguyen, Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Ishan Misra, Ivan Evtimov, Jade Copet, Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay Mahadeokar, Jeet Shah, Jelmer van der Linde, Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park, Joseph Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia, Kalyan Vasuden Alwala, Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth Heafield, Kevin Stone, Khalid El-Arini, Krithika Iyer, Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Lauren Rantala-Yearly, Laurens van der Maaten, Lawrence Chen, Liang Tan, Liz Jenkins, Louis Martin, Lovish Madaan, Lubo Malo, Lukas Blecher, Lukas Landzaat, Luke de Oliveira, Madeline Muzzi, Mahesh Pasupuleti, Mannat Singh, Manohar Paluri, Marcin Kardas, Mathew Oldham, Mathieu Rita, Maya Pavlova, Melanie Kambadur, Mike Lewis, Min Si, Mitesh Kumar Singh, Mona Hassan, Naman Goyal, Narjes Torabi, Nikolay Bashlykov, Nikolay Bogoychev, Niladri Chatterji, Olivier Duchenne, Onur Çelebi, Patrick Alrassy, Pengchuan Zhang, Pengwei Li, Petar Vasic, Peter Weng, Prajjwal Bhargava, Pratik Dubal, Praveen Krishnan, Punit Singh Koura, Puxin Xu, Qing He, Qingxiao Dong, Ragavan Srinivasan, Raj Ganapathy, Ramon Calderer, Ricardo Silveira Cabral, Robert Stojnic, Roberta Raileanu, Rohit Girdhar, Rohit Patel, Romain Sauvestre, Ronnie Polidoro, Roshan Sumbaly, Ross Taylor, Ruan Silva, Rui Hou, Rui Wang, Saghar Hosseini, Sahana Chennabasappa, Sanjay Singh, Sean Bell, Seohyun Sonia Kim, Sergey Edunov, Shaoliang Nie, Sharan Narang, Sharath Raparthy, Sheng Shen, Shengye Wan, Shruti Bhosale, Shun Zhang, Simon Vandenhende, Soumya Batra, Spencer Whitman, Sten Sootla, Stephane Collot, Suchin Gururangan, Sydney Borodinsky, Tamar Herman, Tara Fowler, Tarek Sheasha, Thomas Georgiou, Thomas Scialom, Tobias Speckbacher, Todor Mihaylov, Tong

Xiao, Ujjwal Karn, Vedanuj Goswami, Vibhor Gupta, Vignesh Ramanathan, Viktor Kerkez, Vincent Gouget, Virginie Do, Vish Vogeti, Vladan Petrovic, Weiwei Chu, Wenhan Xiong, Wenyin Fu, Whitney Meers, Xavier Martinet, Xiaodong Wang, Xiaoqing Ellen Tan, Xinfeng Xie, Xuchao Jia, Xuwei Wang, Yaelle Goldschlag, Yashesh Gaur, Yasmine Babaei, Yi Wen, Yiwen Song, Yuchen Zhang, Yue Li, Yuning Mao, Zacharie Delapierre Coudert, Zheng Yan, Zhengxing Chen, Zoe Papakipos, Aaditya Singh, Aaron Grattafiori, Abha Jain, Adam Kelsey, Adam Shajnfeld, Adithya Gangidi, Adolfo Victoria, Ahuva Goldstand, Ajay Menon, Ajay Sharma, Alex Boesenberg, Alex Vaughan, Alexei Baevski, Allie Feinstein, Amanda Kallet, Amit Sangani, Anam Yunus, Andrei Lupu, Andres Alvarado, Andrew Caples, Andrew Gu, Andrew Ho, Andrew Poulton, Andrew Ryan, Ankit Ramchandani, Annie Franco, Aparajita Saraf, Arkabandhu Chowdhury, Ashley Gabriel, Ashwin Bharambe, Assaf Eisenman, Azadeh Yazdan, Beau James, Ben Maurer, Benjamin Leonhardi, Bernie Huang, Beth Loyd, Beto De Paola, Bhargavi Paranjape, Bing Liu, Bo Wu, Boyu Ni, Braden Hancock, Bram Wasti, Brandon Spence, Brani Stojkovic, Brian Gamido, Britt Montalvo, Carl Parker, Carly Burton, Catalina Mejia, Changan Wang, Changkyu Kim, Chao Zhou, Chester Hu, Ching-Hsiang Chu, Chris Cai, Chris Tindal, Christoph Feichtenhofer, Damon Civin, Dana Beaty, Daniel Kreymer, Daniel Li, Danny Wyatt, David Adkins, David Xu, Davide Testuggine, Delia David, Devi Parikh, Diana Liskovich, Didem Foss, Dingkang Wang, Duc Le, Dustin Holland, Edward Dowling, Eissa Jamil, Elaine Montgomery, Eleonora Presani, Emily Hahn, Emily Wood, Erik Brinkman, Esteban Arcaute, Evan Dunbar, Evan Smothers, Fei Sun, Felix Kreuk, Feng Tian, Firat Ozgenel, Francesco Caggioni, Francisco Guzmán, Frank Kanayet, Frank Seide, Gabriela Medina Florez, Gabriella Schwarz, Gada Badeer, Georgia Swee, Gil Halpern, Govind Thattai, Grant Herman, Grigory Sizov, Guangyi, Zhang, Guna Lakshminarayanan, Hamid Shojanazeri, Han Zou, Hannah Wang, Hanwen Zha, Haroun Habeeb, Harrison Rudolph, Helen Suk, Henry Aspegren, Hunter Goldman, Igor Molybog, Igor Tufanov, Irina-Elena Veliche, Itai Gat, Jake Weissman, James Geboski, James Kohli, Japhet Asher, Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang, Jennifer Chan, Jenny Zhen, Jeremy Reizenstein, Jeremy Teboul, Jessica Zhong, Jian Jin, Jingyi Yang, Joe Cummings, Jon Carvill, Jon Shepard, Jonathan McPhie, Jonathan Torres, Josh Ginsburg, Junjie Wang, Kai Wu, Kam Hou U, Karan Saxena, Karthik Prasad, Kartikay Khandelwal, Katayoun Zand, Kathy Matosich, Kaushik Veeraraghavan, Kelly Michelena, Keqian Li, Kun Huang, Kunal Chawla, Kushal Lakhota, Kyle Huang, Lailin Chen, Lakshya Garg, Lavender A, Leandro Silva, Lee Bell, Lei Zhang, Liangpeng Guo, Licheng Yu, Liron Moshkovich, Luca Wehrstedt, Madian Khabsa, Manav Avalani, Manish Bhatt, Maria Tsim-poukelli, Martynas Mankus, Matan Hasson, Matthew Lennie, Matthias Reso, Maxim Groshev, Maxim Naumov, Maya Lathi, Meghan Keneally, Michael L.

- Seltzer, Michal Valko, Michelle Restrepo, Mihir Patel, Mik Vyatskov, Mikayel Samvelyan, Mike Clark, Mike Macey, Mike Wang, Miquel Jubert Her-moso, Mo Metanat, Mohammad Rastegari, Mun-ish Bansal, Nandhini Santhanam, Natascha Parks, Natasha White, Navyata Bawa, Nayan Singhal, Nick Egebo, Nicolas Usunier, Nikolay Pavlovich Laptev, Ning Dong, Ning Zhang, Norman Cheng, Oleg Chernoguz, Olivia Hart, Omkar Salpekar, Ozlem Kalinli, Parkin Kent, Parth Parekh, Paul Saab, Pa-van Balaji, Pedro Rittner, Philip Bontrager, Pierre Roux, Piotr Dollar, Polina Zvyagina, Prashant Ratan-chandani, Pritish Yuvraj, Qian Liang, Rachad Alao, Rachel Rodriguez, Rafi Ayub, Raghotham Murthy, Raghu Nayani, Rahul Mitra, Raymond Li, Rebekkah Hogan, Robin Battey, Rocky Wang, Rohan Mah-eswari, Russ Howes, Ruty Rinott, Sai Jayesh Bondu, Samyak Datta, Sara Chugh, Sara Hunt, Sargun Dhillon, Sasha Sidorov, Satadru Pan, Saurabh Verma, Seiji Yamamoto, Sharadh Ramaswamy, Shaun Lind-say, Shaun Lindsay, Sheng Feng, Shenghao Lin, Shengxin Cindy Zha, Shiva Shankar, Shuqiang Zhang, Shuqiang Zhang, Sinong Wang, Sneha Agar-wal, Soji Sajuyigbe, Soumith Chintala, Stephanie Max, Stephen Chen, Steve Kehoe, Steve Satterfield, Sudarshan Govindaprasad, Sumit Gupta, Sungmin Cho, Sunny Virk, Suraj Subramanian, Sy Choudhury, Sydney Goldman, Tal Remez, Tamar Glaser, Tamara Best, Thilo Kohler, Thomas Robinson, Tianhe Li, Tianjun Zhang, Tim Matthews, Timothy Chou, Tzook Shaked, Varun Vontimitta, Victoria Ajayi, Victoria Montanez, Vijai Mohan, Vinay Satish Kumar, Vishal Mangla, Vlad Ionescu, Vlad Poenaru, Vlad Tiberiu Mihailescu, Vladimir Ivanov, Wei Li, Wenchen Wang, Wenwen Jiang, Wes Bouaziz, Will Constable, Xiaocheng Tang, Xiaofang Wang, Xiaojuan Wu, Xiaolan Wang, Xide Xia, Xilun Wu, Xinbo Gao, Yan-jun Chen, Ye Hu, Ye Jia, Ye Qi, Yenda Li, Yilin Zhang, Ying Zhang, Yossi Adi, Youngjin Nam, Yu, Wang, Yuchen Hao, Yundi Qian, Yuzi He, Zach Rait, Zachary DeVito, Zef Rosnbrick, Zhaoduo Wen, Zhenyu Yang, and Zhiwei Zhao. 2024. [The Llama 3 Herd of Models](#).
- Sergey Edunov, Myle Ott, Marc’ Aurelio Ranzato, and Michael Auli. 2020. [On the evaluation of machine translation systems trained with back-translation](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 2836–2846, Online. Association for Computational Lin-guistics.
- AbdelRahim Elmadany, Ife Adebara, and Muhammad Abdul-Mageed. 2024. [Toucan: Many-to-many trans-lation for 150 african language pairs](#). In *Findings of the Association for Computational Linguistics: ACL 2024*, pages 13189–13206, Bangkok, Thailand. As-sociation for Computational Linguistics.
- Ignatius Ezeani, Paul Rayson, Ikechukwu E. Onyenwe, Chinedu Uchechukwu, and Mark Hepple. 2020. [Igbo-english machine translation: an evaluation benchmark](#). In *1st AfricaNLP Workshop Proceed-ings, AfricaNLP@ICLR 2020, Virtual Conference, Formerly Addis Ababa Ethiopia, 26th April 2020*.
- Angela Fan, Shruti Bhosale, Holger Schwenk, Zhiyi Ma, Ahmed El-Kishky, Siddharth Goyal, Man-deep Baines, Onur Celebi, Guillaume Wenzek, Vishrav Chaudhary, Naman Goyal, Tom Birch, Vi-taliy Liptchinsky, Sergey Edunov, Edouard Grave, Michael Auli, and Armand Joulin. 2020. [Beyond english-centric multilingual machine translation](#).
- Christian Federmann. 2018. [Appraise evaluation frame-work for machine translation](#). In *Proceedings of the 27th International Conference on Computational Lin-guistics: System Demonstrations*, pages 86–88, Santa Fe, New Mexico. Association for Computational Lin-guistics.
- Christian Federmann, Tom Kocmi, and Ying Xin. 2022. [NTREX-128 – news test references for MT evalua-tion of 128 languages](#). In *Proceedings of the First Workshop on Scaling Up Multilingual Evaluation*, pages 21–24, Online. Association for Computational Linguistics.
- Yukun Feng, Feng Li, Ziang Song, Boyuan Zheng, and Philipp Koehn. 2022. [Learn to remember: Trans-former with recurrent memory for document-level machine translation](#). In *Findings of the Association for Computational Linguistics: NAACL 2022*, pages 1409–1420, Seattle, United States. Association for Computational Linguistics.
- Markus Freitag, Isaac Caswell, and Scott Roy. 2019. [APE at scale and its implications on MT evaluation biases](#). In *Proceedings of the Fourth Conference on Machine Translation (Volume 1: Research Papers)*, pages 34–44, Florence, Italy. Association for Com-putational Linguistics.
- Eva Martínez Garcia, Carles Creus, Cristina Espana-Bonet, and Lluís Màrquez. 2017. Using word embed-dings to enforce document-level lexical consistency in machine translation. *The Prague Bulletin of Math-ematical Linguistics*, 108(1):85.
- Eva Martínez Garcia, Cristina Espana-Bonet, and Lluís Màrquez Villodre. 2014. Document-level ma-chine translation as a re-translation process. *Proce-samiento del Lenguaje Natural*, 53:103–110.
- Gemma Team, Morgane Riviere, Shreya Pathak, Pier Giuseppe Sessa, Cassidy Hardin, Surya Bhupati-raj, Léonard Hussenot, Thomas Mesnard, Bobak Shahriari, Alexandre Ramé, Johan Ferret, Peter Liu, Pouya Tafti, Abe Friesen, Michelle Casbon, Sabela Ramos, Ravin Kumar, Charline Le Lan, Sammy Jerome, Anton Tsitsulin, Nino Vieillard, Piotr Stanczyk, Sertan Girgin, Nikola Momchev, Matt Hoffman, Shantanu Thakoor, Jean-Bastien Grill, Behnam Neyshabur, Olivier Bachem, Alanna Wal-ton, Aliaksei Severyn, Alicia Parrish, Aliya Ah-mad, Allen Hutchison, Alvin Abdagic, Amanda Carl, Amy Shen, Andy Brock, Andy Coenen, An-thony Laforge, Antonia Paterson, Ben Bastian, Bilal

- Piot, Bo Wu, Brandon Royal, Charlie Chen, Chintu Kumar, Chris Perry, Chris Welty, Christopher A. Choquette-Choo, Danila Sinopalnikov, David Weinberger, Dimple Vijaykumar, Dominika Rogozińska, Dustin Herbison, Elisa Bandy, Emma Wang, Eric Noland, Erica Moreira, Evan Senter, Evgenii Eltyshv, Francesco Visin, Gabriel Rasskin, Gary Wei, Glenn Cameron, Gus Martins, Hadi Hashemi, Hanna Klimczak-Plucińska, Harleen Batra, Harsh Dhand, Ivan Nardini, Jacinda Mein, Jack Zhou, James Svensson, Jeff Stanway, Jetha Chan, Jin Peng Zhou, Joana Carrasqueira, Joana Iljazi, Jocelyn Becker, Joe Fernandez, Joost van Amersfoort, Josh Gordon, Josh Lipschultz, Josh Newlan, Ju yeong Ji, Kareem Mohamed, Kartikeya Badola, Kat Black, Katie Millican, Keelin McDonnell, Kelvin Nguyen, Kiranbir Sodhia, Kish Greene, Lars Lowe Sjoesund, Lauren Usui, Laurent Sifre, Lena Heuermann, Leticia Lago, Lilly McNealus, Livio Baldini Soares, Logan Kilpatrick, Lucas Dixon, Luciano Martins, Machel Reid, Manvinder Singh, Mark Iverson, Martin Görner, Mat Velloso, Mateo Wirth, Matt Davidow, Matt Miller, Matthew Rahtz, Matthew Watson, Meg Risdal, Mehran Kazemi, Michael Moynihan, Ming Zhang, Minsuk Kahng, Minwoo Park, Mofi Rahman, Mohit Khatwani, Natalie Dao, Nenshad Bardoliwalla, Nesh Devanathan, Neta Dumai, Nilay Chauhan, Oscar Wahltinez, Pankil Botarda, Parker Barnes, Paul Barham, Paul Michel, Pengchong Jin, Petko Georgiev, Phil Culliton, Pradeep Kuppala, Ramona Comanescu, Ramona Merhej, Reena Jana, Reza Ardeshtir Rokni, Rishabh Agarwal, Ryan Mullins, Samaneh Saadat, Sara Mc Carthy, Sarah Perrin, Sébastien M. R. Arnold, Sebastian Krause, Shengyang Dai, Shruti Garg, Shruti Sheth, Sue Ronstrom, Susan Chan, Timothy Jordan, Ting Yu, Tom Eccles, Tom Hennigan, Tomas Kocisky, Tulsee Doshi, Vihan Jain, Vikas Yadav, Vilobh Meshram, Vishal Dharmadhikari, Warren Barkley, Wei Wei, Wenming Ye, Woohyun Han, Woosuk Kwon, Xiang Xu, Zhe Shen, Zhitao Gong, Zichuan Wei, Victor Cotruta, Phoebe Kirk, Anand Rao, Minh Giang, Ludovic Peran, Tris Warkentin, Eli Collins, Joelle Barral, Zoubin Ghahramani, Raia Hadsell, D. Sculley, Jeanine Banks, Anca Dragan, Slav Petrov, Oriol Vinyals, Jeff Dean, Demis Hassabis, Koray Kavukcuoglu, Clement Farabet, Elena Buchatskaya, Sebastian Borgeaud, Noah Fiedel, Armand Joulin, Kathleen Kenealy, Robert Dadashi, and Alek Andreiev. 2024. [Gemma 2: Improving open language models at a practical size](#).
- Naman Goyal, Cynthia Gao, Vishrav Chaudhary, Peng-Jen Chen, Guillaume Wenzek, Da Ju, Sanjana Krishnan, Marc’ Aurelio Ranzato, Francisco Guzmán, and Angela Fan. 2022. [The Flores-101 evaluation benchmark for low-resource and multilingual machine translation](#). *Transactions of the Association for Computational Linguistics*, 10:522–538.
- Christian Herold and Hermann Ney. 2023. [Improving long context document-level machine translation](#). In *Proceedings of the 4th Workshop on Computational Approaches to Discourse (CODI 2023)*, pages 112–125, Toronto, Canada. Association for Computational Linguistics.
- Yuchen Jiang, Tianyu Liu, Shuming Ma, Dongdong Zhang, Jian Yang, Haoyang Huang, Rico Sennrich, Ryan Cotterell, Mrinmaya Sachan, and Ming Zhou. 2022. [BlonDe: An automatic evaluation metric for document-level machine translation](#). In *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 1550–1565, Seattle, United States. Association for Computational Linguistics.
- Tom Kocmi, Eleftherios Avramidis, Rachel Bawden, Ondřej Bojar, Anton Dvorkovich, Christian Federmann, Mark Fishel, Markus Freitag, Thamme Gowda, Roman Grundkiewicz, Barry Haddow, Philipp Koehn, Benjamin Marie, Christof Monz, Makoto Morishita, Kenton Murray, Makoto Nagata, Toshiaki Nakazawa, Martin Popel, Maja Popović, and Mariya Shmatova. 2023. [Findings of the 2023 conference on machine translation \(WMT23\): LLMs are here but not quite there yet](#). In *Proceedings of the Eighth Conference on Machine Translation*, pages 1–42, Singapore. Association for Computational Linguistics.
- Tom Kocmi, Vilém Zouhar, Eleftherios Avramidis, Roman Grundkiewicz, Marzena Karpinska, Maja Popović, Mrinmaya Sachan, and Mariya Shmatova. 2024. [Error span annotation: A balanced approach for human evaluation of machine translation](#). In *Proceedings of the Ninth Conference on Machine Translation*, pages 1440–1453, Miami, Florida, USA. Association for Computational Linguistics.
- Moshe Koppel and Noam Ordan. 2011. [Translationese and its dialects](#). In *Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies*, pages 1318–1326, Portland, Oregon, USA. Association for Computational Linguistics.
- Julia Kreutzer, Isaac Caswell, Lisa Wang, Ahsan Wahab, Daan van Esch, Nasanbayar Ulzii-Orshikh, Allahsera Tapo, Nishant Subramani, Artem Sokolov, Claytone Sikasote, Monang Setyawan, Supheakmungkol Sarin, Sokhar Samb, Benoît Sagot, Clara Rivera, Annette Rios, Isabel Papadimitriou, Salomey Osei, Pedro Ortiz Suarez, Iroko Orife, Kelechi Ogueji, Andre Niyongabo Rubungo, Toan Q. Nguyen, Mathias Müller, André Müller, Shamsuddeen Hassan Muhammad, Nanda Muhammad, Ayanda Mnyakeni, Jamshidbek Mirzakhlov, Tapiwanashe Matangira, Colin Leong, Nze Lawson, Sneha Kudugunta, Yacine Jernite, Mathias Jenny, Orhan Firat, Bonaventure F. P. Dossou, Sakhile Dlamini, Nisansa de Silva, Sakine Çabuk Ballı, Stella Biderman, Alessia Battisti, Ahmed Baruwa, Ankur Bapna, Pallavi Baljekar, Israel Abebe Azime, Ayodele Awokoya, Duygu Ataman, Orevaoghene Ahia, Oghenefego Ahia, Sweta Agrawal, and Mofetoluwa Adeyemi. 2022. [Quality at a glance: An audit of web-crawled multilingual](#)

- datasets. *Transactions of the Association for Computational Linguistics*, 10:50–72.
- Sneha Kudugunta, Isaac Caswell, Biao Zhang, Xavier Garcia, Christopher A. Choquette-Choo, Katherine Lee, Derrick Xin, Aditya Kusupati, Romi Stella, Ankur Bapna, and Orhan Firat. 2023. [Madlad-400: A multilingual and document-level large audited dataset](#).
- Chia-Hsuan Lee, Aditya Siddhant, Viresh Ratnakar, and Melvin Johnson. 2022. [DOCmT5: Document-level pretraining of multilingual language models](#). In *Findings of the Association for Computational Linguistics: NAACL 2022*, pages 425–437, Seattle, United States. Association for Computational Linguistics.
- Jindřich Libovický and Jindřich Helcl. 2017. [Attention strategies for multi-source sequence-to-sequence learning](#). In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 196–202, Vancouver, Canada. Association for Computational Linguistics.
- António Lopes, M. Amin Farajian, Rachel Bawden, Michael Zhang, and André F. T. Martins. 2020. [Document-level neural MT: A systematic comparison](#). In *Proceedings of the 22nd Annual Conference of the European Association for Machine Translation*, pages 225–234, Lisboa, Portugal. European Association for Machine Translation.
- Yinquan Lu, Wenhao Zhu, Lei Li, Yu Qiao, and Fei Yuan. 2024. Llamax: Scaling linguistic horizons of llm by enhancing translation capabilities beyond 100 languages. *arXiv preprint arXiv:2407.05975*.
- Eva Martínez García, Carles Creus, and Cristina España-Bonet. 2019. [Context-aware neural machine translation decoding](#). In *Proceedings of the Fourth Workshop on Discourse in Machine Translation (DiscoMT 2019)*, pages 13–23, Hong Kong, China. Association for Computational Linguistics.
- Sameen Maruf, Fahimeh Saleh, and Gholamreza Haffari. 2021. A survey on document-level neural machine translation: Methods and evaluation. *ACM Computing Surveys (CSUR)*, 54(2):1–36.
- Arya D. McCarthy, Rachel Wicks, Dylan Lewis, Aaron Mueller, Winston Wu, Oliver Adams, Garrett Nicolai, Matt Post, and David Yarowsky. 2020. [The Johns Hopkins University Bible corpus: 1600+ tongues for typological exploration](#). In *Proceedings of the Twelfth Language Resources and Evaluation Conference*, pages 2884–2892, Marseille, France. European Language Resources Association.
- Lesly Miculicich, Dhananjay Ram, Nikolaos Pappas, and James Henderson. 2018. [Document-level neural machine translation with hierarchical attention networks](#). In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, pages 2947–2954, Brussels, Belgium. Association for Computational Linguistics.
- Alireza Mohammadshahi, Vassilina Nikoulina, Alexandre Berard, Caroline Brun, James Henderson, and Laurent Besacier. 2022. [SMaLL-100: Introducing shallow multilingual machine translation model for low-resource languages](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 8348–8359, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Mathias Müller, Annette Rios, Elena Voita, and Rico Sennrich. 2018. [A large-scale test set for the evaluation of context-aware pronoun translation in neural machine translation](#). In *Proceedings of the Third Conference on Machine Translation: Research Papers*, pages 61–72, Brussels, Belgium. Association for Computational Linguistics.
- NLLB Team, Marta R. Costa-jussà, James Cross, Onur Çelebi, Maha Elbayad, Kenneth Heafield, Kevin Hefernan, Elahe Kalbassi, Janice Lam, Daniel Licht, Jean Maillard, Anna Sun, Skyler Wang, Guillaume Wenzek, Al Youngblood, Bapi Akula, Loic Barraut, Gabriel Mejia Gonzalez, Prangthip Hansanti, John Hoffman, Semarley Jarrett, Kaushik Ram Sadagopan, Dirk Rowe, Shannon Spruit, Chau Tran, Pierre Andrews, Necip Fazil Ayan, Shruti Bhosale, Sergey Edunov, Angela Fan, Cynthia Gao, Vedanuj Goswami, Francisco Guzmán, Philipp Koehn, Alexandre Mourachko, Christophe Ropers, Safiyyah Saleem, Holger Schwenk, and Jeff Wang. 2022. [No language left behind: Scaling human-centered machine translation](#).
- NLLB Team et al. 2024. Scaling neural machine translation to 200 languages. *Nature*, 630(8018):841.
- OpenAI. 2024. Introducing ChatGPT. <https://openai.com/index/chatgpt/>. [Accessed 01-06-2024].
- Myle Ott, Sergey Edunov, Alexei Baevski, Angela Fan, Sam Gross, Nathan Ng, David Grangier, and Michael Auli. 2019. fairseq: A fast, extensible toolkit for sequence modeling. In *Proceedings of NAACL-HLT 2019: Demonstrations*.
- Kishore Papineni, Salim Roukos, Todd Ward, and Wei-Jing Zhu. 2002. [Bleu: a method for automatic evaluation of machine translation](#). In *Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics*, pages 311–318, Philadelphia, Pennsylvania, USA. Association for Computational Linguistics.
- Maja Popović. 2015. [chrF: character n-gram F-score for automatic MT evaluation](#). In *Proceedings of the Tenth Workshop on Statistical Machine Translation*, pages 392–395, Lisbon, Portugal. Association for Computational Linguistics.
- Matt Post. 2018. [A call for clarity in reporting BLEU scores](#). In *Proceedings of the Third Conference on Machine Translation: Research Papers*, pages 186–191, Brussels, Belgium. Association for Computational Linguistics.

- Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J. Liu. 2020. [Exploring the limits of transfer learning with a unified text-to-text transformer](#). *Journal of Machine Learning Research*, 21(140):1–67.
- Annette Rios Gonzales, Laura Mascarell, and Rico Sennrich. 2017. [Improving word sense disambiguation in neural machine translation with sense embeddings](#). In *Proceedings of the Second Conference on Machine Translation*, pages 11–19, Copenhagen, Denmark. Association for Computational Linguistics.
- Holger Schwenk, Vishrav Chaudhary, Shuo Sun, Hongyu Gong, and Francisco Guzmán. 2021a. [WikiMatrix: Mining 135M parallel sentences in 1620 language pairs from Wikipedia](#). In *Proceedings of the 16th Conference of the European Chapter of the Association for Computational Linguistics: Main Volume*, pages 1351–1361, Online. Association for Computational Linguistics.
- Holger Schwenk, Guillaume Wenzek, Sergey Edunov, Edouard Grave, Armand Joulin, and Angela Fan. 2021b. [CCMatrix: Mining billions of high-quality parallel sentences on the web](#). 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 6490–6500, Online. Association for Computational Linguistics.
- Yirong Sun, Dawei Zhu, Yanjun Chen, Erjia Xiao, Xinghao Chen, and Xiaoyu Shen. 2025. [Fine-grained and multi-dimensional metrics for document-level machine translation](#). In *Proceedings of the 2025 Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 4: Student Research Workshop)*, pages 1–17, Albuquerque, USA. Association for Computational Linguistics.
- Zewei Sun, Mingxuan Wang, Hao Zhou, Chengqi Zhao, Shujian Huang, Jiajun Chen, and Lei Li. 2022. [Re-thinking document-level neural machine translation](#). In *Findings of the Association for Computational Linguistics: ACL 2022*, pages 3537–3548, Dublin, Ireland. Association for Computational Linguistics.
- Jörg Tiedemann and Yves Scherrer. 2017. [Neural machine translation with extended context](#). In *Proceedings of the Third Workshop on Discourse in Machine Translation*, pages 82–92, Copenhagen, Denmark. Association for Computational Linguistics.
- Sami Ul Haq, Sadaf Abdul Rauf, Arsalan Shaukat, and Abdullah Saeed. 2020. [Document level NMT of low-resource languages with backtranslation](#). In *Proceedings of the Fifth Conference on Machine Translation*, pages 442–446, Online. Association for Computational Linguistics.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. [Attention is all you need](#). In *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc.
- Elena Voita, Rico Sennrich, and Ivan Titov. 2019. [When a good translation is wrong in context: Context-aware machine translation improves on deixis, ellipsis, and lexical cohesion](#). In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 1198–1212, Florence, Italy. Association for Computational Linguistics.
- Elena Voita, Pavel Serdyukov, Rico Sennrich, and Ivan Titov. 2018. [Context-aware neural machine translation learns anaphora resolution](#). In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1264–1274, Melbourne, Australia. Association for Computational Linguistics.
- Jiayi Wang, David Ifeoluwa Adelani, Sweta Agrawal, Marek Masiak, Ricardo Rei, Eleftheria Briakou, Marine Carpuat, Xuanli He, Sofia Bourhim, Andiswa Bukula, Muhidin Mohamed, Temitayo Olatoye, Tosin Adewumi, Hamam Mokayed, Christine Mwase, Wangui Kimotho, Foutse Yuehgoh, Anuoluwapo Aremu, Jessica Ojo, Shamsuddeen Hassan Muhammad, Salomey Osei, Abdul-Hakeem Omotayo, Chiamaka Chukwuneke, Perez Ogayo, Oumaima Hourrane, Salma El Anigri, Lolwethu Ndolela, Thabiso Mangwana, Shafie Abdi Mohamed, Ayinde Hassan, Oluwabusayo Olufunke Awoyomi, Lama Alkhaled, Sana Al-Azzawi, Naome A. Etori, Millicent Ochieng, Clemencia Siro, Samuel Njoroge, Eric Muchiri, Wangari Kimotho, Lyse Naomi Wamba Momo, Daud Abolade, Simbiat Ajao, Iyanuoluwa Shode, Ricky Macharm, Ruqayya Nasir Iro, Saheed S. Abdullahi, Stephen E. Moore, Bernard Opoku, Zainab Akinjobi, Abeeb Afolabi, Nnaemeka Obiefuna, Onyekachi Raphael Ogbu, Sam Brian, Verrah Akinyi Otiende, Chinedu Emmanuel Mbonu, Sakayo Toadoun Sari, Yao Lu, and Pontus Stenetorp. 2024a. [Afrimte and africomet: Enhancing comet to embrace under-resourced african languages](#).
- Longyue Wang, Zefeng Du, Wenxiang Jiao, Chenyang Lyu, Jianhui Pang, Leyang Cui, Kaiqiang Song, Derek Wong, Shuming Shi, and Zhaopeng Tu. 2024b. [Benchmarking and improving long-text translation with large language models](#). In *Findings of the Association for Computational Linguistics ACL 2024*, pages 7175–7187, Bangkok, Thailand and virtual meeting. Association for Computational Linguistics.
- Longyue Wang, Chenyang Lyu, Tianbo Ji, Zhirui Zhang, Dian Yu, Shuming Shi, and Zhaopeng Tu. 2023. [Document-level machine translation with large language models](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 16646–16661, Singapore. Association for Computational Linguistics.
- Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chaumond, Clement Delangue, Anthony Moi, Pierric Cistac, Tim Rault, Remi Louf, Morgan Funtowicz, Joe Davison, Sam Shleifer, Patrick von Platen,

- Clara Ma, Yacine Jernite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gugger, Mariama Drame, Quentin Lhoest, and Alexander Rush. 2020. [Trans-formers: State-of-the-art natural language processing](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 38–45, Online. Association for Computational Linguistics.
- Billy T. M. Wong and Chunyu Kit. 2012. [Extending machine translation evaluation metrics with lexical cohesion to document level](#). In *Proceedings of the 2012 Joint Conference on Empirical Methods in Natural Language Processing and Computational Natural Language Learning*, pages 1060–1068, Jeju Island, Korea. Association for Computational Linguistics.
- Minghao Wu, George Foster, Lizhen Qu, and Gholamreza Haffari. 2023. Document flattening: Beyond concatenating context for document-level neural machine translation. In *Proceedings of the 17th Conference of the European Chapter of the Association for Computational Linguistics*, Dubrovnik, Croatia.
- Minghao Wu, Thuy-Trang Vu, Lizhen Qu, George Foster, and Gholamreza Haffari. 2024. [Adapting large language models for document-level machine translation](#).
- 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.
- Fei Yuan, Yinqun Lu, Wenhao Zhu, Lingpeng Kong, Lei Li, Yu Qiao, and Jingjing Xu. 2023. [Lego-MT: Learning detachable models for massively multilingual machine translation](#). In *Findings of the Association for Computational Linguistics: ACL 2023*, pages 11518–11533, Toronto, Canada. Association for Computational Linguistics.
- Dawei Zhu, Pinzhen Chen, Miaoran Zhang, Barry Haddow, Xiaoyu Shen, and Dietrich Klakow. 2024a. [Fine-tuning large language models to translate: Will a touch of noisy data in misaligned languages suffice?](#) In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 388–409, Miami, Florida, USA. Association for Computational Linguistics.
- Dawei Zhu, Sony Trenous, Xiaoyu Shen, Dietrich Klakow, Bill Byrne, and Eva Hasler. 2024b. [A preference-driven paradigm for enhanced translation with large language models](#). In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pages 3385–3403, Mexico City, Mexico. Association for Computational Linguistics.
- Wenhao Zhu, Hongyi Liu, Qingxiu Dong, Jingjing Xu, Shujian Huang, Lingpeng Kong, Jiajun Chen, and Lei Li. 2024c. [Multilingual machine translation with large language models: Empirical results and analysis](#). In *Findings of the Association for Computational Linguistics: NAACL 2024*, pages 2765–2781, Mexico City, Mexico. Association for Computational Linguistics.
- Ahmet Üstün, Viraat Aryabumi, Zheng-Xin Yong, Wei-Yin Ko, Daniel D’souza, Gbemileke Onilude, Neel Bhandari, Shivalika Singh, Hui-Lee Ooi, Amr Kayid, Freddie Vargus, Phil Blunsom, Shayne Longpre, Niklas Muennighoff, Marzieh Fadaee, Julia Kreutzer, and Sara Hooker. 2024. Aya model: An instruction finetuned open-access multilingual language model. *arXiv preprint arXiv:2402.07827*.

## A More details about AFRIDOC-MT

### A.1 Translation guidelines

The translation guidelines, aside the details shared at the workshop on translation and terminology creation can be found below.

- Thank you for agreeing to work on this project. Below is the link to access the data for translation. The files are in .csv format, and you can open them using Google Sheets or Microsoft Excel (for offline work).
- Each file contains 2500 sentences, and they are named in the format of a serial number followed by your first name.
- Please do not delete double empty rows, as they serve to separate paragraphs. Also, avoid deleting any rows, columns, or provided text.
- Use the language field to input the translations. It is essential not to rely on translation engines, as our quality assurance process can detect this. Depending on such tools may result in potential issues that you would need to address, leading to additional work on your part.
- We will provide a list of extracted terminologies soon so that you can harmonize how terminologies are translated.
- Thank you for your attention to these guidelines. Should you have any questions, concerns, or suggestions, feel free to contact us or reach out to your language coordinator.

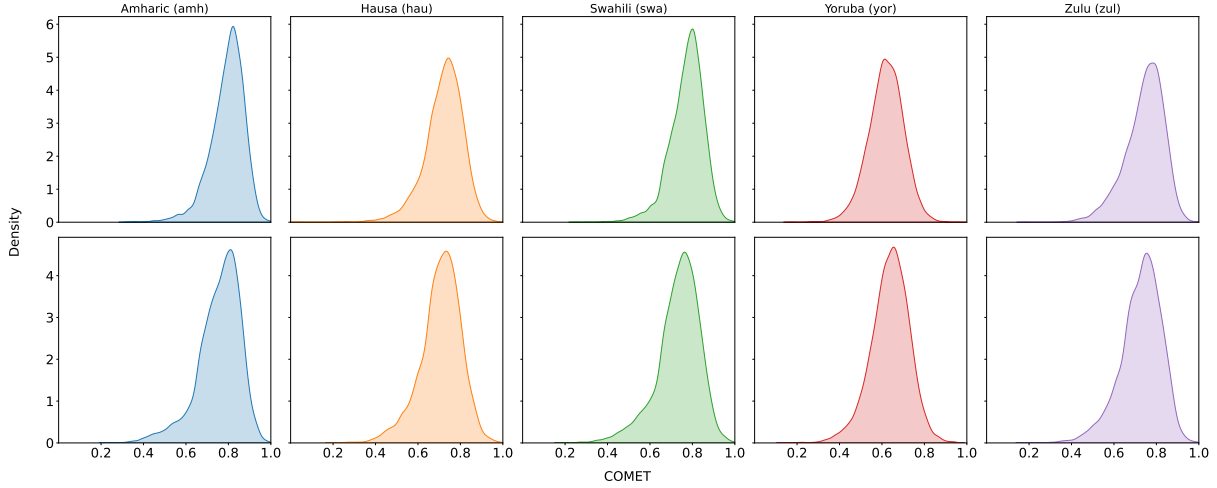


Figure 3: Distribution of the quality estimation of the translated sentences using COMET scores for the *health* (top), *tech* (bottom).

## A.2 Creation of pseudo-documents for AFRIDOC-MT

Given that the translated documents vary in length in terms of sentences and tokens, and considering the maximum token length limitations of the different LLMs used, we adopted a chunking approach for document-level evaluation. In this approach, documents were divided into smaller pseudo-documents that fit within the maximum length constraints of the models. To establish an appropriate chunk size, each document was divided into fixed-size chunks of  $k$  sentences, with the possibility that the final chunk may contain fewer than  $k$  sentences. These sentence groups, referred to as pseudo-documents, were used for document-level translation.

We conducted an initial analysis, testing different values for  $k$  (5, 10, and 25), with  $k=1$  serving as our sentence-level setup. Table 11 presents the resulting number of parallel pseudo-documents, as well as the average number of tokens per pseudo-document per language for the various model tokenizers, including the 95th percentile token count. Our analysis revealed that Amharic and Yorùbá—languages with unique characteristics such as non-Latin scripts and diacritics, respectively—had the largest average token counts across the tokenizers. Additionally, the domain with the highest number of average tokens for pseudo-document varies from language to language.

To accommodate both languages in our experiments, we chose pseudo-documents with  $k=10$ . However, for the SFT models described in Appendix B.2, we used both  $k=5$  and  $k=10$ .

## B Experimental details

### B.1 Evaluated Models

#### B.1.1 Translation Models

**M2M-100** (Fan et al., 2020) is a transformer-based multilingual neural translation model from Meta, trained to translate between 100 languages, including several African languages. It has three variants of different sizes: 400M parameters, 1.2B parameters, and 12B parameters. For our experiments, we evaluated the 400M and 1.2B variants.

**NLLB** (NLLB Team et al., 2024) is a model similar to M2M-100, with broader coverage, trained to translate between just over 200 languages, including more than 50 African languages. It also has different sizes: 600M, 1.3B, 3.3B, and 54B parameters. For this work, we evaluated the first three variants.

**MADLAD-400** (Kudugunta et al., 2023) is a multilingual translation model based on the T5 architecture (Raffel et al., 2020), covering 450 languages, including many African languages. It was trained on data collected from the Common Crawl dataset. The dataset underwent a thorough self-audit to filter out noisy content and ensure its quality for training MT models.

**Toucan** (Elmadany et al., 2024; Adebara et al., 2024) is another multilingual but Afro-centric translation model based on the T5 architecture, covering 150 language pairs of African languages. It was first pre-trained on large multilingual texts covering over 500 African languages and then finetuned on translation task covering over 100 language pairs.

| Languages/Split                                                        | Models      | Full           |                | 25 sent.      |               | 10 sent.      |               | 5 sent.       |              |
|------------------------------------------------------------------------|-------------|----------------|----------------|---------------|---------------|---------------|---------------|---------------|--------------|
|                                                                        |             | Health         | Tech           | Health        | Tech          | Health        | Tech          | Health        | Tech         |
| Sizes of data splits in AFRIDOC-MT pseudo-document                     |             |                |                |               |               |               |               |               |              |
| Train                                                                  |             | 240            | 187            | 402           | 369           | 812           | 789           | 1506          | 1483         |
| Dev                                                                    |             | 33             | 25             | 56            | 48            | 112           | 106           | 209           | 204          |
| Test                                                                   |             | 61             | 59             | 108           | 106           | 224           | 227           | 417           | 418          |
| Statistics of LLM tokens in AFRIDOC-MT pseudo-document training splits |             |                |                |               |               |               |               |               |              |
| en                                                                     | NLLB-200    | 923.7/2017.6   | 941.9/1982.1   | 551.5/951.7   | 477.4/758.8   | 273.0/430.9   | 223.2/343.6   | 147.2/233.8   | 118.8/184.9  |
|                                                                        | MADLAD-400  | 971.0/2095.2   | 991.4/2100.1   | 579.7/1017.1  | 502.4/797.8   | 287.0/449.3   | 235.0/362.0   | 154.7/245.0   | 125.0/196.9  |
|                                                                        | Aya-101     | 1008.2/2183.5  | 1020.5/2184.3  | 601.9/1038.0  | 517.2/820.2   | 298.0/463.4   | 241.9/372.6   | 160.7/255.0   | 128.7/199.0  |
|                                                                        | LLama3      | 801.4/1788.0   | 842.5/1798.4   | 478.5/833.8   | 427.0/664.0   | 236.9/372.9   | 199.7/304.2   | 127.8/203.0   | 106.3/166.0  |
|                                                                        | Gemma-2     | 802.9/1820.1   | 857.9/1857.6   | 479.3/841.0   | 434.8/689.6   | 237.3/375.0   | 203.4/314.0   | 128.0/205.0   | 108.2/169.0  |
|                                                                        | ModernBERT  | 801.0/1800.7   | 862.7/1819.3   | 478.3/837.8   | 437.3/685.6   | 236.8/373.4   | 204.5/311.0   | 127.8/204.0   | 108.9/171.0  |
| am                                                                     | NLLB-200    | 1304.4/2785.8  | 1376.3/2888.7  | 778.8/1329.9  | 697.5/1130.8  | 385.6/592.0   | 326.2/520.0   | 207.9/328.0   | 173.5/282.9  |
|                                                                        | MADLAD-400  | 1624.8/3393.6  | 1685.0/3487.4  | 970.0/1684.2  | 853.9/1380.4  | 480.2/750.0   | 399.4/640.2   | 258.9/413.8   | 212.5/342.9  |
|                                                                        | Aya-101     | 1887.4/3937.9  | 1934.7/4126.9  | 1126.8/1931.8 | 980.5/1598.0  | 557.9/855.4   | 458.5/722.0   | 300.8/477.8   | 244.0/390.0  |
|                                                                        | LLama3      | 6798.0/13986.2 | 6829.6/14750.9 | 4058.5/6971.8 | 3461.1/5584.8 | 2009.3/3084.4 | 1618.7/2560.8 | 1083.3/1716.0 | 861.2/1379.9 |
|                                                                        | Gemma-2     | 2817.9/5857.5  | 2868.4/6227.4  | 1682.1/2896.4 | 1453.2/2342.4 | 832.4/1267.8  | 679.3/1071.6  | 448.5/710.0   | 361.0/575.0  |
|                                                                        | ModernBERT  | 7347.8/15045.1 | 7386.4/15952.3 | 4386.4/7544.1 | 3742.8/6035.8 | 2171.1/3331.2 | 1749.9/2760.4 | 1170.2/1851.0 | 930.6/1485.9 |
| ha                                                                     | NLLB-200    | 1204.4/2713.7  | 1171.4/2463.0  | 719.0/1252.8  | 593.6/962.6   | 356.0/554.0   | 277.6/430.6   | 191.9/306.8   | 147.7/232.0  |
|                                                                        | MADLAD-400  | 1297.1/2849.4  | 1260.5/2643.7  | 774.4/1359.7  | 638.8/1042.0  | 383.4/606.4   | 298.8/465.6   | 206.7/329.0   | 158.9/251.0  |
|                                                                        | Aya-101     | 1614.9/3497.4  | 1535.3/3241.9  | 964.1/1672.3  | 778.0/1254.6  | 477.3/742.6   | 363.9/563.2   | 257.4/410.8   | 193.6/306.0  |
|                                                                        | LLama3      | 1916.7/4012.9  | 1822.6/3917.9  | 1144.3/1988.8 | 923.7/1513.6  | 566.6/882.4   | 432.1/674.6   | 305.5/488.8   | 230.0/365.9  |
|                                                                        | Gemma-2     | 1642.4/3568.9  | 1581.3/3373.4  | 980.6/1716.7  | 801.4/1297.8  | 485.5/757.4   | 374.8/584.0   | 261.8/417.8   | 199.4/317.8  |
|                                                                        | ModernBERT  | 1998.5/4207.5  | 1916.8/4139.7  | 1193.1/2057.8 | 971.5/1575.8  | 590.8/916.9   | 454.4/701.0   | 318.6/510.8   | 241.8/382.9  |
| sw                                                                     | NLLB-200    | 1100.8/2494.8  | 1094.8/2187.5  | 657.2/1145.9  | 554.8/896.4   | 325.4/517.0   | 259.5/409.6   | 175.4/280.0   | 138.1/218.0  |
|                                                                        | MADLAD-400  | 1177.3/2629.9  | 1155.3/2293.9  | 702.8/1227.6  | 585.5/938.6   | 348.0/547.0   | 273.8/436.0   | 187.6/297.0   | 145.7/231.9  |
|                                                                        | Aya-101     | 1345.3/2925.0  | 1311.0/2667.8  | 803.2/1390.9  | 664.4/1076.2  | 397.6/627.9   | 310.7/487.4   | 214.4/339.0   | 165.3/261.0  |
|                                                                        | LLama3      | 1668.1/3605.0  | 1619.4/3364.9  | 995.9/1735.4  | 820.7/1330.0  | 493.1/771.4   | 383.9/599.8   | 266.0/418.0   | 204.3/323.0  |
|                                                                        | Gemma-2     | 1413.3/3097.3  | 1377.1/2770.0  | 843.8/1467.7  | 697.9/1126.2  | 417.8/658.9   | 326.4/513.0   | 225.3/356.8   | 173.7/277.9  |
|                                                                        | ModernBERT  | 1757.9/3753.4  | 1719.7/3594.1  | 1049.5/1822.8 | 871.6/1421.0  | 519.7/810.0   | 407.7/632.0   | 280.2/441.0   | 217.0/342.8  |
| yo                                                                     | NLLB-200    | 1702.6/3854.7  | 1724.8/3577.1  | 1016.5/1857.2 | 874.1/1428.6  | 503.2/814.7   | 408.8/644.6   | 271.3/443.8   | 217.5/348.9  |
|                                                                        | MADLAD-400  | 1983.6/4470.9  | 1990.4/4136.7  | 1184.3/2137.5 | 1008.7/1650.2 | 586.3/939.4   | 471.7/742.2   | 316.1/512.0   | 251.0/401.9  |
|                                                                        | Aya-101     | 2729.2/5832.3  | 2659.8/5549.7  | 1629.4/2956.4 | 1347.9/2211.6 | 806.7/1292.4  | 630.4/988.0   | 434.9/704.0   | 335.4/544.0  |
|                                                                        | LLama3      | 2945.8/6322.4  | 2880.0/5995.5  | 1758.6/3203.9 | 1459.4/2400.4 | 870.5/1406.0  | 682.5/1077.6  | 469.3/767.8   | 363.0/585.9  |
|                                                                        | Gemma-2     | 2620.4/5745.5  | 2593.5/5406.9  | 1564.3/2867.7 | 1314.3/2143.8 | 774.4/1245.4  | 614.6/965.6   | 417.4/678.0   | 327.0/530.0  |
|                                                                        | ModernBERT  | 3648.3/7780.9  | 3595.2/7600.6  | 2178.1/4002.0 | 1822.0/3020.4 | 1078.3/1761.4 | 852.1/1339.8  | 581.4/957.2   | 453.3/733.9  |
| zu                                                                     | NLLB-200    | 1201.8/2513.3  | 1230.4/2555.7  | 717.5/1233.0  | 623.5/1016.6  | 355.2/554.3   | 291.6/461.2   | 191.5/300.0   | 155.1/250.0  |
|                                                                        | MADLAD-400  | 1215.2/2524.0  | 1230.7/2519.6  | 725.5/1284.8  | 623.7/1007.2  | 359.2/557.8   | 291.7/465.6   | 193.7/305.5   | 155.2/251.0  |
|                                                                        | Aya-101     | 1491.3/3012.2  | 1485.2/3180.8  | 890.3/1521.8  | 752.7/1213.0  | 440.8/688.9   | 352.0/554.4   | 237.7/372.8   | 187.3/298.9  |
|                                                                        | LLama3      | 1921.7/3822.6  | 1834.3/3933.4  | 1147.3/1963.9 | 929.7/1512.4  | 568.1/885.4   | 434.9/689.2   | 306.4/475.8   | 231.5/373.0  |
|                                                                        | Gemma-2     | 1787.5/3573.5  | 1703.0/3666.1  | 1067.2/1834.8 | 863.0/1416.2  | 528.3/819.4   | 403.6/637.6   | 284.9/447.8   | 214.8/343.9  |
|                                                                        | MordernBERT | 2073.7/4134.2  | 1965.8/4239.3  | 1238.1/2138.4 | 996.3/1625.6  | 613.0/956.3   | 466.1/737.0   | 330.6/515.8   | 248.0/399.0  |

Table 11: AFRIDOC-MT Pseudo-document statistics. The number of translation instances in the data AFRIDOC-MT pseudo-document splits. average and 95th percentile (average/95 percentile) of the AFRIDOC-MT document train split tokenization statistics using the different LLM tokenizers.

### B.1.2 Large Language Models

**Aya-101** (Üstün et al., 2024) is an instruction-tuned mT5 model (Xue et al., 2021) designed to handle both discriminative and generative multi-lingual tasks. With 13B parameters, it covers 100 languages and is capable of translating between a wide range of languages, including African languages.

**Gemma2** (Gemma Team et al., 2024) is a decoder-only LLM trained on billions of tokens sourced from the web. The training data primarily consists of English-language text, but it also includes code and mathematical content. While Gemma2 has an English-centric focus, it also possesses multilingual capabilities. We evaluate the base Gemma2 model with 9B parameters, as well as its instruction-tuned version.

**LLama3.1** (Dubey et al., 2024) is another decoder-only LLM trained on trillions of tokens across multiple languages. It was fine-tuned using existing instruction datasets as well as synthetically generated instruction data to create its instruction-tuned version. One advantage LLama3.1 has over

other models is its context window of 128K tokens, the largest among all models considered in this work, making it particularly suitable for document-based tasks such as document-level translation. We evaluate the base LLama3.1 model with 8B parameters, as well as its instruction-tuned version.

**LLaMAX3** (Lu et al., 2024) is a multilingual LLM built on the LLama3 with 8B parameters as its base. It was trained on 102 languages, including several African languages, through continued pretraining. Using an English instruction dataset (Alpaca), it was further fine-tuned to create LLaMAX3-Alpaca. We evaluated both models and compared their performance across various tasks.

### B.2 Supervised Finetuning

We perform supervised fine-tuning to tailor LLMs for translation tasks. To train sentence-level MT systems, we use all parallel sentences from AFRIDOC-MT to construct the training set, enabling the LLMs to translate across multiple directions and domains. Following Zhu et al. (2024b), we augment the parallel data with translation in-

structions, which are randomly sampled from a pre-defined set of 31 MT instructions for each training example.<sup>12</sup> To train document-level MT systems, we follow the same process, but train on longer segments formed by concatenating multiple sentences. When fine-tuning, we use a learning rate of  $5e^{-6}$  and an effective batch size of 64. Models are trained for only one epoch, as further training does not result in improvements and may even lead to performance degradation.

Similarly, we fine-tuned the 1.3B version of NLLB-200 for sentence and pseudo-document (with 10 sentences) translation using the Fairseq (Ott et al., 2019) codebase. We used all the training examples from 30 language directions across both domains. The model was fine-tuned for 50k steps using a learning rate of  $5e^{-5}$ , token batch size of 2048 and a gradient accumulation of 2. The checkpoint with the lowest validation loss was selected as the best model for evaluation.

### B.3 Evaluation setup

The models were evaluated using different tools. For example, both the NLLB-200 and M2M-100 models were evaluated with the Fairseq codebase, while Toucan and MADLAD-400 were evaluated using the Hugging Face (HF) codebase. All other LLMs, including LLama3.1 (both instruction-tuned and SFT models), Gemma, and Aya-101, were evaluated using EleutherAI LM Evaluation Harness (lm-eval) tool (Biderman et al., 2024). In all cases, greedy decoding was used.

The models evaluated have different context lengths. For encoder-decoder models, M2M-100 and NLLB have a maximum sequence length of 1024 and 512 respectively. Aya-101 and MADALAD, based on the T5 architecture, do not have a pre-specified maximum sequence length, so we fixed their maximum sequence length to 1024 for all experiments involving encoder-decoder models. However, for decoder-only models, Gemma and LLaMAX3 (based on LLama3) have a maximum sequence length of 8192, while LLama3.1 has a maximum sequence length of 128K. Since all the decoder-only models were evaluated using LM Evaluation Harness, we used a similar setup for them, selecting the maximum length based on the specific needs of each model.

Table 12 shows the maximum number of gen-

<sup>12</sup>We use the same instruction set as described in (Zhu et al., 2024b).

| Setting         | X → eng | eng → X       |
|-----------------|---------|---------------|
| <b>Sentence</b> |         |               |
| sentence        | 512     | 512           |
| <b>Document</b> |         |               |
| 5               | 4096    | 4096          |
| 10              | 4096    | 4096          |
| 25              | 1024    | 8192 (11264)  |
| Full            | 2048    | 16384 (32768) |

Table 12: The maximum number of tokens set for decoder-only LLMs when translating between English and African languages, and vice versa. Special cases for Amharic are indicated in brackets.

eration tokens we set when translating between English and African languages. These numbers were chosen based on the statistics from Table 11. However, for Amharic, when translating pseudo-documents with 25 sentences and full documents, there were instances exceeding the 95th percentile derived from the training statistics. Therefore, we increased the token limit specifically for Amharic.

### B.4 Evaluation prompts

While the translation models we evaluated do not require prompts, MADLAD-400, requires a prefix of the form `<2xx>` token, which is prepended to the source sentence. Here, xx indicates the target language using its language code (e.g., “sw” for Swahili). Similarly, Toucan uses just the target language ISO-639 code as prefix, which is prepended to the source sentence (e.g., “swa” for Swahili). For other models, including Aya-101, we used three different prompts for sentence-level translation and document translation experiments. The main difference between the prompts for these tasks is the explicit mention of “text” or “document” within the prompt, as shown in Table 23. For the base models Gemma2, LLama3.1, LLaMAX3, and Aya-101, we prompted them directly using the respective prompts. However, for the instruction-tuned versions of Gemma2 and LLama3.1, we used their respective chat templates. For all Alpaca-based models, including our SFT models, we used the Alpaca template.

### B.5 Evaluation metrics

We evaluate translation quality with BLEU (Papineni et al., 2002) and ChrF (Popović, 2015) using SacreBLEU<sup>13</sup> (Post, 2018). We run significance tests using bootstrap resampling and report the 95%

<sup>13</sup>case:mixed|eff:no| tok:13a|smooth:exp|v:2.3.1,

confidence interval for the scores, based on a sample size of 1000. We also use AfriCOMET<sup>14</sup> (Wang et al., 2024a) to evaluate the quality of the translation outputs. We report the chrF scores of the best prompt for each model and language direction in the main paper, with all additional results provided in the Appendix C. For document-level experiments, we evaluated the LLMs using the same three prompts as in the sentence-level experiment. For evaluation, we used BLEU and chrF scores but excluded AfriCOMET due to its backbone model, AfroXLM-R-L (Alabi et al., 2022; Adelani et al., 2024a), having a context length of 512 tokens. This made it impractical to compute COMET scores for document-level outputs.

## B.6 GPT-4o as an evaluator for machine translation

We use GPT-4o to assess the quality of translation output, as demonstrated by Sun et al. (2025), which shows a correlation with human judgment. Due to the cost of this task, we limited our evaluation to a few selected models, including Aya-101, GPT-3.5, GPT-4o, and LLaMAX3 fine-tuned on AFRIDOC-MT sentences and pseudo-documents of 10 sentences. We compared translations performed at the sentence-level and pseudo-document level in terms of fluency, content errors, and cohesion errors—specifically lexical (LE) and grammatical (GE) errors—using the same definitions as Sun et al. (2025).

Below are the prompts used to evaluate documents using GPT-4o for fluency, content errors, and cohesion errors—specifically lexical (LE) and grammatical (GE) errors.

- **Fluency:** GPT-4o is prompted to rate the fluency of a document on a scale from 1 to 5, where 5 indicates high fluency and 1 represents low fluency. This evaluation is conducted without providing any reference document. For the final fluency score, we report the average rating across all documents. Below we provide the prompt used.

```
Please evaluate the fluency of the
following text in <<target>>.
```

```
-----
```

```
### **Instructions:**
```

<sup>14</sup><https://huggingface.co/masakhane/africomet-stl-1.1>

```
- **Task**: Evaluate the fluency of
the text.

- Scoring: Provide a score from 1 to
5, where:

- **5**: The text is **highly
fluent**, with no grammatical
errors, unnatural wording, or
stiff syntax.
- **4**: The text is **mostly
fluent**, with minor errors
that do not impede
understanding.
- **3**: The text is **moderately
fluent**, with noticeable
errors that may slightly
affect comprehension.
- **2**: The text has **low
fluency**, with frequent
errors that hinder
understanding.
- **1**: The text is **not fluent
**, with severe errors that
make it difficult to
understand.

- **Explanation**: Support your
score with specific examples to
justify your evaluation.
```

```
-----
```

```
### **Output Format:**
```

```
Provide your evaluation in the
following JSON format:
```

```
““
{
  "Fluency": {
    "Score": "<the score>",
    "Explanation": "<your
explanation on how you made
the decision>"
  }
}
““
```

```
-----
```

```
**Text to Evaluate:**
```

```
<<hypothesis>>
```

```
Answer:
```

- **Accuracy:** GPT-4 is prompted to identify and list the mistakes, such as incorrect translations, omissions, additions, and any other errors, by comparing the model's output to the reference translation. After identifying these errors, we count all of them and compute the average across all documents, reporting that as the content error (CE). Below is the prompt used.

```
Please evaluate the accuracy of the
following translated text in <<
```

```

target>> by comparing it to the
provided reference text.

-----

### **Instructions:**

- **Task**: Compare the text to the
reference text.

- Identify Mistakes: List all
mistakes related to accuracy.

- Mistake Types:

  - **Wrong Translation**:
    Incorrect meaning or
    misinterpretation leading to
    wrong information.
  - **Omission**: Missing words,
    phrases, or information
    present in the reference
    text.
  - **Addition**: Extra words,
    phrases, or information not
    present in the reference
    text.
  - **Others**: Mistakes that are
    hard to define or categorize
    .

- **Note**: If the text expresses
the same information as the
reference text but uses
different words or phrasing, it
is not considered a mistake.

- **Provide a List**: Summarize all
mistakes without repeating the
exact sentences. Provide an
empty list if there are no
mistakes.

-----

### **Output Format:**

Provide your evaluation in the
following JSON format:

'''
{
  "Accuracy": {
    "Mistakes": [
      "<list of all mistakes in the
text with format 'Mistake
Types: summarize the
mistake', provide an empty
list if there are no
mistakes>"
    ]
  }
}
'''

-----

**Reference Text:**

<<reference>>

```

```

**Text to Evaluate:**

<<hypothesis>>

```

- **Cohesion:** GPT-4 is prompted to rate cohesion-related mistakes, including lexical and grammatical errors, in the model's output, comparing it to the reference translation. We count each error individually, compute the average across the documents, and report them as lexical errors (LE) and grammatical errors (GE). Below is the prompt template we used.

```

Please evaluate the cohesion of the
following translated text in
<<target>> by comparing it to
the provided reference text.

-----

### **Instructions:**

- **Task**: Evaluate the cohesion of
the text.

- **Definition**: Cohesion refers to
how different parts of a text
are connected using language
structures like grammar and
vocabulary. It ensures that
sentences flow smoothly and the
text makes sense as a whole.

- Identify Mistakes: List all
mistakes related to cohesion.

- Separate the mistakes into:

  - **Lexical Cohesion Mistakes**:
    Issues with vocabulary
    usage, incorrect or missing
    synonyms, or overuse of
    certain words that disrupt
    the flow.
  - **Grammatical Cohesion
    Mistakes**: Problems with
    pronouns, conjunctions, or
    grammatical structures that
    link sentences and clauses.

- **Provide Lists**: Provide
separate lists for lexical
cohesion mistakes and
grammatical cohesion mistakes.
Provide empty lists if there are
no mistakes.

-----

### **Output Format:**

Provide your evaluation in the
following JSON format:

'''
{

```

```

"Cohesion": {
  "Lexical Cohesion Mistakes": [
    "<list of all mistakes in the
    text one by one, provide
    an empty list if there are
    no mistakes>"
  ],
  "Grammatical Cohesion Mistakes":
  [
    "<list of all mistakes in the
    text one by one, provide
    an empty list if there are
    no mistakes>"
  ]
}
}
'''

-----

**Reference Text:**

<<reference>>

**Text to Evaluate:**

<<hypothesis>>

```

Fluency can only have values between 1 and 5. However, the other metrics, including CE, GE, and LE, do not have a specific range and can take on any value because they are counts. Refer to (Sun et al., 2025) for more details about these metrics.

### B.7 Human Evaluation Setup

Beyond using GPT-4o as a judge, we also conduct human evaluation on a subset of outputs from GPT-3.5, LLaMAX3-SFT<sub>1</sub>, and LLaMAX3-SFT<sub>10</sub> for two domains, focusing specifically on translation into five African languages due to cost constraints. Translation into English was excluded, as existing automatic metrics, including GPT-based evaluations, are already reliable for this direction.

For the human evaluation, three native speakers of the African languages—primarily translators involved in the dataset creation—were recruited. Each annotator was assigned 80 documents to evaluate, tasked with marking as many error spans as possible and rating the overall quality on a scale from 0 to 100. This annotation followed the error span annotation (ESA) (Kocmi et al., 2024) protocol as implemented within the Appraise Evaluation Framework (Federmann, 2018). To assess consistency and inter-annotator agreement, 30 of the 80 documents were shared among all three annotators. Table 13 shows statistics for 80 documents sampled from the models in both domains for each annotator. Each annotator was remunerated with

| Model                     | Setup   | Full   |      | Shared |      |
|---------------------------|---------|--------|------|--------|------|
|                           |         | health | tech | health | tech |
| GPT-3.5                   | Sent.   | 5      | 5    | -      | 5    |
|                           | Pseudo. | 5      | 5    | -      | 5    |
| LLaMAX3-SFT <sub>1</sub>  | Sent.   | 5      | 5    | 5      | -    |
|                           | Pseudo. | 5      | 5    | 5      | -    |
| LLaMAX3-SFT <sub>10</sub> | Sent.   | 5      | 5    | 5      | 5    |
|                           | Pseudo. | 5      | 5    | 5      | 5    |
| Total                     |         | 25     | 25   | 15     | 15   |

Table 13: The number of documents annotated by each annotator for human direct assessment.

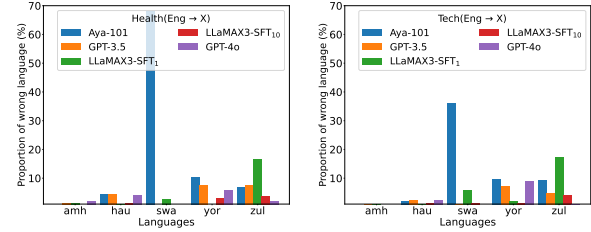


Figure 4: Rate of off-target translation ( $k = 10$ ).

\$55.<sup>15</sup>

### B.8 Qualitative Analysis

Alongside the human direct assessment of the translation outputs, we shared a subset of the outputs with one author per language, each a native speaker. They were tasked with analyzing the outputs to answer two key questions: (1) What common errors or flaws do the models exhibit across different setups? and (2) How fluent are the translation outputs produced by the models across the various settings?

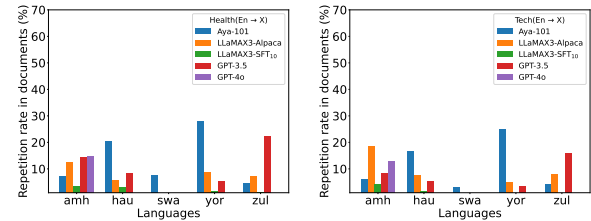


Figure 5: Word repetition rate in the pseudo-document translation ( $k = 10$ ).

## C More experimental results

### C.1 Sentence-level evaluation

Given that AFRIDOC-MT is a document-level translation dataset, and due to the limited context length of most translation models and LLMs, which makes it impossible to translate a full document at once, we opted to translate the sentences within the documents and then merge them back to form the complete document. This also serves as a baseline for document-level translation. In the main paper, we present the results for the best prompt

<sup>15</sup>Annotation protocol.

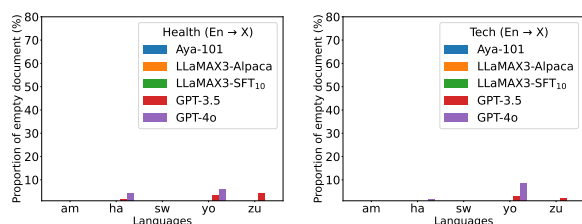


Figure 6: Proportion of empty outputs for pseudo-documents.

for each language pair and model using d-chrF. In this section, we also provide the full results on the merged documents using d-chrF and d-BLEU in Tables 15 and 16. Furthermore, we present results for evaluating just the sentences (without merging them back into documents) using s-BLEU, s-chrF, and s-COMET in Tables 17 and 18. In Figures 16 to 19, we provide plots that summarize some of the results in the table for a few models. Although the main findings are summarized in the main draft, below are some other points we identify.

**M2M-100 is not competitive** Neither version of M2M-100, which was once a state-of-the-art translation model, is competitive with other translation models such as Toucan, NLLB-200, and MADLAD-400, even when compared to models of similar sizes, across all metrics at both the sentence and document levels.

**Base LLMs are not translators for African languages.** Base LLMs without instruction tuning and supervised fine-tuning, such as Gemma2 and LLaMAX3, do not show competitive translation performance either. This can be explained by the fact that they are just language models with limited coverage of African languages. However, LLaMAX3, which was trained on more than 100 languages, including African languages, through continued pre-training, shows improved performance, surpassing LLaMA3.1-IT.

**Amharic and Yorùbá are the worst performing language directions.** When translating from English into African languages, our results show that both Amharic and Yoruba perform the least effectively. This may be attributed to specific properties of these languages, such as the use of non-Latin script in Amharic and the use of diacritics in Yoruba, which in turn increase the tokenization rate of these languages by the different model tokenizers.

## C.2 Document-level evaluation

For document-level evaluation, we split the documents into chunks of 10 sentences and translate these chunks using the different models. In Tables 19 and 20 we provide the full results on the merged pseudo-documents using d-chrF and d-BLEU. And below are some other relevant points from the results. It is important to note that we also trained and evaluated NLLB-200 for pseudo-document translation. However, due to its 512-token maximum sequence length, it is not competitive. Nevertheless, the results still show the influence of fine-tuning. Below are other findings.

**Gemma2-IT shows better translation performance.** Compared to the sentence-level setup, where Gemma2-IT and LLaMAX3-Alpaca achieved similar performance on average, in the pseudo-document setup, Gemma2-IT not only outperforms LLaMAX3-Alpaca but also surpasses GPT-3.5. Although we cannot provide an exact explanation for this performance, we hypothesize that its pre-training setup might be a contributing factor.

**Fine-tuning data has an impact on translation quality.** Our results show that both LLaMA3.1 and LLaMAX3 models, when fine-tuned on sentences, performed significantly worse on pseudo-document evaluations compared to the same models fine-tuned on pseudo-documents for both domains. All these models were trained using a similar setup, with the primary difference being the data used for fine-tuning.

**Language-specific performance trends** Overall, no clear trend is observed in MT performance across language family classes. However, Amharic (a non-Latin script language) and Yorùbá (a heavily diacriticized language) result in the lowest chrF scores, while Swahili—the most widely spoken indigenous African language—performs best.

## C.3 Findings from GPT-4o as a judge

In Tables 21 and 22 we present the average GPT-4o evaluation results for four models. When translating into African languages, there is no clear pattern: for example, GPT-3.5, despite having the lowest fluency score, also had the fewest content, lexical, and grammatical errors, which is counterintuitive. In contrast, when translating into English, the pattern is clear and consistent: translations of pseudo-documents show better fluency and fewer errors

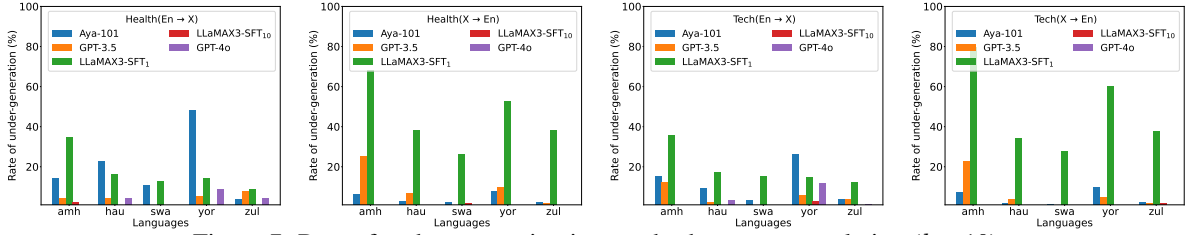


Figure 7: Rate of under-generation in pseudo-document translation ( $k = 10$ )

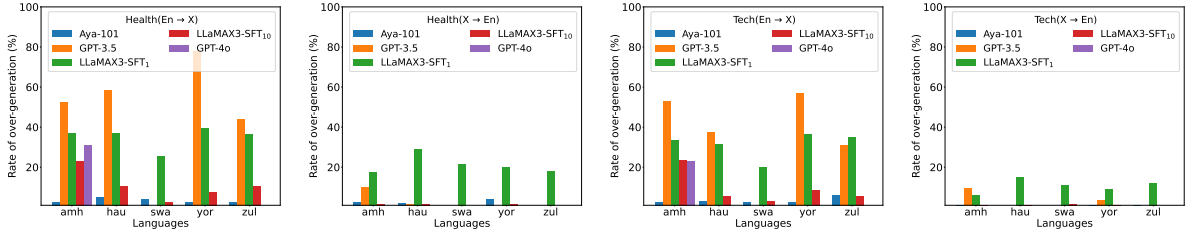


Figure 8: Rate of over-generation in pseudo-document translation ( $k = 10$ )

| Model                     | Setup | amh         | hau         | swa         | yor         | zul         |
|---------------------------|-------|-------------|-------------|-------------|-------------|-------------|
| GPT-3.5                   | Sent  | 18.3        | 42.1        | <b>64.1</b> | 12.6        | 8.9         |
|                           | Doc10 | 4.8         | 32.3        | 58.5        | 6.9         | 17.5        |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | <b>58.1</b> | <b>86.0</b> | 62.1        | <b>66.1</b> | <b>52.0</b> |
|                           | Doc10 | 19.0        | 54.5        | 38.7        | 40.1        | 21.6        |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 54.3        | 83.7        | 61.7        | 62.3        | 37.1        |

Table 14: Average DA score (scale 0–100) from three human evaluators per language in the *tech* domain.

overall. These results suggest that using GPT-4o as a translation judge is not yet well-suited for low-resource languages.

#### C.4 Findings from human evaluation

We were able to obtain DA scores from three annotators for all the languages. For each language, we calculated inter-annotator agreement using Krippendorff’s alpha  $\alpha$  over 30 document instances. We obtained  $\alpha$  scores of 0.46, 0.57, 0.40, and 0.81, and 0.54 for Amharic, Hausa, Swahili, Yorùbá, and Zulu respectively. These are relatively low scores, except for Yorùbá. We present the average DA scores in Tables 10 and 14 for the *health* and *tech* domains, respectively. The results show that annotators rate documents translated at the sentence-level as higher quality than those translated at the pseudo-document level. Additionally, GPT-3.5 received the lowest ratings among the three models. LLaMAX3-SFT<sub>1</sub>, a model trained on sentence-level data, was rated the best across all languages when evaluated on sentences. However, when evaluated on pseudo-documents, its performance was rated lower than that of LLaMAX3-SFT<sub>10</sub>. These findings are consistent with the d-chrF scores for the models, but they do not align with the evaluations from GPT-4o as a judge.

## D More discussion and analysis

**What language benefits more from supervised finetuning?** We focus on the sentence-level task and translated across all 30 directions for which the model was trained, evaluating both NLLB-200 (1.3B) and its fine-tuned version using d-chrF. Figures 13 and 14 show performance improvements after supervised fine-tuning of NLLB-200 for both domains. The results shows that translating into Yorùbá, which is the direction with the lowest d-chrF score from English among all the languages, benefited the most. One major factor contributing to this is the presence of diacritics. Furthermore, looking at their actual performances and not just the differences, our results show that translations into Swahili and English—both relatively high-resource languages—yield higher BLEU and chrF scores (see Figures 11 and 12), even after supervised finetuning. Hence, there is much to be done to improve translation performance between low-resource language pairs.

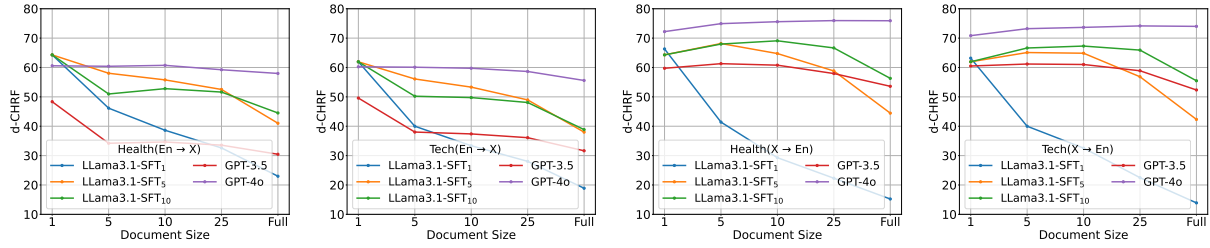


Figure 9: Average chrF score across languages for documents of different sizes.

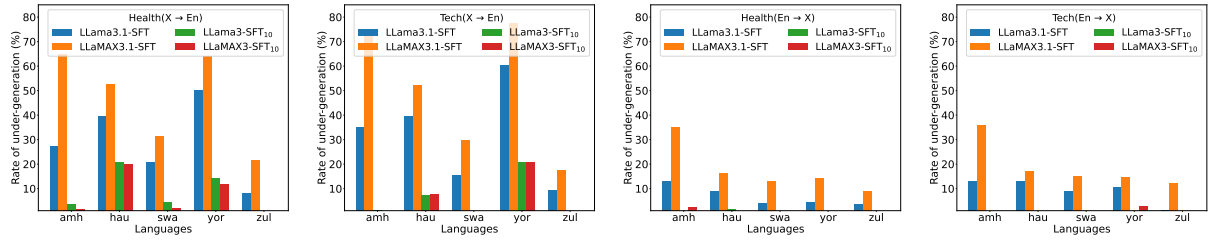


Figure 10: Rate of under-generation in our SFT models.

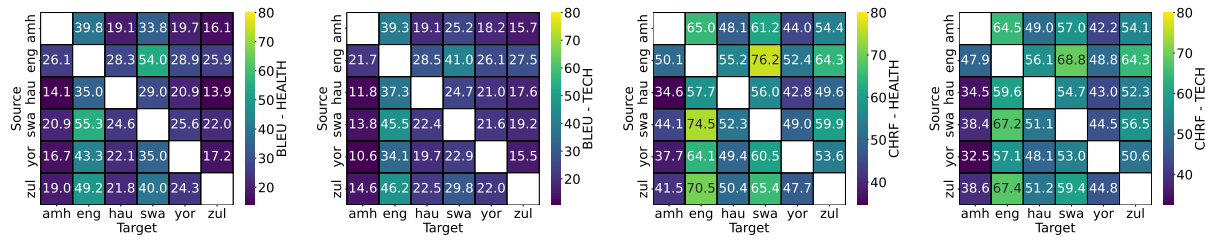


Figure 11: s-BLEU and s-chrF pair-wise comparison of supervised finetuning of NLLB-1.3B on AFRIDOC-MT

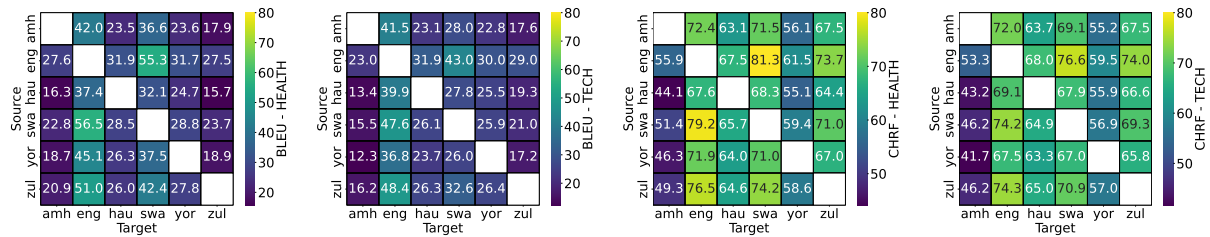


Figure 12: d-BLEU and d-chrF pair-wise comparison of supervised finetuning of NLLB-1.3B on AFRIDOC-MT

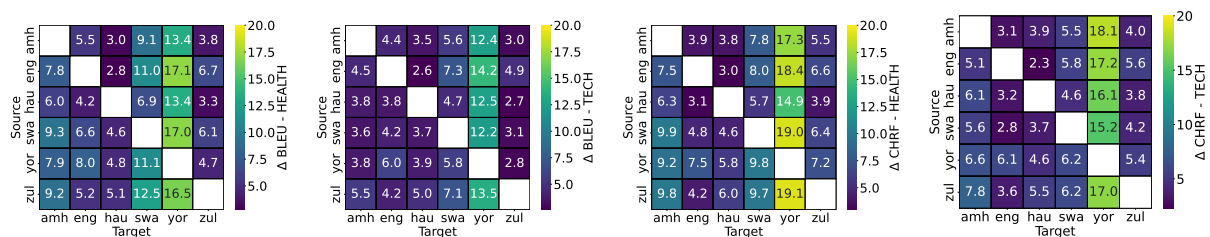


Figure 13: Change ( $\Delta$ ) in s-BLEU and s-chrF for sentence evaluation comparing NLLB1.3B before and after supervised finetuning on AFRIDOC-MT

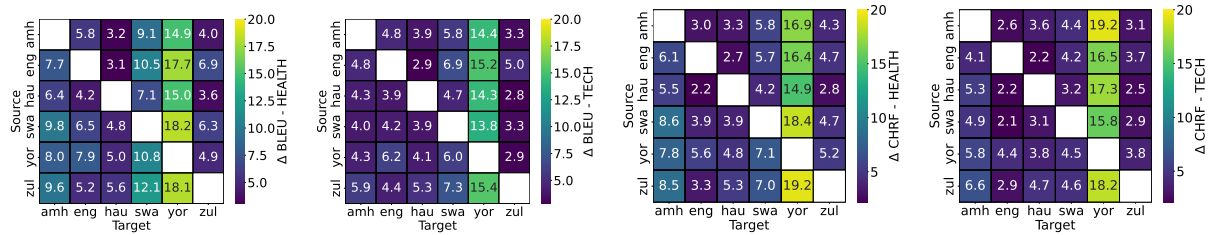


Figure 14: Change ( $\Delta$ ) in d-BLEU and d-chrF for sentence evaluation comparing NLLB1.3B before and after supervised finetuning on AFRiDOC-MT









| Model                     | Setup | $eng \rightarrow X$ |          |      |      |     | $X \rightarrow eng$ |          |      |      |     |
|---------------------------|-------|---------------------|----------|------|------|-----|---------------------|----------|------|------|-----|
|                           |       | d-CHRF↑             | Fluency↑ | CE↓  | LE↓  | GE↓ | d-CHRF↑             | Fluency↑ | CE↓  | LE↓  | GE↓ |
| Amharic                   |       |                     |          |      |      |     |                     |          |      |      |     |
| Aya-101                   | Sent  | 36.6                | 2.8      | 9.0  | 3.8  | 2.5 | 64.6                | 3.0      | 15.8 | 9.6  | 5.1 |
|                           | Doc10 | 28.7                | 2.5      | 8.9  | 2.9  | 2.1 | 61.6                | 3.6      | 12.6 | 8.5  | 3.8 |
| GPT-3.5                   | Sent  | 20.4                | 1.1      | 7.2  | 10.8 | 3.1 | 48.3                | 3.0      | 9.5  | 3.7  | 2.4 |
|                           | Doc10 | 11.6                | 1.0      | 3.3  | 1.8  | 1.6 | 41.6                | 4.2      | 6.6  | 2.0  | 1.3 |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 46.8                | 3.5      | 10.0 | 2.9  | 2.0 | 66.6                | 3.5      | 11.5 | 6.0  | 2.9 |
|                           | Doc10 | 24.1                | 1.7      | 10.4 | 1.9  | 1.8 | 22.6                | 3.0      | 9.0  | 2.5  | 1.6 |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 37.8                | 2.6      | 7.3  | 1.8  | 1.6 | 64.0                | 4.1      | 8.7  | 4.6  | 2.7 |
| Hausa                     |       |                     |          |      |      |     |                     |          |      |      |     |
| Aya-101                   | Sent  | 56.4                | 2.9      | 9.1  | 4.1  | 3.7 | 61.5                | 2.9      | 17.1 | 10.2 | 5.6 |
|                           | Doc10 | 48.5                | 2.9      | 8.3  | 2.3  | 1.9 | 62.3                | 3.2      | 14.6 | 8.4  | 4.5 |
| GPT-3.5                   | Sent  | 44.3                | 1.4      | 11.3 | 5.2  | 5.8 | 52.4                | 2.5      | 13.2 | 6.8  | 3.7 |
|                           | Doc10 | 23.1                | 1.1      | 7.4  | 2.9  | 2.8 | 52.7                | 4.1      | 9.2  | 4.1  | 2.3 |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 62.5                | 3.2      | 9.5  | 4.0  | 3.5 | 58.9                | 2.9      | 9.6  | 4.1  | 3.0 |
|                           | Doc10 | 29.0                | 2.3      | 8.8  | 2.3  | 1.9 | 22.9                | 2.6      | 9.0  | 3.0  | 1.9 |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 51.9                | 3.9      | 14.6 | 2.2  | 1.9 | 66.7                | 4.2      | 9.6  | 4.5  | 2.6 |
| Swahili                   |       |                     |          |      |      |     |                     |          |      |      |     |
| Aya-101                   | Sent  | 44.7                | 1.3      | 9.5  | 5.0  | 3.4 | 70.8                | 3.3      | 17.5 | 9.9  | 4.8 |
|                           | Doc10 | 34.7                | 1.8      | 8.2  | 5.0  | 3.0 | 71.2                | 3.8      | 14.1 | 9.8  | 4.0 |
| GPT-3.5                   | Sent  | 76.7                | 4.9      | 4.2  | 1.8  | 1.2 | 75.0                | 3.6      | 12.4 | 7.6  | 3.7 |
|                           | Doc10 | 76.1                | 4.9      | 3.9  | 0.8  | 0.6 | 77.7                | 4.7      | 7.1  | 4.5  | 2.0 |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 73.1                | 3.6      | 12.3 | 4.8  | 3.4 | 73.1                | 3.9      | 11.8 | 8.4  | 2.7 |
|                           | Doc10 | 42.2                | 3.2      | 9.5  | 3.6  | 2.5 | 33.1                | 3.1      | 8.9  | 3.2  | 1.9 |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 74.4                | 4.4      | 10.5 | 3.6  | 2.2 | 77.8                | 4.6      | 8.9  | 6.2  | 2.7 |
| Yoruba                    |       |                     |          |      |      |     |                     |          |      |      |     |
| Aya-101                   | Sent  | 31.2                | 1.2      | 8.2  | 5.4  | 3.1 | 57.9                | 2.3      | 23.8 | 13.6 | 6.8 |
|                           | Doc10 | 18.7                | 1.4      | 6.7  | 2.7  | 2.2 | 56.1                | 2.9      | 18.7 | 10.1 | 5.1 |
| GPT-3.5                   | Sent  | 21.3                | 1.1      | 7.8  | 5.5  | 3.9 | 52.1                | 2.6      | 12.9 | 5.0  | 3.3 |
|                           | Doc10 | 10.1                | 1.0      | 4.1  | 2.3  | 2.1 | 51.7                | 3.9      | 9.0  | 3.6  | 2.1 |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 57.5                | 4.0      | 12.3 | 3.8  | 2.7 | 64.7                | 3.2      | 12.2 | 6.1  | 3.1 |
|                           | Doc10 | 33.8                | 2.7      | 8.1  | 2.4  | 1.8 | 27.2                | 3.2      | 9.1  | 3.4  | 2.0 |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 52.2                | 4.2      | 10.3 | 2.4  | 1.7 | 71.8                | 4.4      | 10.2 | 5.9  | 2.4 |
| Zulu                      |       |                     |          |      |      |     |                     |          |      |      |     |
| Aya-101                   | Sent  | 58.6                | 2.7      | 13.7 | 6.0  | 3.6 | 67.4                | 2.9      | 19.1 | 12.2 | 8.6 |
|                           | Doc10 | 54.9                | 3.1      | 16.1 | 3.9  | 2.7 | 69.0                | 3.3      | 15.7 | 10.7 | 4.5 |
| GPT-3.5                   | Sent  | 51.1                | 1.5      | 20.9 | 10.0 | 6.1 | 59.5                | 2.6      | 15.9 | 8.8  | 5.7 |
|                           | Doc10 | 29.2                | 1.3      | 10.1 | 4.0  | 3.2 | 61.1                | 4.0      | 10.8 | 5.3  | 2.9 |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 67.5                | 3.3      | 12.4 | 5.5  | 3.6 | 70.5                | 3.6      | 12.5 | 7.2  | 3.8 |
|                           | Doc10 | 33.7                | 2.4      | 8.4  | 2.7  | 2.2 | 31.5                | 3.1      | 8.6  | 3.4  | 2.1 |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 55.0                | 3.6      | 12.1 | 3.0  | 2.0 | 74.1                | 4.4      | 9.4  | 5.5  | 2.5 |

Table 21: Document-level evaluation in the *health* domain, judged by GPT-4o. Compares sentence- vs. document-level outputs on Fluency (1–5 scale), Content Errors (CE), Lexical (LE), and Grammatical Cohesion Errors (GE). Best scores in bold.

| Model                     | Setup | $eng \rightarrow X$ |          |      |     |     | $X \rightarrow eng$ |          |      |      |     |
|---------------------------|-------|---------------------|----------|------|-----|-----|---------------------|----------|------|------|-----|
|                           |       | d-CHRF↑             | Fluency↑ | CE↓  | LE↓ | GE↓ | d-CHRF↑             | Fluency↑ | CE↓  | LE↓  | GE↓ |
| Amharic                   |       |                     |          |      |     |     |                     |          |      |      |     |
| Aya-101                   | Sent  | 37.3                | 3.2      | 9.9  | 3.4 | 3.3 | 65.2                | 2.9      | 19.3 | 11.2 | 5.8 |
|                           | Doc10 | 30.1                | 2.6      | 8.3  | 2.4 | 1.9 | 62.5                | 3.5      | 13.2 | 7.8  | 3.9 |
| GPT-3.5                   | Sent  | 22.6                | 1.3      | 8.5  | 4.2 | 3.3 | 47.4                | 2.4      | 8.0  | 3.6  | 2.4 |
|                           | Doc10 | 13.5                | 1.1      | 4.3  | 2.0 | 1.9 | 38.5                | 3.6      | 7.5  | 2.8  | 1.7 |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 42.8                | 3.2      | 11.6 | 3.0 | 2.4 | 63.0                | 3.3      | 14.2 | 7.9  | 3.4 |
|                           | Doc10 | 21.7                | 1.7      | 7.3  | 2.2 | 1.8 | 24.2                | 3.0      | 8.6  | 2.9  | 1.9 |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 37.7                | 2.9      | 8.5  | 2.5 | 1.7 | 65.4                | 4.0      | 10.6 | 5.4  | 2.5 |
| Hausa                     |       |                     |          |      |     |     |                     |          |      |      |     |
| Aya-101                   | Sent  | 58.9                | 3.1      | 10.4 | 3.5 | 2.9 | 64.8                | 3.3      | 17.1 | 13.8 | 6.5 |
|                           | Doc10 | 55.0                | 3.8      | 11.7 | 2.5 | 2.1 | 65.5                | 3.5      | 14.0 | 10.9 | 4.5 |
| GPT-3.5                   | Sent  | 49.2                | 1.9      | 11.9 | 7.4 | 6.0 | 56.5                | 2.8      | 12.7 | 7.1  | 5.2 |
|                           | Doc10 | 29.7                | 1.6      | 8.8  | 4.5 | 3.7 | 56.3                | 4.1      | 10.3 | 4.9  | 3.9 |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 62.4                | 4.0      | 10.1 | 4.1 | 3.0 | 53.5                | 2.9      | 10.3 | 4.7  | 3.4 |
|                           | Doc10 | 29.9                | 2.7      | 7.9  | 2.9 | 2.1 | 27.6                | 2.8      | 8.2  | 3.0  | 2.0 |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 58.6                | 4.3      | 11.9 | 2.4 | 2.1 | 68.5                | 4.5      | 8.7  | 4.3  | 2.0 |
| Swahili                   |       |                     |          |      |     |     |                     |          |      |      |     |
| Aya-101                   | Sent  | 42.4                | 1.5      | 9.0  | 5.0 | 3.3 | 69.1                | 3.2      | 18.4 | 10.9 | 5.4 |
|                           | Doc10 | 51.7                | 2.9      | 9.3  | 2.9 | 1.9 | 68.8                | 3.4      | 15.3 | 10.4 | 4.6 |
| GPT-3.5                   | Sent  | 72.6                | 4.7      | 5.0  | 0.8 | 0.5 | 71.5                | 3.4      | 15.3 | 9.4  | 4.6 |
|                           | Doc10 | 72.1                | 4.8      | 5.7  | 0.5 | 0.3 | 73.5                | 4.3      | 11.1 | 7.9  | 4.0 |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 67.6                | 3.6      | 12.3 | 4.8 | 3.4 | 67.5                | 3.8      | 12.2 | 6.2  | 3.0 |
|                           | Doc10 | 37.0                | 2.8      | 7.7  | 3.0 | 2.5 | 32.3                | 2.9      | 8.7  | 2.9  | 2.3 |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 68.3                | 4.2      | 9.5  | 3.2 | 2.1 | 73.1                | 4.4      | 9.1  | 6.1  | 2.6 |
| Yoruba                    |       |                     |          |      |     |     |                     |          |      |      |     |
| Aya-101                   | Sent  | 31.4                | 1.9      | 9.0  | 4.8 | 3.7 | 58.5                | 2.8      | 23.2 | 13.5 | 5.4 |
|                           | Doc10 | 22.3                | 2.3      | 7.1  | 2.5 | 2.1 | 55.7                | 3.2      | 13.8 | 8.0  | 4.4 |
| GPT-3.5                   | Sent  | 23.0                | 1.8      | 11.0 | 7.6 | 5.0 | 54.0                | 2.7      | 13.4 | 6.6  | 4.5 |
|                           | Doc10 | 12.7                | 1.2      | 6.9  | 2.9 | 2.6 | 53.0                | 3.8      | 9.8  | 4.5  | 2.4 |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 55.2                | 4.1      | 9.5  | 3.8 | 2.3 | 57.3                | 3.1      | 9.7  | 4.8  | 3.2 |
|                           | Doc10 | 30.5                | 2.7      | 7.1  | 2.4 | 1.8 | 28.5                | 3.2      | 8.4  | 2.4  | 1.8 |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 49.3                | 4.1      | 9.0  | 1.5 | 1.2 | 67.7                | 4.4      | 10.6 | 4.4  | 2.1 |
| Zulu                      |       |                     |          |      |     |     |                     |          |      |      |     |
| Aya-101                   | Sent  | 58.9                | 2.9      | 12.2 | 4.7 | 3.7 | 67.1                | 3.2      | 17.0 | 13.0 | 5.5 |
|                           | Doc10 | 55.0                | 3.4      | 10.4 | 3.7 | 2.4 | 68.4                | 3.6      | 15.4 | 10.0 | 4.6 |
| GPT-3.5                   | Sent  | 53.6                | 1.8      | 22.7 | 7.5 | 5.9 | 59.9                | 2.8      | 16.1 | 11.1 | 7.1 |
|                           | Doc10 | 35.1                | 1.6      | 19.1 | 4.7 | 3.7 | 61.2                | 3.9      | 11.3 | 5.8  | 2.8 |
| LLaMAX3-SFT <sub>1</sub>  | Sent  | 66.0                | 3.2      | 12.0 | 5.3 | 3.9 | 66.8                | 3.6      | 11.1 | 5.6  | 3.0 |
|                           | Doc10 | 31.7                | 2.4      | 7.7  | 3.1 | 2.4 | 29.8                | 3.2      | 8.0  | 3.1  | 2.2 |
| LLaMAX3-SFT <sub>10</sub> | Doc10 | 60.9                | 3.5      | 11.0 | 4.0 | 2.4 | 71.6                | 4.4      | 8.7  | 5.9  | 3.4 |

Table 22: Document-level evaluation in the *tech* domain, judged by GPT-4o. Compares sentence- vs. document-level outputs on Fluency (1–5 scale), Content Errors (CE), Lexical (LE), and Grammatical Cohesion Errors (GE). Best scores in bold.

---

**Prompt 1**

{system\_prompt}  
Translate the following {source\_language} text to {target\_language}:  
Provide only the translation.  
{source\_language} text: {{source\_sentence}}  
{target\_language} text:

**Prompt 2**

{system\_prompt}  
Translate the following {domain} text from {source\_language} to {target\_language}:  
Provide only the translation.  
{source\_language} document: {{source\_document}}  
{target\_language} document:

**Prompt 3**

{system\_prompt}  
Please provide the {target\_language} translation for the following {source\_language} text:{{source\_document}}  
Provide only the translation.

---

**Prompt 1**

{system\_prompt}  
Translate the following {source\_language} document to {target\_language}:  
Provide only the translation.  
{source\_language} document: {{source\_document}}  
{target\_language} document:

**Prompt 2**

{system\_prompt}  
Translate the following {domain} document from {source\_language} to {target\_language}:  
Provide only the translation.  
{source\_language} document: {{source\_document}}  
{target\_language} document:

**Prompt 3**

{system\_prompt}  
Please provide the {target\_language} translation for the following {source\_language} document:{{source\_document}}  
Provide only the translation.

---

Table 23: The task prompts used for evaluating LLMs are applied to both sentence-level and document-level translation tasks.

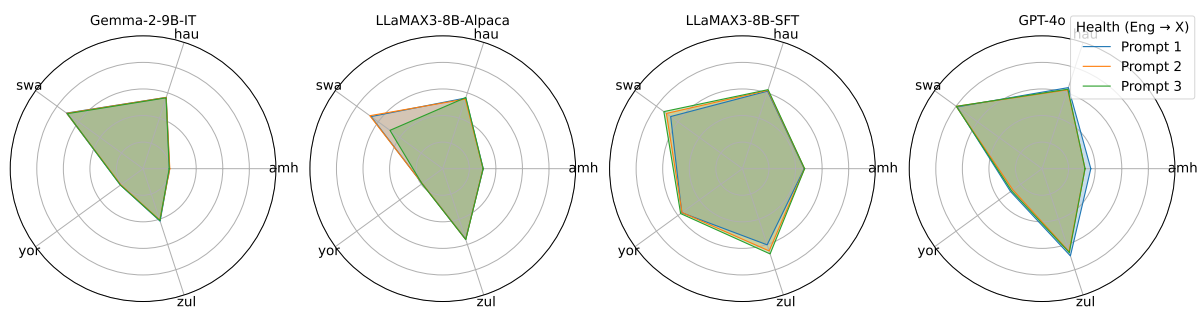


Figure 15: d-chrF scores for some LLMs for sentence-level translation using different prompts when translating into African languages

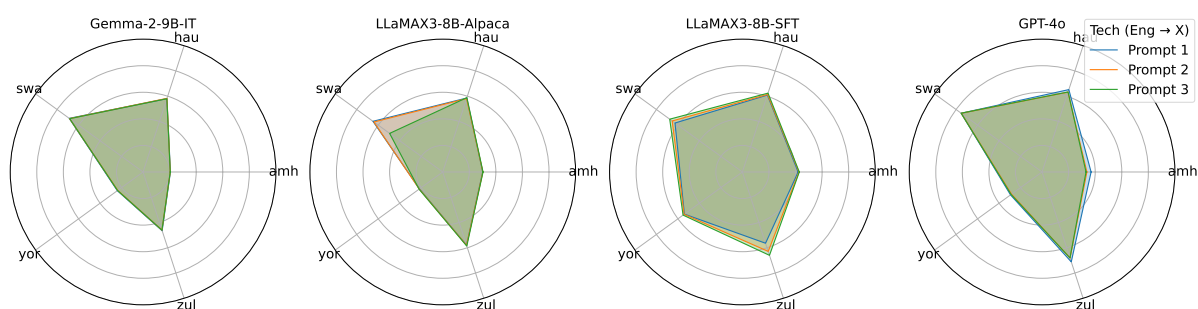


Figure 16: d-chrF scores for some LLMs for sentence-level translation using different prompts when translating into African languages

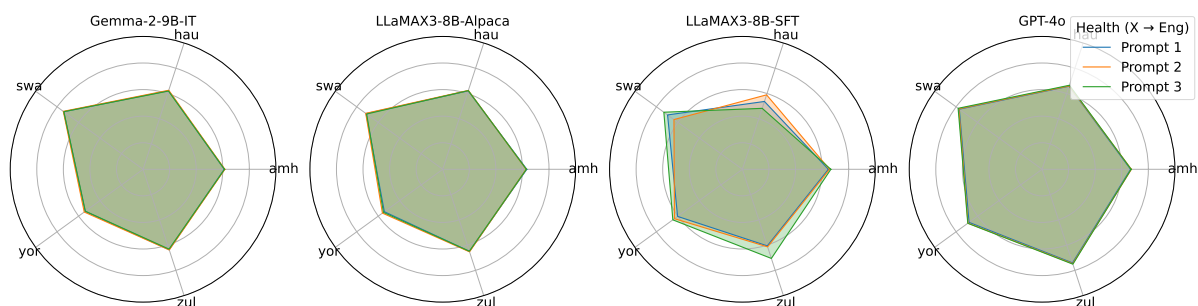


Figure 17: d-chrF scores for some LLMs for sentence-level translation using different prompts when translating into English

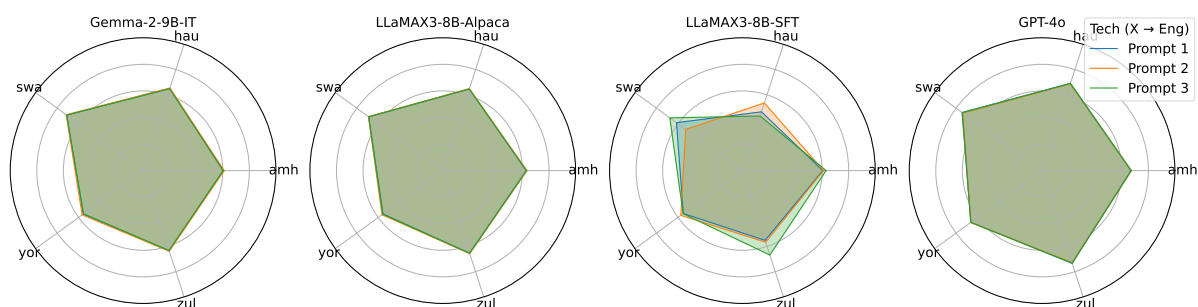


Figure 18: d-chrF scores for some LLMs for sentence-level translation using different prompts when translating into English

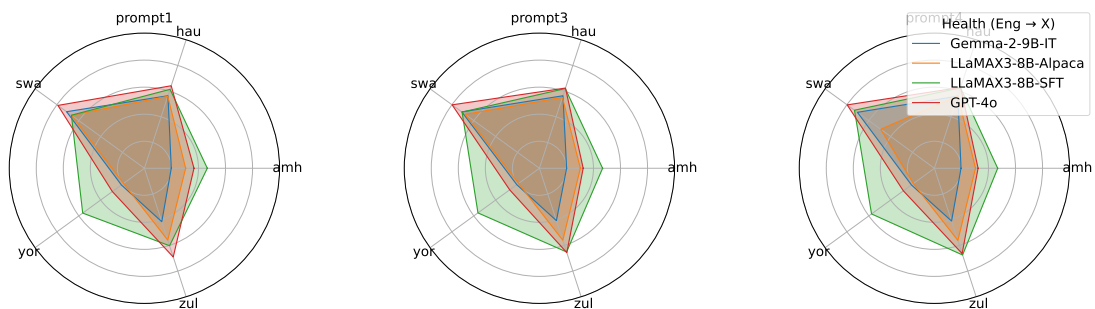


Figure 19: d-chrF scores for some LLMs for sentence-level translation using different prompts when translating into African languages

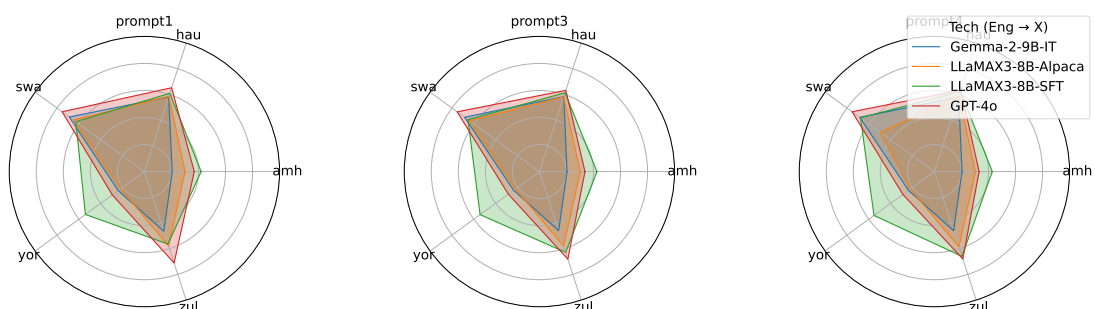


Figure 20: d-chrF scores for some LLMs for sentence-level translation using different prompts when translating into African languages

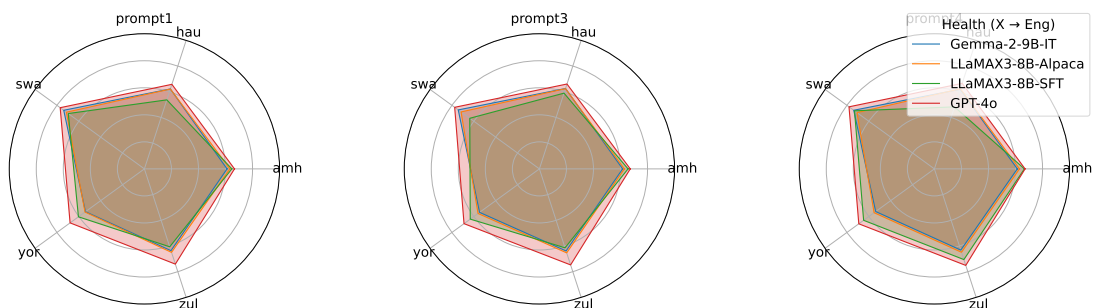


Figure 21: d-chrF scores for some LLMs for sentence-level translation using different prompts when translating into English

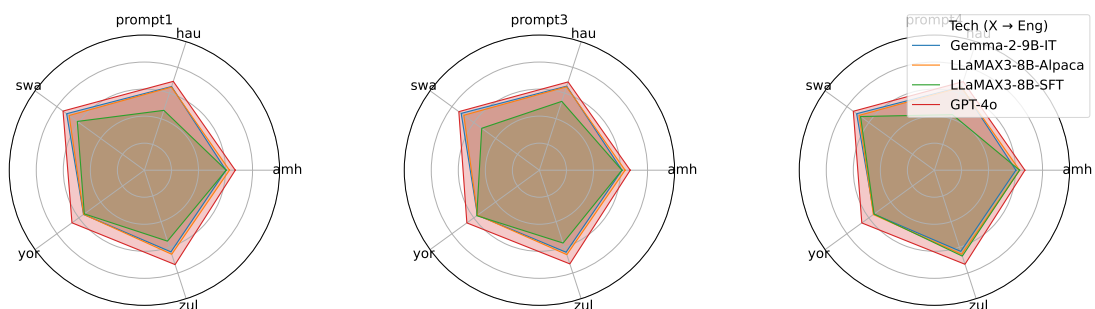


Figure 22: d-chrF scores for some LLMs for sentence-level translation using different prompts when translating into English