

Extending Automatic Machine Translation Evaluation to Book-Length Documents

Kuang-Da Wang^{1†}, Shuoyang Ding^{2†}, Chao-Han Huck Yang²,
Ping-Chun Hsieh¹, Wen-Chih Peng¹, Vitaly Lavruchin², Boris Ginsburg²

¹National Yang Ming Chiao Tung University, ²NVIDIA
gdwang.cs10@nycu.edu.tw shuoyangd@nvidia.com

Abstract

Despite Large Language Models (LLMs) demonstrating superior translation performance and long-context capabilities, evaluation methodologies remain constrained to sentence-level assessment due to dataset limitations, token number restrictions in metrics, and rigid sentence boundary requirements. We introduce SEGALe, an evaluation scheme that extends existing automatic metrics to long-document translation by treating documents as continuous text and applying sentence segmentation and alignment methods. Our approach enables previously unattainable document-level evaluation, handling translations of arbitrary length generated with document-level prompts while accounting for under-/over-translations and varied sentence boundaries. Experiments show our scheme significantly outperforms existing long-form document evaluation schemes, while being comparable to evaluations performed with groundtruth sentence alignments. Additionally, we apply our scheme to book-length texts and newly demonstrate that many open-weight LLMs fail to effectively translate documents at their reported maximum context lengths.

1 Introduction

Since the inception of Large Language Models (LLMs), the paradigm of machine translation (MT) has been shifting toward an LLM-based approach. In the WMT 2024 general translation shared task (Kocmi et al., 2024), LLM-based systems demonstrated strong performance, ultimately dominating submissions across all language pairs. Additionally, because of their long context windows, LLM-based systems may potentially be able to generate translations that better capture discourse-level phenomena and maintain coherence across longer spans of text. This development aligns with the long-standing

trend in MT research to move beyond sentence-level processing toward *paragraph-level* (Deutsch et al., 2023), *discourse-level* (Bawden et al., 2018), and *document-level* (Zhu et al., 2024) translation.

However, despite claims that modern LLMs can process inputs of up to 1M tokens (Yang et al., 2025), evaluations of LLM-based translations remain largely confined to *sentence-level* or *segment-level*, in that they are only prompted to translate one sentence or segment at a time. This forms a significant gap between what LLMs can generate and what existing metrics can evaluate. This limitation stems from the following two challenges:

1. Commonly used model-based evaluation metrics have relatively low maximum token length limitations (e.g., 512 tokens for COMET).
2. Current automatic evaluation metrics require adherence to pre-defined sentence boundaries. This forces evaluators to either use sentence-level prompting or add artificial boundaries.

In this paper, we propose SEGALe (SEGment, ALign, and Evaluate), a scheme that solves these challenges by extending existing automatic evaluation metrics to long documents. To summarize, our approach applies to arbitrarily long documents by using sentence segmenters and aligners to create appropriate sentence-level alignments. We treat those automatically aligned sentence pairs in two different ways. In the case where a valid sentence alignment is found, we apply the existing evaluation metric to the sentence pair. In the case where translation errors occur - either when content from the source text is missing in the translation (*under-translation*) or when there is hallucinated content that is not present in the source (*over-translation*) - we detect these as null alignments and assign a fixed penalty. At the end, along with metric scores, we also report the ratio of null alignments (NA ratio) to help track these over-translation and under-

[†] Equal contribution.

translation errors, which current MT evaluation metrics are having trouble detecting reliably.

Our experiments demonstrate that this scheme evaluates translations with comparable performance to existing sentence-level metrics when applied to cases with over-translation and under-translation. In addition, it handles cases where LLMs are liberal with sentence boundaries, which create many-to-one and one-to-many sentence alignments. Lastly, we newly demonstrate that we can successfully apply this scheme to evaluate translations of book-length texts, and reveal that many open-weight LLMs cannot translate documents of their reported context length, because the number of under- and over-translation errors rise sharply as the input length gets longer. Our code and artifacts are available at <https://github.com/nvmlabs/SEGALE>.

2 Related Work

2.1 Document-Level Translation

Document-level MT extends translation beyond isolated sentences by leveraging broader context for coherence. Existing approaches include simply concatenating adjacent sentences as a larger input to a standard MT model (Scherrer et al., 2019; Junczys-Dowmunt, 2019; Sun et al., 2022), as well as more advanced architectures that introduce context-specific modules: multi-encoder models encode previous sentences with separate encoders and hybrid attention mechanisms (Jean et al., 2017; Bawden et al., 2018; Voita et al., 2019; Miculicich et al., 2018; Maruf et al., 2019; Herold and Ney, 2023). Recent work has also focused on improving the quality of document-level translation by utilizing larger-scale document-level corpus (Thai et al., 2022; Al Ghussin et al., 2023; Post and Junczys-Dowmunt, 2023; Pal et al., 2024), as well as leveraging large language models (LLMs) (Karpinska and Iyyer, 2023; Wang et al., 2023).

Despite the progress, document-level translation still has a few limitations. First, a lot of work stick to a relatively small number of maximum input/output length. For example, Scherrer et al. (2019); Post and Junczys-Dowmunt (2023) both have maximum context length of 250 tokens, while Al Ghussin et al. (2023); Pal et al. (2024) have 512. Besides, some work (Junczys-Dowmunt, 2019; Post and Junczys-Dowmunt, 2023) introduce artificial sentence boundaries to the input, which provides native sentence segmentations for eval-

uation. This requires specialized training data or prompt, and there is no guarantee that the system will generate matching sentence boundaries as the input document.

2.2 Machine Translation Evaluation

Machine translation evaluation has shifted from string-based metrics (e.g., BLEU (Papineni et al., 2002), chrF (Popović, 2015)) to model-based metrics (e.g., COMET (Rei et al., 2020), MetricX (Juraska et al., 2024), GEMBA (Kocmi and Federmann, 2023)). Human evaluations like direct assessment (DA) and multi-dimensional quality metrics (MQM) played a crucial role in this paradigm shift by providing meta-evaluations and training data for model-based metrics.

Most model-based metrics are trained and evaluated on the segment-level. For example, COMET limits each input (source, target, and reference) to 512 tokens, while MetricX has a combined limit of 1,536 tokens across all inputs. In contrast, Qwen-2.5 (Yang et al., 2024), a recent open-source LLM, can handle input of 131,072 tokens and generate up to 8,192 tokens. Prior efforts have explored extending MT evaluation metrics beyond sentence-level. Vernikos et al. (2022) proposed adding prior sentences as context when training model-based metrics. Deutsch et al. (2023) trained metrics on paragraph-level data but found limited benefits. These studies are orthogonal to ours – they focus on building new model-based metrics with longer maximum input length, while we focus on applying existing metrics to long-form text.

Closest to the spirit of our work is mwerSegmenter (Matusov et al., 2005). It is a joint sentence segmentation and alignment scheme that has been the long-standing evaluation standard for unsegmented speech translation¹. The high-level idea is to jointly segment and align long-form model output by minimizing the word error rate (WER) between the segmented text and the already-segmented reference text. The assumption behind the idea is that perfectly segmented and aligned sentences are more likely to be translated well, and thus should have a low WER. Similar to mwerSegmenter, Wang et al. (2023) implemented a segmentation and alignment scheme based on Bleualign (Sennrich and Volk, 2010a), but there was no extensive discussion regarding the validity of the scheme.

¹Specifically, mwerSegmenter has been the evaluation standard for the IWSLT speech translation shared tasks (<https://iwslt.org/>)

Apart from that, [Raunak et al. \(2024\)](#) proposed to extend existing metrics based on running evaluations on aligned sliding windows over sentences in a document, but the algorithm is still limited to the sentence-level prompting paradigm.

A few recent investigations ([Salesky et al., 2023](#); [Sperber et al., 2024](#)) of mwerSegmenter in the context of long-form audio data raised concerns about the segmentation quality. The reader shall see that our results corroborate the concerns. Contemporary to our work, [Post and Hoang \(2025\)](#) resolved a few issues with the current mwerSegmenter implementation such as lack of standardized tokenization and empty translation hypotheses.

2.3 Long-Context LLM Evaluation

Recent progress in extending LLM architectures to handle longer contexts has spurred considerable research interest. Parallel to architectural advances, there has been growing attention toward systematically evaluating the capabilities of LLMs on long-context tasks. [Kamradt \(2023\)](#) developed an evaluation focusing on models’ abilities to retrieve deeply nested information. Similarly, [Bai et al. \(2024\)](#) introduced a long-context bilingual benchmark for assessing models’ comprehension and reasoning abilities, while [An et al. \(2024\)](#) shows that standardized evaluation criteria across multiple long-context scenarios are essential for comprehensive model assessment. Furthermore, [Hsieh et al. \(2024\)](#) highlights that there are discrepancies between theoretical capabilities and effective usable context lengths of contemporary LLMs.

Despite these advances, the evaluation methodologies have predominantly focused on general comprehension tasks rather than specialized applications like long-context machine translation. Existing metrics face limitations such as fixed maximum token lengths and rigid assumptions about sentence boundaries, which hinder effective evaluation of extensive, continuous texts, like books.

3 Preliminaries

Our ultimate goal is to find a way to evaluate the translation quality in the following scenarios:

- For translations of documents of arbitrary length generated with document-level prompts, while handling under-/over-translations and varied sentence boundaries
- For both reference-based and reference-free evaluations, thus enabling broader applications

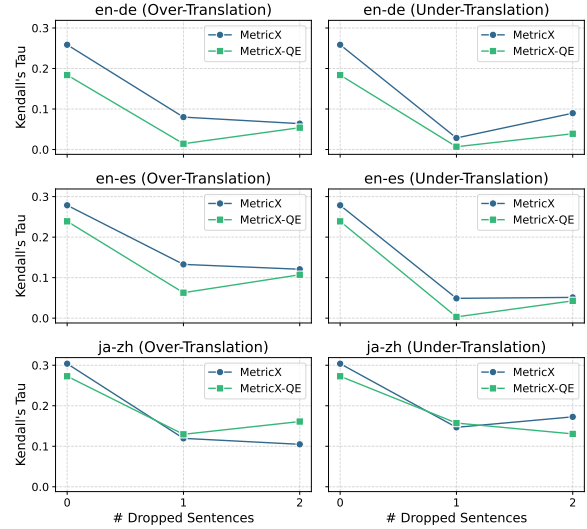


Figure 1: The Kendall’s τ correlations of MetricX-24 and MetricX-24-QE show limited sensitivity to over- and under-translation; sentences with more than three drops are insufficient to estimate correlations reliably.

like curating high-quality training data for LLMs (similar to [Finkelstein et al., 2024](#))

We start by justifying why extensions of existing metrics are required for document-level MT evaluation, rather than directly feeding concatenated sentence pairs from a document into existing metrics. A key limitation of the direct concatenation approach is that commonly used model-based evaluation metrics have relatively low maximum token length limits. Additionally, even for documents that are within the maximum token length limit, we show with a preliminary study in this section that directly applying state-of-the-art MT metrics to concatenated sentences is actually not able to reliably detect under- and over-translation errors.²

We evaluate the performance of MetricX-24 ([Juraska et al., 2024](#)) on such concatenation approach. To avoid going over the maximum token length limit of MetricX-24, we filter out cases where the concatenation of source and target inputs exceed 1024 tokens in length. We compute both MetricX-24 and its reference-free variant, MetricX-24-QE, across three language pairs: English–German (en-de), English–Spanish (en-es), and Japanese–Chinese (ja-zh). We use the dataset from the WMT 2024 Metrics Shared Task ([Freitag](#)

²The conclusion may seem different from a prior study [Deutsch et al. \(2023\)](#), but it’s actually not a direct contradiction, because the evaluation in [Deutsch et al. \(2023\)](#) focuses only on cases where one-to-one mapping between source, target, and reference exists.

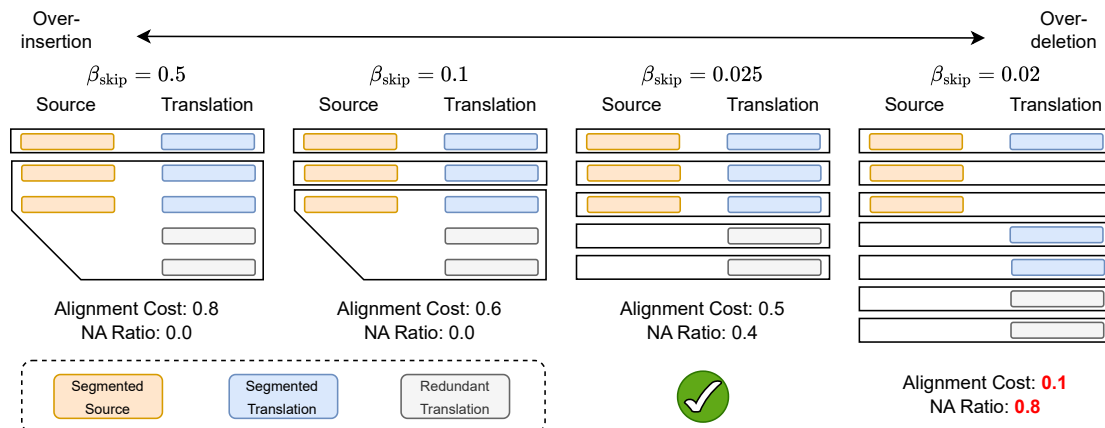


Figure 2: The effect of the skip cost (β_{skip}) on alignment behavior for over-translation. Higher skip costs increase the risk of over-insertion by allowing loose semantic matches to align, while lower skip costs enforce stricter alignment, leading to over-deletion. Over-deletion is indicated by spikes in the null alignment ratio (NA ratio) and low alignment costs, both shown in red.

et al., 2024), which contains: (1) source texts and reference translations in WMT 2024 General Translation Shared Task, (2) translations of the source texts from the system submissions, and (3) MQM human evaluations of the system translations.

To simulate translation errors, we manipulate the texts in two ways: for over-translation, we remove one or two sentences from both source and reference texts, while for under-translation, we remove sentences only from the system translations.³ We measure the performance of the metrics by computing Kendall’s τ correlations between the metric outputs and human evaluation scores.⁴

The results of this preliminary study (Figure 1) confirm that MetricX-24 and MetricX-24-QE have limited sensitivity to over- and under-translations, even within token length limits. This empirical evidence shows that the direct concatenation approach is inadequate for document-level MT evaluation. To apply existing MT metrics to document-level MT evaluation, we need an extension scheme that can properly handle these translation errors while working within the constraints of existing metrics.

4 Method

We now introduce SEGALe, our proposed extension scheme. SEGALe consists of three steps:

1. SEGment the document-level system translation into individual sentences

³We limit ourselves to removing at most two sentences since removing more would leave too few documents for reliable Kendall’s τ correlation calculations.

⁴For more details on meta-evaluation, see Section 5.1.

2. ALign the source and system translations
3. Evaluate the aligned sentence blocks with existing metrics, then average sentence-level scores to obtain a document-level score

4.1 Sentence Segmentation

We use off-the-shelf sentence segmentation models to segment documents into sentences. We experimented with ersatz (Wicks and Post, 2021) and spaCy (Honnibal et al., 2020).

4.2 Sentence Alignment

Given a sentence-segmented source document $\mathbf{S} = \{s_1, \dots, s_N\}$ and its translation $\mathbf{T} = \{t_1, \dots, t_M\}$, the goal of sentence alignment is to identify a minimal-cost alignment path that maps contiguous spans of source sentences to contiguous spans of target sentences. We use Vecalign (Thompson and Koehn, 2019) to perform such alignment, with a few changes detailed below.

Adaptive Penalty Search Ideally, all over- and under-translation errors will result in *null alignments*. In Vecalign, null alignments are modeled via a skip cost, which is parameterized by a percentile-based threshold β_{skip} . If the skip cost is set too high, many alignments are forced between unrelated sentence blocks – essentially reverting to the scenario in our preliminary experiment. Conversely, if the skip cost becomes too low, the aligner will start to assign null alignments even to semantically related sentence pairs. Figure 2 illustrates an example of over-translation and demonstrates how different β_{skip} values impact the alignment result.

Given that the optimal value of β_{skip} can vary depending on the severity of over- or under-translation in each individual document, we implement an adaptive search strategy to enhance robustness. We leverage the insight that over-deletion is often signaled by sudden spikes in the null alignment ratio and abnormally low average alignment costs. Since the optimal alignment typically occurs just before over-deletion sets in, our approach starts with a relatively high β_{skip} and progressively decreases it in small steps.

At each step, alignment quality is monitored using two heuristics to determine whether to terminate the search: (a) when the average alignment cost drops below a threshold, indicating excessive skipping, or (b) when the null alignment ratio exceeds a predefined limit at a step. Both patterns typically suggest that the skip penalty has become too lenient. In such cases, we revert to the previous step and treat it as the final alignment result. See Appendix B for implementation details and heuristic settings.

Building Better Text Embeddings Text embedding models are crucial for sentence alignment. We observe that sentence segmentation granularities vary across languages, which strains existing text embedding models. For example, suppose we have a long source sentence s that should align to smaller target sentences $\{t_1, \dots, t_M\}$. In Vecalign, the scoring function calculates similarity between s and all consecutive blocks of $\{t_1, \dots, t_M\}$. This is not what text embedding models are trained for, leading to suboptimal alignments.

Motivated by this observation, we build our own text embedding model that is specifically designed to handle the sentence segmentation granularities we described above. Our model is fine-tuned from BGE-M3 (Chen et al., 2024), which achieves high performance on bitext-mining task with only 568M parameters and without relying on instructions. The fine-tuning is performed on a synthetic dataset with query, positive and negative example triplets built from the News Commentary v18⁵ dataset (see more details in Appendix B.3), with the FlagEmbedding toolkit⁶. The readers shall see that our fine-tuned text embedding model outperforms both LASER (Artetxe and Schwenk, 2019) and BGE-M3 text embedding models in our experiments.

⁵<https://data.statmt.org/news-commentary/v18.1/>

⁶<https://github.com/FlagOpen/FlagEmbedding>

4.3 Evaluation via Existing Metrics

Once the target translation is segmented and aligned with the source document, we calculate the segment-level translation quality using existing metrics, then average the scores to obtain a document-level score. In the case when a non-one-to-one alignment is established, we concatenate the aligned sentence blocks and evaluate them with the metric. Two numbers are reported per document: the average segment score and ratio of null alignments over aligned sentence pairs ("NA ratio"). For each null alignment, we assign the worst possible score (0 for COMET, 25 for MetricX) and include it in the average calculation.

5 Experiments

5.1 Setup

Our experiments are aimed to test if the proposed evaluation scheme can achieve the two goals stated in Section 3 – in other words, whether it is (1) robust to all kinds of anomalies in system translations and (2) effective with both reference-based and reference-free metrics. Our data and metric setup reflect the above goals.

To establish meaningful comparisons, we compare with two baselines. One is calculating metric scores using the groundtruth sentence boundaries and alignments provided by the dataset ("Gold"), which serves as a performance upper bound.⁷ The other is calculating metric scores using the sentence boundaries and alignments derived by mwerSegmenter (Matusov et al., 2005).

Dataset We use the same dataset from preliminary experiments in Section 3. In all experiments, we merge existing sentence boundaries in the system translation to simulate system translations generated at document-level. We adhere to the same sentence boundaries on the source and reference sides during evaluation.

There are significant limitations if we only conduct meta-evaluation on the original test set, because the original test set is always guaranteed to have perfect sentence alignments (i.e., no null alignments). Hence, in addition to the original test set ("*original*" case), we create three synthetic test sets by introducing anomalies into the original

⁷As the reader shall see, there are times when our score is higher than the upper bound performance. This is likely caused by the sentence segmentation variations between the sentence segmenter and boundaries in the test set. It shouldn't be interpreted as our method being better in a meaningful way.

	COMET	COMET-QE	MetricX	MetricX-QE	NA Ratio ($ \Delta_{\text{Gold}} $)
Original					
Gold	0.3110	0.2800	0.3085	0.2683	0.0% (–)
SEGALE	0.3085	0.2768	0.3074	0.2630	0.7% (0.7%)
mwerSegmenter	0.2874	0.2547	0.2799	0.2360	0.0% (0.0%)
Over-Translate					
Gold	0.3848	0.3547	0.3598	0.3001	10.0% (–)
SEGALE	0.3532	0.3368	0.3499	0.3082	11.2% (1.2%)
mwerSegmenter	0.3506	0.3213	0.3251	0.2729	0.0% (10.0%)
Under-Translate					
Gold	0.3521	0.3174	0.3102	0.2832	10.0% (–)
SEGALE	0.3493	0.3295	0.3268	0.2852	5.8% (4.2%)
mwerSegmenter	0.2183	0.1903	0.1770	0.1441	2.1% (7.9%)
Flex-Boundary					
Gold	0.3096	0.2788	0.3066	0.2665	0.0% (–)
SEGALE	0.3067	0.2746	0.3033	0.2597	1.2% (1.2%)
mwerSegmenter	0.2611	0.2325	0.2524	0.2131	0.0% (0.0%)

Table 1: Correlation between document-level scores and human judgments under different evaluation settings. The first four columns are Kendall’s τ correlation coefficients ($\uparrow$), with the last column being the average NA ratio and the absolute difference from the groundtruth ($|\Delta_{\text{Gold}}|$). All numbers are averaged across three language pairs (en-de, en-es, ja-zh), and all reported numbers of our method are calculated with ersatz sentence segmenter and our fine-tuned BGE-M3 text embedding model.

test set, namely “*over-translate*”, “*under-translate*”, and “*flex-boundary*” cases. The first two cases are created by randomly removing 10% of the sentences from the source/reference sides and system translations, respectively. The last case is created by merging 10% of neighboring sentences in the source texts with GPT-4o, using carefully designed prompts to preserve the original semantic content and word choices. Since the system translations were generated from the source texts before merging, this introduces sentence boundary variations without modifying the system translations and their accompanying human judgments. For more details, please refer to Appendix A.

Underlying Metrics Our experiments cover both reference-based and reference-free (“QE”) variants of COMET⁸ and MetricX⁹.

Meta-Evaluation Similar to previous work and preliminary experiments, we use correlation between document-level scores and human judgments as the primary metric. Although both system translation and human judgments are performed at segment-level, previous work (Deutsch et al., 2023) has shown that MQM annotations are done with context of surrounding sentences, and sen-

tences appear in document order. Hence, they are a good proxy for document translation quality. For cases with introduced null alignments, we assign 25 as the human-annotated MQM score for each null alignment, which is then converted into z-score in accordance with each human annotator’s scoring distribution. Like previous work, we average the segment-level human judgment scores as the document-level scores.

We also report NA ratio for each method as the auxiliary metric. Ideally, we would like to achieve the same NA ratio as the groundtruth ($|\Delta_{\text{Gold}}| = 0$), but the reader should note that perfect NA ratio on its own doesn’t necessarily imply a good evaluation scheme.¹⁰ The correlation with human judgments should still be treated as the ultimate meta-evaluation metric.

5.2 Results

Main Results Table 1 shows a concise version of our main results (averaged across three language pairs). In terms of correlation with human judgments, SEGALE achieves near-ideal performance, outperforming mwerSegmenter while maintaining comparable correlation with human judgments to Gold. The gap is especially significant

⁸Reference-based: [wmt22-comet-da](#). Reference-free: [wmt22-cometwiki-da](#)

⁹[metricx-24-hybrid-large-v2p6](#)

¹⁰For example, in the “original” case, a very bad hypothetical evaluation scheme that aligns a random segment to the source can achieve the same 0% NA ratio as groundtruth.

	COMET	MetricX	NA Ratio ($ \Delta_{\text{Gold}} $)
Original			
ersatz+LASER	0.3072	0.3051	1.2% (1.2%)
ersatz+BGE-m3	0.3037	0.3055	1.3% (1.3%)
ersatz+BGE-m3-ft	0.3085	0.3074	0.7% (0.7%)
spacy+BGE-m3-ft	0.3028	0.3049	1.4% (1.4%)
Over-Translate			
ersatz+LASER	0.3331	0.3386	8.6% (1.4%)
ersatz+BGE-m3	0.3233	0.3252	9.8% (0.2%)
ersatz+BGE-m3-ft	0.3532	0.3499	11.2% (1.2%)
spacy+BGE-m3-ft	0.3554	0.3505	10.1% (0.1%)
Under-Translate			
ersatz+LASER	0.3405	0.3176	6.3% (3.7%)
ersatz+BGE-m3	0.3381	0.3176	4.2% (5.8%)
ersatz+BGE-m3-ft	0.3493	0.3268	5.8% (4.2%)
spacy+BGE-m3-ft	0.3481	0.3258	6.0% (4.0%)
Flex-Boundary			
ersatz+LASER	0.3041	0.3017	1.9% (1.9%)
ersatz+BGE-m3	0.3001	0.3012	2.2% (2.2%)
ersatz+BGE-m3-ft	0.3067	0.3033	1.2% (1.2%)
spacy+BGE-m3-ft	0.3066	0.3033	1.5% (1.5%)

Table 2: Ablation study on different sentence embeddings and segmenters. Numbers are calculated similarly to Table 1 but only include reference-based metrics due to space limits. Boldface numbers indicate the highest correlation for the first two columns, and the NA ratio with the smallest $|\Delta_{\text{Gold}}|$ for the last column.

for the "under-translate" case, where mwerSegmenter suffers significant performance drops while SEGALe remains robust. For a detailed version with per-language-pair breakdown, please refer to Appendix C. The readers shall see that the trend shown in Table 1 is generally consistent across all language pairs.

For NA ratio, we can observe that mwerSegmenter perfectly matches the groundtruth in two settings ("original" and "flex-boundary"). However, upon closer inspection, we conclude that this is not because mwerSegmenter can detect over-/under-translation errors more accurately, but rather because mwerSegmenter was not designed to account for certain translation anomalies. For example, one crucial problem with mwerSegmenter is that it does not have an explicit mechanism to handle null alignments at the sentence level, nor does it have the semantic features that allows it to flag semantically irrelevant sentence pairs. Hence, in the event of a system generating a hallucinated sentence, mwerSegmenter simply merges it to an arbitrary neighboring sentence, which can cause the underlying metric to miss over-translation errors (shown in Section 3). On the other hand, the mwerSegmenter version we experimented with does not allow for empty translation hypotheses, which can cause chunks of other well-translated sentences

to be chopped up as an attempt to minimize deletion errors. Obviously, this also greatly interferes with the scores generated by the underlying metric. As for SEGALe, while it is also not perfect with inferring null alignments, the deviations from the groundtruth are more modest and are less likely to lead to catastrophic performance drops like mwerSegmenter in the "under-translate" case.

While mwerSegmenter doesn't provide sufficient logging that would allow us to pinpoint the exact reason for its performance drop, looking at the segmented text, it is clear that mwerSegmenter's algorithm struggles to handle synonym substitution and paraphrasing, particularly when translation quality is poor and significant structural changes are present. This exemplifies the limitation of using WER, instead of a semantic-based score, as the scoring function for segmentation and alignment.

The readers should also note that while the performance trend remains consistent between reference-based and reference-free metrics in our experiments, mwerSegmenter does require a reference translation as input, which in practice limits its applicability in reference-free evaluation settings. SEGALe, on the other hand, directly performs cross-lingual sentence alignment without requiring a reference translation.

Impact of Sentence Embedding Table 2 shows a comparison of SEGALe with different sentence embeddings. It can be observed that our fine-tuned BGE-M3 embedding consistently outperforms LASER and the original BGE-M3 embedding in all data configurations. The benefits of fine-tuning are especially significant for the "over-translate" and "under-translate" cases, which shows that our fine-tuning process successfully specializes the embedding model for capturing over- and under-translation errors during the coarse-to-fine search process of the sentence alignment step.

Impact of Sentence Segmenter Most of the numbers reported in this paper are calculated with the ersatz sentence segmenter. We also experimented with spaCy as the sentence segmenter as another ablation study, also shown in Table 2. Although we observed that spaCy tends to segment sentences into smaller units than ersatz (which does not align well with the long segments in WMT test sets), the impact on the performance seems to be minimal as we did not observe a consistent trend.

Model ID	Reference
utter-project/EuroLLM-9B-Instruct	Martins et al. (2024)
Qwen/Qwen2.5-14B-Instruct	Yang et al. (2024)
Qwen/Qwen2.5-72B-Instruct	Yang et al. (2024)
meta-llama/Llama-3.1-8B-Instruct	Dubey et al. (2024)
meta-llama/Llama-3.1-70B-Instruct	Dubey et al. (2024)
CohereForAI/aya-expense-8b	Dang et al. (2024)
CohereForAI/aya-expense-32b	Dang et al. (2024)

Table 3: List of LLMs evaluated for book-length translation capability.

6 Evaluation of Book-Length Translation Capability of Existing LLMs

Now that we have validated our evaluation method on WMT 2024 metrics shared task dataset, we briefly demonstrate that our method can be applied to assess the book-length translation capability of existing LLMs by conducting a similar experiment as Wang et al. (2024a). Our dataset comes from the Chinese-English (zh-en) section of the WMT 2024 Discourse-Level Literary Translation task (Wang et al., 2024b). Because the test set only contains book chapters instead of full books, we randomly pick a book with ID 2-xz1tq from the training split of the dataset and use it as our test set.

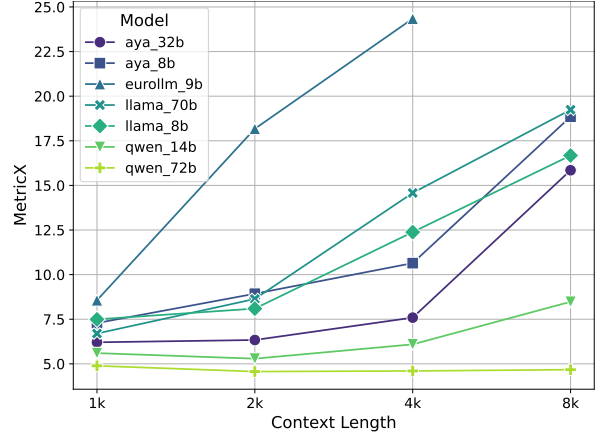
The LLMs evaluated are listed in Table 3. For simplicity, we adopted the same prompt and translation extraction procedure as used in WMT 2024 general machine translation shared task¹¹ for all the LLMs. Since current LLMs are constrained by maximum generation lengths and cannot translate the entire book in a single pass, we divide the content into segments of 1k, 2k, 4k, and 8k tokens, using tokenization from the tiktoken tokenizer^{12,13}. Most of these models have a maximum generation length of 8k tokens, except for EuroLLM, which is capped at 4k.

Figure 3 shows the translation quality and NA ratio of the LLMs at different context lengths. Most models exhibit a sharp degradation in translation quality at context length of 4k or 8k. For example, at 4k context length, EuroLLM refuses to translate as instructed, but rather resorts to summarizing the input document in the target language. Comparing with the trend in NA ratio, it is also clear that under-

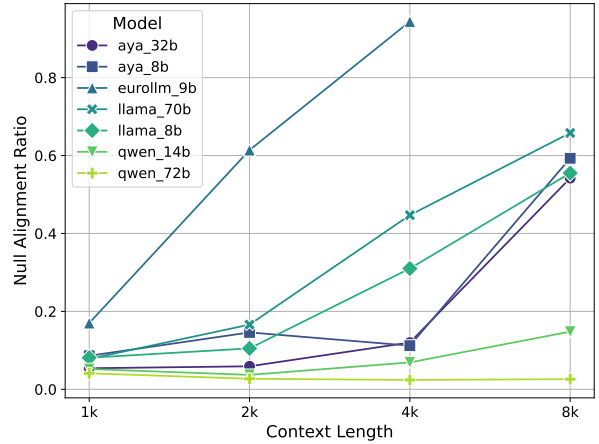
¹¹<https://github.com/wmt-conference/wmt-collect-translations/blob/704b3825730f93a3ee3a0fda44af9414937b6d5a/tools/prompts.py#L23>

¹²<https://github.com/openai/tiktoken>

¹³Note that tiktoken tokenizer is only used to count the number of tokens in the segments and LLM-specific tokenization will still be applied during inference.



(a) Translation Quality (SEGALE-MetricX ↓)



(b) Null Alignment Ratio (↓)

Figure 3: LLM Translation Performance at Different Context Lengths

translation/over-translation errors played a significant role in such degradation. The only noteworthy exception is Qwen2.5-72B-Instruct, which shows a much more stable performance across different context lengths. In fact, with the increasing context length up to 4k, there is a small improvement in translation quality, which shows its ability to utilize long-context information to obtain better translations.

This benchmark shows a significant gap between claimed max generation length and the actual capability of LLMs to translate long-context documents. As future work keeps improving LLM’s long-context processing capabilities, we call on the community to adopt this evaluation practice to gain better insights into such capabilities in downstream applications such as machine translation.

7 Conclusion

We propose SEGALE, a novel extension scheme that enables evaluation metrics to evaluate unseg-

mented document-level translations of arbitrary lengths. SEGAL works with any existing evaluation metric and eliminates the dependency on sentence-level prompting and pre-segmented reference translations. Experimental results show that our extension scheme achieves strong correlation with human judgments while demonstrating robustness to common LLM translation anomalies like over- and under-translation. Through this work, we aim to facilitate machine translation research in its ongoing shift away from sentence-level paradigm, while also offering new perspectives for evaluating LLMs’ long-context generation capabilities.

Limitations

Underlying Metrics

We acknowledge that an LLM-based metric based on long-context, open-source LLMs is a promising (and probably the eventual) solution to the problem of long-context MT evaluation. While previous work has shown that LLM-based metrics such as GEMBA (Kocmi and Federmann, 2023) or AutoMQM (Fernandes et al., 2023) can perform on-par as state-of-the-art BERT-based metrics such as COMET, they have to rely on GPT-4 or GPT-4o as the underlying LLM and are currently prohibitively expensive for MT evaluation of book-length documents. We plan to explore the potential of the open-source LLM counterparts of these metrics for long-context MT evaluation in a separate study.

All underlying metrics explored in our experiments are sentence-level and do not explicitly incorporate discourse-level information. There are existing methods that extend sentence-level metrics with surrounding context, such as Vernikos et al. (2022) and Raunak et al. (2024). We did not choose to go down that path because it will further complicate our experimental setup, and adding extra context to the underlying metrics is an orthogonal problem to our proposed extension scheme. It is worth noting that the modular nature of our proposed extension scheme makes it fully compatible with these extensions, and we expect that such combination will yield better results.

The current design of our adaptive penalty search strategy is built upon the observation that existing metrics have limited sensitivity to over- and under-translation errors (Section 3). Right now, such errors are penalized by a fixed penalty over null alignments, which could appear crude as the underlying metrics continue to improve. In such

case, the hyperparameters of the adaptive penalty search strategy may need to be revisited to delegate more of these errors to the underlying metrics. The downside of this approach is that the NA ratio may become a less reliable indicator of over- and under-translation errors.

Meta-Evaluation and Experimental Setup

We did not include a direct comparison to other sentence alignment methods like Bleualign (Sennrich and Volk, 2010b). We focus on Vecalign, whose superior performance to Bleualign is well-documented (Thompson and Koehn, 2019; Steingrimsson et al., 2023). Besides, the adaptive penalty search mechanism is an enhancement designed specifically for Vecalign and is crucial for our approach’s effectiveness. This mechanism is incompatible with other aligners, which lack the fine-grained control over null alignments that is required by adaptive penalty search.

Another aspect worth considering is that the human judgments used in our evaluation do not explicitly instruct the annotators to consider discourse-level information. Hence, whether those extra information will show benefits in meta-evaluations based on current human judgments remains unclear. Since WMT 2025¹⁴ has recently shifted to include multi-line, long-form texts in its human evaluation setup, exploring discourse-level effects would be more meaningful using this updated dataset.

Our synthetic evaluation setup approximates real-world LLM translations but does not capture behaviors such as discourse-level content reordering. We plan to conduct more extensive real-world testing of our method and compare it against a broader range of document-level MT evaluation paradigms in future work.

Fine-tuned Embedding Model

The fine-tuned BGE-m3 embedding model used in our proposed extension scheme is largely a proof-of-concept, due to the fact that it is trained on a small dataset, covering only languages of our interest. We believe that a specialized text embedding model for sentence alignment is not only useful for our proposed extension scheme, but also for its more traditional use cases, such as curation of web-crawled data. In the future, we plan to explore extending the volume of the training data and sup-

¹⁴<https://www2.statmt.org/wmt25/translation-task.html>

ported languages to improve the usability of our proposed extension scheme.

Acknowledgments

We thank Brian Thompson for discussions regarding Vecalign, Rachel Wicks for discussions regarding ersatz, Miguel Romero Calvo for discussions regarding contrastive fine-tuning of embedding models, and Oleksii Hrinchuk for discussions regarding mwerSegmenter. We also thank Simeng Sun, Elizabeth Salesky, and anonymous reviewers for their valuable comments and suggestions on earlier drafts of this paper. This project has been supported by NVIDIA Taiwan Research & Dev Center (TRDC) Grant.

References

- Yusser Al Ghussin, Jingyi Zhang, and Josef van Genabith. 2023. [Exploring paracrawl for document-level neural machine translation](#). In *Proceedings of the 17th Conference of the European Chapter of the Association for Computational Linguistics*, pages 1304–1310, Dubrovnik, Croatia. Association for Computational Linguistics.
- Chenxin An, Shansan Gong, Ming Zhong, Xingjian Zhao, Mukai Li, Jun Zhang, Lingpeng Kong, and Xipeng Qiu. 2024. [L-eval: Instituting standardized evaluation for long context language models](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 14388–14411, Bangkok, Thailand. Association for Computational Linguistics.
- Mikel Artetxe and Holger Schwenk. 2019. [Massively multilingual sentence embeddings for zero-shot cross-lingual transfer and beyond](#). *Transactions of the Association for Computational Linguistics*, 7:597–610.
- Yushi Bai, Xin Lv, Jiajie Zhang, Hongchang Lyu, Jiankai Tang, Zhidian Huang, Zhengxiao Du, Xiao Liu, Aohan Zeng, Lei Hou, Yuxiao Dong, Jie Tang, and Juanzi Li. 2024. [LongBench: A bilingual, multi-task benchmark for long context understanding](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 3119–3137, Bangkok, Thailand. Association for Computational Linguistics.
- 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.
- Jianlv Chen, Shitao Xiao, Peitian Zhang, Kun Luo, Defu Lian, and Zheng Liu. 2024. [Bge m3-embedding: Multi-lingual, multi-functionality, multi-granularity text embeddings through self-knowledge distillation](#). *Preprint*, arXiv:2402.03216.
- John Dang, Shivalika Singh, Daniel D’souza, Arash Ahmadian, Alejandro Salamanca, Madeline Smith, Aidan Peppin, Sungjin Hong, Manoj Govindassamy, Terrence Zhao, Sandra Kublik, Meor Amer, Viraat Aryabumi, Jon Ander Campos, Yi Chern Tan, Tom Kocmi, Florian Strub, Nathan Grinsztajn, Yannis Flet-Berliac, and 26 others. 2024. [Aya expand: Combining research breakthroughs for a new multilingual frontier](#). *CoRR*, abs/2412.04261.
- Daniel Deutsch, Juraj Juraska, Mara Finkelstein, and Markus Freitag. 2023. [Training and meta-evaluating machine translation evaluation metrics at the paragraph level](#). In *Proceedings of the Eighth Conference on Machine Translation*, pages 996–1013, Singapore. Association for Computational Linguistics.
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, and 82 others. 2024. [The llama 3 herd of models](#). *CoRR*, abs/2407.21783.
- Patrick Fernandes, Daniel Deutsch, Mara Finkelstein, Parker Riley, André Martins, Graham Neubig, Ankush Garg, Jonathan Clark, Markus Freitag, and Orhan Firat. 2023. [The devil is in the errors: Leveraging large language models for fine-grained machine translation evaluation](#). In *Proceedings of the Eighth Conference on Machine Translation*, pages 1066–1083, Singapore. Association for Computational Linguistics.
- Mara Finkelstein, David Vilar, and Markus Freitag. 2024. [Introducing the NewsPaLM MBR and QE dataset: LLM-generated high-quality parallel data outperforms traditional web-crawled data](#). In *Proceedings of the Ninth Conference on Machine Translation*, pages 1355–1372, Miami, Florida, USA. Association for Computational Linguistics.
- Markus Freitag, Nitika Mathur, Daniel Deutsch, Chiklu Lo, Eleftherios Avramidis, Ricardo Rei, Brian Thompson, Frederic Blain, Tom Kocmi, Jiayi Wang, David Ifeoluwa Adelani, Marianna Buchicchio, Chrysoula Zerva, and Alon Lavie. 2024. [Are LLMs breaking MT metrics? results of the WMT24 metrics shared task](#). In *Proceedings of the Ninth Conference on Machine Translation*, pages 47–81, Miami, Florida, USA. Association for Computational Linguistics.
- 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.
- Matthew Honnibal, Ines Montani, Sofie Van Landeghem, and Adriane Boyd. 2020. [spaCy: Industrial-strength Natural Language Processing in Python](#).
- Cheng-Ping Hsieh, Simeng Sun, Samuel Kriman, Shantanu Acharya, Dima Rekesh, Fei Jia, Yang Zhang, and Boris Ginsburg. 2024. [RULER: what’s the real context size of your long-context language models?](#) *CoRR*, abs/2404.06654.
- Sébastien Jean, Stanislas Lauly, Orhan Firat, and Kyunghyun Cho. 2017. [Does neural machine translation benefit from larger context?](#) *CoRR*, abs/1704.05135.
- Marcin Junczys-Dowmunt. 2019. [Microsoft translator at WMT 2019: Towards large-scale document-level neural machine translation](#). In *Proceedings of the Fourth Conference on Machine Translation (Volume 2: Shared Task Papers, Day 1)*, pages 225–233, Florence, Italy. Association for Computational Linguistics.
- Juraj Juraska, Daniel Deutsch, Mara Finkelstein, and Markus Freitag. 2024. [MetricX-24: The Google submission to the WMT 2024 metrics shared task](#). In *Proceedings of the Ninth Conference on Machine Translation*, pages 492–504, Miami, Florida, USA. Association for Computational Linguistics.
- Gregory Kamradt. 2023. [Needle in a haystack - pressure testing llms](#). *Github*.
- Marzena Karpinska and Mohit Iyyer. 2023. [Large language models effectively leverage document-level context for literary translation, but critical errors persist](#). In *Proceedings of the Eighth Conference on Machine Translation*, pages 419–451, Singapore. 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, Marzena Karpinska, Philipp Koehn, Benjamin Marie, Christof Monz, Kenton Murray, Masaaki Nagata, Martin Popel, Maja Popović, and 3 others. 2024. [Findings of the WMT24 general machine translation shared task: The LLM era is here but MT is not solved yet](#). In *Proceedings of the Ninth Conference on Machine Translation*, pages 1–46, Miami, Florida, USA. Association for Computational Linguistics.
- Tom Kocmi and Christian Federmann. 2023. [Large language models are state-of-the-art evaluators of translation quality](#). In *Proceedings of the 24th Annual Conference of the European Association for Machine Translation*, pages 193–203, Tampere, Finland. European Association for Machine Translation.
- Pedro Henrique Martins, Patrick Fernandes, João Alves, Nuno Miguel Guerreiro, Ricardo Rei, Duarte M. Alves, José Pombal, M. Amin Farajian, Manuel Faysse, Mateusz Klimaszewski, Pierre Colombo, Barry Haddow, José G. C. de Souza, Alexandra Birch, and André F. T. Martins. 2024. [Eurollm: Multilingual language models for europe](#). *CoRR*, abs/2409.16235.
- Sameen Maruf, André F. T. Martins, and Gholamreza Haffari. 2019. [Selective attention for context-aware neural machine translation](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 3092–3102, Minneapolis, Minnesota. Association for Computational Linguistics.
- Evgeny Matusov, Gregor Leusch, Oliver Bender, and Hermann Ney. 2005. [Evaluating machine translation output with automatic sentence segmentation](#). In *2005 International Workshop on Spoken Language Translation, IWSLT 2005, Pittsburgh, PA, USA, October 24-25, 2005*, pages 138–144. ISCA.
- 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.
- Proyag Pal, Alexandra Birch, and Kenneth Heafield. 2024. [Document-level machine translation with large-scale public parallel corpora](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 13185–13197, Bangkok, Thailand. Association for Computational Linguistics.
- 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 and Hieu Hoang. 2025. [Effects of automatic alignment on speech translation metrics](#). In *Proceedings of the 22nd International Conference on Spoken Language Translation (IWSLT 2025)*, pages 84–92, Vienna, Austria (in-person and online). Association for Computational Linguistics.
- Matt Post and Marcin Junczys-Dowmunt. 2023. [Escaping the sentence-level paradigm in machine translation](#). *CoRR*, abs/2304.12959.

- Vikas Raunak, Tom Kocmi, and Matt Post. 2024. [SLIDE: Reference-free evaluation for machine translation using a sliding document window](#). In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 2: Short Papers)*, pages 205–211, Mexico City, Mexico. Association for Computational Linguistics.
- Ricardo Rei, Craig Stewart, Ana C Farinha, and Alon Lavie. 2020. [COMET: A neural framework for MT evaluation](#). In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 2685–2702, Online. Association for Computational Linguistics.
- Elizabeth Salesky, Kareem Darwish, Mohamed Al-Badrashiny, Mona Diab, and Jan Niehues. 2023. [Evaluating multilingual speech translation under realistic conditions with resegmentation and terminology](#). In *Proceedings of the 20th International Conference on Spoken Language Translation (IWSLT 2023)*, pages 62–78, Toronto, Canada (in-person and online). Association for Computational Linguistics.
- Yves Scherrer, Jörg Tiedemann, and Sharid Loáigiga. 2019. [Analysing concatenation approaches to document-level NMT in two different domains](#). In *Proceedings of the Fourth Workshop on Discourse in Machine Translation (DiscoMT 2019)*, pages 51–61, Hong Kong, China. Association for Computational Linguistics.
- Thibault Sellam, Dipanjan Das, and Ankur P Parikh. 2020. [Bleurt: Learning robust metrics for text generation](#). In *Proceedings of ACL*.
- Rico Sennrich and Martin Volk. 2010a. [Mt-based sentence alignment for ocr-generated parallel texts](#). In *Proceedings of the 9th Conference of the Association for Machine Translation in the Americas: Research Papers, AMTA 2010, Denver, Colorado, USA, October 31 - November 4, 2010*. Association for Machine Translation in the Americas.
- Rico Sennrich and Martin Volk. 2010b. [MT-based sentence alignment for OCR-generated parallel texts](#). In *Proceedings of the 9th Conference of the Association for Machine Translation in the Americas: Research Papers*, Denver, Colorado, USA. Association for Machine Translation in the Americas.
- Matthias Sperber, Ondřej Bojar, Barry Haddow, Dávid Javorský, Xutai Ma, Matteo Negri, Jan Niehues, Peter Polák, Elizabeth Salesky, Katsuhito Sudoh, and Marco Turchi. 2024. [Evaluating the IWSLT2023 speech translation tasks: Human annotations, automatic metrics, and segmentation](#). In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 6484–6495, Torino, Italia. ELRA and ICCL.
- Steinthor Steingrímsson, Hrafn Loftsson, and Andy Way. 2023. [SentAlign: Accurate and scalable sentence alignment](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 256–263, Singapore. Association for Computational Linguistics.
- Zewei Sun, Mingxuan Wang, Hao Zhou, Chengqi Zhao, Shujian Huang, Jiajun Chen, and Lei Li. 2022. [Rethinking document-level neural machine translation](#). In *Findings of the Association for Computational Linguistics: ACL 2022*, pages 3537–3548, Dublin, Ireland. Association for Computational Linguistics.
- Katherine Thai, Marzena Karpinska, Kalpesh Krishna, Bill Ray, Moira Inghilleri, John Wieting, and Mohit Iyyer. 2022. [Exploring document-level literary machine translation with parallel paragraphs from world literature](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 9882–9902, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Brian Thompson and Philipp Koehn. 2019. [Vecalign: Improved sentence alignment in linear time and space](#). In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 1342–1348, Hong Kong, China. Association for Computational Linguistics.
- Aäron van den Oord, Yazhe Li, and Oriol Vinyals. 2018. [Representation learning with contrastive predictive coding](#). *CoRR*, abs/1807.03748.
- Giorgos Vernikos, Brian Thompson, Prashant Mathur, and Marcello Federico. 2022. [Embarrassingly easy document-level MT metrics: How to convert any pretrained metric into a document-level metric](#). In *Proceedings of the Seventh Conference on Machine Translation (WMT)*, pages 118–128, Abu Dhabi, United Arab Emirates (Hybrid). Association for Computational Linguistics.
- 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.
- Longyue Wang, Zefeng Du, Wenxiang Jiao, Chenyang Lyu, Jianhui Pang, Leyang Cui, Kaiqiang Song, Derek Wong, Shuming Shi, and Zhaopeng Tu. 2024a. [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. Association for Computational Linguistics.
- Longyue Wang, Siyou Liu, Chenyang Lyu, Wenxiang Jiao, Xing Wang, Jiahao Xu, Zhaopeng Tu, Yan Gu, Weiyu Chen, Minghao Wu, Liting Zhou, Philipp Koehn, Andy Way, and Yulin Yuan. 2024b. [Findings of the WMT 2024 shared task on discourse-level](#)

- literary translation. In *Proceedings of the Ninth Conference on Machine Translation*, pages 699–700, Miami, Florida, USA. 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.
- Rachel Wicks and Matt Post. 2021. [A unified approach to sentence segmentation of punctuated text in many languages](#). 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 3995–4007, Online. Association for Computational Linguistics.
- An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang, and 22 others. 2024. [Qwen2.5 technical report](#). *CoRR*, abs/2412.15115.
- An Yang, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoyan Huang, Jiandong Jiang, Jianhong Tu, Jianwei Zhang, Jingren Zhou, Junyang Lin, Kai Dang, Kexin Yang, Le Yu, Mei Li, Minmin Sun, Qin Zhu, Rui Men, Tao He, and 9 others. 2025. [Qwen2.5-1m technical report](#). *CoRR*, abs/2501.15383.
- Dawei Zhu, Sony Trenous, Xiaoyu Shen, Dietrich Klakow, Bill Byrne, and Eva Hasler. 2024. [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.

A Synthetic Test Data Creation

To evaluate the robustness of our evaluation framework, we construct synthetic test data simulating three common alignment challenges: under-/over-translation and varied sentence boundaries.

A.1 Synthetic Under- and Over-Translations

We simulate under- and over-translation scenarios by randomly removing 10% of the segments from either the source/reference sides or the system translations, respectively. To maintain meaningful context, we avoid sampling from documents that contain only a single segment.

A.2 Synthetic Sentence Boundary Variation

To simulate sentence boundary variation, we generate synthetic test data by merging 10% of adjacent segments in the source side. This process is conducted using GPT-4o, with the following constraints:

- Only segments that are neither the first nor the last in a document are eligible for merging.
- For each eligible candidate, an attempt is made to merge it with its subsequent segment.
- Merging is only accepted if the semantic difference between the merged and original segments is minimal. We use BLEURT (Sellam et al., 2020) to assess the semantic similarity, accepting only merges with a BLEURT score greater than 0.85. If a candidate pair fails this criterion, another eligible pair is sampled.

GPT-4o is instructed to merge adjacent segments into a single, fluent sentence without changing the original meaning, vocabulary, or the order of information. Figure 4 shows the prompt template used to guide the model.

Figure 5 provides an example of the merging process before and after GPT-4o rewriting. The sentences initially presented separately are transformed into a single sentence using appropriate transitional phrases.

This procedure enables us to test our evaluation method’s robustness in conditions reflecting realistic variations in sentence boundary alignments while ensuring that human annotations remain valid and can be directly reused without recalibration.

B Implementation Details

In this section, we provide detailed implementation information on extending existing evaluation metrics to book-length documents. Our proposed approach is designed as a flexible framework where the underlying models can be readily substituted. Here, we specifically outline the experimental configurations used in this study.

B.1 Sentence Segmentation

For sentence segmentation, our experiments employ two different models: spaCy and ersatz. spaCy requires specification of the target language, whereas ersatz is language-agnostic, making it suitable for multilingual segmentation tasks.

The experiments in this paper cover five languages: English, German, Spanish, Japanese, and Chinese. Corresponding spaCy models for each language are as follows:

- English (en): en_core_web_sm
- German (de): de_core_news_sm
- Spanish (es): es_core_news_sm
- Japanese (ja): ja_core_news_sm
- Chinese (zh): zh_core_web_sm

B.2 Sentence Alignment

We adopt a robust sentence alignment strategy based on Vecalign (Thompson and Koehn, 2019), leveraging multilingual sentence embeddings and an efficient dynamic programming approximation to identify many-to-many alignments between sentence segments.

In our experiments, we set the maximum number of allowed overlap size to 16. This allows us to search for source-target sentence alignments of size $N-M$, where $N + M \leq 16$. While we use a fixed overlap size in our experiments, it can be estimated from reference documents. Detailed explanations are provided in Appendix B.4.

Adaptive Penalty Search. In Vecalign, null alignments are handled via a skip cost, parameterized by a percentile-based threshold β_{skip} , representing a quantile in the empirical distribution of 1-1 alignment costs. To identify the optimal β_{skip} , we perform a search starting from 0.2, which is the default value in Vecalign, and progressively decrease it in small steps of 0.005. At each step, we detect signals of over-deletion. Upon detection, we revert to the previous step and treat it as the final alignment result.

System:
You are a helpful assistant.

User :
Please merge these two segments into one sentence while preserving their original meaning, word choice, and order. Instead of simply concatenating them, use appropriate transitional expressions so that the segments are naturally connected without merely inserting a period or extra whitespace. Ensure the final result flows coherently and no important information is omitted. Return only the merged text on a single line, with no additional commentary or extraneous text.
First segment: <first segment>
Second segment: <second segment>

Figure 4: Prompt used for instructing GPT-4o to merge sentences.

First segment:
Move will also help transform land at the derelict Gartshore Works site.

Second segment:
A Scottish recycling business that has already processed more than a million tonnes of construction waste has opened a second plant following a multi-million pound investment.

After GPT-4o merging:
Move will also help transform land at the derelict Gartshore Works site **as** a Scottish recycling business that has already processed more than a million tonnes of construction waste has opened a second plant following a multi-million pound investment.

Figure 5: Example of merged sentences before and after GPT-4o rewriting.

Heuristic Termination Conditions. The search is terminated based on two heuristic signals that indicate over-deletion:

- (a) The average alignment cost falls below 0.3.
- (b) The null alignment ratio at a step exceeds 0.15.

Both patterns suggest that the skip cost has become too lenient, leading to excessive null alignments. In addition, two rare edge cases are handled with early stopping:

- (a) If the average alignment cost increases rather than decreases at a step, indicating misalignments with semantically distant content.
- (b) If the average alignment cost exceeds 0.7, indicating poor alignment quality overall.

Parameter Tuning. These heuristic termination conditions were tuned empirically on the test set from the WMT24 Discourse-Level Literary Translation shared task, using only the Chinese→English portion. To remain within the context length of our selected LLMs, we segment each instance from the training and validation sets into chunks of up to 1024 tokens. Translations are generated using meta-llama/Llama-3.1-8B-Instruct and GPT-4o₂₀₂₄₋₀₈₋₀₆, and sentence embeddings are computed using LASER. We empirically validate alignment quality by manually inspecting whether translation errors are consistently marked as null

alignments. Once verified, the same configuration is used for all experiments in this paper.

While the heuristic termination conditions are robust in our experiments, they may vary depending on the translation direction and source/reference sentence boundaries. These parameters can be further refined using reference translations, which are typically assumed to be perfectly aligned – i.e., with a null alignment (NA) ratio of zero. This provides a basis for estimating how much the NA ratio increases when the skip cost becomes too lenient, as well as the expected alignment cost under ideal conditions. These estimates can then inform the selection of appropriate heuristic parameters when evaluating future system translations.

B.3 Details for Text Embedding Model Fine-Tuning

We used News Commentary 18.1 parallel data from any language pairs that is a combination of Chinese (zh), English (en), German (de), Japanese (ja), and Spanish (es), which are the languages of interest in our evaluation. For each language pair, we use either the full dataset or only first 10,000 lines of the dataset, whichever is smaller. We build the example triplets with the following: we concatenate each of the two neighboring sentence pairs in the parallel corpus and use the source/target side as the query/positive example, respectively. As for the negative example, we always construct two

variants: (1) we randomly drop one of the two sentences on the target side (2) we retrieve a nearby sentence by randomly looking forward 1 to 3 sentences in the dataset and use it to substitute the second sentence on the target side. The intuition behind these example triplets is for the model to better distinguish a good translation from (1) an incomplete translation, and (2) a sentence that has a similar topic but is not a translation. The resulting synthetic dataset has 130,436 examples.

We fine-tune the BGE-M3 embedding on the synthetic dataset with InfoNCE loss (van den Oord et al., 2018) for two epochs. We did not conduct an extensive hyperparameter search, but simply use the setup in the fine-tuning tutorial in the FlagEmbedding package¹⁵. We directly used the checkpoint at the end of the training without using a validation set to select the best checkpoint.

B.4 Setting Overlap Size for Alignment

Beyond the text embedding model, we also observe that the choice of overlap size in Vecalign can significantly impact the alignment quality. The overlap size defines the size of the blocks that are compared to each other in the alignment search. A small overlap size causes some ideal alignments to fall out of search space. For example, if the overlap size is set to 8, but a long source sentence should be aligned to a sentence block of size 16 on the target side, such groundtruth alignment will never be considered by the search algorithm. On the other hand, a large overlap size will increase the computational cost.

With datasets that come with human-segmented sentence boundaries and alignment (which covers the vast majority of use cases), one can easily and accurately estimate the appropriate overlap size by re-segmenting the reference documents with sentence segmenter, calculate the maximum ratio of sentences as segmented by human and by the sentence segmenter. With datasets that does not come with human-segmented sentence boundaries or references, one would have to first conduct a pilot study with sentence-level translation to estimate the appropriate overlap size. However, the good news is that with a given sentence segmenter, the overlap size only needs to be estimated once per language, and can be reused for different datasets. Besides, in the case where estimation is not very

accurate, one can always err on the safe side and set a larger overlap size.

B.5 Details for Evaluation Setup

For budgetary reasons, WMT datasets from recent years exclude some segments from MQM annotations. To ensure meaningful comparison between the automatic metrics and human judgments at the document level, we exclude documents from the evaluation if more than 20% of the segments are missing MQM annotations.

Running mwerSegmenter requires a tokenization preprocessing step. For ja-zh, we follow the anonymous reviewer’s recommendation to tokenize the output into individual characters. For other languages, we use the tokenizer implemented in Sacremoses¹⁶.

C Supplementary Results

Since the results in the main paper are condensed versions with averaging across different language pairs or showing only a subset of the metrics evaluated, we attach the full breakdown of the results in Table 4 for readers’ reference.

D Licensing of Artifacts

Almost all code, model, and data artifacts we used in this paper are publicly available with permissive licenses (MIT/Apache 2.0/CC-BY-4.0). The only exceptions are Aya models (CC-BY-NC 4.0), and GPT-4o (OpenAI API Terms of Use), which still allows research use. We also plan to release all created code, model, and data artifacts under a permissive license.

E Use of AI Assistants

We used a code editor with generative AI functionalities during code development and paper writing (in the latter case, it only assists with LaTeX code completion and minor text editing). We also used various AI assistants for creating miscellaneous single-use data processing scripts, as well as all the figures in this paper. All AI-generated artifacts were carefully reviewed and accepted by the authors.

¹⁵https://github.com/FlagOpen/FlagEmbedding/blob/024e789d599eb4cf9a208e98d27508ad455f5ecb/Tutorials/7_Fine-tuning/7.1.2_Fine-tune.ipynb

¹⁶<https://pypi.org/project/sacremoses/>

	COMET			COMET-QE			METRIX			METRIX-QE			NA Ratio		
	en-de	en-es	ja-zh	en-de	en-es	ja-zh	en-de	en-es	ja-zh	en-de	en-es	ja-zh	en-de	en-es	ja-zh
Original															
Gold	0.3239	0.2320	0.3770	0.2667	0.2248	0.3486	0.3468	0.2352	0.3434	0.2807	0.1995	0.3247	0%	0%	0%
mwerSegmenter	0.2551	0.2376	0.3696	0.2045	0.2228	0.3368	0.2666	0.2413	0.3318	0.2009	0.2050	0.3022	0%	0%	0%
SEGALE-ersatz-BGE-M3-ft	0.3229	0.2410	0.3616	0.2662	0.2291	0.3351	0.3458	0.2454	0.3310	0.2776	0.2050	0.3063	0.25%	0.75%	1.32%
SEGALE-spaCy-BGE-M3-ft	0.3222	0.2350	0.3513	0.2652	0.2261	0.3314	0.3457	0.2401	0.3288	0.2775	0.2035	0.3047	0.99%	0.81%	2.34%
SEGALE-ersatz-LASER	0.3208	0.2340	0.3669	0.2637	0.2258	0.3380	0.3452	0.2392	0.3310	0.2782	0.2029	0.3069	0.23%	0.82%	2.40%
SEGALE-ersatz-BGE-M3	0.3193	0.2300	0.3617	0.2635	0.2205	0.3344	0.3448	0.2339	0.3377	0.2779	0.1977	0.3108	0.36%	0.81%	2.84%
Over-Translate															
Gold	0.3590	0.3975	0.3979	0.3145	0.3962	0.3533	0.3727	0.3430	0.3637	0.3004	0.2620	0.3380	10%	10%	10%
mwerSegmenter	0.2799	0.3816	0.3904	0.2410	0.3747	0.3483	0.2704	0.3533	0.3515	0.2020	0.2957	0.3211	10%	10%	10%
SEGALE-ersatz-BGE-M3-ft	0.3318	0.3488	0.3789	0.3012	0.3628	0.3464	0.3588	0.3400	0.3508	0.2973	0.3003	0.3269	9.29%	13.28%	11.09%
SEGALE-spaCy-BGE-M3-ft	0.3383	0.3495	0.3783	0.3007	0.3625	0.3446	0.3618	0.3405	0.3492	0.3002	0.3005	0.3255	9.47%	9.55%	11.35%
SEGALE-ersatz-LASER	0.3327	0.3146	0.3521	0.2994	0.3361	0.3129	0.3584	0.3212	0.3361	0.2956	0.2752	0.3126	8.16%	9.07%	8.66%
SEGALE-ersatz-BGE-M3	0.3104	0.3102	0.3494	0.2783	0.3278	0.3135	0.3381	0.3030	0.3344	0.2794	0.2574	0.3092	9.26%	8.61%	11.37%
Under-Translate															
Gold	0.3488	0.3349	0.3725	0.3050	0.3416	0.3056	0.3542	0.2564	0.3199	0.2881	0.2406	0.3208	10%	10%	10%
mwerSegmenter	0.1198	0.1935	0.3415	0.0947	0.1826	0.2937	0.1169	0.1266	0.2876	0.0754	0.0950	0.2619	2.65%	3.64%	0.02%
SEGALE-ersatz-BGE-M3-ft	0.3378	0.3499	0.3601	0.3000	0.3619	0.3266	0.3523	0.2944	0.3338	0.2909	0.2566	0.3082	5.10%	6.62%	5.72%
SEGALE-spaCy-BGE-M3-ft	0.3378	0.3481	0.3583	0.2998	0.3610	0.3266	0.3528	0.2923	0.3322	0.2916	0.2556	0.3088	5.47%	6.61%	6.01%
SEGALE-ersatz-LASER	0.3392	0.3415	0.3409	0.2934	0.3524	0.3132	0.3488	0.2850	0.3189	0.2891	0.2529	0.2926	5.55%	6.34%	7.13%
SEGALE-ersatz-BGE-M3	0.3412	0.3225	0.3507	0.3001	0.3441	0.3142	0.3504	0.2726	0.3299	0.2903	0.2447	0.3042	2.98%	4.14%	5.54%
Flex-Boundary															
Gold	0.3220	0.2321	0.3746	0.2632	0.2270	0.3463	0.3433	0.2372	0.3394	0.2768	0.2011	0.3216	0%	0%	0%
mwerSegmenter	0.2301	0.1853	0.3680	0.1803	0.1831	0.3342	0.2365	0.1929	0.3279	0.1741	0.1652	0.2999	0%	0%	0%
SEGALE-ersatz-BGE-M3-ft	0.3213	0.2299	0.3688	0.2620	0.2227	0.3392	0.3414	0.2379	0.3306	0.2721	0.2010	0.3060	0.98%	1.59%	1.05%
SEGALE-spaCy-BGE-M3-ft	0.3211	0.2313	0.3673	0.2616	0.2237	0.3389	0.3415	0.2391	0.3292	0.2720	0.2022	0.3053	1.43%	1.64%	1.35%
SEGALE-ersatz-LASER	0.3182	0.2300	0.3641	0.2609	0.2232	0.3357	0.3416	0.2373	0.3263	0.2746	0.2003	0.3033	1.31%	1.81%	2.55%
SEGALE-ersatz-BGE-M3	0.3169	0.2280	0.3554	0.2596	0.2183	0.3312	0.3406	0.2339	0.3292	0.2729	0.1971	0.3035	1.19%	2.03%	3.33%

Table 4: Full breakdown of our results on the WMT 2024 Metrics Shared Task dataset.