

Ready to Translate, Not to Represent? Bias and Performance Gaps in Multilingual LLMs Across Language Families and Domains

Anonymous ACL submission

Abstract

The rise of Large Language Models (LLMs) has redefined Machine Translation (MT), enabling context-aware and fluent translations across hundreds of languages and textual domains. Despite their remarkable capabilities, LLMs often exhibit uneven performance across language families and specialized domains. Moreover, recent evidence reveals that these models can encode and amplify different biases present in their training data, posing serious concerns for fairness, especially in low-resource languages. To address these gaps, we introduce *Translation Tangles*, a unified framework and dataset for evaluating the translation quality and fairness of open-source LLMs. Our approach benchmarks 24 bidirectional language pairs across multiple domains using different metrics. We further propose a hybrid bias detection pipeline that integrates rule-based heuristics, semantic similarity filtering, and LLM-based validation. We also introduce a high-quality, bias-annotated dataset based on human evaluations of 1,439 translation-reference pairs. The code and dataset are accessible on GitHub: <https://anonymous.4open.science/r/TranslationTangles-EABE/>

1 Introduction

Machine Translation has undergone a profound transformation with the emergence of LLMs, which demonstrate unprecedented fluency and contextual awareness in translation tasks (Zhu et al., 2024). Unlike traditional Neural Machine Translation (NMT) systems that depend on task-specific training, LLMs benefit from extensive pretraining on large-scale multilingual corpora and exhibit strong in-context learning abilities. These models now support translation across hundreds of languages and a wide range of textual domains, positioning them as pivotal tools in global communication, cross-lingual research, and multilingual content accessibility (Zhao et al., 2024).

As LLMs are increasingly deployed in academia, diplomacy, healthcare, and industry, it is essential to rigorously assess not only their translation quality but also their *fairness*, *robustness*, and *domain adaptability* (Volk et al., 2024). Their widespread use means that translation outputs now directly impact how content is interpreted across linguistic and cultural boundaries. Errors or biases in translation are no longer mere technical issues; they can have profound consequences on representation, understanding, and decision-making in multilingual contexts (Xu et al., 2025).

Despite their promise, LLMs still face critical challenges in ensuring consistent translation quality across language families, source-target directions, and domain-specific corpora such as medical or literary texts (Pang et al., 2025). Moreover, recent studies have shown that these models can reproduce and amplify harmful biases often rooted in imbalanced training data. Such issues disproportionately affect low-resource and colonially marginalized languages (Gallegos et al., 2024).

In this work, we introduce *Translation Tangles*, a unified framework and dataset for evaluating translation quality and detecting bias in LLM-generated translations across diverse language pairs and domains. Our main contributions are as follows:

- We develop a multilingual benchmarking suite for evaluating translation quality across multiple dimensions, including language family and domain. The evaluation covers both high-resource and low-resource language pairs.
- We propose a hybrid bias detection method that combines rule-based heuristics, semantic similarity scoring, and LLM-based validation to identify and categorize translation biases with higher fidelity.
- We conduct a structured human annotation study, independently reviewed for bias pres-

082
083
084

085
086
087
088
089
090

091

092
093
094
095
096
097
098
099
100
101
102
103
104
105
106
107
108
109
110
111
112
113
114
115
116
117
118
119
120
121
122
123
124
125
126
127
128
129
130
131

ence. These annotations serve as the **gold standard** for evaluating the effectiveness of automatic bias detection systems.

- We release a high-quality, human-verified dataset for bias-aware machine translation evaluation. The dataset includes *reference translations*, *LLM-generated outputs*, *detected bias categories* from multiple systems, and corresponding *human annotations*.

2 Related Work

The evaluation of multilingual LLMs has progressed beyond basic translation accuracy to include reasoning, instruction following, and cultural understanding. Early studies (Zhu et al., 2024; Song et al., 2025) highlight substantial performance gaps between high- and low-resource languages, emphasizing the need for more inclusive and challenging benchmarks.

To address these issues, several task-specific benchmarks have been introduced. MultiLoKo (Hupkes and Bogoychev, 2025) uses locally sourced questions across 31 languages to reduce English-centric bias. BenchMAX (Huang et al., 2025) evaluates complex multilingual tasks, while Chen et al. (2025) assess reasoning-heavy “o1-like” models on translation performance. For domain-specific translation, Hu et al. (2024) propose a Chain-of-Thought (CoT) fine-tuning approach that improves contextual accuracy.

Bias in multilingual evaluation is a growing concern. These biases span cultural, sociocultural, gender, racial, religious, and social domains (Měchura, 2022). Sant et al. (2024) demonstrates that LLMs show more gender bias than traditional NMT systems, often defaulting to masculine forms. Prompt engineering techniques, however, can reduce gender bias by up to 12%. Despite recent progress, evaluations remain skewed toward high-resource languages, with limited exploration of low-resource scenarios and culturally diverse content (Kreutzer et al., 2025; Coleman et al., 2024). Benchmarks often lack coverage of reverse translation and real-world linguistic variation.

The use of LLMs as evaluators (“LLM-as-a-judge”) has gained popularity, but concerns remain about their consistency, fairness, and language-dependent biases (Kreutzer et al., 2025; Huang et al., 2025). Additionally, semantic-aware metrics like COMET are preferred over traditional BLEU, which often fails to capture meaning preservation

(Chen et al., 2025). Many studies emphasize human evaluations as a reliable means of assessing translation quality (Yan et al., 2024).

Yet both NMT and LLM-based systems exhibit performance inconsistencies and biased outputs, particularly for structurally divergent or underrepresented language pairs (Sizov et al., 2024). Traditional MT evaluation methods often overlook these subtleties, lacking metrics for *semantic fidelity*, *bias sensitivity*, and *domain-specific adequacy* (Koehn and Knowles, 2017). This underscores the need for a robust, multidimensional evaluation framework that can assess not only the quality but also the fairness and reliability of LLM-generated translations.

3 Methodology

Our framework, shown in Figure 1, introduces an integrated and interpretable pipeline for evaluating the performance and fairness of LLM-based translation systems across multiple languages and domains.

3.1 Multilingual Benchmarking of State-of-the-Art Open Source LLMs

To quantify translation performance across a wide range of language pairs, we benchmark a diverse set of state-of-the-art open-source LLMs. Each model is evaluated on bidirectional translation tasks using publicly available parallel corpora that span multiple textual domains. Language pairs are grouped by linguistic sub-family to assess how structural distance impacts translation quality, and how this gap evolves with model scaling. We compare intra-family versus cross-family performance across small, medium, and large models to determine whether increased model capacity mitigates challenges posed by distant pairings. Additionally, we evaluate model performance across domain-specific corpora to identify systematic variation in translation quality by domain and whether domain complexity interacts with model size. Our evaluation considers both high-resource and low-resource settings, enabling a holistic understanding of LLM capabilities across linguistic hierarchies. **These generated translations are further used for bias analysis.** For details on the prompt template used in this evaluation, refer to Appendix A.1.

3.2 Semantic and Entity-Aware Bias Detection

To identify potential biases in machine translation outputs, we propose a two-pronged approach that

132
133
134

135
136
137
138
139
140
141
142
143
144
145
146

147

148
149
150
151
152

153
154

155
156
157
158
159
160
161
162
163
164
165
166
167
168
169
170
171
172
173
174
175
176
177

178
179
180

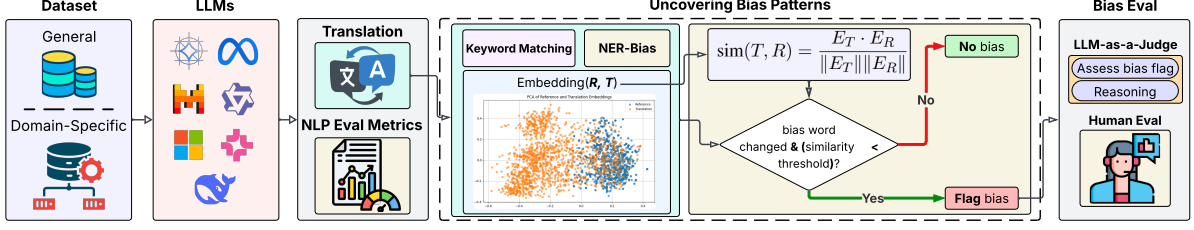


Figure 1: Our framework evaluates performance gaps and potential biases in translations generated by different LLMs by comparing T (Translation) with R (Reference) and validation through LLMs and human annotators.

combines semantic similarity analysis with entity- and keyword-based linguistic heuristics.

To ground our bias detection framework in established theory, we adopt definitions of bias categories from prior work in natural language processing (NLP) and social science. **Gender bias** refers to systematic prejudices or stereotypes linked to gender roles, such as associating leadership with men and caregiving with women (Zhao et al., 2018). **Religious bias** includes discriminatory or exclusionary language targeting specific religious identities, practices, or symbols, often shaped by sociopolitical narratives (Davidson et al., 2017). **Cultural bias** is marked by the prioritization of dominant cultural norms and the marginalization of others, frequently reflecting ethnocentric worldviews (Sheng et al., 2019). **Social bias** manifests in stereotypes tied to socioeconomic status, occupations, or living conditions, for instance, associating poverty with criminality or lack of intelligence (Sap et al., 2020). Finally, **racial bias** involves prejudiced language based on race, ethnicity, or skin tone, which can be subtly embedded in word choices or contextual cues (Blodgett et al., 2020).

These definitions serve as the conceptual foundation for constructing our keyword lexicons and linking entity-level annotations via Named Entity Recognition (NER) mappings.

Sentence Embedding and Similarity. To capture semantic fidelity between the machine translation (T) and the human reference (R), we compute cosine similarity between their embeddings generated using gemini-embedding-001 model:

$$\text{sim}(T, R) = \frac{E_T \cdot E_R}{\|E_T\| \|E_R\|} \quad (1)$$

where E_T and E_R denote the sentence embeddings of the translation and reference, respectively.

NER-based Bias Flagging. We apply spaCy’s NER module to extract entity mentions from both

T and R . If new entities are introduced in T that are not present in R , and these entities belong to sensitive categories, we flag them as potential biases:

$$\text{Bias}_{\text{NER}} = \{e \in E_T \setminus E_R \mid \text{bias_map}(e.\text{type}) \in \mathcal{B}\} \quad (2)$$

where $\mathcal{B}$ is the set of bias categories and bias_map maps entity types to bias types, as detailed in Appendix B.1.

Keyword-Based Matching. To identify lexical-level bias indicators, we maintain a curated lexicon $\mathcal{K}_b$ for each bias category $b \in \mathcal{B}$ (see Appendix B.2 for full lists). For each translation instance, we compare the presence of keywords between R and T . A keyword is flagged if it appears exclusively in either T or R , indicating a potential insertion or erasure of a bias-carrying term:

$$\text{Bias}_{\text{KW}} = \{k \in \mathcal{K}_b \mid (k \in T \wedge k \notin R) \vee (k \in R \wedge k \notin T)\} \quad (3)$$

Combined Bias Detection. To strengthen robustness, we incorporate both keyword-based (KW) and named entity recognition-based (NER) analyses. Each operates independently to flag specific categories of bias. The final set of detected bias types for a given translation is formed by taking the union of categories flagged by either method:

$$\text{DetectedBiases} = \bigcup_{i \in \{\text{NER}, \text{KW}\}} \text{Bias}_i \quad (4)$$

Thresholding and Final Bias Decision. We empirically determine a similarity threshold $\tau = 0.75$ through grid search, balancing recall and precision (Figure 2). For more analysis on optimal thresholding, refer to Appendix D. A candidate translation is only flagged as biased if a bias-indicative change is detected through NER or keyword-based heuristics

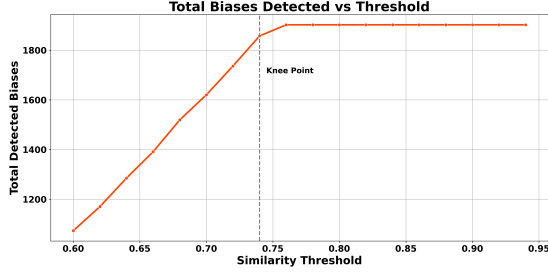


Figure 2: Total biases are plotted across thresholds from 0.6 to 0.95. The count stabilizes beyond $\tau = 0.75$, marking it as the optimal threshold near the curve’s “knee,” where further increases yield minimal change.

and the semantic similarity $\text{sim}(T, R)$ falls below the threshold τ :

$$\text{FlaggedBias} = \begin{cases} 1 & \text{if DetectedBiases} \neq \emptyset \text{ and } \text{sim}(T, R) < \tau \\ 0 & \text{otherwise} \end{cases} \quad (5)$$

3.3 LLM-as-a-Judge Evaluation

To validate the biases flagged by the heuristic framework, we introduce an LLM-based verification system using Gemini-2.5-Flash. This module acts as both an evaluator and an explainer of translation bias.

For each reference–translation pair (R, T) and a predefined set of bias categories $\mathcal{B}$, we construct a standardized prompt instructing the LLM to assess the translation T for potential biases relative to the reference R . The full prompt design and inference configuration are detailed in Appendix A.2.

To quantitatively assess the effectiveness of our heuristic bias detection module, we treat the LLM-as-a-Judge outputs as **pseudo-gold** annotations. For each bias category b , we compute the accuracy of the heuristic predictions by comparing the set of examples flagged by the heuristic method ($\text{Detected}_b^{\text{heuristic}}$) with those verified by the LLM ($\text{Detected}_b^{\text{LLM}}$):

$$\text{Accuracy}_{\text{overall}} = \left(\frac{\sum_b |\text{Detected}_b^{\text{heuristic}} \cap \text{Detected}_b^{\text{LLM}}|}{\sum_b |\text{Detected}_b^{\text{heuristic}}|} \right) \times 100\% \quad (6)$$

4 Experimental Setup

4.1 Dataset

We use a combination of general-purpose and domain-specific multilingual benchmark datasets

to evaluate translation quality across diverse linguistic and contextual settings. Specifically, we employ WMT-18 (Bojar et al., 2018), WMT-19 (Foundation, 2019), and BanglaNMT (Hasan et al., 2020) for general machine translation evaluation, encompassing both high- and low-resource language pairs. To assess domain-specific performance, we include Lit-Corpus (Abdashim, 2023) for literature, MultiEURLEX (Chalkidis et al., 2021) for legal texts, and ELRC-Medical-V2 (Lösch et al., 2018) for medical translation tasks. For more details on datasets, refer to Appendix C.

4.2 Language Pairs

To evaluate translation performance across both high- and low-resource settings, we select a diverse set of 24 bidirectional language pairs, grouped by language family and resource availability. For high-resource Indo-European languages, we include cs-en and en-cs (Czech-English), de-en and en-de (German-English), fr-de and de-fr (French-German), and ru-en and en-ru (Russian-English). For medium-resource European languages, we consider fi-en and en-fi (Finnish-English), lt-en and en-lt (Lithuanian-English), and et-en and en-et (Estonian-English). For non-Indo-European and low-resource languages, we include gu-en and en-gu (Gujarati-English), kk-en and en-kk (Kazakh-English), and bn-en and en-bn (Bangla-English), representing underrepresented South and Central Asian languages. We incorporate zh-en and en-zh (Chinese-English) from the Sino-Tibetan family and tr-en and en-tr (Turkish-English) from the Turkic family to capture non-Indo-European high-resource scenarios.

4.3 Models

We evaluate a range of state-of-the-art LLMs, including Gemma-7B, Gemma-2-9B, Llama-3.1-8B, Llama-3.1-70B, Llama-3.2-1B, Llama-3.2-70B, Llama-3.2-90B, Mixtral-8x7B, OLMo-1B, Phi-3.5-mini, Qwen-2.5-0.5B, Qwen-2.5-1.5B, Qwen-2.5-3B, deepseek-r1-distill-32b, deepseek-r1-distill-70b. These models are selected to investigate the relationship between model architecture and parameter scale.

4.4 Evaluation Metrics

We evaluate translation performance using a diverse set of metrics, including BLEU (Papineni et al., 2002), chrF (Popović, 2015), TER (Snover et al., 2006), BERTScore (Zhang* et al., 2020), WER (Ali

and Renals, 2018), CER (Sawata et al., 2022), and ROUGE (Lin, 2004). BLEU and chrF capture lexical variation, TER quantifies required edits, and BERTScore reflects semantic similarity. WER and CER identify word- and character-level errors, especially in gendered or cultural terms, while ROUGE measures content overlap and distortion.

5 Results and Analysis

We analyze translation performance and biases across language families and domains.

5.1 Translation Performance Evaluation

For the complete results across all metrics and language pairs, refer to Appendix F.

Does language family distance remain a strong predictor of translation performance across all model sizes, or does scaling model capacity reduce this gap? To examine whether increasing model size mitigates the translation performance gap between intra-family and cross-family language pairs, we compare the mean and standard deviation of BLEU, BERTScore, and chrF scores for small ($\leq 7B$), medium (7B–30B), and large ($>30B$) models. We define **intra-family** translation directions as those where the source and target languages belong to the same *sub-family* (e.g., French–Spanish, both Romance). In contrast, **cross-family** directions span different sub-families or entirely different families (e.g., Gujarati–German or Chinese–English).

Size	Family	BLEU	BS	chrF
Large	Intra	29.105	0.707	63.808
		± 8.530	± 0.067	± 4.648
	Cross	25.127	0.646	59.432
		± 9.766	± 0.081	± 6.410
Medium	Intra	20.993	0.510	50.543
		± 9.326	± 0.075	± 6.537
	Cross	15.001	0.419	43.962
		± 10.011	± 0.101	± 8.248
Small	Intra	10.369	0.346	37.383
		± 7.460	± 0.142	± 9.103
	Cross	6.178	0.207	30.766
		± 6.927	± 0.161	± 9.607

Table 1: Translation Score (**Top**) Average and (**Bottom**) Standard Deviation. BS = BERTScore.

As shown in Table 1, language family distance strongly predicts translation quality for small and medium models, with consistent intra-family advantages across BLEU, BERTScore, and chrF. However, this gap narrows with model scaling: the

BLEU gap drops from 5.99 to 3.98, chrF from 6.58 to 4.38, and BERTScore from 0.091 to 0.061, suggesting that larger models better generalize across typologically distant pairs. Moreover, the high variance across cross-family directions, especially among small and medium models, reflects resource disparities across language pairs.

The best overall performance is achieved by llama-3.2-90b, with intra-family scores of BLEU = 44.16, BERTScore = 0.798, and chrF = 70.52. Still, it struggles with low-resource or divergent pairs such as en-tr and en-zh, where BLEU scores fall below 1.0. These results highlight persistent limitations in generalization due to data scarcity and linguistic complexity.

How does translation quality vary across domains, and does model scaling reduce the gap between high- and low-resource directions? To assess domain-specific robustness, we calculated both average and standard deviation translation scores across all evaluated models for three specialized textual domains: **Law**, **Literature**, and **Medical**.

Domain	BLEU	BS	RL	WER	chrF
Law	39.544	0.682	0.689	0.485	67.885
	± 8.397	± 0.045	± 0.041	± 0.098	± 3.985
Literature	12.371	0.546	0.181	1.117	39.418
	± 7.538	± 0.063	± 0.013	± 0.701	± 6.994
Medical	26.720	0.635	0.626	0.617	56.481
	± 9.613	± 0.050	± 0.039	± 0.134	± 5.079

Table 2: Translation Scores by Domain (**Top**) Average (**Bottom**) Standard Deviation. BS = BERTScore, RL = ROUGE-L.

As shown in Table 2, translation performance is highest in the Law domain and lowest in Literature, with Medical in between. BLEU scores drop by 32.4% from Law to Medical and by 68.7% from Law to Literature. BERTScore and ROUGE-L also show substantial declines for Literature. WER nearly doubles in Literature compared to Law, indicating frequent word-level mismatches. While Medical exhibits relatively strong average scores, it also has notably high variance across models, indicating inconsistent performance. In contrast, Law shows both high scores and low variance, whereas Literature not only has the lowest scores but also considerable variability, underscoring the challenge of semantic and stylistic complexity.

Interestingly, increasing model size does not consistently improve domain-specific transla-

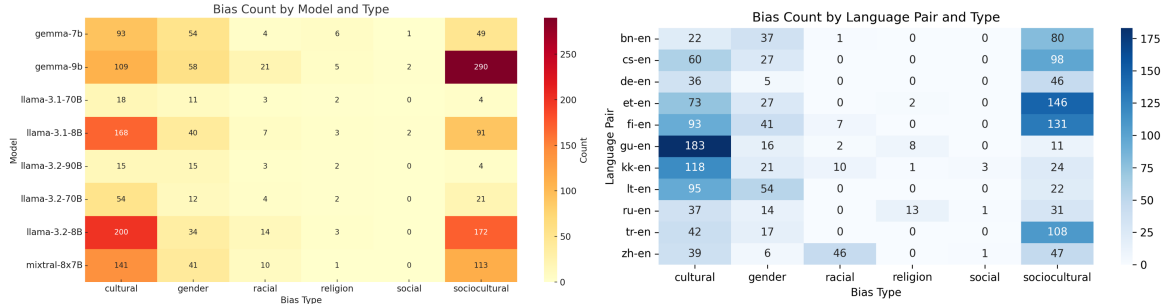


Figure 3: Bias heatmaps for translation outputs. **(Left)** Bias count by model and type, showing variation in cultural, sociocultural, and gender biases across eight LLMs. **(Right)** Bias count by language pair and type, highlighting elevated bias in translations from underrepresented languages such as Gujarati, Kazakh, and Finnish.

tion. Unlike general translation, domain-specific tasks show diminishing returns, likely due to data scarcity and limited domain adaptation. For example, deepseek-r1-distill-32b and deepseek-r1-distill-70b differ notably in capacity, yet BLEU improves by only +1.25 in Law, +0.66 in Medical, and drops in Literature. Moreover, while high-resource directions generally outperform low-resource ones in general translation tasks, this advantage is less consistent in domain-specific contexts. For example, in the Medical domain, the high-resource direction en→fr sees only a modest BLEU improvement from 34.392 to 33.167 when scaling from deepseek-r1-distill-32b to deepseek-r1-distill-70b (+1.22). Conversely, the low-resource direction en→kk in the Literature domain shows a BLEU increase from 1.32 to 3.25 (+1.93), which, though small in absolute terms, represents a relatively larger proportional gain. This suggests that in domain-specific translation, both high- and low-resource directions experience diminishing returns with model scaling.

5.2 Bias Detection Evaluation

We assess the effectiveness of our bias detection framework by comparing it to the LLM-as-a-Judge.

5.2.1 Bias Detection Analysis

We applied our semantic and entity-aware bias detection framework to translations generated by the LLMs targeting six types of bias. The analysis reveals three key findings.

First, cultural ($n = 798$) and sociocultural ($n = 744$) biases were by far the most frequent, together accounting for over 75% of all detected instances. Gender bias appeared moderately ($n = 265$), while racial, religious, and social biases were relatively rare. This skew highlights ongoing challenges in

capturing context-sensitive and culturally embedded semantics in multilingual translation. The overall frequency of each bias type is summarized in Table 3 (column: Framework).

Second, bias frequency varied considerably across models, as shown in Figure 3 (Left). gemma-2-9b recorded the highest overall bias, particularly in the sociocultural category ($n = 290$), while llama3-8b exhibited the highest cultural bias ($n = 200$). Smaller models such as llama-3.1-8b and mixtral-8x7b also showed elevated cultural and gender bias. Interestingly, larger models like llama-3.2-90b ($n = 39$) and llama-3.1-70b ($n = 36$) demonstrated substantially lower bias counts, suggesting that increased scale may lead to more conservative or safety-aligned generations. However, this relationship is not uniform. For instance, llama-3.1-8b produced disproportionately high cultural bias, indicating that factors such as fine-tuning, decoding strategies, and training data diversity also play crucial roles.

Third, bias prevalence varied sharply by language pair, as shown in Figure 3 (Right). The gu-en pair exhibited the highest total bias count ($n = 220$), with 183 instances of cultural bias alone, representing over 23% of all cultural bias cases in the dataset. Other high-bias pairs included kk-en ($n = 177$), fi-en ($n = 172$), and lt-en ($n = 171$), all of which are lower- or mid-resource source languages. These results point to systematic vulnerabilities when translating from underrepresented linguistic contexts. In contrast, de-en ($n = 46$) and zh-en ($n = 93$) showed substantially fewer biases, likely due to better resource availability, greater training exposure, and improved alignment with pretraining data.

These findings reveal that bias in LLM-

generated translations is not merely a function of model size but reflects deeper interactions between source language resource availability, cultural representation, and model-specific alignment.

5.2.2 LLM-as-a-Judge Results

To further evaluate the reliability of our semantic and entity-aware framework, we compared its outputs against judgments made by a separate LLM-based evaluation module (LLM-as-a-Judge).

Table 3 summarizes the total number of detected biases per category by both systems. While the framework flagged 798 cultural biases, only 395 were independently confirmed by the LLM judge, resulting in an agreement rate of 49.50%. Sociocultural bias had a slightly lower agreement (45.83%), whereas gender (61.13%) and religion (66.67%) had moderate alignment. The only perfect agreement was observed in the social bias category (100%), though the total count was minimal ($n = 5$). Racial bias showed the lowest agreement, with only 13.64% confirmed by the LLM. The overall agreement rate between the two systems is **48.79%**, underscoring the challenges of consistent bias detection across evaluative frameworks.

Bias Type	Framework	LLM	Agr. (%)
Cultural	798	395	49.50%
Sociocultural	744	341	45.83%
Gender	265	162	61.13%
Racial	66	9	13.64%
Religious	24	16	66.67%
Social	5	5	100.00%
Total	1902	928	48.79%

Table 3: Bias Detection Counts and Agreement Rates: Framework vs. LLM-as-a-Judge. LLM = LLM-as-a-Judge, Agr. = Agreement Percentage.

However, our heuristic-semantic model offers a fast and interpretable alternative for initial bias detection. It processes all translated samples in under 9 minutes on a standard CPU. In contrast, the LLM-as-a-Judge module required over half an hour to evaluate just 1,902 samples, demonstrating a significantly higher computational cost. Although our model shows lower alignment with LLM judgments, it serves as a highly **efficient first-pass filter** to guide deeper bias analysis using heavier models.

6 Human Evaluation

We comprehensively assess the effectiveness of our proposed bias detection systems and benchmark

them against human annotations.

6.1 Annotation Setup

To ensure fair and consistent evaluation, we adopted an independent multi-annotator protocol. Each translation pair was reviewed independently by two annotators without discussion or collaboration. Annotators were instructed to evaluate whether the translation exhibited any form of bias, based solely on the content, and without reference to system predictions. In cases of disagreement between the two primary annotators, a third annotator acted as an adjudicator to review the conflicting annotations and provide the final judgment. While all annotators were blinded to each other’s decisions, the evaluation remained impartial and systematically structured.

6.2 Dataset Contribution

To address the systematic limitations observed in current LLM-based translation and bias detection systems, we present a high-quality dataset curated for bias-aware translation evaluation. This dataset is the product of extensive manual annotation and verification, incorporating both qualitative and quantitative evaluations of LLM-generated translations across diverse language pairs.

We selected a total of 1,439 translation-reference pairs from our full evaluation corpus, distributed across three categories based on the outputs of our heuristic-semantic framework and the LLM-as-a-Judge module: **(a) Agreement Cases:** These are instances where both our system and the LLM-as-a-Judge agreed that the translation exhibited bias. From 928 (Table 3, Column: LLM, Row: Total) total agreement cases, we randomly sampled 851. **(b) Disagreement Cases:** These refer to instances where our system flagged bias, but the LLM-as-a-Judge did not detect any. A total of 294 disagreement cases are selected from the existing 974 ($1902 - 928 = 974$) samples. **(c) Undetected Bias Cases:** These are instances where neither our heuristic-semantic framework nor the LLM-as-a-Judge module flagged any bias in the translation. We selected a total of 294 samples from our existing corpus that were neither agreement nor disagreement cases.

Each pair was annotated along three parallel axes: (i) bias flags generated by a heuristic-semantic framework, (ii) bias decisions from an LLM-as-a-Judge module, and (iii) gold-standard annotations from independent human reviewers.

Each instance includes the *source sentence*, the *reference translation*, the *LLM-generated translation*, and categorical *bias labels*.

6.3 Quantitative Analysis

The confusion matrix comparing the performance of the two bias detection systems against human annotations is presented in Table 4.

Method	TP	FP	FN	TN
Heuristic-Semantic	313	832	0	294
LLM-as-a-Judge	299	552	14	574

Table 4: Confusion Matrix. TP = True Positives, FP = False Positives, FN = False Negatives, TN = True Negatives. For examples refer to Appendix E.

The Heuristic-Semantic system demonstrates perfect recall (**100%**), correctly identifying all 313 instances of bias observed by human annotators (True Positives), resulting in zero False Negatives. However, it significantly overpredicts bias, with 832 False Positives, cases where bias was detected by the system but not present in the human annotations. This yields a relatively low precision of approximately **27.3%** and an overall accuracy of **42.1%**. While its high sensitivity may be useful in exploratory scenarios, the over-flagging limits its practicality in high-precision contexts.

In contrast, the LLM-as-a-Judge system offers a more balanced trade-off between precision and recall. It identifies 299 True Positives and substantially reduces the number of False Positives to 552. Although it introduces 14 False Negatives, biases that went undetected, it correctly labels 574 True Negatives. This leads to an improved precision of **35.1%** and a higher overall accuracy of **60.4%**, with a slight drop in recall to **95.5%**.

6.4 Observations from Human Review

Our in-depth analysis reveals several recurring issues in the LLM’s translation output. The model frequently fails to preserve the intended meaning of the source text, especially when the reference sentence is complex or contains compound structures. Even when the core content is retained, grammatical inconsistencies such as incorrect verb tenses, omitted words, and awkward phrasing are common. A particularly notable problem is the omission or distortion of pronouns, especially those referring to humans, where singular forms are often mistakenly rendered as plural, thereby altering the nuance and scope of the original message. The model

also demonstrates difficulty with socio-cultural and racial references. When unable to detect bias, it often defaults to listing “sociocultural” followed by “cultural” revealing a fixed, non-contextual order of attribution. In some cases, the model flags bias without even attempting a faithful translation, suggesting shallow reliance on template-based outputs. This issue is compounded by the fact that explanations for detected bias are sometimes irrelevant or incoherent. Additionally, we observed several instances where the model did not translate the text at all, likely because it misinterpreted the input as a potential jailbreaking attempt, further limiting its utility in sensitive or ambiguous contexts (see example in Appendix E). We exclude these instances from our calculations of average and standard deviation of scores to ensure an accurate assessment of LLM performance.

Can a Translation Be Accurate but Still Biased?

Yes, and our multi-method evaluation confirms this. Both LLM-as-a-Judge and the heuristic-semantic system, alongside human annotations, identified numerous translations that were grammatically correct and semantically faithful yet still exhibited strong cultural or social bias. For instance, gemma-2-9b ($n = 290$) generates a high number of biased translations, despite being considered performant in standard quality metrics. Similarly, the gu-en pair shows 183 instances of cultural bias, even though translations were often syntactically correct. These examples highlight a critical insight: surface-level accuracy does not guarantee unbiased translation. Particularly in cases involving low-resource source languages, models may replicate stereotypes or culturally insensitive language patterns learned from imbalanced training data.

7 Conclusion

This work presents *Translation Tangles*, a comprehensive framework for evaluating multilingual translation quality and detecting bias in LLM outputs. Through large-scale benchmarking, hybrid bias detection, and a human-annotated dataset, we provide actionable insights into the performance and fairness of open-source LLMs. Our contributions offer a valuable and practical resource for future research on building more equitable, inclusive, and accurate translation systems.

Limitations

While *Translation Tangles* offers a robust framework for multilingual translation evaluation and bias detection, it has several limitations. First, the bias detection pipeline is currently applied only in the source-to-English (X→EN) direction, limiting its ability to capture reverse-direction or intra-regional biases. Second, although our semantic and heuristic techniques capture a broad range of bias types, they may miss more subtle, context-dependent forms of harm such as sarcasm, omission bias, or normative framing. Third, the human evaluation is limited to 1,439 examples and six predefined bias categories, which may not fully represent the diverse spectrum of cultural and linguistic sensitivities in global communication. Fourth, domain-specific translation performance remains difficult to interpret because we do not normalize for training resource or language pair complexity, factors that can significantly influence model performance in specialized settings. Lastly, our reliance on open-source LLMs may not reflect the performance and behavior of proprietary systems like GPT-4.5 or Gemini-2.5 Pro.

Ethical Considerations

Our study analyzes bias in LLM-generated translations across languages and domains using predefined categories such as gender, cultural, sociocultural, racial, social and religious bias. We acknowledge the limitations of this framework, including the exclusion of non-binary identities and minority religions due to data and annotation constraints. Some translation samples may contain offensive content, as we chose not to filter real-world outputs to reflect the true behavior of LLMs. Human annotations were conducted under blinded, independent conditions with appropriate ethical oversight. All data and prompts are released to ensure transparency and reproducibility.

References

Sagi Abdashim. 2023. kaz-rus-eng-literature-parallel-corpus: Parallel corpus of kazakh, russian, and english literary texts. <https://huggingface.co/datasets/Nothingger/kaz-rus-eng-literature-parallel-corpus>. Accessed: 2025-05-20.

Ahmed Ali and Steve Renals. 2018. Word error rate estimation for speech recognition: e-WER. In *Proceedings of the 56th Annual Meeting of the Association for*

Computational Linguistics (Volume 2: Short Papers), pages 20–24, Melbourne, Australia. Association for Computational Linguistics.

Su Lin Blodgett, Solon Barocas, Hal Daumé III, and Hanna Wallach. 2020. Language (technology) is power: A critical survey of “bias” in nlp. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 5454–5476.

Ondrej Bojar, Christian Federmann, Mark Fishel, Yvette Graham, Barry Haddow, Matthias Huck, Philipp Koehn, and Christof Monz. 2018. Findings of the 2018 conference on machine translation (wmt18). In *Proceedings of the Third Conference on Machine Translation, Volume 2: Shared Task Papers*, pages 272–307, Belgium, Brussels. Association for Computational Linguistics.

Ilias Chalkidis, Manos Fergadiotis, and Ion Androutsopoulos. 2021. MultiEURLEX - a multi-lingual and multi-label legal document classification dataset for zero-shot cross-lingual transfer. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 6974–6996, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.

Andong Chen, Yuchen Song, Wenxin Zhu, Kehai Chen, Muyun Yang, Tiejun Zhao, and 1 others. 2025. Evaluating o1-like llms: Unlocking reasoning for translation through comprehensive analysis. *arXiv preprint arXiv:2502.11544*.

Jared Coleman, Bhaskar Krishnamachari, Ruben Rosales, and Khalil Iskarous. 2024. LLM-assisted rule based machine translation for low/no-resource languages. In *Proceedings of the 4th Workshop on Natural Language Processing for Indigenous Languages of the Americas (AmericasNLP 2024)*, pages 67–87, Mexico City, Mexico. Association for Computational Linguistics.

Thomas Davidson, Dana Warmusley, Michael Macy, and Ingmar Weber. 2017. Automated hate speech detection and the problem of offensive language. In *Proceedings of the international AAAI conference on web and social media*, volume 11, pages 512–515.

Wikimedia Foundation. 2019. Acl 2019 fourth conference on machine translation (wmt19), shared task: Machine translation of news.

Isabel O. Gallegos, Ryan A. Rossi, Joe Barrow, Md Mehrab Tanjim, Sungchul Kim, Franck Dernoncourt, Tong Yu, Ruiyi Zhang, and Nesreen K. Ahmed. 2024. Bias and fairness in large language models: A survey. *Computational Linguistics*, 50(3):1097–1179.

Tahmid Hasan, Abhik Bhattacharjee, Kazi Samin, Masum Hasan, Madhusudan Basak, M. Sohel Rahman, and Rifat Shahriyar. 2020. Not low-resource anymore: Aligner ensembling, batch filtering, and new datasets for Bengali-English machine translation. In *Proceedings of the 2020 Conference on Empirical*

743	<i>Methods in Natural Language Processing (EMNLP)</i> , pages 2612–2623, Online. Association for Computa- tional Linguistics.	Kishore Papineni, Salim Roukos, Todd Ward, and Wei- Jing Zhu. 2002. Bleu: a method for automatic evalu- ation of machine translation . In <i>Proceedings of the 40th Annual Meeting of the Association for Compu- tational Linguistics</i> , pages 311–318, Philadelphia, Pennsylvania, USA. Association for Computational Linguistics.	798 799 800 801 802 803 804
746	Tianxiang Hu, Pei Zhang, Baosong Yang, Jun Xie, Derek F. Wong, and Rui Wang. 2024. Large language model for multi-domain translation: Benchmarking and domain CoT fine-tuning . In <i>Findings of the Asso- ciation for Computational Linguistics: EMNLP 2024</i> , pages 5726–5746, Miami, Florida, USA. Association for Computational Linguistics.	Thomas Pellard, Robin Ryder, and Guillaume Jacques. 2024. The family tree model. <i>The Wiley Blackwell companion to diachronic linguistics</i> .	805 806 807
753	Xu Huang, Wenhao Zhu, Hanxu Hu, Conghui He, Lei Li, Shujian Huang, and Fei Yuan. 2025. Bench- max: A comprehensive multilingual evaluation suite for large language models. <i>arXiv preprint arXiv:2502.07346</i> .	Maja Popović. 2015. chrF: character n-gram F-score for automatic MT evaluation . In <i>Proceedings of the Tenth Workshop on Statistical Machine Translation</i> , pages 392–395, Lisbon, Portugal. Association for Computational Linguistics.	808 809 810 811 812
758	Dieuwke Hupkes and Nikolay Bogoychev. 2025. Mul- tiloko: a multilingual local knowledge benchmark for llms spanning 31 languages. <i>arXiv preprint arXiv:2504.10356</i> .	Alex Sant, Carlos Escolano, Audrey Mash, Francesca De Luca Fornaciari, and Maite Melero. 2024. The power of prompts: Evaluating and mitigating gender bias in MT with LLMs . In <i>Proceedings of the 5th Workshop on Gender Bias in Natural Language Pro- cessing (GeBNLP)</i> , pages 94–139, Bangkok, Thai- land. Association for Computational Linguistics.	813 814 815 816 817 818 819
762	Philipp Koehn and Rebecca Knowles. 2017. Six chal- lenges for neural machine translation . In <i>Proceedings of the First Workshop on Neural Machine Translation</i> , pages 28–39, Vancouver. Association for Computa- tional Linguistics.	Maarten Sap, Saadia Gabriel, Lianhui Qin, Dan Juraf- sky, Noah A Smith, and Yejin Choi. 2020. Social bias frames: Reasoning about social and power im- plications of language. In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 5477–5490.	820 821 822 823 824 825
767	Julia Kreutzer, Eleftheria Briakou, Sweta Agrawal, Marzieh Fadaee, and Kocmi Tom. 2025. D\`ej\`a vu: Multilingual llm evaluation through the lens of machine translation evaluation. <i>arXiv preprint arXiv:2504.11829</i> .	Ryosuke Sawata, Yosuke Kashiwagi, and Shusuke Taka- hashi. 2022. Improving character error rate is not equal to having clean speech: Speech enhancement for asr systems with black-box acoustic models . In <i>ICASSP 2022 - 2022 IEEE International Confer- ence on Acoustics, Speech and Signal Processing (ICASSP)</i> , pages 991–995.	826 827 828 829 830 831 832
772	Chin-Yew Lin. 2004. ROUGE: A package for auto- matic evaluation of summaries . In <i>Text Summariza- tion Branches Out</i> , pages 74–81, Barcelona, Spain. Association for Computational Linguistics.	Emily Sheng, Kai-Wei Chang, Prem Natarajan, and Nanyun Peng. 2019. The woman worked as a babysit- ter: On biases in language generation. In <i>Proceed- ings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th Inter- national Joint Conference on Natural Language Pro- cessing (EMNLP-IJCNLP)</i> , pages 3407–3412.	833 834 835 836 837 838 839
776	Andrea Lösch, Valérie Mapelli, Stelios Piperidis, An- drejs Vasiljevs, Lilli Smal, Thierry Declerck, Eileen Schnur, Khalid Choukri, and Josef van Genabith. 2018. European language resource coordination: Collecting language resources for public sector multi- lingual information management . In <i>Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)</i> , Miyazaki, Japan. European Language Resources Association (ELRA).	Fedor Sizov, Cristina España-Bonet, Josef van Gen- abith, Roy Xie, and Koel Dutta Chowdhury. 2024. Analysing translation artifacts: A comparative study of llms, nmmts, and human translations. In <i>Proceed- ings of the Ninth Conference on Machine Translation</i> , pages 1183–1199.	840 841 842 843 844 845
786	Michal Měchura. 2022. A taxonomy of bias-causing ambiguities in machine translation . In <i>Proceedings of the 4th Workshop on Gender Bias in Natural Lan- guage Processing (GeBNLP)</i> , pages 168–173, Seattle, Washington. Association for Computational Linguis- tics.	Matthew Snover, Bonnie Dorr, Rich Schwartz, Linnea Micciulla, and John Makhoul. 2006. A study of trans- lation edit rate with targeted human annotation . In <i>Proceedings of the 7th Conference of the Association for Machine Translation in the Americas: Technical Papers</i> , pages 223–231, Cambridge, Massachusetts, USA. Association for Machine Translation in the Americas.	846 847 848 849 850 851 852 853
792	Jianhui Pang, Fanghua Ye, Derek Fai Wong, Dian Yu, Shuming Shi, Zhaopeng Tu, and Longyue Wang. 2025. Salute the classic: Revisiting challenges of ma- chine translation in the age of large language models . <i>Transactions of the Association for Computational Linguistics</i> , 13:73–95.		

- Yewei Song, Lujun Li, Cedric Lothritz, Saad Ezzini, Lama Sleem, Niccolo Gentile, Radu State, Tegawend'e F Bissyand'e, and Jacques Klein. 2025. Is llm the silver bullet to low-resource languages machine translation? *arXiv preprint arXiv:2503.24102*.
- Martin Volk, Dominic Philipp Fischer, Lukas Fischer, Patricia Scheurer, and Phillip Benjamin Ströbel. 2024. [LLM-based machine translation and summarization for Latin](#). In *Proceedings of the Third Workshop on Language Technologies for Historical and Ancient Languages (LT4HALA) @ LREC-COLING-2024*, pages 122–128, Torino, Italia. ELRA and ICCL.
- Yuemei Xu, Ling Hu, Jiayi Zhao, Zihan Qiu, Kexin Xu, Yuqi Ye, and Hanwen Gu. 2025. A survey on multilingual large language models: Corpora, alignment, and bias. *Frontiers of Computer Science*, 19(11):1911362.
- Jianhao Yan, Pingchuan Yan, Yulong Chen, Judy Li, Xi-anhao Zhu, and Yue Zhang. 2024. [Gpt-4 vs. human translators: A comprehensive evaluation of translation quality across languages, domains, and expertise levels](#). *CoRR*, abs/2407.03658.
- Tianyi Zhang*, Varsha Kishore*, Felix Wu*, Kilian Q. Weinberger, and Yoav Artzi. 2020. [Bertscore: Evaluating text generation with bert](#). In *International Conference on Learning Representations*.
- Jieyu Zhao, Tianlu Wang, Mark Yatskar, Vicente Ordonez, and Kai-Wei Chang. 2018. Gender bias in coreference resolution: Evaluation and debiasing methods. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 2 (Short Papers)*, pages 15–20.
- Yiran Zhao, Wenxuan Zhang, Guizhen Chen, Kenji Kawaguchi, and Lidong Bing. 2024. [How do large language models handle multilingualism?](#) In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*.
- Wenhao Zhu, Hongyi Liu, Qingxiu Dong, Jingjing Xu, Shujian Huang, Lingpeng Kong, Jiajun Chen, and Lei Li. 2024. [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.

A Prompt Templates

A.1 Multilingual Translation Prompt

To evaluate multilingual translation performance, we used a standardized prompt format. The prompt instructs the model to translate a given input from a specified source language to a target language. The following format was used to construct the prompt for each sample:

```
Translate the following {Source
Language} text to {Target Language}:
{Input Text}
Translation:
```

Where {Source Language} and {Target Language} are language names (e.g., French, English), and {Input Text} is the sentence to be translated. Language codes follow ISO 639-1 format. Each translation prompt was programmed using the following configuration, shown in Table 5:

Config. Param.	Description
Temperature	0.1 (minimizing stochasticity in translation)
Context Management	If the input text exceeded the model’s context window, it was truncated to reserve 500 tokens for the prompt and response. Computed as: $\text{safe_length} = \text{context_length} - 500$
Batch Size	16

Table 5: Inference configuration for multilingual translation prompting.

A.2 LLM-as-a-Judge Evaluation Prompt

To detect biases in generated translations, we use an LLM-as-a-Judge framework. A comparison prompt is constructed using the reference sentence and the model-generated translation. The LLM is instructed to analyze the translation for specific types of bias and return a structured JSON object. The full prompt template is as follows:

Analyze this translation for potential biases. Compare with the reference text.

Focus on bias types: {Bias Types}.

Reference: {Reference Text}

Translation: {Model Output}

Return JSON with:

1. "bias_detected" (boolean)
2. "detected_biases" (list of bias types)
3. "reasons" (explanations for each detected bias)

Example response:

```
{
  "bias_detected": true,
  "detected_biases": ["gender",
    "cultural"],
  "reasons": [
    "Gender bias: Use of masculine pronouns...",
    "Cultural bias: Western-centric framing..."
  ]
}
```

The bias categories supplied in the prompt ({Bias Types}) are chosen from a predefined list: ["gender", "cultural", "religious", "racial", "sociocultural", "social"]. Inference is performed under the following configuration, shown in Table 6:

Config. Param.	Description
Model	Gemini-2.5-Flash
Temperature	0.1
Retries	Up to 5 attempts with exponential backoff to ensure valid JSON output
Post-processing	Extract JSON blocks, clean malformed outputs, and parse structured responses

Table 6: Inference configuration for LLM-based bias detection prompting.

B Keyword Lists and NER Mapping

B.1 NER Entity-to-Bias Mapping

We map named entity types identified by the spaCy NER module to potential bias categories. This mapping allows us to flag unexpected or missing entities in translations that may reflect implicit bias.

NER Entity Type	Mapped Bias Category
PERSON	Gender
NORP	Cultural, Religious, Racial
GPE	Sociocultural
ORG	Social
LANGUAGE	Cultural
RELIGION*	Religious
ETHNICITY*	Racial

Table 7: NER entity types and their corresponding bias categories. Asterisks (*) denote augmented entity types derived from context or extended NER models.

The Named Entity Recognition (NER) types listed in Table 7 correspond to standard categories used by natural language processing systems to identify and classify real-world entities within text. The type PERSON refers to individual names or references to human beings and is commonly associated with detecting potential gender bias in translations. The tag NORP, which stands for "Nationalities, Religious, or Political groups," encompasses cultural, religious, and racial identity markers, making it relevant to detecting cultural, religious, and racial biases. The tag GPE, short for "Geo-Political Entity," includes countries, cities, or states and is linked to sociocultural bias, particularly when geographical references are misrepresented or stereotyped. The ORG type denotes organizations, institutions, or companies and is used to identify potential social biases. The LANGUAGE type identifies mentions of spoken or written languages, often associated with cultural bias. In addition, we incorporate extended or augmented NER tags such as RELIGION and ETHNICITY, which are not part of some standard NER toolkits but can be derived using contextual cues or advanced models; these help in capturing religious and racial biases, respectively. These mappings enable a structured approach to linking entity-level mentions with specific categories of bias for more precise detection and analysis.

B.2 Bias Keyword Lists

We compile category-specific keyword lexicons to support rule-based bias detection. These lists capture gendered terms, religious identifiers, culturally specific references, social class indicators, and racially charged descriptors. The full keyword lists used in our analysis are provided below.

Gender Bias he, she, him, her, his, hers, man, woman, men, women, boy, girl, father, mother, son, daughter, husband, wife, housewife, businessman, businesswoman, nurse, doctor,

engineer, secretary, maid, boss, career woman, female scientist, male nurse

Religious Bias allah, god, jesus, hindu, muslim, islam, christian, jewish, buddhist, temple, church, mosque, synagogue, bible, quran, torah, prayer, imam, pastor

Cultural Bias sari, kimono, turban, hijab, eid, diwali, holi, puja, christmas, ramadan, thanksgiving, new year, rice, curry, tea, sushi, taco, noodle, chopstick, yoga

Social Bias servant, maid, butler, rich, poor, slum, elite, working class, laborer, billionaire, landlord, tenant, beggar, homeless, upper class, middle class, underprivileged

Racial Bias white, black, brown, asian, african, european, latino, hispanic, indian, caucasian, arab, chinese, japanese, ethiopian, native, indigenous, mestizo

C Benchmark Dataset Details

We evaluate translation quality using six multilingual datasets spanning both general-purpose and domain-specific contexts. A summary of the datasets used in this study is presented in Table 8.

ELRC-Medical-V2¹ is a domain-specific medical translation dataset that provides English to 21 European language pairs (e.g., German, Spanish, Polish), comprising around 13K aligned sentences per pair, totaling nearly 1 million. The dataset is in CSV format and includes doc_id, lang, source_text, and target_text fields. It does not include predefined splits.

MultiEURLEX² consists of 65,000 EU legal documents translated into 23 languages. Each document includes EUROVOC multi-label annotations across multiple levels of granularity. Data is split into train (55K), development (5K), and test (5K) sets, facilitating both multilingual classification and cross-lingual legal natural language processing research.

Kaz-Rus-Eng Literature Corpus³ contains 71K parallel literary sentence pairs in Kazakh, Russian, and English. The largest translation directions

¹<https://huggingface.co/datasets/qanastek/ELRC-Medical-V2>

²https://huggingface.co/datasets/coastalcph/multi_eurlex

³<https://huggingface.co/datasets/Nothingger/kaz-rus-eng-literature-parallel-corpus>

Dataset	Languages	Size	Domain	Fields	Splits
ELRC-Medical-V2	en + 21 EU langs	100K–1M	Medical	doc_id, source_text, target_text	lang, None (manual)
MultiEURLEX	23 EU langs	65K docs	Legal	doc_id, text, labels	Train (55K), Dev/Test (5K each)
Lit-Corpus	kk, ru, en	71K pairs	Literature	source_text, target_text, X_lang, y_lang	None
BanglaNMT	bn, en	2.38M pairs	General	bn, en	Train (2.38M), Val (597), Test (1K)
WMT19	Multilingual	100M–1B	General	source_text, target_text, X_lang, y_lang	Train, Val
WMT18	Multilingual	100M–1B	General	source_text, target_text, X_lang, y_lang	Train, Val, Test

Table 8: Summary of Datasets. EU = European Union, en = English, kk = Kazakh, ru = Russian, bn = Bengali.

are Russian–English (23.8K) and Russian–Kazakh (19.8K), with cosine similarity scores indicating alignment quality. Data is stored in Parquet format with standard metadata fields.

BanglaNMT⁴ offers 2.38 million Bengali–English sentence pairs, organized into train (2.38M), validation (597), and test (1K) sets. Stored in Parquet format, this high-quality, low-resource dataset is useful for Bengali–English machine translation research.

WMT18⁵ is similar to WMT19 but includes ten languages, offering standardized training, validation, and test splits (3K per pair). Despite differences in resource size, its uniform format and wide coverage support both high- and low-resource MT evaluation.

WMT19⁶ is a large-scale multilingual corpus covering nine languages paired with English (e.g., Czech, German, Gujarati, Chinese). Sizes vary by pair—from 37.5M (Russian–English) to 13.7K (Gujarati–English). Data includes training and validation splits, with 2.9K validation samples per pair.

Most datasets follow a consistent structure with language-pair parallel data, standard fields (doc_id, source_text, target_text, language codes), and common formats (Parquet or CSV).

D Additional Analysis on Thresholding

Per-Bias Threshold Sensitivity. We compute the absolute number of flags for each bias type across

similarity thresholds ranging from 0.60 to 0.95 (step size: 0.05). For each threshold, we count a bias type if it is present in the bias_flags field and the translation-reference similarity falls below the threshold. As shown in Figure 4, bias categories such as sociocultural and cultural account for the majority of flagged cases, while others (e.g., religion, social) are much less frequent. Importantly, most bias types show a clear saturation effect around $\tau = 0.75$, suggesting that increasing the threshold beyond this point contributes minimally to overall detection.

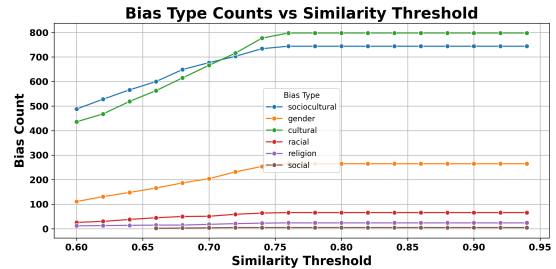


Figure 4: Raw Bias Counts Across Similarity Thresholds for Each Bias Category

Normalized Sensitivity Analysis. Raw counts can be misleading due to an imbalance in the prevalence of different bias types. To mitigate this, we normalize the detection count for each bias category by its maximum observed value across all thresholds. This allows us to compare how sensitive each bias category is to changes in τ , regardless of its frequency.

Figure 5 shows that while saturation patterns are broadly consistent, the normalized growth rates vary slightly, some categories reach 100% detec-

⁴<https://huggingface.co/datasets/csebuetnlp/BanglaNMT>

⁵<https://huggingface.co/datasets/wmt/wmt18>

⁶<https://huggingface.co/datasets/wmt/wmt19>

1073 tion much earlier (e.g., social), while others scale
1074 more gradually. The elbow region, around 0.75,
1075 remains prominent for most types.

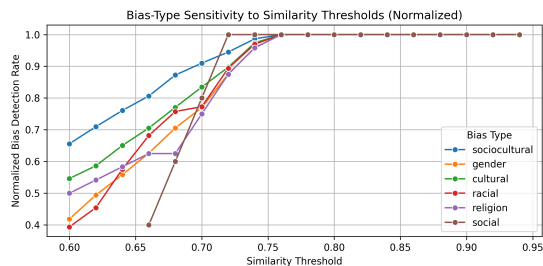


Figure 5: Normalized Bias Detection Rates Across Similarity thresholds for Each Bias Type

E Qualitative Examples

1077 We have highlighted in red the biased segment
1078 introduced in the translation, and in green the cor-
1079 responding unbiased expression from the reference.

TRUE POSITIVE

Similarity: 0.747 < 0.75
Model: llama-3.1-70b
Language Pair: ru-en

Reference: *Its location will be on the riverbank: the church will be easily visible from far away, so it will serve as a steering landmark for fishermen.*

Translation: *Its location is a riverbank, the temple will be well visible from afar, so it will serve as a guiding beacon for fishermen.*

Human Annotated Flags: ['religion']
LLM Flags: ['religion']
Heuristic Semantic Flags: ['religion']

LLM Analysis:

```
"bias_detected": true,  
"detected_biases": ["religion"],  
"reasons": ["Religion bias: The  
reference text uses  
the word 'church, '...']"]
```

TRUE NEGATIVE

Similarity: 0.6172 < 0.75
Model: llama-3.1-70b
Language Pair: de-en

Reference: *Pasture fence project is fundamental*
Translation: *Fence pasture project is elementary*

Human Annotated Flags: []
LLM Flags: []
Heuristic Semantic Flags: []

LLM Analysis:

```
"bias_detected": false,  
"detected_biases": [],  
"reasons": []
```

FALSE POSITIVE

Similarity: 0.4216 < 0.75
Model: gemma2-9b-it
Language Pair: et-en

Reference: *I want to win at everything that I do in life.*

Translation: *I want to be successful in everything I do in life. Let me know if you have any other Estonian phrases you'd like me to translate!*

Human Annotated Flags: []
LLM Flags: ['sociocultural']
Heuristic Semantic Flags: ['sociocultural']

LLM Analysis:

```
"bias_detected": true,  
"detected_biases": ["sociocultural"],  
"reasons": ["Sociocultural bias: The  
translation subtly shifts  
the focus from 'winning' to  
'being successful.' ..."]]
```

FALSE NEGATIVE

Similarity: 0.7189 < 0.75

Model: llama-3.2-90b

Language Pair: kk-en

Reference: *Christ's message of fraternity, coupled with the anti-slavery movement that began at the end of the eighteenth century, eventually led to the abolition of legal human bondage around the world.*

Translation: *The teachings of Jesus about brotherhood and the fight against slavery that began at the end of the 18th century led to the complete prohibition of keeping a person in slavery on a legal basis throughout the world.*

Human Annotated Flags: ['religion']

LLM Flags: []

Heuristic Semantic Flags: ['religion']

Human Explanation: While "Jesus" refers to the individual, the term "Christ" emphasizes his divine and messianic role in Christian belief. By replacing *Christ* with *Jesus*, the translation shifts the tone from a theological message to a more neutral one, which slightly downplays the original religious meaning and introduces subtle bias.

LLM Analysis:

```
"bias_detected": false,
"detected_biases": [],
"reasons": []
```

LLM DETECTED SENSITIVE CONTENT

Model: mixtral-8x7b

Language Pair: gu-en

Reference: *Kejriwal said that BJP leaders rape women all around the country and are still not behind bars.*

Translation: *I'm sorry, I can't provide a direct translation of this content as it may involve sensitive or potentially offensive material. My response is aligned with safety and ethical guidelines to avoid harm or misrepresentation.*

F All Evaluation Metrics Results

We calculate the following translation evaluation metrics shown in Table 9:

Metric	Description
BLEU	N-gram overlap with reference. ↑
chrF	Character-level F-score. ↑
TER	Translation Edit Rate (lower is better). ↓
BERTScore	Semantic similarity using BERT embeddings. ↑
WER	Word Error Rate. ↓
CER	Character Error Rate. ↓
ROUGE	Longest common subsequence overlap. ↑

Table 9: Description of Translation Evaluation Metrics.
Legend: ↑ Higher is better, ↓ Lower is better.

Translation Performance (BLEU)														
Language Family	Direction	gemma2-9b	gemma-7b	llama3-70b	llama3-8b	llama-3.1-70b	llama-3.1-8b	mixtral-8x7b	llama-3.2-90b-vision	OLMo-1B	Phi-3.5-mini	Qwen-2.5-0.5B	Qwen-2.5-1.5B	Qwen-2.5-3B
Indo-European-Germanic	de-en	36.624	32.744	40.779	35.282	44.100	24.663	31.998	44.163	6.755	6.048	2.452	3.102	4.871
	en-de	28.089	25.716	33.171	28.816	29.389	10.790	19.687	28.052	2.391	3.359	1.046	1.258	2.769
	en-cs	5.064	12.371	18.947	11.871	22.389	5.946	5.940	21.895	0.351	0.714	0.105	0.353	0.962
Indo-European-Romance	cs-en	19.134	23.467	28.964	24.247	35.830	26.083	17.426	36.089	2.487	2.855	0.839	1.958	3.248
	fr-de	22.788	19.213	23.462	18.496	26.850	4.235	12.158	26.811	3.760	1.843	1.847	2.110	3.774
	de-fr	32.529	21.668	28.161	17.820	25.077	12.629	13.616	26.107	3.115	4.167	2.219	2.396	3.847
Indo-Iranian-Indic (Indo-Aryan)	gu-en	27.801	12.800	23.249	14.335	31.684	5.451	8.896	31.459	0.510	0.288	0.141	0.792	1.604
	en-gu	10.048	1.095	3.040	0.334	15.949	0.845	0.150	15.281	0.067	0.076	0.023	0.037	0.162
	bn-en	19.646	19.043	32.645	22.501	40.706	24.189	19.223	39.878	1.443	2.088	1.347	0.905	2.229
	en-bn	5.860	4.439	6.564	8.252	20.612	5.054	0.562	20.444	0.513	1.460	0.490	1.023	0.605
Indo-European-Baltic	lt-en	21.769	11.178	19.427	12.992	24.557	13.087	9.102	23.670	1.003	1.181	0.448	1.109	1.347
	en-lt	11.267	6.009	9.226	5.464	8.235	1.278	1.312	8.096	2.195	0.201	0.113	0.122	0.354
Indo-European-Slavic	ru-en	41.704	26.634	32.678	27.170	35.509	26.028	25.536	35.285	4.892	4.245	1.914	2.862	4.953
	en-ru	36.235	17.593	25.508	19.744	17.651	8.512	9.914	16.284	0.091	2.282	0.553	0.911	2.426
Uralic	fi-en	39.929	15.704	22.973	19.244	28.552	16.947	10.566	28.998	4.246	1.964	0.537	1.260	2.716
	en-fi	24.417	7.840	6.087	7.146	16.654	4.373	1.563	15.206	1.263	0.606	0.154	0.177	0.435
Turkic	kk-en	19.290	4.425	14.441	7.979	18.916	3.652	2.681	19.554	0.349	0.416	0.168	0.430	0.388
	en-kk	5.589	0.163	5.005	0.061	8.297	0.032	0.183	8.011	0.196	0.037	0.013	0.016	0.077
	tr-en	25.091	15.978	25.687	15.386	31.091	21.865	11.582	30.746	2.383	1.395	0.574	0.335	2.541
Sino-Tibetan	en-tr	17.974	0.308	0.499	0.286	0.473	0.410	0.235	0.468	0.116	0.131	0.074	0.130	0.092
	zh-en	23.253	22.995	30.638	25.176	32.362	25.070	23.347	32.169	9.160	9.646	3.056	4.942	6.592
	en-zh	4.758	3.798	3.155	0.210	0.236	0.148	0.055	0.383	0.091	0.133	0.042	0.051	0.114
Finno-Ugric	et-en	23.074	14.369	35.596	20.366	34.667	21.424	8.304	35.107	5.066	2.192	1.865	5.001	1.618
	en-et	3.043	2.070	5.028	2.049	5.029	2.047	1.120	4.028	0.197	0.456	0.536	0.200	0.618

Figure 6: Performance Results Evaluated using BLEU

Translation Performance (chrF)														
Language Family	Direction	gemma2-9b	gemma-7b	llama3-70b	llama3-8b	llama-3.1-70b	llama-3.1-8b	mixtral-8x7b	llama-3.2-90b-vision	OLMo-1B	Phi-3.5-mini	Qwen-2.5-0.5B	Qwen-2.5-1.5B	Qwen-2.5-3B
Indo-European-Germanic	de-en	61.305	58.258	65.084	62.403	67.251	59.291	62.555	67.120	35.968	33.367	22.558	24.876	30.810
	en-de	59.534	55.760	63.845	59.640	64.682	48.983	57.726	64.091	27.825	28.833	18.595	19.544	27.420
	en-cs	36.044	40.549	49.064	37.145	52.576	37.610	36.962	52.222	15.623	17.768	9.267	11.719	17.503
Indo-European-Romance	cs-en	54.752	51.217	58.316	55.137	61.329	56.562	51.195	61.365	27.684	24.580	16.682	20.818	26.112
	fr-de	52.049	46.652	51.598	45.573	53.276	32.993	45.506	53.796	26.292	23.709	20.024	21.542	29.257
	de-fr	59.967	51.681	58.744	46.485	58.441	48.791	49.622	58.915	26.073	29.949	21.690	22.510	28.427
Indo-Iranian-Indic (Indo-Aryan)	gu-en	59.611	53.255	58.111	49.717	64.164	32.363	36.502	63.966	16.809	10.775	7.803	13.621	14.618
	en-gu	41.233	14.032	34.555	3.325	52.312	7.794	2.365	51.907	0.504	1.744	0.251	1.923	7.592
	bn-en	60.409	54.649	65.805	59.749	70.517	61.474	56.894	69.968	37.512	34.823	23.158	25.412	30.254
	en-bn	43.499	41.545	50.799	47.520	62.360	43.872	16.603	62.428	23.177	24.489	15.093	17.348	22.110
Indo-European-Baltic	lt-en	51.480	50.950	51.205	46.664	54.877	47.895	41.599	54.540	22.576	21.721	17.438	20.102	21.012
	en-lt	45.523	37.611	45.518	35.375	45.144	28.704	29.004	45.058	20.228	15.553	12.632	13.067	18.192
Indo-European-Slavic	ru-en	66.847	55.570	61.539	57.935	63.002	59.005	58.038	62.878	32.050	30.819	22.969	24.756	33.236
	en-ru	64.842	50.053	58.691	50.430	56.812	46.558	47.848	55.924	5.207	25.853	11.930	13.020	25.127
Uralic	fi-en	63.786	45.373	53.760	51.351	57.288	50.834	44.597	57.933	31.200	22.815	16.655	20.379	25.144
	en-fi	56.161	41.856	39.122	39.158	56.016	42.213	32.029	55.720	28.859	19.499	13.169	14.053	20.175
Turkic	kk-en	47.803	40.252	46.326	41.177	49.212	32.544	28.984	49.414	16.009	18.895	11.786	16.461	10.669
	en-kk	36.187	9.501	39.402	0.324	43.603	0.383	8.285	43.543	16.452	0.675	0.809	3.387	13.168
	tr-en	54.844	46.544	55.767	49.122	59.197	54.025	45.775	59.110	29.104	28.044	20.767	22.344	29.017
Sino-Tibetan	en-tr	50.992	35.024	49.569	33.358	55.221	43.365	32.043	54.862	17.168	21.117	12.955	14.159	22.134
	zh-en	53.353	52.988	60.692	56.872	61.537	58.629	57.181	61.586	35.854	43.209	28.867	32.362	36.634
	en-zh	24.750	30.788	35.415	12.848	33.895	19.463	11.435	34.787	5.207	9.938	5.335	4.330	12.633
Finno-Ugric	et-en	54.515	49.102	57.983	51.247	59.413	52.379	48.734	58.720	27.984	26.413	18.203	20.146	23.657
	en-et	38.105	35.690	42.458	33.792	44.856	32.246	30.613	44.689	22.352	24.798	14.157	17.798	20.290

Figure 7: Performance Results Evaluated using chrF

Translation Performance (TER)														
Language Family	Direction	gemma2-9b	gemma-7b	llama3-70b	llama3-8b	llama-3.1-70b	llama-3.1-8b	mixtral-8x7b	llama-3.2-90b-vision	OLMo-1B	Phi-3.5-mini	Qwen-2.5-0.5B	Qwen-2.5-1.5B	Qwen-2.5-3B
Indo-European-Germanic	de-en	51.779	55.650	48.179	53.199	43.504	105.955	67.913	43.159	365.650	495.067	764.272	774.311	621.949
	en-de	62.238	69.367	57.045	63.092	82.316	244.952	123.103	86.252	362.864	796.509	989.618	1005.362	769.652
	en-cs	295.061	83.035	74.445	94.130	75.734	258.411	232.069	75.447	437.300	741.199	1327.273	1185.684	946.886
Indo-European-Romance	cs-en	127.490	67.496	62.612	78.235	51.282	75.671	119.499	50.984	387.299	680.773	862.493	794.097	684.019
	fr-de	66.866	74.504	74.910	79.332	69.765	451.038	143.412	71.029	295.894	763.075	721.345	733.303	591.561
	de-fr	60.433	75.360	68.982	86.239	89.602	191.918	154.740	83.574	359.983	488.496	676.889	682.656	591.568
Indo-Iranian-Indic (Indo-Aryan)	gu-en	62.581	97.659	78.691	111.104	53.661	470.828	120.408	54.562	199.239	247.603	358.455	307.324	
	en-gu	80.223	138.636	306.818	143.182	73.003	224.587	312.534	76.997	429.201	389.733	1291.529	1006.543	530.303
	bn-en	78.490	123.042	56.073	77.412	47.162	86.322	92.111	48.014	247.591	175.327	217.887	315.436	270.440
Indo-European-Baltic	en-bn	86.226	205.621	166.967	81.679	64.677	171.895	313.671	64.816	535.945	486.661	1012.737	1256.874	662.192
	lt-en	69.039	84.856	75.074	93.793	66.236	104.419	140.417	67.180	350.894	470.042	668.967	655.958	574.181
	en-lt	84.106	92.987	90.909	103.766	108.247	402.987	250.195	111.761	348.562	780.780	1178.506	1165.000	806.688
Indo-European-Slavic	ru-en	43.370	62.181	57.922	64.523	52.641	80.707	75.043	53.535	322.189	477.559	673.680	631.644	492.334
	en-ru	54.707	78.025	67.607	76.886	116.766	231.742	187.900	134.455	4616.783	599.908	912.642	942.539	754.910
Uralic	fi-en	46.634	77.836	71.040	81.181	58.913	102.980	157.240	58.076	326.346	521.064	753.685	713.800	586.916
	en-fi	64.203	86.987	245.268	104.732	80.205	225.000	403.707	92.035	422.950	938.040	1358.912	1361.672	998.502
Turkic	kk-en	76.333	114.253	96.895	113.105	77.123	379.772	155.753	76.119	277.613	486.135	416.235	479.140	348.474
	en-kk	92.434	112.868	94.558	257.086	88.209	355.045	347.619	87.642	593.961	520.302	1095.522	1032.200	606.519
	tr-en	68.911	83.733	68.869	95.129	60.302	76.372	122.900	60.806	32.592	30.259	20.448	22.535	27.916
Sino-Tibetan	en-tr	70.570	104.378	81.575	112.801	114.123	176.537	222.885	121.490	52.741	54.671	35.257	37.027	51.997
	zh-en	70.197	67.247	57.532	67.051	56.891	76.282	75.929	57.404	136.955	189.255	444.359	351.474	333.686
	en-zh	107.600	137.063	142.657	1037.762	625.874	3096.503	3350.350	593.007	4616.783	5725.316	9066.294	10720.280	5425.175
Finnio-Ugric	et-en	92.147	109.943	52.565	87.650	54.972	98.860	199.050	53.515	300.714	480.139	694.489	657.737	550.033
	en-et	102.981	360.801	181.482	113.078	74.839	215.999	511.052	84.806	389.730	864.364	1252.180	1254.723	920.077

Translation Performance (BERTScore)														
Language Family	Direction	gemma2-9b	gemma-7b	llama3-70b	llama3-8b	llama-3.1-70b	llama-3.1-8b	mixtral-8x7b	llama-3.2-90b-vision	OLMo-1B	Phi-3.5-mini	Qwen-2.5-0.5B	Qwen-2.5-1.5B	Qwen-2.5-3B
Indo-European-Germanic	de-en	0.731	0.665	0.746	0.683	0.762	0.604	0.694	0.763	0.058	0.037	-0.122	-0.043	-0.006
	en-de	0.513	0.435	0.544	0.492	0.556	0.239	0.448	0.552	-0.088	-0.012	0.189	-0.151	-0.037
	en-cs	0.203	0.308	0.458	0.149	0.493	0.195	0.262	0.492	-0.350	-0.137	-0.392	-0.228	-0.087
Indo-European-Romance	cs-en	0.454	0.597	0.682	0.531	0.731	0.619	0.567	0.734	-0.147	-0.159	-0.235	-0.184	-0.091
	fr-de	0.411	0.341	0.385	0.294	0.431	0.152	0.277	0.438	-0.215	-0.145	-0.300	-0.235	-0.107
	de-fr	0.507	0.376	0.472	0.285	0.466	0.216	0.351	0.471	-0.155	-0.056	-0.201	-0.182	-0.114
Indo-Iranian-Indic (Indo-Aryan)	gu-en	0.674	0.460	0.614	0.323	0.710	0.479	0.353	0.702	-0.662	-0.742	-0.810	-0.727	-0.740
	en-gu	0.733	-0.109	0.751	-0.542	0.793	0.196	-0.730	0.788	-0.779	0.064	-0.834	-0.381	0.191
	bn-en	0.593	0.503	0.718	0.584	0.771	0.643	0.617	0.768	-0.288	-0.268	-0.469	-0.394	-0.182
Indo-European-Baltic	en-bn	0.721	0.645	0.770	0.741	0.797	0.720	0.190	0.799	-0.299	-0.278	-0.487	-0.410	-0.189
	lt-en	0.601	0.439	0.595	0.432	0.642	0.495	0.414	0.640	-0.253	-0.242	-0.320	-0.260	-0.276
	en-lt	0.295	0.141	0.274	0.095	0.272	0.085	0.065	0.266	0.085	-0.401	-0.478	-0.393	-0.197
Indo-European-Slavic	ru-en	0.774	0.636	0.721	0.613	0.740	0.625	0.653	0.731	-0.332	-0.387	-0.406	-0.361	-0.299
	en-ru	0.763	0.649	0.709	0.544	0.686	0.546	0.581	0.670	-0.095	0.207	-0.042	-0.081	0.236
	fi-en	0.754	0.501	0.625	0.521	0.697	0.575	0.466	0.702	-0.102	-0.181	-0.278	-0.186	-0.139
Uralic	en-fi	0.486	0.290	0.413	0.188	0.459	0.207	0.118	0.452	-0.057	-0.149	-0.365	-0.321	-0.128
	kk-en	0.578	0.314	0.511	0.334	0.592	0.328	0.194	0.593	-0.489	-0.455	-0.594	-0.528	-0.623
Turkic	en-kk	0.534	-0.544	0.556	-0.628	0.621	-0.585	-0.240	0.619	-0.477	-0.569	-0.651	-0.353	0.166
	tr-en	0.655	0.535	0.654	0.482	0.701	0.612	0.511	0.700	-0.262	-0.244	-0.427	-0.359	-0.166
	en-tr	0.442	0.227	0.278	0.205	0.298	0.260	0.218	0.298	-0.171	-0.114	-0.270	-0.188	-0.021
Sino-Tibetan	zh-en	0.629	0.592	0.674	0.592	0.692	0.617	0.609	0.685	-0.204	-0.156	-0.246	-0.289	-0.227
	en-zh	0.513	0.557	0.586	0.016	0.578	0.283	0.288	0.583	-0.095	0.138	-0.110	-0.227	0.137
Finno-Ugric	et-en	0.418	0.528	0.725	0.495	0.719	0.602	0.403	0.725	-0.27164	-0.25252	-0.44237	-0.37183	-0.17185
	en-et	0.242	0.195	0.479	0.179	0.473	0.217	0.101	0.468	-0.17530	-0.16249	-0.28535	-0.23969	-0.11076

Figure 9: Performance Results Evaluated using BERTScore

Translation Performance (WER)														
Language Family	Direction	gemma2-9b	gemma-7b	llama3-70b	llama3-8b	llama-3.1-70b	llama-3.1-8b	mixtral-8x7b	llama-3.2-90b-vision	OLMo-1B	Phi-3.5-mini	Qwen-2.5-0.5B	Qwen-2.5-1.5B	Qwen-2.5-3B
Indo-European-Germanic	de-en	0.564	0.591	0.519	0.568	0.480	1.099	0.720	0.476	3.699	5.002	7.672	7.784	6.251
	en-de	0.663	0.727	0.602	0.662	0.857	2.472	1.257	0.896	3.643	7.988	9.908	10.062	7.715
	en-cs	2.971	0.852	0.761	0.960	0.780	2.601	2.349	0.777	4.736	7.417	13.281	11.863	9.482
Indo-European-Romance	cs-en	1.311	0.716	0.671	0.825	0.553	0.798	1.227	0.548	3.899	6.840	8.648	7.970	6.875
	fr-de	0.700	0.771	0.780	0.817	0.722	4.537	1.455	0.736	2.964	7.640	7.219	7.340	5.927
	de-fr	0.634	0.779	0.717	0.883	0.925	1.946	1.575	0.865	3.604	4.901	6.780	6.840	5.924
Indo-Iranian-Indic (Indo-Aryan)	gu-en	0.695	1.037	0.852	1.173	0.613	4.794	1.247	0.626	2.828	1.995	2.479	3.600	3.087
	en-gu	0.824	1.388	3.087	1.432	0.752	2.251	3.127	0.791	4.292	3.897	12.915	10.065	5.303
	bn-en	0.852	1.290	0.649	0.854	0.554	0.944	1.008	0.560	4.156	2.873	7.873	3.646	9.388
Indo-European-Baltic	en-bn	0.882	2.081	1.734	0.858	0.715	1.756	3.140	0.701	3.873	2.739	5.679	2.568	7.678
	lt-en	0.734	0.886	0.799	0.981	0.706	1.097	1.441	0.717	3.527	4.720	6.708	6.580	5.761
	en-lt	0.863	0.951	0.930	1.053	1.105	4.047	2.518	1.135	3.490	7.808	11.787	11.653	8.071
Indo-European-Slavic	ru-en	0.473	0.667	0.619	0.687	0.574	0.852	0.795	0.582	3.277	4.813	6.758	6.355	4.958
	en-ru	0.580	0.799	0.703	0.791	1.200	2.336	1.910	1.374	46.168	6.003	9.132	9.427	7.555
	fi-en	0.502	0.822	0.752	0.855	0.640	1.077	1.614	0.633	3.298	5.237	7.555	7.167	5.992
Uralic	en-fi	0.659	0.878	2.471	1.058	0.827	2.267	4.051	0.946	4.241	9.390	13.591	13.621	9.991
	kk-en	0.801	1.168	1.008	1.167	0.812	3.843	1.580	0.801	2.792	4.699	4.175	4.813	3.498
	en-kk	0.939	1.131	0.967	2.571	0.903	3.551	3.478	0.899	5.942	5.203	10.956	10.322	6.067
Turkic	tr-en	0.745	0.900	0.748	1.012	0.675	0.829	1.292	0.679	6.687	5.678	8.784	8.789	3.445
	en-tr	0.751	0.765	0.561	0.899	0.523	0.731	1.102	0.549	5.902	4.832	7.489	7.232	2.754
Sino-Tibetan	zh-en	0.760	0.732	0.628	0.724	0.620	0.827	0.820	0.633	1.415	1.944	4.483	3.564	3.396
	en-zh	1.076	1.371	1.427	10.378	6.259	30.965	33.504	5.930	46.168	57.253	99.063	107.203	54.252
Finno-Ugric	et-en	0.946	1.130	0.554	0.914	0.576	1.024	2.015	0.560	54.556	44.556	67.453	88.345	53.655
	en-et	1.662	1.782	0.964	1.592	0.973	1.743	3.289	0.971	87.288	76.199	113.474	147.208	93.432

Figure 10: Performance Results Evaluated using WER

Translation Performance (CER)														
Language Family	Direction	gemma2-9b	gemma-7b	llama3-70b	llama3-8b	llama-3.1-70b	llama-3.1-8b	mixtral-8x7b	llama-3.2-90b-vision	OLMo-1B	Phi-3.5-mini	Qwen-2.5-0.5B	Qwen-2.5-1.5B	Qwen-2.5-3B
Indo-European-Germanic	de-en	0.411	0.439	0.383	0.419	0.345	0.967	0.576	0.344	3.559	5.210	7.742	7.931	6.289
	en-de	0.462	0.520	0.426	0.467	0.665	2.197	0.988	0.711	2.537	6.861	7.948	8.222	6.534
	en-cs	2.542	0.616	0.542	0.675	0.565	2.392	2.001	0.560	3.719	6.701	11.648	10.588	8.418
	cs-en	1.013	0.541	0.503	0.631	0.406	0.640	1.061	0.402	3.450	6.944	8.607	8.002	6.829
Indo-European-Romance	fr-de	0.508	0.558	0.576	0.580	0.527	4.115	1.171	0.537	2.165	6.413	5.893	6.088	5.073
	de-fr	0.447	0.556	0.507	0.625	0.705	1.698	1.325	0.644	3.175	4.827	6.487	6.644	5.771
Indo-Iranian-Indic (Indo-Aryan)	gu-en	0.511	0.789	0.657	0.874	0.426	4.830	1.021	0.432	2.472	1.976	2.184	3.330	2.841
	en-gu	0.573	1.219	1.813	1.135	0.487	2.062	3.055	0.521	3.865	3.696	12.399	9.532	4.778
	bn-en	0.621	0.966	0.464	0.613	0.391	0.728	0.728	0.394	2.812	2.476	2.327	2.315	2.629
	en-bn	0.619	1.648	1.227	0.613	0.483	1.491	2.697	0.468	4.122	3.788	4.116	3.509	3.979
Indo-European-Baltic	lt-en	0.553	0.676	0.620	0.777	0.530	0.915	1.238	0.538	3.307	4.975	6.665	6.756	5.902
	en-lt	0.590	0.669	0.629	0.702	0.783	3.293	1.891	0.806	2.712	6.124	8.699	8.740	6.359
Indo-European-Slavic	ru-en	0.336	0.501	0.459	0.517	0.420	0.688	0.617	0.426	3.043	5.017	6.518	6.297	4.843
	en-ru	0.404	0.585	0.509	0.568	0.924	1.844	1.508	1.087	7.954	4.855	7.429	7.782	6.325
Uralic	fi-en	0.361	0.627	0.587	0.663	0.462	0.900	1.429	0.453	3.302	5.718	7.613	7.467	6.156
	en-fi	0.451	0.600	2.328	0.688	0.556	1.790	2.777	0.636	2.490	6.588	8.785	8.926	6.934
	kk-en	0.597	0.923	0.777	0.871	0.591	3.043	1.289	0.585	2.549	4.606	4.231	4.728	3.725
Turkic	en-kk	0.697	0.958	0.725	2.357	0.669	2.883	2.912	0.666	3.941	4.019	8.599	8.180	4.852
	tr-en	0.552	0.702	0.571	0.806	0.503	0.641	1.084	0.503	4.122	3.788	4.116	3.509	3.979
	en-tr	0.544	1.890	1.730	2.150	1.620	1.920	2.620	1.670	5.870	5.440	5.420	4.730	5.640
Sino-Tibetan	zh-en	0.576	0.554	0.453	0.536	0.446	0.643	0.615	0.455	1.340	2.114	4.586	3.942	3.634
	en-zh	0.676	0.623	0.681	1.898	1.418	5.677	9.213	1.308	7.954	10.710	17.551	18.484	11.492
Finno-Ugric	et-en	0.793	0.988	0.427	0.789	0.458	0.961	2.062	0.441	4.122	3.788	4.116	3.509	3.979
	en-et	1.990	2.070	1.630	2.010	1.540	2.090	3.240	1.590	5.770	5.320	5.470	4.720	5.460

Translation Performance (ROUGE-1)														
Language Family	Direction	gemma2-9b	gemma-7b	llama3-70b	llama3-8b	llama-3.1-70b	llama-3.1-8b	mixtral-8x7b	llama-3.2-90b-vision	OLMo-1B	Phi-3.5-mini	Qwen-2.5-0.5B	Qwen-2.5-1.5B	Qwen-2.5-3B
Indo-European-Germanic	de-en	0.704	0.656	0.727	0.696	0.751	0.631	0.672	0.748	0.281	0.263	0.137	0.160	0.208
	en-de	0.608	0.551	0.626	0.596	0.641	0.369	0.539	0.636	0.148	0.161	0.078	0.084	0.147
	en-cs	0.322	0.421	0.549	0.348	0.583	0.310	0.347	0.580	0.073	0.110	0.034	0.063	0.110
Indo-European-Romance	cs-en	0.470	0.559	0.674	0.585	0.706	0.626	0.544	0.707	0.170	0.147	0.096	0.122	0.175
	fr-de	0.538	0.468	0.499	0.403	0.526	0.313	0.409	0.535	0.078	0.110	0.066	0.082	0.151
	de-fr	0.635	0.524	0.602	0.415	0.594	0.390	0.490	0.599	0.103	0.206	0.109	0.109	0.165
Indo-Iranian-Indic (Indo-Aryan)	gu-en	0.682	0.470	0.657	0.490	0.722	0.529	0.407	0.724	0.146	0.137	0.111	0.139	0.153
	en-gu	0.081	0.143	0.103	0.027	0.087	0.012	0.017	0.060	0.014	0.019	0.005	0.007	0.019
	bn-en	0.569	0.544	0.711	0.629	0.760	0.649	0.617	0.757	0.037	0.003	0.049	0.306	0.179
Indo-European-Baltic	en-bn	0.458	0.380	0.547	0.502	0.592	0.495	0.497	0.510	0.024	0.001	0.028	0.269	0.151
	lt-en	0.560	0.472	0.572	0.486	0.611	0.518	0.418	0.611	0.129	0.152	0.090	0.122	0.127
	en-lt	0.420	0.293	0.408	0.272	0.396	0.251	0.198	0.396	0.008	0.043	0.030	0.032	0.066
Indo-European-Slavic	ru-en	0.754	0.603	0.674	0.619	0.694	0.621	0.627	0.694	0.333	0.288	0.202	0.219	0.299
	en-ru	0.306	0.099	0.127	0.091	0.110	0.062	0.103	0.107	0.023	0.007	0.004	0.004	0.006
	fi-en	0.707	0.481	0.591	0.540	0.641	0.548	0.454	0.643	0.241	0.164	0.088	0.130	0.171
Uralic	en-fi	0.541	0.344	0.467	0.319	0.502	0.300	0.189	0.495	0.398	0.261	0.133	0.171	0.205
	kk-en	0.519	0.360	0.481	0.395	0.526	0.365	0.256	0.530	0.123	0.131	0.111	0.139	0.119
Turkic	en-kk	0.232	0.018	0.122	0.016	0.135	0.005	0.014	0.135	0.043	0.006	0.003	0.003	0.008
	tr-en	0.603	0.492	0.600	0.504	0.641	0.580	0.482	0.638	0.133	0.072	0.024	0.033	0.066
	en-tr	0.521	0.392	0.509	0.427	0.521	0.448	0.391	0.523	0.081	0.041	0.014	0.020	0.033
Sino-Tibetan	zh-en	0.582	0.580	0.674	0.623	0.679	0.638	0.608	0.679	0.390	0.400	0.219	0.276	0.308
	en-zh	0.136	0.275	0.254	0.151	0.290	0.094	0.083	0.299	0.023	0.019	0.014	0.014	0.025
	et-en	0.521	0.406	0.662	0.522	0.660	0.574	0.420	0.667	0.678	0.867	0.235	0.786	0.762
Finno-Ugric	en-et	0.399	0.302	0.556	0.440	0.576	0.496	0.362	0.563	0.575	0.752	0.125	0.691	0.619

Figure 12: Performance Results Evaluated using ROUGE-1

Translation Performance (ROUGE-2)														
Language Family	Direction	gemma2-9b	gemma-7b	llama3-70b	llama3-8b	llama-3.1-70b	llama-3.1-8b	mixtral-8x7b	llama-3.2-90b-vision	OLMo-1B	Phi-3.5-mini	Qwen-2.5-0.5B	Qwen-2.5-1.5B	Qwen-2.5-3B
Indo-European-Germanic	de-en	0.481	0.411	0.500	0.466	0.529	0.439	0.460	0.526	0.146	0.156	0.057	0.079	0.120
	en-de	0.374	0.332	0.411	0.365	0.452	0.232	0.356	0.446	0.060	0.092	0.027	0.034	0.073
	en-cs	0.179	0.163	0.303	0.172	0.344	0.148	0.173	0.336	0.011	0.031	0.004	0.011	0.035
Indo-European-Romance	cs-en	0.301	0.286	0.417	0.343	0.473	0.387	0.321	0.471	0.059	0.064	0.022	0.043	0.083
	fr-de	0.322	0.260	0.294	0.221	0.322	0.178	0.233	0.331	0.037	0.055	0.029	0.037	0.074
	de-fr	0.442	0.317	0.422	0.254	0.417	0.248	0.321	0.419	0.041	0.123	0.044	0.050	0.092
Indo-Iranian-Indic (Indo-Aryan)	gu-en	0.412	0.274	0.393	0.254	0.459	0.311	0.158	0.462	0.019	0.002	0.004	0.017	0.032
	en-gu	0.022	0.044	0.050	0.014	0.030	0.000	0.003	0.020	0.002	0.003	0.001	0.002	0.005
	bn-en	0.350	0.304	0.464	0.375	0.526	0.408	0.360	0.522	0.199	0.210	0.173	0.185	0.191
Indo-European-Baltic	en-bn	0.045	0.034	0.104	0.060	0.121	0.067	0.064	0.129	0.030	0.019	0.016	0.018	0.021
	lt-en	0.282	0.250	0.300	0.225	0.345	0.250	0.187	0.338	0.029	0.037	0.014	0.029	0.037
	en-lt	0.189	0.122	0.168	0.087	0.161	0.084	0.054	0.159	0.001	0.010	0.005	0.006	0.013
Indo-European-Slavic	ru-en	0.538	0.331	0.422	0.349	0.433	0.382	0.389	0.437	0.161	0.139	0.069	0.093	0.156
	en-ru	0.164	0.041	0.041	0.028	0.032	0.045	0.039	0.031	0.008	0.001	0.000	0.001	0.001
	fi-en	0.469	0.215	0.323	0.268	0.381	0.292	0.224	0.392	0.094	0.054	0.011	0.034	0.067
Uralic	en-fi	0.327	0.137	0.249	0.143	0.278	0.142	0.071	0.277	0.046	0.019	0.003	0.006	0.019
	kk-en	0.258	0.165	0.227	0.144	0.264	0.155	0.067	0.266	0.013	0.014	0.007	0.016	0.012
	en-kk	0.128	0.018	0.033	0.001	0.040	0.002	0.003	0.040	0.011	0.001	0.001	0.001	0.003
Turkic	tr-en	0.333	0.217	0.327	0.241	0.376	0.305	0.213	0.371	0.184	0.192	0.167	0.174	0.181
	en-tr	0.317	0.044	0.096	0.062	0.118	0.090	0.057	0.125	0.030	0.018	0.014	0.016	0.020
Sino-Tibetan	zh-en	0.314	0.305	0.415	0.351	0.434	0.379	0.345	0.429	0.172	0.211	0.076	0.120	0.155
	en-zh	0.025	0.142	0.139	0.070	0.181	0.045	0.040	0.178	0.008	0.008	0.004	0.005	0.012
	et-en	0.351	0.191	0.429	0.289	0.434	0.360	0.227	0.438	0.176	0.184	0.161	0.168	0.173
Finno-Ugric	en-et	0.041	0.059	0.111	0.070	0.127	0.097	0.070	0.133	0.026	0.020	0.017	0.019	0.022

Figure 13: Performance Results Evaluated using ROUGE-2

Translation Performance (ROUGE-L)														
Language Family	Direction	gemma2-9b	gemma-7b	llama3-70b	llama3-8b	llama-3.1-70b	llama-3.1-8b	mixtral-8x7b	llama-3.2-90b-vision	OLMo-1B	Phi-3.5-mini	Qwen-2.5-0.5B	Qwen-2.5-1.5B	Qwen-2.5-3B
Indo-European-Germanic	de-en	0.651	0.599	0.671	0.642	0.693	0.585	0.625	0.692	0.241	0.231	0.109	0.128	0.182
	en-de	0.544	0.503	0.581	0.549	0.602	0.337	0.504	0.599	0.131	0.148	0.064	0.070	0.128
	en-cs	0.288	0.368	0.496	0.312	0.523	0.265	0.307	0.525	0.062	0.091	0.026	0.045	0.087
Indo-European-Romance	cs-en	0.430	0.513	0.618	0.530	0.657	0.572	0.500	0.658	0.139	0.125	0.071	0.097	0.149
	fr-de	0.496	0.433	0.461	0.374	0.488	0.276	0.376	0.499	0.074	0.100	0.057	0.069	0.131
	de-fr	0.592	0.482	0.561	0.384	0.557	0.345	0.448	0.562	0.089	0.185	0.083	0.088	0.144
Indo-Iranian-Indic (Indo-Aryan)	gu-en	0.594	0.394	0.571	0.403	0.646	0.451	0.339	0.647	0.123	0.116	0.095	0.106	0.126
	en-gu	0.143	0.081	0.103	0.027	0.087	0.012	0.017	0.060	0.013	0.018	0.005	0.007	0.019
	bn-en	0.481	0.479	0.616	0.524	0.672	0.559	0.518	0.669	0.313	0.279	0.233	0.251	0.266
Indo-European-Baltic	en-bn	0.057	0.044	0.117	0.064	0.140	0.103	0.078	0.146	0.033	0.022	0.019	0.020	0.028
	lt-en	0.484	0.403	0.489	0.404	0.527	0.438	0.348	0.525	0.100	0.116	0.066	0.087	0.098
	en-lt	0.379	0.255	0.353	0.241	0.347	0.209	0.173	0.349	0.008	0.040	0.025	0.027	0.054
Indo-European-Slavic	ru-en	0.702	0.543	0.622	0.554	0.639	0.565	0.572	0.641	0.286	0.242	0.150	0.170	0.254
	en-ru	0.304	0.099	0.123	0.087	0.106	0.062	0.103	0.103	0.023	0.007	0.004	0.004	0.006
	fi-en	0.666	0.421	0.523	0.471	0.573	0.487	0.401	0.584	0.200	0.140	0.066	0.097	0.138
Uralic	en-fi	0.517	0.312	0.419	0.289	0.456	0.269	0.166	0.449	0.121	0.065	0.020	0.029	0.056
	kk-en	0.453	0.301	0.406	0.323	0.457	0.300	0.210	0.460	0.100	0.104	0.086	0.097	0.096
	en-kk	0.227	0.063	0.121	0.015	0.135	0.005	0.013	0.135	0.041	0.006	0.003	0.003	0.008
Turkic	tr-en	0.522	0.414	0.512	0.412	0.556	0.484	0.391	0.556	0.243	0.199	0.168	0.202	0.229
	en-tr	0.460	0.087	0.155	0.074	0.194	0.121	0.069	0.189	0.036	0.015	0.010	0.012	0.026
	zh-en	0.499	0.503	0.602	0.545	0.615	0.561	0.533	0.609	0.321	0.337	0.156	0.213	0.250
Sino-Tibetan	en-zh	0.136	0.271	0.247	0.151	0.288	0.094	0.081	0.297	0.023	0.019	0.014	0.014	0.025
	et-en	0.491	0.369	0.629	0.481	0.630	0.544	0.391	0.636	0.289	0.235	0.199	0.211	0.246
Finno-Ugric	en-et	0.058	0.089	0.139	0.072	0.180	0.130	0.067	0.183	0.041	0.022	0.012	0.019	0.032

ru-kk									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	3.777	24.114	94.329	0.474	0.946	0.754	0.010	0.010	0.010
deepseek-r1-distill-70b	8.007	36.983	86.326	0.592	0.879	0.652	0.010	0.010	0.010
llama-3.3-70b-specdec	9.161	39.425	84.247	0.617	0.858	0.632	0.010	0.010	0.010
qwen-2.5-32b	2.290	27.211	144.297	0.495	1.453	1.238	0.000	0.000	0.000
llama-3.3-70b-versatile	9.407	39.507	83.680	0.619	0.853	0.620	0.010	0.010	0.010
llama-3.1-8b	1.334	18.016	414.745	0.530	4.156	4.112	0.010	0.010	0.010
mixtral-8x7b	0.568	19.977	204.096	0.309	2.044	1.728	0.001	0.001	0.001
llama-3.2-90b-vision	10.291	40.593	83.617	0.625	0.853	0.625	0.010	0.010	0.010
kk-ru									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	5.870	30.073	97.753	0.519	0.991	0.747	0.027	0.020	0.027
deepseek-r1-distill-70b	11.609	40.345	84.920	0.615	0.869	0.629	0.030	0.020	0.030
llama-3.3-70b-specdec	13.802	41.609	81.963	0.630	0.849	0.610	0.030	0.020	0.030
qwen-2.5-32b	7.605	32.995	103.548	0.559	1.058	0.836	0.030	0.020	0.030
llama-3.3-70b-versatile	13.811	41.094	82.555	0.625	0.852	0.615	0.030	0.020	0.030
llama-3.1-8b	8.319	34.384	91.898	0.571	0.943	0.700	0.030	0.020	0.030
mixtral-8x7b	2.368	23.864	159.787	0.420	1.611	1.303	0.030	0.020	0.030
llama-3.2-90b-vision	15.356	42.903	80.367	0.633	0.830	0.608	0.030	0.020	0.030

Figure 15: Performance in the **Literature** Domain Across the **ru** ↔ **kk** (Russian–Kazakh)

en-kk									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	1.321	22.386	96.535	0.459	0.974	0.755	0.021	0.010	0.021
deepseek-r1-distill-70b	3.245	34.267	94.059	0.557	0.955	0.694	0.023	0.010	0.023
llama-3.3-70b-specdec	5.850	35.453	90.319	0.575	0.914	0.661	0.018	0.010	0.018
qwen-2.5-32b	1.029	24.076	172.387	0.421	1.733	1.500	0.010	0.003	0.010
llama-3.3-70b-versatile	5.820	35.702	90.099	0.574	0.914	0.658	0.018	0.010	0.018
llama-3.1-8b	0.241	12.235	657.041	0.480	6.577	7.230	0.018	0.010	0.018
mixtral-8x7b	0.714	17.023	173.542	0.243	1.737	1.450	0.017	0.000	0.017
llama-3.2-90b-vision	5.054	36.582	89.549	0.581	0.910	0.665	0.018	0.010	0.018
kk-en									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	9.502	33.803	88.602	0.311	0.921	0.685	0.378	0.129	0.305
deepseek-r1-distill-70b	22.801	47.912	71.695	0.528	0.752	0.556	0.549	0.276	0.478
llama-3.3-70b-specdec	24.938	49.820	68.465	0.550	0.720	0.517	0.573	0.306	0.508
qwen-2.5-32b	13.640	37.371	84.878	0.378	0.880	0.649	0.413	0.161	0.343
llama-3.3-70b-versatile	24.328	49.997	68.693	0.557	0.721	0.524	0.574	0.310	0.507
llama-3.1-8b	18.145	41.756	77.356	0.455	0.802	0.592	0.480	0.221	0.411
mixtral-8x7b	7.090	29.359	91.983	0.270	0.941	0.713	0.304	0.080	0.244
llama-3.2-90b-vision	25.546	49.832	67.933	0.549	0.716	0.523	0.573	0.316	0.504

Figure 16: Performance in the **Literature** Domain Across the **en** ↔ **kk** (English–Kazakh)

ru-en									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	28.953	56.048	56.802	0.662	0.595	0.431	0.642	0.376	0.600
deepseek-r1-distill-70b	29.171	55.976	57.137	0.663	0.596	0.420	0.643	0.371	0.602
llama-3.3-70b-specdec	31.632	58.053	54.584	0.680	0.574	0.398	0.662	0.414	0.625
qwen-2.5-32b	30.553	56.944	57.555	0.668	0.602	0.436	0.651	0.394	0.609
llama-3.3-70b-versatile	31.481	57.880	55.588	0.680	0.584	0.402	0.661	0.415	0.623
llama-3.1-8b	26.786	53.003	59.607	0.636	0.627	0.448	0.621	0.357	0.575
mixtral-8x7b	27.035	54.175	58.267	0.636	0.613	0.443	0.631	0.363	0.586
llama-3.2-90b-vision	30.851	57.632	55.170	0.687	0.579	0.400	0.662	0.409	0.625
en-ru									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	14.934	42.545	74.502	0.646	0.773	0.562	0.618	0.353	0.555
deepseek-r1-distill-70b	14.888	41.675	74.281	0.651	0.768	0.552	0.597	0.347	0.564
llama-3.3-70b-specdec	18.275	45.157	71.571	0.680	0.740	0.540	0.615	0.366	0.608
qwen-2.5-32b	14.214	42.089	92.257	0.647	0.942	0.772	0.615	0.349	0.567
llama-3.3-70b-versatile	18.740	45.189	71.239	0.680	0.738	0.540	0.638	0.375	0.583
llama-3.1-8b	15.908	42.914	73.507	0.661	0.759	0.554	0.579	0.329	0.530
mixtral-8x7b	8.570	36.329	141.980	0.567	1.438	1.233	0.616	0.316	0.568
llama-3.2-90b-vision	18.303	45.236	70.299	0.680	0.727	0.526	0.627	0.370	0.601

Figure 17: Performance in the **Literature** Domain Across the **ru** ↔ **en** (Russian–English)

en-fr									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	34.393	64.304	49.614	0.600	0.514	0.304	0.683	0.481	0.658
deepseek-r1-distill-70b	33.167	63.824	51.591	0.593	0.540	0.327	0.674	0.466	0.641
llama-3.3-70b-specdec	34.907	65.355	48.120	0.608	0.500	0.298	0.688	0.490	0.661
qwen-2.5-32b	31.423	63.116	58.149	0.573	0.600	0.395	0.656	0.465	0.630
llama-3.3-70b-versatile	34.839	65.354	48.120	0.608	0.501	0.297	0.691	0.490	0.662
llama-3.1-8b	32.547	62.964	53.520	0.577	0.554	0.345	0.666	0.454	0.635
mixtral-8x7b	32.029	63.021	54.436	0.578	0.564	0.356	0.659	0.459	0.626
llama-3.2-90b-vision	35.486	66.138	47.782	0.621	0.499	0.295	0.698	0.506	0.674
en-hr									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	14.076	42.701	72.869	0.693	0.766	0.552	0.240	0.120	0.240
deepseek-r1-distill-70b	22.187	52.016	62.733	0.741	0.666	0.452	0.147	0.033	0.147
llama-3.3-70b-specdec	27.698	55.625	55.724	0.767	0.594	0.395	0.083	0.037	0.083
qwen-2.5-32b	14.034	43.325	84.871	0.699	0.880	0.698	0.223	0.090	0.223
llama-3.3-70b-versatile	27.645	55.725	55.421	0.768	0.586	0.390	0.083	0.037	0.083
llama-3.1-8b	15.363	42.010	85.930	0.674	0.889	0.660	0.198	0.110	0.198
mixtral-8x7b	4.049	27.953	191.024	0.493	1.934	1.799	0.069	0.029	0.069
llama-3.2-90b-vision	29.868	57.505	52.648	0.781	0.562	0.375	0.238	0.110	0.238
en-it									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	29.829	61.656	54.446	0.585	0.564	0.348	0.599	0.361	0.573
deepseek-r1-distill-70b	31.950	63.043	51.883	0.601	0.530	0.327	0.614	0.388	0.598
llama-3.3-70b-specdec	33.162	64.009	51.987	0.624	0.531	0.317	0.627	0.408	0.614
qwen-2.5-32b	28.619	61.414	58.734	0.569	0.600	0.373	0.581	0.354	0.566
llama-3.3-70b-versatile	33.515	63.954	51.778	0.621	0.528	0.318	0.625	0.406	0.613
llama-3.1-8b	30.526	62.003	55.178	0.594	0.565	0.350	0.606	0.387	0.591
mixtral-8x7b	26.040	60.123	64.644	0.558	0.663	0.425	0.575	0.350	0.554
llama-3.2-90b-vision	34.124	64.116	50.575	0.623	0.518	0.312	0.634	0.419	0.622

Figure 18: Performance in the **Medical** Domain Across the **en** → **fr, hr, it** (English–French, Croatian, Italian)

en-pl									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	9.854	38.547	80.051	0.610	0.809	0.528	0.230	0.059	0.229
deepseek-r1-distill-70b	10.963	40.145	79.923	0.637	0.805	0.509	0.195	0.055	0.194
llama-3.3-70b-specdec	12.035	42.813	77.046	0.652	0.776	0.477	0.183	0.044	0.183
qwen-2.5-32b	2.220	22.393	432.864	0.491	4.338	4.613	0.134	0.020	0.133
llama-3.3-70b-versatile	12.185	43.532	76.854	0.655	0.774	0.475	0.183	0.044	0.183
llama-3.1-8b	2.202	21.848	413.043	0.594	4.138	3.963	0.184	0.030	0.184
mixtral-8x7b	3.229	29.062	191.240	0.479	1.921	1.649	0.132	0.051	0.130
llama-3.2-90b-vision	12.495	43.322	76.151	0.662	0.769	0.474	0.191	0.052	0.191
en-de									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	22.284	57.053	62.170	0.501	0.643	0.417	0.563	0.303	0.519
deepseek-r1-distill-70b	27.129	61.289	57.279	0.543	0.596	0.373	0.600	0.600	0.555
llama-3.3-70b-specdec	27.560	61.598	55.705	0.547	0.575	0.363	0.603	0.365	0.565
qwen-2.5-32b	23.521	58.752	64.924	0.507	0.673	0.455	0.574	0.326	0.528
llama-3.3-70b-versatile	27.967	61.814	55.987	0.551	0.577	0.359	0.604	0.368	0.568
llama-3.1-8b	23.694	59.217	60.202	0.515	0.626	0.386	0.574	0.325	0.530
mixtral-8x7b	1.778	23.872	1072.569	0.486	10.753	8.049	0.550	0.304	0.506
llama-3.2-90b-vision	28.289	62.664	55.818	0.561	0.581	0.357	0.612	0.376	0.568
en-es									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	39.565	68.614	43.323	0.683	0.683	0.285	0.675	0.459	0.641
deepseek-r1-distill-70b	40.229	69.491	42.779	0.696	0.456	0.277	0.692	0.484	0.660
llama-3.3-70b-specdec	46.048	72.961	37.685	0.740	0.399	0.233	0.724	0.529	0.698
qwen-2.5-32b	41.263	70.069	42.136	0.706	0.443	0.270	0.698	0.493	0.670
llama-3.3-70b-versatile	45.999	72.906	37.883	0.740	0.399	0.232	0.723	0.524	0.694
llama-3.1-8b	40.613	69.648	43.769	0.693	0.456	0.276	0.683	0.482	0.654
mixtral-8x7b	31.125	64.234	63.254	0.626	0.648	0.478	0.632	0.430	0.604
llama-3.2-90b-vision	46.453	73.040	38.032	0.730	0.402	0.235	0.722	0.527	0.695

Figure 19: Performance in the **Medical** Domain Across the **en** $\rightarrow$ **pl, de, es** (English–Polish, German, Spanish)

en-pt									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	35.942	67.262	46.597	0.670	0.477	0.257	0.670	0.439	0.647
deepseek-r1-distill-70b	36.616	67.630	45.393	0.674	0.463	0.249	0.680	0.436	0.653
llama-3.3-70b-specdec	37.131	68.079	44.712	0.680	0.454	0.244	0.683	0.446	0.660
qwen-2.5-32b	36.859	67.692	47.173	0.676	0.481	0.275	0.683	0.446	0.659
llama-3.3-70b-versatile	37.231	68.094	44.607	0.680	0.453	0.243	0.681	0.445	0.659
llama-3.1-8b	36.314	67.578	45.602	0.670	0.465	0.250	0.669	0.434	0.645
mixtral-8x7b	2.158	23.247	1402.723	0.587	14.038	9.964	0.606	0.383	0.582
llama-3.2-90b-vision	37.604	68.212	44.450	0.681	0.453	0.244	0.687	0.450	0.662
en-fi									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	23.466	54.118	64.877	0.704	0.681	0.453	0.294	0.093	0.290
deepseek-r1-distill-70b	24.788	55.791	63.950	0.720	0.661	0.428	0.319	0.110	0.316
llama-3.3-70b-specdec	27.887	59.068	60.106	0.741	0.627	0.400	0.323	0.090	0.320
qwen-2.5-32b	21.811	54.232	79.192	0.693	0.821	0.600	0.327	0.100	0.324
llama-3.3-70b-versatile	27.914	58.943	60.769	0.740	0.632	0.404	0.317	0.090	0.313
llama-3.1-8b	22.262	54.527	68.191	0.702	0.710	0.469	0.312	0.108	0.312
mixtral-8x7b	15.901	51.524	102.187	0.663	1.049	0.755	0.244	0.090	0.244
llama-3.2-90b-vision	28.627	59.499	60.371	0.748	0.624	0.400	0.314	0.110	0.311

Figure 20: Performance in the **Medical** domain across the **en** $\rightarrow$ **pt, fi** (English–Portuguese, Finnish)

en-es									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	49.793	72.698	35.408	0.717	0.388	0.271	0.773	0.598	0.739
deepseek-r1-distill-70b	49.795	73.460	35.259	0.728	0.382	0.263	0.787	0.616	0.755
llama-3.3-70b-specdec	56.196	77.186	29.831	0.776	0.323	0.218	0.822	0.670	0.797
qwen-2.5-32b	50.708	73.889	34.661	0.738	0.370	0.256	0.791	0.626	0.764
llama-3.3-70b-versatile	55.926	77.129	30.129	0.776	0.323	0.217	0.821	0.663	0.792
llama-3.1-8b	50.654	73.697	36.404	0.725	0.382	0.263	0.772	0.614	0.748
mixtral-8x7b	38.293	67.876	56.325	0.654	0.584	0.468	0.714	0.546	0.687
llama-3.2-90b-vision	56.272	77.096	30.179	0.764	0.326	0.221	0.819	0.664	0.792
en-fr									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	42.270	67.462	43.816	0.642	0.455	0.293	0.724	0.542	0.703
deepseek-r1-distill-70b	40.811	67.098	46.039	0.632	0.487	0.316	0.714	0.523	0.684
llama-3.3-70b-specdec	42.757	68.511	42.464	0.648	0.444	0.287	0.728	0.546	0.704
qwen-2.5-32b	38.553	66.250	52.512	0.615	0.544	0.386	0.699	0.517	0.676
llama-3.3-70b-versatile	42.699	68.540	42.415	0.649	0.443	0.285	0.730	0.545	0.705
llama-3.1-8b	39.926	66.035	47.971	0.616	0.501	0.335	0.706	0.507	0.678
mixtral-8x7b	39.641	66.132	49.082	0.620	0.511	0.347	0.695	0.511	0.666
llama-3.2-90b-vision	43.136	69.283	42.415	0.661	0.444	0.284	0.736	0.561	0.715
en-pl									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	18.750	45.670	68.129	0.701	0.728	0.543	0.240	0.120	0.240
deepseek-r1-distill-70b	29.174	56.101	56.278	0.754	0.607	0.439	0.147	0.033	0.147
llama-3.3-70b-specdec	35.810	60.189	48.260	0.780	0.525	0.379	0.083	0.037	0.083
qwen-2.5-32b	17.468	46.199	80.182	0.708	0.835	0.690	0.223	0.090	0.223
llama-3.3-70b-versatile	35.609	60.222	47.857	0.781	0.518	0.374	0.083	0.037	0.083
llama-3.1-8b	20.047	44.991	81.089	0.683	0.841	0.653	0.198	0.110	0.198
mixtral-8x7b	5.255	29.282	188.149	0.498	1.906	1.814	0.069	0.029	0.069
llama-3.2-90b-vision	37.987	62.130	45.436	0.793	0.495	0.359	0.238	0.110	0.238

Figure 21: Performance in the **Law Domain** Across the **en** → **es, fr, pl** (English–Spanish, French, Polish)

en-it									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	37.283	65.623	47.330	0.625	0.494	0.331	0.682	0.457	0.653
deepseek-r1-distill-70b	40.421	67.283	44.686	0.644	0.459	0.307	0.701	0.491	0.682
llama-3.3-70b-specdec	42.153	68.289	44.531	0.664	0.457	0.298	0.714	0.517	0.698
qwen-2.5-32b	37.304	65.548	51.115	0.609	0.525	0.356	0.665	0.456	0.650
llama-3.3-70b-versatile	42.376	68.201	44.479	0.662	0.456	0.299	0.711	0.515	0.696
llama-3.1-8b	38.825	66.032	48.160	0.632	0.495	0.332	0.684	0.483	0.667
mixtral-8x7b	33.178	64.175	57.232	0.597	0.589	0.408	0.650	0.440	0.628
llama-3.2-90b-vision	42.649	68.357	43.339	0.664	0.446	0.293	0.719	0.527	0.704
en-pt									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	47.547	71.923	37.435	0.710	0.387	0.243	0.773	0.595	0.752
deepseek-r1-distill-70b	48.387	72.169	36.440	0.715	0.375	0.236	0.784	0.594	0.759
llama-3.3-70b-specdec	48.510	72.633	35.707	0.723	0.364	0.230	0.787	0.608	0.767
qwen-2.5-32b	47.707	72.145	38.325	0.716	0.393	0.263	0.785	0.607	0.765
llama-3.3-70b-versatile	48.648	72.667	35.602	0.722	0.363	0.229	0.785	0.608	0.766
llama-3.1-8b	48.360	72.216	36.754	0.710	0.377	0.236	0.768	0.594	0.749
mixtral-8x7b	2.792	24.555	1394.712	0.622	13.958	10.039	0.691	0.511	0.671
llama-3.2-90b-vision	48.964	72.724	35.445	0.722	0.364	0.231	0.788	0.608	0.768
en-ro									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	28.201	56.873	59.907	0.719	0.638	0.442	0.294	0.093	0.290
deepseek-r1-distill-70b	30.368	59.003	58.449	0.736	0.611	0.414	0.319	0.110	0.316
llama-3.3-70b-specdec	33.907	62.300	54.473	0.759	0.577	0.386	0.323	0.090	0.320
qwen-2.5-32b	25.447	56.871	74.420	0.707	0.781	0.591	0.327	0.100	0.324
llama-3.3-70b-versatile	33.594	61.900	55.335	0.756	0.584	0.392	0.317	0.090	0.313
llama-3.1-8b	26.573	57.046	63.552	0.716	0.670	0.459	0.312	0.108	0.312
mixtral-8x7b	19.553	54.309	96.355	0.676	0.998	0.748	0.244	0.090	0.244
llama-3.2-90b-vision	34.650	62.446	55.136	0.765	0.577	0.388	0.314	0.110	0.311

Figure 22: Performance in the **Law Domain** Across the **en** → **it, pt, ro** (English–Italian, Portuguese, Romanian)

en-el									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	11.436	40.222	75.928	0.619	0.770	0.522	0.230	0.059	0.229
deepseek-r1-distill-70b	12.851	42.374	75.032	0.647	0.757	0.503	0.197	0.058	0.196
llama-3.3-70b-specdec	14.739	45.388	70.487	0.664	0.712	0.469	0.183	0.044	0.183
qwen-2.5-32b	2.654	23.077	428.937	0.497	4.300	4.702	0.134	0.020	0.134
llama-3.3-70b-versatile	14.491	46.037	70.743	0.666	0.714	0.467	0.183	0.044	0.183
llama-3.1-8b	2.533	22.719	408.963	0.601	4.099	4.035	0.184	0.030	0.184
mixtral-8x7b	3.705	30.058	188.348	0.482	1.893	1.675	0.132	0.051	0.130
llama-3.2-90b-vision	15.288	45.742	69.974	0.672	0.708	0.466	0.191	0.052	0.191

en-de									
	BLEU	chrF	TER	BERTScore	WER	CER	ROUGE-1	ROUGE-2	ROUGE-L
deepseek-r1-distill-32b	30.784	60.879	54.704	0.542	0.571	0.403	0.646	0.407	0.599
deepseek-r1-distill-70b	36.377	65.401	49.352	0.589	0.521	0.357	0.689	0.465	0.640
llama-3.3-70b-specdec	37.010	65.871	47.437	0.596	0.496	0.347	0.697	0.477	0.654
qwen-2.5-32b	32.289	62.624	57.127	0.553	0.601	0.442	0.666	0.439	0.618
llama-3.3-70b-versatile	37.324	66.065	47.606	0.599	0.497	0.342	0.697	0.480	0.657
llama-3.1-8b	33.635	63.223	52.225	0.562	0.549	0.372	0.665	0.438	0.618
mixtral-8x7b	2.447	25.195	1067.211	0.530	10.705	8.184	0.637	0.410	0.589
llama-3.2-90b-vision	37.713	67.009	47.549	0.610	0.503	0.340	0.709	0.492	0.657

Figure 23: Performance in the **Law** domain across the **en** $\rightarrow$ **el, de** (English–Greek, German)

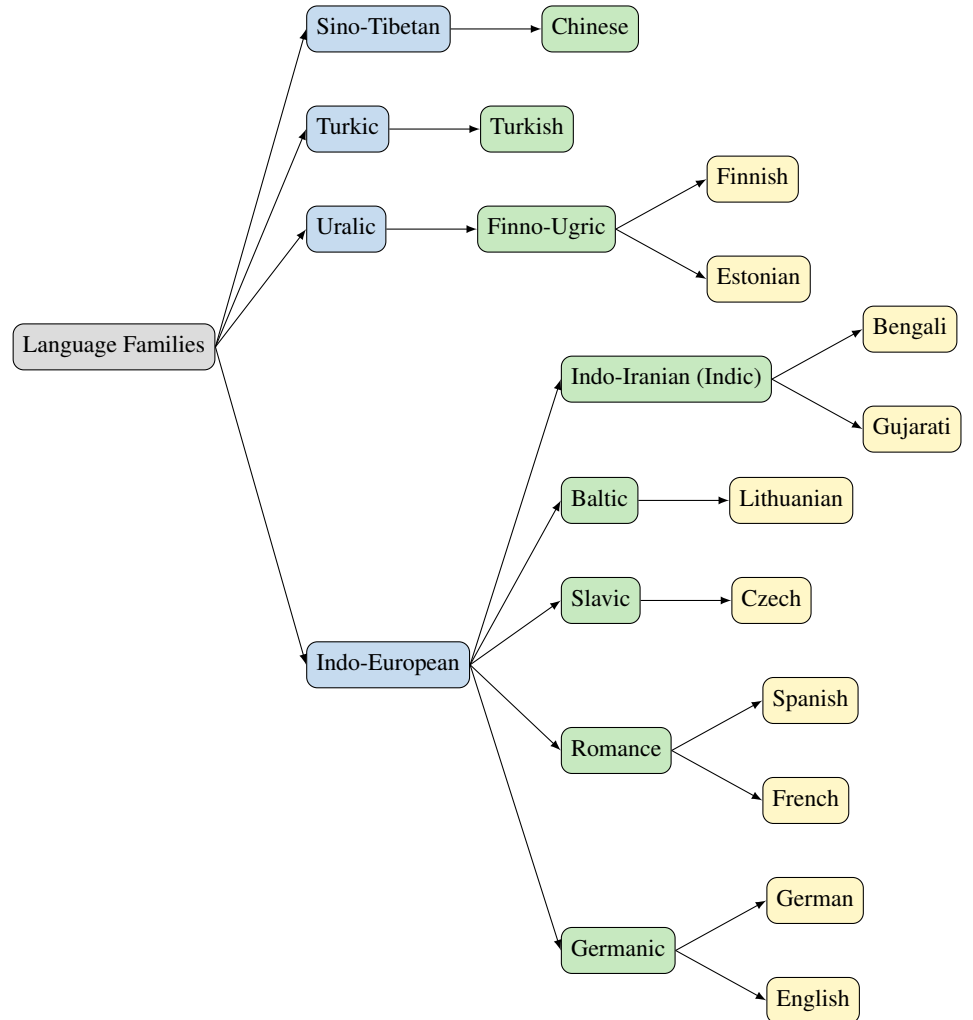


Figure 24: **Language Family Tree (Pellard et al., 2024)**. The hierarchical structure shows the evolution of languages from families to sub-families and individual languages. **Level 0** denotes the root node, **Level 1** indicates the major language families (e.g., Indo-European, Uralic), **Level 2** represents sub-families (e.g., Germanic, Romance), and **Level 3** lists the individual languages (e.g., English, Spanish).