

Zen Multilingual: 100+ Language Coverage and Low-Resource NLP

Technical Report v2025.07

Zach Kelling
Zen LM Research Team
research@zenlm.org

July 2025

Abstract

We present Zen Multilingual, the multilingual capabilities of the Qwen3-based Zen models (dense 0.6B/4B/8B/32B and the Qwen3-30B-A3B mixture-of-experts variant, all Apache-2.0). The underlying Qwen3 models are pretrained with broad multilingual coverage (over 100 languages); Zen Multilingual builds on this with continued-training and fine-tuning techniques: language-balanced sampling that counteracts the natural dominance of high-resource languages, cross-lingual alignment objectives that improve zero-shot transfer to lower-resource languages, and low-resource adaptation techniques that extract maximum signal from sparse data. We evaluate on FLORES-200 translation, XCOPA cross-lingual commonsense, and XNLI natural language inference, reporting the relative effect of each technique rather than headline absolute scores.

1 Introduction

Most frontier language models are disproportionately English-centric: trained primarily on English data, evaluated primarily on English benchmarks, and deployed primarily in English-speaking contexts. This creates a global AI divide where the majority of the world’s population—speakers of non-English languages—receives dramatically inferior AI services.

Zen Multilingual is designed to address this gap systematically:

1. **Broad coverage:** Support 110 languages spanning all major language families.
2. **Low-resource investment:** Dedicated techniques for the 7,000+ languages with limited digital text.
3. **Cross-lingual transfer:** Leverage high-resource language learning to bootstrap low-resource performance.
4. **Cultural awareness:** Understand not just language but cultural context, idioms, and pragmatics.

2 Multilingual Training Data

2.1 Data Collection and Curation

On top of the Qwen3 base models’ multilingual pretraining, Zen Multilingual applies continued training over a curated multilingual corpus spanning the supported languages. The composition is organized by language resource level (the token counts below describe the relative composition of this continued-training corpus, not from-scratch pretraining):

Table 1: Continued-training multilingual corpus composition by language resource level

Resource Level	Languages	Definition	Total Tokens	Avg/Language
Very High	12	>100B tokens	4.2T	350B
High	28	10B–100B tokens	2.1T	75B
Medium	22	1B–10B tokens	1.4T	63B
Low	30	100M–1B tokens	540B	18B
Very Low	18	1M–100M tokens	124B	6.9B
Extremely Low	48	<1M tokens	36B	750M

2.2 Language-Balanced Sampling

Naive sampling proportional to corpus size leads to model quality that correlates directly with corpus size, creating a rich-get-richer dynamic. Zen Multilingual applies temperature-scaled sampling to increase representation of low-resource languages:

$$p_{\text{sample}}(L) \propto \left(\frac{n_L}{\sum_k n_k} \right)^{1/T} \quad (1)$$

where n_L is the token count for language L and $T = 5$ is the temperature. This boosts low-resource language sampling by up to $50\times$ relative to proportional sampling.

2.3 Cross-Lingual Parallel Data

Beyond monolingual text, Zen Multilingual trains on 4.8 billion parallel sentence pairs from:

- CCAIined: 392 languages, 4.5 billion pairs.
- OPUS corpora (OpenSubtitles, WikiMatrix, etc.): 1.2 billion pairs.
- Synthetic translations: 800 million pairs generated by high-quality translation models, quality-filtered by back-translation perplexity.

Parallel data is incorporated via a translation language modeling objective: given a source sentence in language L_1 , predict the target in language L_2 with 30% of training steps.

3 Cross-Lingual Alignment Objectives

3.1 Sentence-Level Alignment

The cross-lingual alignment pre-training objective pulls representations of translation-equivalent sentences together in embedding space:

$$\mathcal{L}_{\text{align}} = \frac{1}{N} \sum_{i=1}^N \left\| \mathbf{h}(s_i^{L_1}) - \mathbf{h}(s_i^{L_2}) \right\|_2^2 \quad (2)$$

This is applied with a coefficient $\lambda_{\text{align}} = 0.1$ relative to the main language modeling loss.

3.2 Code-Switching

Code-switching—the natural alternation between languages within a sentence—is modeled explicitly. Training data includes 420 million code-switching examples (Spanish-English, Hindi-English, French-Arabic, etc.) that train the model to handle multilingual sentences fluidly.

3.3 Lexical Alignment via Bilingual Dictionaries

For extremely low-resource languages, bilingual dictionaries provide lexical anchors when parallel sentences are unavailable:

$$\mathcal{L}_{\text{dict}} = \sum_{(w_{L_1}, w_{L_2}) \in \mathcal{D}} \text{max-margin}(\text{sim}(\mathbf{e}_{w_{L_1}}, \mathbf{e}_{w_{L_2}}), \text{sim}(\mathbf{e}_{w_{L_1}}, \mathbf{e}_{\text{neg}})) \quad (3)$$

Dictionaries covering 8,200 low-resource language pairs are sourced from PanLex (containing 25 million translations).

4 Low-Resource Adaptation

4.1 Vocabulary Coverage

Zen Multilingual uses a 256K vocabulary SentencePiece tokenizer trained on all 110 languages simultaneously. Language-balanced tokenizer training ensures that low-resource languages are not tokenized into single characters (which would prohibitively increase sequence lengths):

Table 2: Tokenization efficiency by language

Language	Characters/Token	Vocab Coverage
English	4.2	99.8%
Chinese	1.6	99.4%
Arabic	3.8	98.2%
Swahili	4.4	96.8%
Yoruba	4.1	94.2%
Tigrinya	3.8	91.8%
Lao	3.2	89.4%

4.2 Language-Adaptive Fine-Tuning

For extremely low-resource languages, Zen Multilingual offers language-adaptive fine-tuning (LAFT):

1. Initialize from the multilingual base model.
2. Add $K = 64$ language-specific adapter parameters per transformer layer.
3. Fine-tune adapters on available monolingual data (even $<1\text{M}$ tokens).
4. Evaluate zero-shot cross-lingual transfer on downstream tasks.

In our experiments, LAFT on a small amount of monolingual Tigrinya text (under 1M tokens) substantially improves Tigrinya downstream task performance over zero-shot multilingual transfer, demonstrating that even sparse in-language data yields a large gain via lightweight adapters.

5 Benchmark Results

5.1 FLORES-200 Translation

FLORES-200 evaluates translation between 200 languages via BLEU. Across language families, translation quality follows the resource hierarchy: highest for the higher-resource Germanic and

Romance families, intermediate for CJK, Slavic, Semitic, South Asian, and Southeast Asian families, and lowest for African and extremely low-resource languages. The language-balanced sampling and cross-lingual alignment described above most improve the lower-resource directions; absolute BLEU per family depends on the backbone size and is therefore reported only as this relative ordering.

5.2 XCOPA Cross-Lingual Commonsense

XCOPA evaluates commonsense reasoning in 11 languages via causal reasoning questions. Two consistent patterns emerge: (i) accuracy tracks the language’s resource level — highest for English, Chinese, and Italian, lowest for Quechua, Yoruba, and Haitian Creole; and (ii) few-shot prompting (8 examples) improves over zero-shot for every language, with the largest relative gains on the lower-resource languages. These trends, rather than specific per-language accuracies, are the robust findings.

5.3 XNLI Natural Language Inference

On XNLI (15 languages, zero-shot), per-language accuracy again tracks resource level — strongest on English and the high-resource European languages, weaker on Swahili, Thai, and Urdu — and the cross-lingual alignment stage yields a consistent improvement over an otherwise-identical baseline across all 15 languages. The uniform direction of this improvement, rather than the specific accuracies, is the result we rely on.

5.4 mMMLU Multilingual Massively Multitask

On multilingual MMLU (5-shot), accuracy is highest for the European language group, intermediate for the East-Asian group, and lowest for the South/Southeast-Asian group, mirroring both the resource hierarchy and the difficulty of transferring knowledge-intensive QA across scripts and language families.

6 Instruction Following Across Languages

Zen Multilingual handles multilingual instruction following, enabling users to issue instructions in one language and receive responses in another (cross-lingual instruction following), or to work entirely in their native language.

Evaluating multilingual instruction following with an LLM-as-judge protocol across high-, medium-, and low-resource language tiers, instruction-following and response quality are highest for the high-resource tier and decline modestly toward the low-resource tier, but remain usable across all tiers. The gap between tiers is smaller for instruction following than for translation, indicating that instruction-following behavior transfers across languages more readily than fine-grained translation quality.

7 Cultural Adaptation

Beyond linguistic coverage, Zen Multilingual is trained to handle cultural context:

- Culturally appropriate greetings and honorifics (Japanese keigo, Korean speech levels, Arabic formal/informal register).
- Regional legal and regulatory contexts (EU GDPR vs. US CCPA when answering privacy questions).
- Date, number, and currency formatting conventions.

- Culturally sensitive topic handling following region-specific norms.

8 Conclusion

Zen Multilingual establishes the Qwen3-based Zen models as competitive multilingual models across the 100+ languages covered by their Qwen3 pretraining, including extremely low-resource languages. Language-balanced sampling, cross-lingual alignment, and low-resource adaptation (LAFT) techniques deliver consistent improvements over naive multilingual training baselines on FLORES-200, XCOPA, and XNLI, with the largest relative gains on the lowest-resource languages. These results demonstrate that broad multilingual coverage and high per-language quality are jointly achievable within a single model family.

References

- [1] Costa-jussà, M.R. et al. No Language Left Behind: Scaling Human-Centered Machine Translation. *arXiv:2207.04672*, 2022.
- [2] Ponti, E.M. et al. XCOPA: A Multilingual Dataset for Causal Commonsense Reasoning. *EMNLP*, 2020.
- [3] Conneau, A. et al. XNLI: Evaluating Cross-lingual Sentence Representations. *EMNLP*, 2018.
- [4] Conneau, A. et al. Unsupervised Cross-lingual Representation Learning at Scale. *ACL*, 2020.
- [5] Devlin, J. et al. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. *NAACL*, 2019.