

The Multilingual Mind: A Survey of Multilingual Reasoning in Language Models

Anonymous ACL submission

Abstract

While reasoning and multilingual capabilities in Language Models (LMs) have achieved remarkable progress in recent years, their integration into a unified paradigm—multilingual reasoning—is at a nascent stage. Multilingual reasoning requires language models to handle logical reasoning across languages while addressing misalignment, biases, and challenges in low-resource settings. This survey provides the first in-depth review of multilingual reasoning in LMs. In this survey, we provide a systematic overview of existing methods that leverage LMs for multilingual reasoning, specifically outlining the challenges, motivations, and foundational aspects of applying language models to reason across diverse languages. We provide an overview of the standard data resources used for training multilingual reasoning in LMs and the evaluation benchmarks employed to assess their multilingual capabilities. Next, we analyze various state-of-the-art methods and their performance on these benchmarks. Finally, we explore future research opportunities to improve multilingual reasoning in LMs, focusing on enhancing their ability to handle diverse languages and complex reasoning tasks.

1 Introduction

If we spoke a different language, we would perceive a somewhat different world.

Ludwig Wittgenstein

Large Language Models (LLMs) (Vaswani, 2017) have emerged as transformative tools in natural language processing, demonstrating state-of-the-art performance in language generation, translation, and summarization. These models, trained on vast corpora, excel in generating human-like text and understanding diverse linguistic contexts. Despite their success in language generation, LLMs often face significant challenges in addressing *underrepresented languages* and *reasoning*.

While the development of Multilingual LLMs (Qin et al., 2024; Huang et al., 2024a) extends LLM’s capabilities in addressing multiple languages and catering to the needs of linguistically diverse communities, their proficiency in generation stems from training on large-scale corpora optimized for next-word prediction rather than logical inference (Ramji and Ramji, 2024). Consequently, while they produce fluent and contextually appropriate responses, they frequently struggle with complex reasoning tasks, particularly those requiring multi-step logic or nuanced understanding (Patel et al., 2024). These limitations become even more pronounced in multilingual settings due to key technical problems like cross-lingual misalignment, biases in training data, and the scarcity of resources for low-resource languages.

Reasoning is formally defined as the process of drawing logical conclusions, enabling individuals and systems to solve problems and make complex decisions. Recent advancements have sought to enhance the reasoning capabilities of LLMs using Chain-of-Thought (CoT) (Wei et al., 2022), fine-tuning (Lobo et al., 2024), and hybrid modeling (Yao et al., 2024), especially in high-resource languages like English. However, reasoning in multilingual contexts remains a relatively **unexplored** domain, where existing efforts predominantly focus on a handful of high-resource languages, leaving low-resource and typologically distant languages **underrepresented**. The lack of robust benchmarks, diverse training corpora, and alignment strategies further impede progress in this vital area.

Multilingual reasoning, which combines logical inference with multilingual capabilities, is essential for creating AI systems that effectively operate across diverse linguistic and cultural contexts (Shi et al., 2022). Such systems hold immense potential for global applications, from multilingual education to culturally adaptive healthcare, ensuring inclusivity and fairness. The

motivation for this survey arises from the urgent need to address these challenges and provide a systematic exploration of methods, resources, and future directions for multilingual reasoning in LLMs. The key contributions of our work are:

1) Comprehensive Overview: We systematically review existing methods that leverage LLMs for multilingual reasoning, outlining challenges, motivations, and foundational aspects of applying reasoning to diverse languages.

2) Training Corpora and Evaluation Benchmarks: We analyze the strengths, limitations, and suitability of existing multilingual corpora and evaluation benchmarks in assessing the reasoning capabilities of LLMs for diverse linguistic tasks.

3) Analysis of State-of-the-Art Methods: We evaluate the performance of various state-of-the-art techniques, including CoT prompting, instruction tuning, and cross-lingual adaptations, on multilingual reasoning benchmark tasks.

4) Future Research Directions: We identify key challenges and provide actionable insights for advancing multilingual reasoning, focusing on adaptive alignment strategies, culturally aware benchmarks, and methods for low-resource languages.

2 Multilingual Reasoning in LLMs

Recent advancements in LLMs have improved their reasoning capabilities. However, extending them across languages introduces several challenges, including consistency, low-resource adaptation, and cultural integration. Below, we describe the preliminaries and key characteristics of multilingual reasoning, focusing on challenges and desiderata for cross-lingual inference.

2.1 Preliminaries

Large Language Models (LLMs). LLMs are transformer-based neural network architectures designed to model the probability of a sequence of tokens. Formally, LLMs are trained to predict the likelihood of a word (or sub-word token) given the preceding words in a sequence $X = \{x_1, \dots, x_n\}$, *i.e.*, $P(X) = \prod_{i=1}^n P(x_i | x_1, \dots, x_{i-1})$, where $P(X)$ is the probability of the entire sequence and $P(x_i | x_1, \dots, x_{i-1})$ is the conditional probability of the i^{th} token given the preceding tokens.

Reasoning. One of the key reasons behind the success of LLMs in mathematical and logical tasks is their reasoning capabilities. Formally, reasoning enables LLMs to draw logical conclusions C from

premises P using a mapping function: $C = f(P)$. To this end, there are different types of reasoning strategies that an LLM can employ:

a) Deductive Reasoning: It derives specific conclusions from general premises. If a given set of premises P_i is true, the conclusion C must be true, *i.e.*, $P_1, P_2, \dots, P_n \Rightarrow C$,

b) Inductive Reasoning: Generalizes patterns from specific instances, leading to probabilistic conclusions, *i.e.*, $P_1, P_2, \dots, P_n \Rightarrow C_{\text{probabilistic}}$

c) Abductive Reasoning: Infers the most plausible explanation (H_{best}) for given observation O , *i.e.*, $O \Rightarrow H_{\text{best}}$

d) Analogical Reasoning: Identifies relationships between domains and transfers knowledge, *i.e.*, $A : B \approx C : D$

e) Commonsense Reasoning: Uses real-world knowledge for intuitive decision-making.

2.2 Desiderata in Multilingual Reasoning

Here, we describe desiderata that lay the foundation for multilingual reasoning in LLMs. Let $L = \{l_1, l_2, \dots, l_m\}$ represent a set of m languages, and let P_l and C_l denote the premise and conclusion in a given language l_i . For a multilingual reasoning model M , the task can be defined as: $M(P_{l_i}) \rightarrow C_{l_i}, \forall l_i \in L$, where M must satisfy the following key desiderata:

1. Consistency: A model should make logically equivalent conclusions across languages for semantically equivalent premises, *i.e.*, $C_{l_i} \approx C_{l_j}$, if $P_{l_i} \equiv P_{l_j}$, $\forall l_i, l_j \in L$, where $\equiv$ indicates semantic equivalence of premises across languages. Consistency ensures that logical conclusions remain invariant of the input language.

2. Adaptability: For languages $l_k \in L_{\text{low-resource}}$, the model must generalize effectively using cross-lingual transfer from high-resource languages and perform robust reasoning, *i.e.*, $\forall l_k \in L_{\text{low-resource}}, M(P_{l_k}) \rightarrow C_{l_k}$.

3. Cultural Contextualization: Reasoning should consider cultural and contextual differences inherent to each language, *i.e.*, for a context c_{l_i} specific to language l_i , the conclusion C_{l_i} should adapt accordingly: $C_{l_i} = f(P_{l_i}, c_{l_i})$, $\forall l_i \in L$, where f is a mapping function that integrates linguistic reasoning with cultural nuances.

4. Cross-Lingual Alignment: The model must align reasoning processes across typologically diverse languages, where typology refers to linguistic differences in syntax, morphology, and structure (*e.g.*, word order variations between

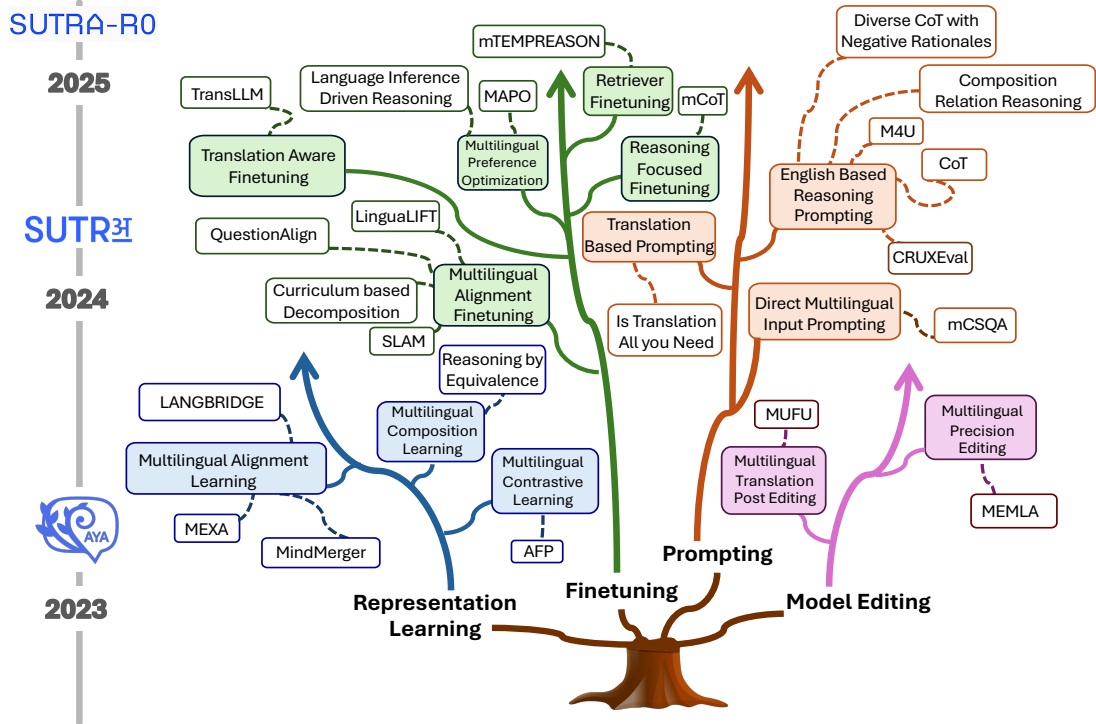


Figure 1: **Taxonomy tree of current Multilingual Reasoning Research.** The thrusts for improving multilingual reasoning mainly include representation learning, fine-tuning, prompting, and model editing. With the emergence of multilingual LLMs, while initial research focused on naive prompting, recent works propose several alignment, editing, and fine-tuning strategies to improve reasoning in multilingual LLMs.

English and Japanese). Given the typological variations T_{l_i} and T_{l_j} for languages l_i and l_j , alignment ensures that reasoning remains consistent and coherent across languages, *i.e.*, if $P_{l_i} \equiv P_{l_j}$, $M(P_{l_i}) \approx M(P_{l_j})$, $\forall l_i, l_j \in L$. Next, we highlight existing works that propose different training corpora and benchmarks for multilingual reasoning in Sec. 3 and then describe previously proposed techniques to improve the multilingual reasoning of LLMs in Sec. 4.

3 Multilingual Reasoning Datasets

Models trained on english corpora exhibit language biases (Lyu et al., 2024), limiting their reasoning capability on non-English languages. Training an LM to solve math problems across languages requires multilingual understanding and mathematical reasoning (Son et al., 2024). Hence, multilingual datasets and benchmarks play a key role in training multilingual LMs and evaluating the effectiveness of various LMs and techniques in handling domain-specific reasoning queries across low- and high-resource languages (Xu et al., 2024; Rasiah et al., 2024; Xue et al., 2024). Below, we detail training datasets (Sec. 3.1) and benchmarks (Sec. 3.2), comprising domains, tasks, and language distribution in current multilingual reasoning datasets.

3.1 Training Corpus

The best strategy to equip an LM with a specific type of reasoning is to train the model on it. However, the training objective differs based on the use case, domain, and language in which the model needs to be adapted. For example, to perform mathematical reasoning (Cobbe et al., 2021; Amini et al., 2019) in a particular language, it needs to be trained with mathematical reasoning datasets, which will differ if we want to adapt the model for legal reasoning.

While most training corpora are predominantly based on mathematical reasoning, XCSQA (Zhu et al., 2024b) and MultiNLI (Williams et al., 2017) are used for enhancing logical and coding reasoning, and sPhinX (Ahuja et al., 2024) is developed to translate instruction-response pairs into 50 languages for fine-tuning. In addition, there are cases where translation datasets like OPUS (Tiedemann, 2012), FLORES-200 (Goyal et al., 2022), and LegoMT (Yuan et al., 2022) are used to map the multilingual representation into the LM’s representation space. Further, Ponti et al. (2020) introduced XCOPA to show that multilingual pre-training and zero-shot fine-tuning underperform compared to translation-based transfer. We argue that, moving

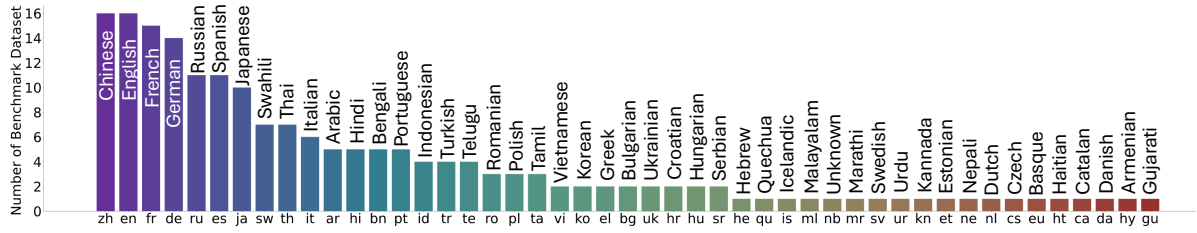


Figure 2: **Language distribution across training corpora and benchmarks for multilingual reasoning.** The y-axis denotes the number of training corpora/benchmark datasets that include a given language (x-axis). We observe a long-tail distribution, denoting that current datasets predominantly cover languages like Chinese, English, French, and German, highlighting the need for benchmarks that represent long-tail languages.

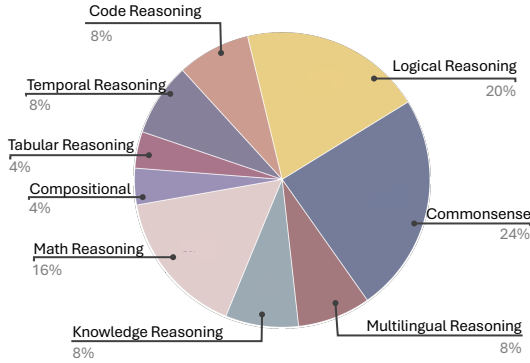


Figure 3: **Distribution of multilingual reasoning datasets.** We find that datasets predominantly comprise logical, commonsense, and math reasoning, and the community needs benchmarks to include compositional and tabular reasoning.

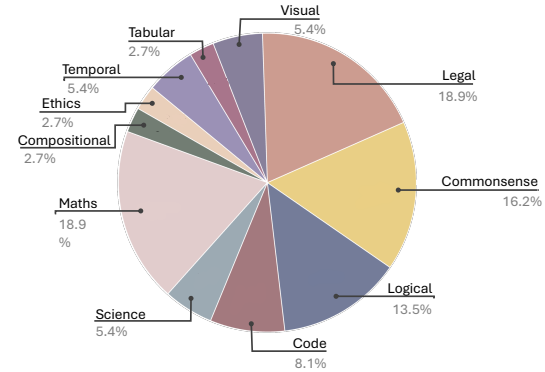


Figure 4: **Distribution of domains in multilingual reasoning datasets.** While legal, commonsense, and math domain dataset cover up to 54% of current multilingual reasoning research, other under-explored domains include ethics, science, visual, and compositional.

forward, selecting the appropriate dataset and training methodology is crucial for optimizing a model’s performance in specialized reasoning tasks.

3.2 Evaluation Benchmark

Benchmarks are key to advancing the field of multilingual reasoning as they provide a systematic framework to assess the performance of models across diverse reasoning tasks. Each reasoning task and domain presents unique challenges, making it crucial to have tailored benchmarks that reflect specific requirements and complexities of those tasks. Below, we analyze the evaluation benchmarks on three key aspects, namely languages (Fig. 2), domain (Fig. 3), and task (Fig. 4).

3.2.1 Domains and Tasks Covered

Multilingual reasoning in LMs spans multiple domains, each with its complexities and requirements, and understanding these differences is essential for developing LMs that can effectively adapt to various applications. For instance, Cobbe et al. (2021) highlighted that mathematical reasoning requires structured multi-step logic and datasets. While Ponti et al. (2020) showed that causal reasoning in XCOPA relies on cross-lingual consis-

tency and commonsense inference, Östling and Tiedemann (2016) noted that multilingual reasoning introduces typological challenges. These studies emphasize the need for tailored approaches to address the specific demands of each task and domain. Hence, it is crucial to **build reliable and robust benchmarks** for developing more robust techniques tailored to handle the complexity of a particular domain and task. Figs. 3-4 show the distribution of datasets across various domains and tasks, highlighting the need to develop more comprehensive benchmarks across multiple domains. Currently, tasks such as math, legal, and commonsense reasoning dominate multilingual benchmarks, collectively accounting for **54%** of the total (Fig. 4). In contrast, domains like science, ethics, and visual, tabular, and temporal reasoning are underrepresented, covering only **35%**. Notably, **crucial domains such as finance and healthcare still lack dedicated evaluation benchmarks** for multilingual reasoning, highlighting a significant gap in the field.

3.2.2 Languages Covered

Comprehensive language coverage is vital for multilingual reasoning, ensuring inclusivity and bal-

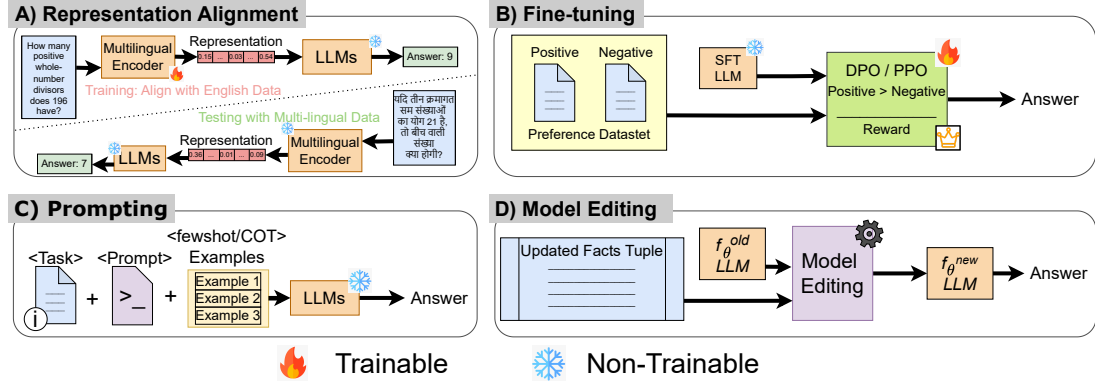


Figure 5: **Taxonomy of Multilingual Reasoning Methods.** A taxonomy of approaches for enhancing multilingual reasoning in models, covering (A) Representation Alignment, (B) Finetuning, (C) Prompting, and (D) Model Editing.

anced performance across low- and high-resource linguistic communities. Based on languages, current benchmarks can be primarily classified into human and coding languages. Benchmarks like XNLI (Conneau et al., 2018), mCSQA (Sakai et al., 2024), and m-ARC (Lai et al., 2023) predominantly focus on high-resource languages like English, Chinese, French, and Spanish. While some efforts include low-resource languages like Swahili (XCOPA (Ponti et al., 2020)), Haitian (M4U (Wang et al., 2024)), and Nepali (mMMLU (Hendrycks et al., 2020)), **their representation remains minimal** and research in these languages remains at a nascent stage. Typologically distant and under-represented languages, such as Kannada, Gujarati (xSTREET (Li et al., 2024a)), and Quechua, are rarely included, **further widening linguistic inequalities**. Datasets like FLORES-200 attempt to balance low- and high-resource languages but fail to achieve comprehensive coverage. To ensure effective LLM performance across diverse linguistic and cultural contexts, it is critical to include a broader range of low-resource and endangered languages (Goyal et al., 2022; Amini et al., 2019) (see the complete distribution of human languages across benchmarks in Fig. 2). Finally, only four benchmarks (Luo et al., 2024; Xu et al., 2024; Zhang et al., 2024b; Li et al., 2024a) incorporate coding languages across multiple languages.

4 Methods

Multilingual reasoning within LMs has garnered significant attention in recent years, leading to the development of diverse techniques for enhancing their capabilities across diverse languages. Prior works have explored various directions to improve multilingual reasoning. Building upon this body of work (see Fig. 5), we identify four primary thrusts,

viz. representation alignment, fine-tuning, prompting, and model editing, collectively contributing to advancing multilingual reasoning in LMs.

a) Representation Alignment. Multilingual reasoning requires consistent representations across languages, but LMs often struggle due to imbalanced training data. Representation alignment ensures that equivalent concepts share similar embeddings, reducing inconsistencies in multilingual inference, vital for reasoning and multilingual generalization. Li et al. (2024b) employs contrastive learning to align multilingual sentence representations by treating translation pairs as positive samples and pulling their embeddings closer, bridging language representation gaps and enhancing model’s cross-lingual reasoning and generation capabilities. Multilingual Alignment Learning is another technique that ensures semantic consistency across languages by aligning their representations for improved multilingual performance (Huang et al., 2024b), bridging multilingual encoders with LLMs using minimal parameters to achieve effective alignment without supervision (Yoon et al., 2024; Kargaran et al., 2024). Similarly, Ruan et al. (2025) integrates all encoder layer representations and employs adaptive fusion-enhanced attention to enable layer-wise alignment between the LLM and multilingual encoder, ensuring consistent cross-lingual representations and improving the model’s multilingual reasoning capabilities. Finally, an exciting new direction is multilingual compositional learning, which constructs compositional representations by combining equivalent token embeddings across multiple languages (Arora et al., 2024) and formalizing problems in an abstract space and solving them step-by-step using self-training for improved alignment across languages (Ranaldi and Pucci, 2025).

b) Finetuning. It leverages cross-lingual data and tasks to fine-tune models for enhanced reasoning and comprehension, leading to numerous innovative approaches. For instance, LinguaLIFT (Zhang et al., 2024a) uses code-switched fine-tuning along with language alignment layers to effectively bridge the gap between English and low-resource languages, helping maintain the nuance and context across linguistic boundaries. Similarly, QuestionAlign (Zhu et al., 2024b) aligns questions and responses in multiple languages, thereby enhancing cross-lingual understanding and consistency in reasoning and Ko et al. (2025) introduces a strategic fine-tuning approach that anchors reasoning in English and then translates results, significantly reducing cross-lingual performance gaps. Strategic fine-tuning using a small but high-quality bilingual dataset can enhance both the reasoning capabilities and non-English language proficiency of LLMs (Ha, 2025). While these methods have leaned towards extensive fine-tuning, SLAM (Fan et al., 2025) introduces a more parameter-efficient strategy and selectively tunes layers critical for multilingual comprehension, **significantly lowering the computational demands** while still maintaining or even enhancing the model’s reasoning capabilities. Translation has also been harnessed as a powerful tool for knowledge transfer in multilingual settings, where TransLLM (Geng et al., 2024) focuses on translation-aware fine-tuning to align different languages, enhancing language understanding but also adapting the model for various cross-lingual tasks. For those aiming at more complex reasoning tasks, **reasoning-focused fine-tuning** has proven beneficial. The Multilingual CoT (mCoT) instruction tuning method (Lai and Nissim, 2024) utilizes a dataset specifically curated for reasoning across languages and combines CoT reasoning with instruction tuning to boost consistency and logical problem-solving in multiple languages. In addition, preference-based techniques to align reasoning outputs across languages emphasize the use of language imbalance as a reward signal in models like Direct Preference and Proximal Policy Optimization (She et al., 2024). Recent research has demonstrated that Process Reward Modeling offers fine-grained feedback at each step of the reasoning process, only Wang et al. (2025) has shown its application on non-English language. Finally, an interesting direction moving forward is curriculum-based and retriever-based fine-tuning techniques

to enhance multilingual reasoning (Anand et al., 2024; Bajpai and Chakraborty, 2024), where models must not only retrieve relevant information but also compare them to evaluate relationships between them (Agrawal et al., 2024; Ranaldi et al., 2025b; Shao et al., 2024; Yang et al., 2025).

c) Prompting. Prompting has emerged as a key technique for enhancing how LLMs adapt and reason across different languages. By guiding the model through specific strategies, prompting facilitates dynamic language adaptation and addresses the data imbalance challenge, thereby enhancing cross-lingual consistency, logical alignment, and the robustness of reasoning. For instance, an effective method is Direct Multilingual Input Prompting (Sakai et al., 2024), where the model directly processes inputs in various native languages without translation, preserving the original linguistic nuances. This approach was notably applied in the paper “Do Moral Judgements” (Khandelwal et al., 2024), where moral scenarios were directly presented in their native languages to assess the model’s reasoning capabilities. Another strategy, Translation-based prompting (Liu et al., 2024) uses translation to convert multilingual inputs into a target language for processing, where tasks are translated into English for reasoning and translated back to the target language for evaluation (Wang et al., 2024; Zhao and Zhang, 2024b). This is also used to generate diverse CoT with Negative Rationales by incorporating both correct and incorrect reasoning paths to refine multilingual reasoning capabilities (Payoungkhamdee et al., 2024). While in-context learning with natural language can be ambiguous and less effective in low-resource languages, program-based demonstrations offer clearer, structured reasoning that transfers better across languages (Ranaldi et al., 2025a). In addition to the above strategies, Dictionary Insertion Prompting (DIP) offers a lightweight and practical alternative by inserting English translations of keywords into non-English prompts, bridging linguistic gaps without full translation and enabling clearer reasoning and improved performance in multilingual tasks (Lu et al., 2024).

d) Model Editing. Model editing is a growing and exciting research area that aims to modify/update the information stored in a model. Formally, model editing strategies update pre-trained models for specific input-output pairs without retraining them and impacting the baseline model performance on other inputs. Multilingual Precision Editing in-

involves making updates to model knowledge while ensuring minimal impact on unrelated information. Multilingual knowledge Editing with neuron-Masked Low-Rank Adaptation (MEMLA) (Xie et al., 2024) enhances multilingual reasoning by leveraging neuron-masked LoRA-based edits to **integrate knowledge across languages and improve multi-hop reasoning** capabilities. Further, Multilingual Translation Post-editing refines translations by correcting errors in multilingual outputs for better alignment, where we can enhance multilingual reasoning by incorporating auxiliary translations into the post-editing process, enabling LLMs to improve semantic alignment and translation quality across languages (Lim et al., 2024). **An emerging complementary direction** investigates inference-time (test-time) compute scaling in enhancing multilingual reasoning. Recent work shows that scaling up compute for English-centric reasoning language models (RLMs) can significantly improve performance across many languages, including low-resource ones, even surpassing larger models (Yong et al., 2025). While most test-time techniques, such as CoT prompting with trial and error, have primarily focused on English, methods like English-Pivoted CoT training (Tran et al., 2025) exploit the model’s strong English reasoning capabilities to support multilingual tasks, offering a promising path to bridge alignment gaps for underrepresented languages.

5 Evaluation Metrics and Benchmarks

Evaluating multilingual reasoning in LLMs requires standardized metrics to ensure logical consistency and cross-lingual coherence. Unlike traditional NLP, it must address inference errors, translation drift, and reasoning stability across languages.

5.1 Metrics

Here, we detail key metrics for evaluating multilingual reasoning, along with their formal definitions:

1) Accuracy. These metrics assess overall correctness in reasoning and multilingual benchmarks: i) *General Accuracy* measures the proportion of correct outputs over total samples, and ii) *Zero-Shot Accuracy*, which evaluates model performance on unseen tasks or categories without fine-tuning.

2) Reasoning and Consistency. These metrics evaluate logical inference and multi-step reasoning ability: i) *Reasoning Accuracy* assesses correctness in logical and step-by-step reasoning tasks and ii)

Path Consistency measures coherence between reasoning steps in CoT prompting.

3) Translation and Cross-Lingual. To ensure multilingual reasoning consistency, models must preserve meaning across languages: i) *Translation Success Rate* measures correctness and semantic preservation in multilingual translations as the ratio of accurate translations and total translations and ii) *Cross-Lingual Consistency* evaluates whether logically equivalent statements yield *consistent reasoning outputs* across different languages.

4) Perplexity and Alignment. They quantify *semantic alignment* and measure whether embeddings across languages remain consistent: i) *Perplexity-Based Alignment* (P_{align})

$$P_{align} = \exp \left(-\frac{1}{N} \sum_{i=1}^N \log P(x_i) \right), \quad (1)$$

where $P(x_i)$ is the model’s probability of predicting token x_i (lower perplexity means better alignment) and ii) *Semantic Alignment* measures the cosine similarity between multilingual sentence embeddings: $S_{align} = \frac{E_l \cdot E_t}{\|E_l\| \|E_t\|}$, where E_l and E_t are sentence embeddings in different languages.

5.2 Performance on Benchmarks

Here, we discuss the performance of the aforementioned methods on standard mathematical (MGSM (Shi et al., 2022), MSVAMP (Chen et al., 2023)), commonsense (xCSQA (Lin et al., 2021)), logical (xNLI (Conneau et al., 2018)) reasoning benchmarks¹. Next, we describe the four most popular benchmarks and detail the performance of reasoning techniques, highlighting existing model gaps that limit their reasoning performance.

MGSM tests multilingual arithmetic reasoning in LMs with 250 translated math problems in ten diverse languages. While recent trends suggest that advanced post-training techniques like MAPO are key for strong performance, fine-tuning strategies may be more impactful than stronger reasoning architectures or relying on the model’s English expertise to improve multilingual performance.

MSVAMP is an out-of-domain multilingual mathematical reasoning dataset comprising 10k problems across ten languages and serves as a comprehensive test bed to evaluate LMs’ generalization in multilingual mathematical contexts. We find that advanced preference optimization achieves much stronger

¹We only cover benchmarks analyzed by more than four papers.

performance than CoT-based fine-tuning, suggesting advanced fine-tuning techniques are a better direction to beat the current best in this benchmark. **xCSQA** is a multilingual extension of the CommonsenseQA dataset, encompassing 12,247 multiple-choice questions translated into 15 languages, designed to assess LMs’ cross-lingual commonsense reasoning capabilities. The current trend shows that stronger fine-tuning strategies like two-step fine-tuning or preference optimization show better performance than selectively fine-tuning specific layers as in SLAM.

xNLI evaluates cross-lingual inference across 15 languages. Recent studies suggest that LM integration with external models (Huang et al., 2024b) and multilingual alignment followed by fine-tuning (Zhang et al., 2024a) outperform contrastive learning methods like TCC (Chia et al., 2023), highlighting the need for more structured multilingual adaptation strategies.

6 Future Directions

With the rapid development of reasoning models, our community must ensure that models remain unbiased towards low-resource languages. Looking forward, we call on the community to put their collective efforts into the following directions:

1. Multilingual Alignment and Reasoning Transfer. A key challenge in multilingual reasoning is the lack of data in different languages. One promising solution is to leverage existing large datasets and transfer/distill their knowledge in the representation space (Yoon et al., 2024; Huang et al., 2024b). Future research should develop cross-lingual knowledge transfer techniques, enabling models to use high-resource languages as a bridge to enhance reasoning in *low-resource languages*. Another direction is to generate synthetic datasets using techniques like back-translation and data augmentation, tailored specifically for reasoning tasks.

2. Explainable and Interpretable Reasoning. Ensuring faithful reasoning in multilingual LLMs is challenging due to linguistic diversity, translation ambiguities, and reasoning inconsistencies. Studies on English CoT reasoning (Tanneru et al., 2024; Lobo et al., 2024) highlight faithfulness issues, which become more severe when extended to low-resource languages. Causal reasoning can enhance cross-lingual alignment, improving interpretability by uncovering cause-and-effect relationships across languages. Future research should focus

on integrating causal reasoning and multilingual CoT frameworks to ensure logical coherence, transparency, and trust in multilingual AI systems.

3. Advanced Training and Inference Techniques. While recent advancements in multilingual reasoning have introduced reasoning-aware fine-tuning and multilingual preference optimization techniques, further efforts are needed to improve training paradigms. Some exciting techniques in this direction includes post-training RL methods that improve reasoning in low-resource languages (Wu et al., 2024) and efficient inference-time scaling and Agentic frameworks (Khanov et al., 2024; Chakraborty et al., 2024). Preliminary post-training works (Xuan et al., 2025) show that they yield mixed results across languages, with effectiveness depending on the base model and required degree of linguistic diversity, highlighting the need for language inclusive training approaches.

4. Unified Evaluation Metrics. A comprehensive evaluation framework is a crucial missing component for assessing multilingual reasoning capabilities. Metrics should measure logical consistency, cultural adaptability, and robustness, considering real-world and adversarial multilingual settings.

5. Multimodal Multilingual Reasoning. While there are a few works on visual reasoning in the multilingual context (Das et al., 2024; Gao et al., 2025), multimodal reasoning (integrating tables, text, image, audio, and video) remains largely unexplored. Advancing this area could enable models to handle complex tasks in low-resource languages and incorporate cross-modal reasoning. Refer to Appendix A for additional directions.

7 Conclusion

Multilingual reasoning in LLMs is a rapidly evolving field, addressing critical challenges like cross-lingual alignment, low-resource language gaps, and cultural adaptation. Our survey highlights advancements in fine-tuning, prompting, and representation learning while identifying gaps in scalability and domain-specific applications. It serves as a call to action for the LLM and reasoning community to focus on advanced alignment techniques, culturally aware reasoning, and scalable architectures. By breaking language barriers and fostering inclusivity, multilingual reasoning can create globally impactful AI systems. Our survey provides a foundation for advancing research in this transformative domain.

8 Limitations

This is the first survey dedicated to the important and emerging topic of multilingual reasoning. We have made every effort to include key studies and recent advancements in this area; however, we acknowledge that some relevant work may have been unintentionally missed. As the field is still in its early stages, this survey does not aim to provide definitive solutions for improving multilingual reasoning. Instead, our goal is to analyze existing approaches and offer a comprehensive evaluation of which techniques demonstrate stronger performance across current benchmarks.

References

Ameeta Agrawal, Andy Dang, Sina Bagheri Nezhad, Rhitabrat Pokharel, and Russell Scheinberg. 2024. Evaluating multilingual long-context models for retrieval and reasoning. *arXiv preprint arXiv:2409.18006*.

Sanchit Ahuja, Kumar Tanmay, Hardik Hansrajbhai Chauhan, Barun Patra, Kriti Aggarwal, Luciano Del Corro, Arindam Mitra, Tejas Indulal Dhamecha, Ahmed Awadallah, Monojit Choudhary, Vishrav Chaudhary, and Sunayana Sitaram. 2024. sphinx: Sample efficient multilingual instruction fine-tuning through n-shot guided prompting. *arXiv*.

Aida Amini, Saadia Gabriel, Peter Lin, Rik Koncel-Kedziorski, Yejin Choi, and Hannaneh Hajishirzi. 2019. Mathqa: Towards interpretable math word problem solving with operation-based formalisms. In *NAACL*.

Avinash Anand, Kritarth Prasad, Chhavi Kirtani, Ashwin R Nair, Manvendra Kumar Nema, Raj Jaiswal, and Rajiv Ratn Shah. 2024. Multilingual mathematical reasoning: Advancing open-source llms in hindi and english. *arXiv*.

Avinash Anand, Kritarth Prasad, Chhavi Kirtani, Ashwin R Nair, Manvendra Kumar Nema, Raj Jaiswal, and Rajiv Ratn Shah. 2025. Multilingual mathematical reasoning: Advancing open-source llms in hindi and english. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 23415–23423.

Gaurav Arora, Srujana Merugu, Shreya Jain, and Vaibhav Saxena. 2024. Towards robust knowledge representations in multilingual llms for equivalence and inheritance based consistent reasoning. *arXiv*.

Ashutosh Bajpai and Tanmoy Chakraborty. 2024. Multilingual llms inherently reward in-language time-sensitive semantic alignment for low-resource languages. *arXiv*.

Ashutosh Bajpai and Tanmoy Chakraborty. 2025. Multilingual llms inherently reward in-language time-sensitive semantic alignment for low-resource languages. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 23469–23477.

Linzheng Chai, Jian Yang, Tao Sun, Hongcheng Guo, Jiaheng Liu, Bing Wang, Xiannian Liang, Jiaqi Bai, Tongliang Li, Qiyao Peng, and 1 others. 2024. xcot: Cross-lingual instruction tuning for cross-lingual chain-of-thought reasoning, 2024. URL <https://arxiv.org/abs/2401.7037>.

Souradip Chakraborty, Soumya Suvra Ghosal, Ming Yin, Dinesh Manocha, Mengdi Wang, Amrit Singh Bedi, and Furong Huang. 2024. Transfer q star: Principled decoding for llm alignment. *arXiv*.

Nuo Chen, Zinan Zheng, Ning Wu, Ming Gong, Dongmei Zhang, and Jia Li. 2023. Breaking language barriers in multilingual mathematical reasoning: Insights and observations.

Yew Ken Chia, Guizhen Chen, Luu Anh Tuan, Soujanya Poria, and Lidong Bing. 2023. Contrastive chain-of-thought prompting. *arXiv*.

Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. 2021. Training verifiers to solve math word problems. *arXiv*.

Alexis Conneau, Guillaume Lample, Ruty Rinott, Adina Williams, Samuel R Bowman, Holger Schwenk, and Veselin Stoyanov. 2018. Xnli: Evaluating cross-lingual sentence representations. *arXiv*.

Rocktim Jyoti Das, Simeon Emilov Hristov, Haonan Li, Dimitar Iliyanov Dimitrov, Ivan Koychev, and Preslav Nakov. 2024. Exams-v: A multi-discipline multilingual multimodal exam benchmark for evaluating vision language models. *arXiv*.

Yuchun Fan, Yongyu Mu, Yilin Wang, Lei Huang, Junhao Ruan, Bei Li, Tong Xiao, Shujian Huang, Xiaocheng Feng, and Jingbo Zhu. 2025. Slam: Towards efficient multilingual reasoning via selective language alignment. *arXiv*.

Junyuan Gao, Jiahe Song, Jiang Wu, Runchuan Zhu, Guanlin Shen, Shasha Wang, Xingjian Wei, Haote Yang, Songyang Zhang, Weijia Li, and 1 others. 2025. Pm4bench: A parallel multilingual multimodal multi-task benchmark for large vision language model. *arXiv preprint arXiv:2503.18484*.

Xiang Geng, Ming Zhu, Jiahuan Li, Zhejian Lai, Wei Zou, Shuaijie She, Jiabin Guo, Xiaofeng Zhao, Yinglu Li, Yuang Li, and 1 others. 2024. Why not transform chat large language models to non-english? *arXiv*.

764	Naman Goyal, Cynthia Gao, Vishrav Chaudhary, Peng-	Viet Dac Lai, Chien Van Nguyen, Nghia Trung Ngo,	816
765	Jen Chen, Guillaume Wenzek, Da Ju, Sanjana Kr-	Thuat Nguyen, Franck Dernoncourt, Ryan A Rossi,	817
766	ishnan, Marc Aurelio Ranzato, Francisco Guzmán,	and Thien Huu Nguyen. 2023. Okapi: Instruction-	818
767	and Angela Fan. 2022. The flores-200 evaluation	tuned large language models in multiple languages	819
768	benchmark for low-resource and multilingual ma-	chine translation. In <i>EMNLP. ACL</i> .	820
769		<i>arXiv</i> .	821
770	Srishti Gureja, Lester James V Miranda, Shayekh Bin	Bryan Li, Tamer Alkhoul, Daniele Bonadiman, Niko-	822
771	Islam, Rishabh Maheshwary, Drishti Sharma, Gusti	laos Pappas, and Saab Mansour. 2024a. Eliciting	823
772	Winata, Nathan Lambert, Sebastian Ruder, Sara	better multilingual structured reasoning from llms	824
773	Hooker, and Marzieh Fadaee. 2024. M-rewardbench:	through code. <i>arXiv</i> .	825
774	Evaluating reward models in multilingual settings.		
775	<i>arXiv preprint arXiv:2410.15522</i> .	Chong Li, Shaonan Wang, Jiajun Zhang, and Chengqing	826
776	Huy Hoang Ha. 2025. Pensez: Less data, better	Zong. 2024b. Improving in-context learning of	827
777	reasoning—rethinking french llm. <i>arXiv preprint</i>	multilingual generative language models with cross-	828
778	<i>arXiv:2503.13661</i> .	lingual alignment. In <i>NAACL</i> .	829
779	Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou,	Zihao Li, Yucheng Shi, Zirui Liu, Fan Yang, Ninghao	830
780	Mantas Mazeika, Dawn Song, and Jacob Steinhardt.	Liu, and Mengnan Du. 2024c. Quantifying multilin-	831
781	2020. Measuring massive multitask language under-	gual performance of large language models across	832
782	standing. <i>arXiv</i> .	languages. <i>arXiv e-prints</i> , pages arXiv–2404.	833
783	Amey Hengle, Prasoon Bajpai, Soham Dan, and Tan-	Zheng Wei Lim, Nitish Gupta, Honglin Yu, and Trevor	834
784	moy Chakraborty. 2025. Can llms reason over ex-	Cohn. 2024. Mufu: Multilingual fused learning for	835
785	tended multilingual contexts? towards long-context	low-resource translation with llm. <i>arXiv</i> .	836
786	evaluation beyond retrieval and haystacks. <i>arXiv</i>	Yankai Lin, Jiapeng Zhou, Yiming Shen, Wenxuan	837
787	<i>preprint arXiv:2504.12845</i> .	Zhou, Zhiyuan Liu, Peng Li, Maosong Sun, and Jie	838
788	Kaiyu Huang, Fengran Mo, Hongliang Li, You Li,	Zhou. 2021. Xcsqa: A benchmark for cross-lingual	839
789	Yuanchi Zhang, Weijian Yi, Yulong Mao, Jincheng	conversational question answering. In <i>EMNLP</i> .	840
790	Liu, Yuzhuang Xu, Jinan Xu, and 1 others. 2024a. A	Chaoqun Liu, Wenxuan Zhang, Yiran Zhao, Anh Tuan	841
791	survey on large language models with multilingual-	Luu, and Lidong Bing. 2024. Is translation all you	842
792	ism: Recent advances and new frontiers. <i>arXiv</i> .	need? a study on solving multilingual tasks with	843
793	Zixian Huang, Wenhao Zhu, Gong Cheng, Lei Li, and	large language models. <i>arXiv</i> .	844
794	Fei Yuan. 2024b. Mindmerger: Efficient boosting	Elita Lobo, Chirag Agarwal, and Himabindu Lakkaraju.	845
795	llm reasoning in non-english languages. <i>arXiv</i> .	2024. On the impact of fine-tuning on chain-of-	846
796	Amir Hossein Kargaran, Ali Modarressi, Nafiseh	thought reasoning. <i>arXiv</i> .	847
797	Nikeghbal, Jana Diesner, François Yvon, and Hin-	Hongyuan Lu, Zixuan Li, and Wai Lam. 2024. Dic-	848
798	rich Schütze. 2024. Mexa: Multilingual evaluation	tionary insertion prompting for multilingual reason-	849
799	of english-centric llms via cross-lingual alignment.	ing on multilingual large language models. <i>arXiv</i>	850
800	<i>arXiv</i> .	<i>preprint arXiv:2411.01141</i> .	851
801	Aditi Khandelwal, Utkarsh Agarwal, Kumar Tanmay,	Xianzhen Luo, Qingfu Zhu, Zhiming Zhang, Libo	852
802	and Monojit Choudhury. 2024. Do moral judgment	Qin, Xuanyu Zhang, Qing Yang, Dongliang Xu, and	853
803	and reasoning capability of llms change with lan-	Wanxiang Che. 2024. Python is not always the best	854
804	guage? a study using the multilingual defining issues	choice: Embracing multilingual program of thoughts.	855
805	test. <i>arXiv</i> .	<i>arXiv preprint arXiv:2402.10691</i> .	856
806	Maxim Khanov, Jirayu Burapachee, and Yixuan Li.	Jiachen Lyu, Katharina Dost, Yun Sing Koh, and Jörg	857
807	2024. Args: Alignment as reward-guided search.	Wicker. 2024. Regional bias in monolingual english	858
808	<i>arXiv</i> .	language models. <i>Machine Learning</i> .	859
809	Hyunwoo Ko, Guijin Son, and Dasol Choi. 2025. Un-	Xuefei Ning, Zifu Wang, Shiyao Li, Zinan Lin, Peiran	860
810	derstand, solve and translate: Bridging the multilin-	Yao, Tianyu Fu, Matthew B Blaschko, Guohao Dai,	861
811	gual mathematical reasoning gap. <i>arXiv preprint</i>	Huazhong Yang, and Yu Wang. 2024. Can llms learn	862
812	<i>arXiv:2501.02448</i> .	by teaching for better reasoning? a preliminary study.	863
813	Huiyuan Lai and Malvina Nissim. 2024. mcot: Multi-	<i>arXiv</i> .	864
814	lingual instruction tuning for reasoning consistency	Robert Östling and Jörg Tiedemann. 2016. Continuous	865
815	in language models. <i>arXiv</i> .	multilinguality with language vectors. <i>arXiv</i> .	866

867	Nisarg Patel, Mohith Kulkarni, Mihir Parmar, Aashna	Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu,	922
868	Budhiraja, Mutsumi Nakamura, Neeraj Varshney, and	Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan	923
869	Chitta Baral. 2024. Multi-logieval: Towards eval-	Zhang, Y. K. Li, Y. Wu, and Daya Guo. 2024.	924
870	uating multi-step logical reasoning ability of large	Deepseekmath: Pushing the limits of mathemati-	925
871	language models. <i>arXiv</i> .	cal reasoning in open language models . <i>Preprint</i> ,	926
		arXiv:2402.03300.	927
872	Patomporn Payoungkhamdee, Peerat Limkonchotiwat,	Shuaijie She, Wei Zou, Shujian Huang, Wenhao Zhu, Xi-	928
873	Jinheon Baek, Potsawee Manakul, Can Udom-	ang Liu, Xiang Geng, and Jiajun Chen. 2024. Mapo:	929
874	charoenchaikit, Ekapol Chuangsuwanich, and Sarana	Advancing multilingual reasoning through multilin-	930
875	Nutanong. 2024. An empirical study of multilingual	gual alignment-as-preference optimization. <i>arXiv</i> .	931
876	reasoning distillation for question answering. In <i>Con-</i>		
877	<i>ference on Empirical Methods in Natural Language</i>	Freda Shi, Mirac Suzgun, Markus Freitag, Xuezhi Wang,	932
878	<i>Processing</i> .	Suraj Srivats, Soroush Vosoughi, Hyung Won Chung,	933
879	Edoardo Maria Ponti, Goran Glavaš, Olga Majewska,	Yi Tay, Sebastian Ruder, Denny Zhou, and 1 others.	934
880	Qianchu Liu, Ivan Vulić, and Anna Korhonen. 2020.	2022. Language models are multilingual chain-of-	935
881	Xcopa: A multilingual dataset for causal common-	thought reasoners. <i>arXiv</i> .	936
882	sense reasoning. In <i>EMNLP</i> .		
883	Libo Qin, Qiguang Chen, Yuhang Zhou, Zhi Chen,	Guijin Son, Dongkeun Yoon, Juyoung Suk, Javier Aula-	937
884	Yinghui Li, Lizi Liao, Min Li, Wanxiang Che, and	Blasco, Mano Aslan, Vu Trong Kim, Shayekh Bin	938
885	Philip S Yu. 2024. Multilingual large language	Islam, Jaume Prats-Cristià, Lucía Tormo-Bañuelos,	939
886	model: A survey of resources, taxonomy and fron-	and Seungone Kim. 2024. Mm-eval: A multilingual	940
887	tiers. <i>arXiv</i> .	meta-evaluation benchmark for llm-as-a-judge and	941
		reward models. <i>arXiv preprint arXiv:2410.17578</i> .	942
888	Raghav Ramji and Keshav Ramji. 2024. Inductive lin-	Yueqi Song, Simran Khanuja, and Graham Neu-	943
889	guistic reasoning with large language models. <i>arXiv</i> .	big. 2024. What is missing in multilingual vi-	944
		sual reasoning and how to fix it. <i>arXiv preprint</i>	945
890	Leonardo Ranaldi, Barry Haddow, and Alexandra Birch.	<i>arXiv:2403.01404</i> .	946
891	2025a. When natural language is not enough: The		
892	limits of in-context learning demonstrations in mul-	Sree Harsha Tanneru, Dan Ley, Chirag Agarwal, and	947
893	tilingual reasoning. In <i>Findings of the Association</i>	Himabindu Lakkaraju. 2024. On the hardness of	948
894	<i>for Computational Linguistics: NAACL 2025</i> , pages	faithful chain-of-thought reasoning in large language	949
895	7369–7396.	models. <i>arXiv</i> .	950
896	Leonardo Ranaldi and Giulia Pucci. 2025. Multilin-	Jörg Tiedemann. 2012. Opus: An open source parallel	951
897	gual reasoning via self-training. In <i>Proceedings of</i>	corpus .	952
898	<i>the 2025 Conference of the Nations of the Americas</i>		
899	<i>Chapter of the Association for Computational Lin-</i>	Khanh-Tung Tran, Barry O’Sullivan, and Hoang D	953
900	<i>guistics: Human Language Technologies (Volume 1:</i>	Nguyen. 2025. Scaling test-time compute for low-	954
901	<i>Long Papers)</i> , pages 11566–11582.	resource languages: Multilingual reasoning in llms.	955
		<i>arXiv preprint arXiv:2504.02890</i> .	956
902	Leonardo Ranaldi, Federico Ranaldi, Fabio Massimo	A Vaswani. 2017. Attention is all you need. <i>NeurIPS</i> .	957
903	Zanzotto, Barry Haddow, and Alexandra Birch.		
904	2025b. Improving multilingual retrieval-augmented	Hongyu Wang, Jiayu Xu, Senwei Xie, Ruiping Wang,	958
905	language models through dialectic reasoning argu-	Jialin Li, Zhaojie Xie, Bin Zhang, Chuyan Xiong,	959
906	mentations. <i>arXiv preprint arXiv:2504.04771</i> .	and Xilin Chen. 2024. M4u: Evaluating multilingual	960
		understanding and reasoning for large multimodal	961
907	Vishvakshen Rasiyah, Ronja Stern, Veton Matoshi,	models. <i>arXiv</i> .	962
908	Matthias Stürmer, Ilias Chalkidis, Daniel E Ho, and		
909	Joel Niklaus. 2024. One law, many languages:	Weixuan Wang, Minghao Wu, Barry Haddow, and	963
910	Benchmarking multilingual legal reasoning for ju-	Alexandra Birch. 2025. Demystifying multilingual	964
911	dicial support.	chain-of-thought in process reward modeling. <i>arXiv</i>	965
		<i>preprint arXiv:2502.12663</i> .	966
912	Zhiwen Ruan, Yixia Li, He Zhu, Longyue Wang, Wei-	Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten	967
913	hua Luo, Kaifu Zhang, Yun Chen, and Guanhua	Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou,	968
914	Chen. 2025. Layalign: Enhancing multilingual rea-	and 1 others. 2022. Chain-of-thought prompting elic-	969
915	soning in large language models via layer-wise adap-	its reasoning in large language models. <i>NeurIPS</i> .	970
916	tive fusion and alignment strategy. <i>arXiv preprint</i>		
917	<i>arXiv:2502.11405</i> .	Zihao Wei, Jingcheng Deng, Liang Pang, Hanxing Ding,	971
918	Yusuke Sakai, Hidetaka Kamigaito, and Taro Watanabe.	Huawei Shen, and Xueqi Cheng. 2024. Mlake: Mul-	972
919	2024. mcsqa: Multilingual commonsense reasoning	tilingual knowledge editing benchmark for large lan-	973
920	dataset with unified creation strategy by language	guage models. <i>arXiv preprint arXiv:2404.04990</i> .	974
921	models and humans. <i>arXiv</i> .		

975	Adina Williams, Nikita Nangia, and Samuel R Bowman. 2017. A broad-coverage challenge corpus for sentence understanding through inference. <i>arXiv</i> .	
976		
977		
978	Zhaofeng Wu, Ananth Balashankar, Yoon Kim, Jacob Eisenstein, and Ahmad Beirami. 2024. Reuse your rewards: Reward model transfer for zero-shot cross-lingual alignment. <i>arXiv</i> .	
979		
980		
981		
982	Jiakuan Xie, Pengfei Cao, Yuheng Chen, Yubo Chen, Kang Liu, and Jun Zhao. 2024. Memla: Enhancing multilingual knowledge editing with neuron-masked low-rank adaptation. <i>arXiv</i> .	
983		
984		
985		
986	Ruiyang Xu, Jialun Cao, Yaojie Lu, Hongyu Lin, Xi-anpei Han, Ben He, Shing-Chi Cheung, and Le Sun. 2024. Cruxeval-x: A benchmark for multilingual code reasoning, understanding and execution. <i>arXiv</i> .	
987		
988		
989		
990	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. <i>Frontiers of Computer Science</i> , 19(11):1911362.	
991		
992		
993		
994		
995	Weihaio Xuan, Rui Yang, Heli Qi, Qingcheng Zeng, Yunze Xiao, Yun Xing, Junjue Wang, Huitao Li, Xin Li, Kunyu Yu, and 1 others. 2025. Mmlu-prox: A multilingual benchmark for advanced large language model evaluation. <i>arXiv preprint arXiv:2503.10497</i> .	
996		
997		
998		
999		
1000	Siqiao Xue, Tingting Chen, Fan Zhou, Qingyang Dai, Zhixuan Chu, and Hongyuan Mei. 2024. Famma: A benchmark for financial domain multilingual multimodal question answering. <i>arXiv preprint arXiv:2410.04526</i> .	
1001		
1002		
1003		
1004		
1005	Wen Yang, Junhong Wu, Chen Wang, Chengqing Zong, and Jiajun Zhang. 2024. Language imbalance driven rewarding for multilingual self-improving. <i>arXiv preprint arXiv:2410.08964</i> .	
1006		
1007		
1008		
1009	Yahan Yang, Soham Dan, Shuo Li, Dan Roth, and Insup Lee. 2025. Mr. guard: Multilingual reasoning guardrail using curriculum learning. <i>arXiv preprint arXiv:2504.15241</i> .	
1010		
1011		
1012		
1013	Wenlin Yao, Haitao Mi, and Dong Yu. 2024. Hdflow: Enhancing llm complex problem-solving with hybrid thinking and dynamic workflows. <i>arXiv</i> .	
1014		
1015		
1016	Zheng-Xin Yong, M Farid Adilazuarda, Jonibek Mansurov, Ruochen Zhang, Niklas Muennighoff, Carsten Eickhoff, Genta Indra Winata, Julia Kreutzer, Stephen H Bach, and Alham Fikri Aji. 2025. Crosslingual reasoning through test-time scaling. <i>arXiv preprint arXiv:2505.05408</i> .	
1017		
1018		
1019		
1020		
1021		
1022	Dongkeun Yoon, Joel Jang, Sungdong Kim, Seungone Kim, Sheikh Shafayat, and Minjoon Seo. 2024. Lang-bridge: Multilingual reasoning without multilingual supervision. <i>arXiv</i> .	
1023		
1024		
1025		
1026	Fei Yuan, Yinquan Lu, WenHao Zhu, Lingpeng Kong, Lei Li, Yu Qiao, and Jingjing Xu. 2022. Lego-mt: Learning detachable models for massively multilingual machine translation. <i>arXiv</i> .	
1027		
1028		
1029		
	Hongbin Zhang, Kehai Chen, Xuefeng Bai, Yang Xiang, and Min Zhang. 2024a. Lingualift: An effective two-stage instruction tuning framework for low-resource language tasks. <i>arXiv</i> .	1030
		1031
		1032
		1033
	Yidan Zhang, Boyi Deng, Yu Wan, Baosong Yang, Haoran Wei, Fei Huang, Bowen Yu, Junyang Lin, and Jingren Zhou. 2024b. P-mmeval: A parallel multilingual multitask benchmark for consistent evaluation of llms. <i>arXiv preprint arXiv:2411.09116</i> .	1034
		1035
		1036
		1037
		1038
	Jinman Zhao and Xueyan Zhang. 2024a. Exploring the limitations of large language models in compositional relation reasoning. <i>arXiv preprint arXiv:2403.02615</i> .	1039
		1040
		1041
	Jinman Zhao and Xueyan Zhang. 2024b. Large language model is not a (multilingual) compositional relation reasoner. In <i>First Conference on Language Modeling</i> .	1042
		1043
		1044
		1045
	Shaolin Zhu, Shaoyang Xu, Haoran Sun, Leiyu Pan, Menglong Cui, Jiangcun Du, Renren Jin, António Branco, Deyi Xiong, and 1 others. 2024a. Multilingual large language models: A systematic survey. <i>arXiv preprint arXiv:2411.11072</i> .	1046
		1047
		1048
		1049
		1050
	Wenhao Zhu, Shujian Huang, Fei Yuan, Cheng Chen, Jiajun Chen, and Alexandra Birch. 2024b. The power of question translation training in multilingual reasoning: Broadened scope and deepened insights. <i>arXiv</i> .	1051
		1052
		1053
		1054
	Wenhao Zhu, Shujian Huang, Fei Yuan, Shuaijie She, Jiajun Chen, and Alexandra Birch. 2024c. Question translation training for better multilingual reasoning. <i>arXiv preprint arXiv:2401.07817</i> .	1055
		1056
		1057
		1058

A Appendix

Related Surveys The earliest surveys (Qin et al., 2024; Xu et al., 2025)—both from April 2024 focus on laying foundational taxonomies of Multilingual LLMs (MLLMs): (Qin et al., 2024) survey resources, taxonomy, and emerging frontiers in MLLMs, while (Xu et al., 2025) delve deeply into multilingual corpora, alignment techniques, and bias issues. Huang et al. (2024a) broadens the scope to multiple perspectives—training/inference, security, cultural domains, and datasets—framing “new frontiers” in multilingual LLM research. Finally, survey by (Zhu et al., 2024a) provides the most comprehensive “systematic” treatment: it covers architectures, pre-training objectives, alignment datasets, a detailed evaluation roadmap (including safety, interpretability, reasoning), and real-world applications across domains. *This survey is the first survey dedicated specifically to multilingual reasoning, drilling deeply into logical inference across languages, its unique challenges (misalignment, bias, low-resource gaps), and the benchmarks and methods tailored to evaluate and improve reasoning capabilities.*

Additional Future Directions. Below, are some additional future directions to advance multilingual reasoning in language models.

1. New Benchmarks: As multilingual reasoning advances, robust evaluation benchmarks are essential because reasoning is highly domain-specific in nature, developing targeted benchmarks is crucial, especially in high-stakes fields like healthcare, law, and finance, where accuracy directly affects decision-making. For instance, Xue et al. (2024) introduces FAMMA which shows significant challenges in the field of Financial Question Answering.

2. Efficient Reasoning Models. An emerging direction in reasoning research is enhancing resource efficiency in reasoning-aware models. Recent works like (Ning et al., 2024) propose strategies for more efficient reasoning, reducing computational costs while maintaining logical consistency. However, this area remains largely unexplored in multilingual settings, offering a key opportunity to develop scalable reasoning models that generalize across languages with minimal resources.

3. Miscellaneous Tasks. LLMs have given extraordinary performance in many tasks yet they struggle with complex compositional reasoning (Zhao and Zhang, 2024a), performing only slightly better than random guessing. Models also struggle to reason over long texts, especially in low-resource languages (Hengle et al., 2025), often failing to combine information or recognize what’s missing—even when they can retrieve facts.

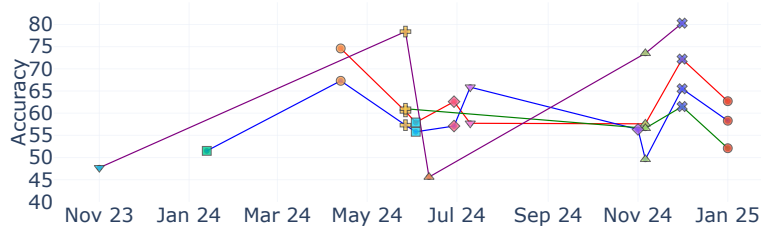


Figure 6: Accuracy trends of various methods on multilingual reasoning benchmarks, including MGSM, MSVAMP, XNLI, and XCSQA. The x -axis represents the arXiv paper submission date, and the y -axis indicates percentage accuracy.

We show a detailed tabular format of the languages used in different reasoning datasets along with their languages.

af	Afrikaans	ar	Arabic	be	Belarusian	bg	Bulgarian
bn	Bengali	ca	Catalan	cs	Czech	da	Danish
de	German	el	Greek	en	English	es	Spanish
et	Estonian	eu	Basque	fa	Persian	fi	Finnish
fr	French	ha	Hausa	he	Hebrew	hi	Hindi
hr	Croatian	ht	Haitian	hu	Hungarian	hy	Armenian
id	Indonesian	id	Indonesian	is	Icelandic	it	Italian
ja	Japanese	kn	Kannada	ko	Korean	lb	Luxembourgish
mk	Macedonian	ml	Malayalam	mr	Marathi	nb	Norwegian Bokmal
ne	Nepali	nl	Dutch	pl	Polish	pt	Portuguese
qu	Quechua	ro	Romanian	ru	Russian	sk	Slovak
sl	Slovenian	sr	Serbian	sv	Swedish	tr	Turkish
uk	Ukrainian	ur	Urdu	vi	Vietnamese	zh	Chinese

Table 1: Language Codes and Their Corresponding Languages

Table 2: Multilingual Datasets and their respective papers, domains, and languages.

Dataset	Paper	Domain	Languages
MSVAMP	(She et al., 2024; Yoon et al., 2024; Zhu et al., 2024c,b; Lai and Nissim, 2024; Chai et al., 2024; Huang et al., 2024b; Zhang et al., 2024a; Fan et al., 2025)	Maths	zh, th, ja, en, de, fr, es, bn, sw
MGSM	(She et al., 2024; Yoon et al., 2024; Zhu et al., 2024c,b; Lai and Nissim, 2024; Chai et al., 2024; Huang et al., 2024b; Liu et al., 2024; Zhang et al., 2024a; Fan et al., 2025)	Maths	zh, th, ja, en, de, fr, es, ru, bn, sw, te
MNumGLUESub	(She et al., 2024)	Maths	bn, th, sw, ja, zh, ru, de, es, fr, en
MetaMathQA	(Yoon et al., 2024; Zhu et al., 2024c,b; Lai and Nissim, 2024; Huang et al., 2024b)	Maths	en
Proof-Pile 2	(Yoon et al., 2024)	Maths	en
Exams Dataset	(Payoungkhamdee et al., 2024)	Science and Humanities	ar, de, fr, es, it, pl, vi, pt, sr, hu, tr, bg, hr, mk, sq
M4U Benchmark	(Wang et al., 2024)	Science	zh, en, de
XCSQA	(Zhu et al., 2024b; Zhang et al., 2024a; Fan et al., 2025)	Common Sense	zh, en, de, fr, es, ru, hi
XNLI	(Zhu et al., 2024b; Liu et al., 2024; Zhang et al., 2024a)	Logical	zh, th, ur, en, de, fr, es, ru, el, tr, bg, hi, sw
MultiNLI	(Zhu et al., 2024b), (Huang et al., 2024b)	Logical	en
BBH-Hard	(Luo et al., 2024)	Temporal, Tabular, Spatial	Python, R, C++, Java, Javascript
NLVR2	(Song et al., 2024)	Visual	en
MARVL	(Song et al., 2024)	Visual	id, sw, ta, tr, zh
xSTREET	(Li et al., 2024a)	Logical	ar, zh, ja, en, es, ru
Translated Code Comments (TCC)	(Li et al., 2024a)	Code	Java, JavaScript, Python
mCoT-MATH	(Lai and Nissim, 2024)	Maths	zh, th, ja, en, de, fr, es, ru, bn, hi, te
Reasoning by Equivalence Dataset	(Arora et al., 2024)	Logical	en, fr, es, de, pt, hi
Reasoning by Inheritance Dataset	(Arora et al., 2024)	Logical	en, fr, es, de, pt, hi
XCOT	(Chai et al., 2024)	Maths	de, fr, es, ru, zh, ja, th, te, bn, sw, en
mCSQA	(Sakai et al., 2024)	Common Sense	zh, ja, en, fr, de, pt, ru

Dataset	Paper	Domain	Languages
Rulings, Legislation, Court View Generation, Critically Prediction, Law Area Prediction, Judgment Prediction Datasets	(Rasiah et al., 2024)	Legal	de, fr, it, ro, en
mRewardBench	(Gureja et al., 2024)	Logical and CommonSense	ar, cs, de, el, es, fa, fr, he, hi, id, it, ja, ko, nl, pl, pt, ro, ru, tr, uk, vi, zh
Moral Judgement Dataset	(Khandelwal et al., 2024)	Moral	en, zh, hi, ru, es, sw
MCR	(Zhao and Zhang, 2024b)	Compositional	ja, ko, fr
mTEMPREASON	(Bajpai and Chakraborty, 2025)	Temporal	ro, de, fr
XCOPA	(Liu et al., 2024)	Common Sense	zh, it, vi, tr, id, sw, th, et, ta, ht, qu
mARC	(Kargaran et al., 2024)	Common Sense	zh, ja, en, de, fr, es
IndiMathQA	(Anand et al., 2025)	Maths	en, hi
CRUXEval	(Xu et al., 2024)	Code	C#, C++, D, Go, Java, JavaScript, Julia, Luca, Perl, PHP, R, Racket, Ruby, Rust, Scala, Shell, Swift, TypeScript
Dataset	Paper	Domain	Languages
mMMLU	(Kargaran et al., 2024)	Common Sense	ar, zh, vi, id, en, de, fr, it, nl, eu, es, pt, ca, da, ru, hr, hy, hu, ro, ne, kn, uk, sr, sv, mr, nb, ml, is, bn, hi, ta, te, gu
MMWP Benchmark	(Zhang et al., 2024a)	Maths	af, ar, be, bn, eu, gu, ha, hi, hy, is, kn, lb, mk, ml, mr, ne, sk, sw, ta, te, th, bg, ca, cs, da, fi, hr, hu, id, ko, nb, pl, pt, ro, sl, sr, uk, vi, de, en, es, fr, it, ja, nl, ru, sv, zh

Reasoning Type	Papers
Deductive	Lai and Nissim (2024), Chai et al. (2024), Huang et al. (2024b), Zhang et al. (2024a), Huang et al. (2024b), Fan et al. (2025), Payoungkhamdee et al. (2024), Luo et al. (2024), Song et al. (2024), Li et al. (2024a), Arora et al. (2024), Rasiah et al. (2024), Sakai et al. (2024), Khandelwal et al. (2024), Kargaran et al. (2024), Anand et al. (2025), Xu et al. (2024), She et al. (2024), Zhu et al. (2024b), Li et al. (2024c), Lim et al. (2024), Bajpai and Chakraborty (2025), Li et al. (2024b)
Inductive	Chai et al. (2024), Huang et al. (2024b), Zhang et al. (2024a), Huang et al. (2024b), Fan et al. (2025), Payoungkhamdee et al. (2024), Luo et al. (2024), Song et al. (2024), Li et al. (2024a), Arora et al. (2024), Rasiah et al. (2024), Sakai et al. (2024), Khandelwal et al. (2024), Kargaran et al. (2024), Anand et al. (2025), Xu et al. (2024), She et al. (2024), Zhu et al. (2024b), Li et al. (2024c), Lim et al. (2024), Bajpai and Chakraborty (2025), Li et al. (2024b), Wei et al. (2024), Xie et al. (2024), Yang et al. (2024), Geng et al. (2024), Yang et al. (2025), Ko et al. (2025), Ruan et al. (2025), Lu et al. (2024), Agrawal et al. (2024), Ranaldi et al. (2025b), Ha (2025), Ranaldi et al. (2025a), Ranaldi and Pucci (2025), Xuan et al. (2025), Yoon et al. (2024), Zhu et al. (2024c), Lai and Nissim (2024), Chai et al. (2024), Huang et al. (2024b), Zhang et al. (2024a), Huang et al. (2024b), Fan et al. (2025), Payoungkhamdee et al. (2024), Luo et al. (2024), Song et al. (2024), Li et al. (2024a), Arora et al. (2024), Rasiah et al. (2024), Sakai et al. (2024), Khandelwal et al. (2024), Kargaran et al. (2024), Anand et al. (2025), Xu et al. (2024), She et al. (2024), Zhu et al. (2024b), Li et al. (2024c), Lim et al. (2024), Bajpai and Chakraborty (2025), Li et al. (2024b), Wei et al. (2024), Xie et al. (2024), Yang et al. (2024), Geng et al. (2024), Yang et al. (2025), Ko et al. (2025), Ruan et al. (2025), Lu et al. (2024), Agrawal et al. (2024), Ranaldi et al. (2025b), Ha (2025), Ranaldi et al. (2025a), Ranaldi and Pucci (2025)
Abductive	Huang et al. (2024b), Zhang et al. (2024a)
Analogical	Zhang et al. (2024a), Huang et al. (2024b), Fan et al. (2025), Payoungkhamdee et al. (2024), Luo et al. (2024), Song et al. (2024), Li et al. (2024a), Arora et al. (2024), Rasiah et al. (2024), Sakai et al. (2024), Khandelwal et al. (2024), Kargaran et al. (2024), Anand et al. (2025), Xu et al. (2024), She et al. (2024), Zhu et al. (2024b), Li et al. (2024c), Lim et al. (2024), Bajpai and Chakraborty (2025), Li et al. (2024b), Wei et al. (2024), Xie et al. (2024), Yang et al. (2024), Geng et al. (2024), Yang et al. (2025), Ko et al. (2025), Ruan et al. (2025), Lu et al. (2024), Agrawal et al. (2024), Ranaldi et al. (2025b), Ha (2025), Ranaldi et al. (2025a), Ranaldi and Pucci (2025)
Commonsense	Huang et al. (2024b), Fan et al. (2025), Payoungkhamdee et al. (2024), Luo et al. (2024), Song et al. (2024), Li et al. (2024a), Arora et al. (2024), Rasiah et al. (2024), Sakai et al. (2024), Khandelwal et al. (2024), Kargaran et al. (2024), Anand et al. (2025), Xu et al. (2024), She et al. (2024), Zhu et al. (2024b), Li et al. (2024c), Lim et al. (2024), Bajpai and Chakraborty (2025), Li et al. (2024b), Wei et al. (2024), Xie et al. (2024)

Table 3: Categorization of Papers by Reasoning Type