

# Estimating Uncertainty from Reasoning: A Large-Scale Study of Multi- and Crosslingual MCQA Performance in LLMs

Andrea Alfarano<sup>1,†\*</sup> Andrea Bacciu<sup>2,†</sup> Saab Mansour<sup>2</sup> Amin Mantrach<sup>2</sup> Marcello Federico<sup>2</sup>

<sup>1</sup>INSAIT, Sofia <sup>2</sup>Amazon

andrea.alfarano@insait.ai

{andbac, saabm, mantrach, marcfe}@amazon.com

## Abstract

Uncertainty estimation (UE) enables LLM-powered systems to recognize when to abstain, yet existing research has predominantly focused on English. We present the first large-scale evaluation of UE methods across 22 languages, spanning high-, mid-, and low-resource settings. Using two human-curated Q&A datasets, we compare open and closed box UE methods (nine in total) across different model sizes and architectures while eliciting long-form reasoning, avoiding LLM-as-a-judge and embedding-based scoring, which can introduce evaluation noise. We report three main actionable findings. First, we find that prompting models to reason in English while keeping questions in low-resource languages substantially improves UE performance, suggesting that comprehension of low-resource languages is largely intact, and that the reliability bottleneck lies in generation rather than understanding. Second, prompting models to reason in English closes the UE performance gap between low and high-resource languages, demonstrating that generation language matters more than the question language. Third, the choice of UE method should depend on model scale: at smaller scales, open-box probability-based methods outperform alternatives; at larger scales, closed-box self-verbalized uncertainty becomes superior. Finally, we provide an analysis of threshold selection for selective prediction, offering guidance on calibrating abstention in multilingual settings.

## 1 Introduction

Large Language Models (LLMs) have changed how people access and interact with information, supporting tasks from everyday planning to complex question-answering (Bommasani, 2021). Recent research has primarily focused on improving

task performance, for example through chain-of-thought prompting (Wei et al., 2022) and instruction tuning (Ouyang et al., 2022), enabling models to solve increasingly complex problems. However, high accuracy alone is insufficient: downstream systems must recognize when an LLM’s answer is not grounded in its knowledge, enabling them to abstain, defer to humans, or fall back to safer behavior. This has motivated a parallel line of work on uncertainty estimation (UE), which seeks to determine when models know the answer, enabling LLM-powered systems to recognize and communicate their lack of knowledge (Kuhn et al., 2023). Existing research on LLM uncertainty has predominantly focused on English, leaving limited evidence on whether UE methods maintain their efficacy in other languages, especially in low-resource settings (Kuhn et al. (2023); Kossen et al. (2024); Santilli et al. (2025); Cecere et al. (2025), inter alia). The only dedicated multilingual UE study we are aware of (Xue et al., 2025) evaluates just three methods on five languages, relies on machine-translated data without human post-editing, and uses a multiple-choice, short-answer format rather than open-ended generation. Their dataset yields a median answer length of just one word, meaning this setup primarily evaluates how language affects question comprehension while offering limited evidence about uncertainty during longer, more linguistically rich generation. Establishing multilingual trustworthiness through UE is also methodologically challenging because standard metrics such as AUROC rely on ground truth labels. When ground truth is approximated via LLM-as-judge, BERTScore, or n-gram overlap, these proxies introduce noise that can misrank uncertainty methods (Santilli et al., 2025). In multilingual settings, this problem is compounded: neural judges may behave inconsistently across languages, introducing language-dependent bias into UE comparisons. As a result, existing English-centric evaluations do not yet pro-

\*Work done during internship at Amazon.

†Equal contribution.

vide decisive evidence that UE methods transfer reliably to low-resource languages and long-form generation. To address these challenges, we evaluate nine UE methods on two human-curated QA datasets across 22 languages, covering high-, mid-, and low-resource settings and representing 80% of the world’s native speakers. We elicit long-form, language-rich reasoning (on average 150 words) while preserving objective correctness from QA labels, thereby avoiding LLM-as-judge and embedding-based proxies.

Our results reveal three key patterns: (1) UE method performance varies largely across language families, consistent with effects of resource availability and linguistic distance; (2) generation language has a larger impact on UE than question language, a factor largely overlooked by prior work; and (3) while the closed-box Self Verbalized method achieves the best overall performance, sampling-based closed-box methods degrade substantially on low-resource languages, whereas open-box probability-based methods maintain more consistent performance across all language resource levels.

We make the following contributions: (i) we propose a label-grounded framework for evaluating UE on open-ended multilingual generation without model-based correctness metrics; (ii) a large-scale study of UE across 9 models and 22 languages; and (iii) practical guidance for deploying uncertainty-aware systems in multilingual settings. Specifically, we address: (RQ1) which UE methods are robust across models and languages, (RQ2) how model scale affects UE efficacy, (RQ3) how the language of reasoning shapes UE quality, (RQ4) UE robustness under cross-lingual answer settings, and (RQ5) threshold selection strategies for selective prediction.

## 2 Related Work and Background

Uncertainty estimation (UE) is critical for deploying large language models (LLMs) in high-stakes environments, enabling systems to flag unreliable predictions. Existing literature is commonly categorized by the degree of model access available: *open-box methods*, which utilize internal model representations, and *closed-box methods*, which rely solely on textual outputs.

### 2.1 Open-Box Methods

Open-box approaches derive uncertainty directly from model internals, such as hidden states or token probabilities.

Inference method approaches require no additional training and aggregate token-level logits into sequence-level scores. Standard metrics include length-normalized log-probabilities, the geometric mean of token probabilities, and minimum token probability (Fomicheva et al., 2020; Vazhentsev et al., 2023). Inference-only methods remain the standard for open-weights models due to their computational efficiency and lack of supervision requirements.

### 2.2 Closed-Box Methods

Closed-box techniques are essential when access to logits or gradients is restricted (e.g., commercial APIs). These methods rely exclusively on the generated text.

**Self verbalized uncertainty** This approach prompts the LLM to explicitly state its confidence as a score or linguistic expression (Kadavath et al., 2022; Tian et al., 2023b). While accessible, self-verbalization is frequently miscalibrated and often struggles to discriminate correctness effectively compared to logit-based baselines (Xiong et al., 2024; Groot and Valdenegro Toro, 2024).

**Consistency and Sampling Strategies** Consistency-based methods rely on the intuition that correct reasoning is stable across stochastic samples. *Semantic Entropy* and its variants (Kuhn et al., 2023; Farquhar et al., 2024; Cecere et al., 2025) cluster multiple generations by meaning and compute entropy over these clusters. Other works extend this method using graph theory, modeling generations as nodes in a similarity graph to derive uncertainty from structural properties (Lin et al., 2023; Vashurin et al., 2025).

### 2.3 Evaluation and Multilingual Challenges

UE performance is typically evaluated via *discrimination* AUROC but reliable evaluation remains a challenge.

**Ground Truth Matching** Santilli et al. (2025) demonstrate that automated ground truth matching like ROUGE or BERTScore correlate strongly with superficial features (e.g., length) rather than factual accuracy, inflating reported UE performance.

To address this, we avoid automated ground truth matching and rely on exact matching of multiple-choice answers against human-curated labels. Prior work has also used MCQA exact matching to evaluate model reliability. For example, [Madhusudhan et al. \(2025\)](#) study whether models select an explicit “I Don’t Know” option as a binary abstention behavior. Our framework differs in three ways: we evaluate continuous-valued uncertainty signals (rather than binary abstention), extract them from long-form reasoning traces (rather than the answer choice itself), and operate across 22 languages (rather than English only).

**Multilingual Uncertainty Estimation** While UE is well-studied in English, multilingual evaluation remains underexplored. To the best of our knowledge, [Xue et al. \(2025\)](#) present the only dedicated multilingual UE study, introducing *Mling-Conf* to benchmark three uncertainty methods across five languages. However, their evaluation relies on machine-translated data without human post-editing and employs a short-answer format with a median response length of just one word. This setup primarily assesses how language affects question comprehension, offering limited insight into uncertainty during extended generation. Our benchmark addresses this gap by eliciting long-form reasoning in the target language while preserving objective correctness labels.

### 3 Methodology

Our goal is to evaluate LLM UE methods across multiple languages, including low-resource ones, while supporting long text generation and using model-free metrics to maintain an unbiased setup.

To achieve this, we leverage human-curated Multiple-Choice Question Answer (MCQA) datasets where the set of possible labels is a fixed set of options (such as A, B, C, and D). MCQA labels can be interpreted without relying on string-matching, LLM-as-a-judge, or embedding similarity techniques, thus providing ground truth without any approximation or language bias. MCQA labels are too short to derive uncertainty estimates from the model; for this reason, we apply UE methods to the LLM’s reasoning text, similarly to the approach of [Podolak and Verma \(2025\)](#).

Specifically, we prompt the model to generate a reasoning explanation before producing its answer (e.g., Chain of Thoughts ([Wei et al., 2022](#))). We then apply uncertainty estimation methods to this

reasoning text (on average 150 words), while using exact matching against MCQA labels to determine answer correctness. This allows us to retain the reliability of MCQA labels as ground truth while extending UE evaluation to long text generation. We deliberately choose MCQA over open-ended QA datasets for two reasons: (i) open-ended answers are extremely short factoids (median 1–2 words), so they do not provide longer text for UE than MCQA labels do, and (ii) their evaluation does not ensure a language- and correctness-unbiased setup. A detailed justification is provided in Appendix B. Finally, we evaluate the performance of each UE method using the Area Under the Receiver Operating Characteristic (AUROC) curve, which is the standard metric in the UE literature for measuring how well uncertainty scores correlate with answer correctness. The complete process is illustrated in Figure 1.

**Datasets:** We use two human-curated MCQA datasets, namely Global-MMLU ([Singh et al., 2025](#)) (with 4 possible choices) and MMLU-ProX ([Xuan et al., 2025](#)) (with 10 possible choices). We focus on the 22 languages that are common to both datasets (the complete list of languages is available in Appendix C). These two datasets are parallel, meaning that every question-answer pair is available in every language, making our multilingual evaluation comparable across languages. These two datasets cover a broad set of categories: Global-MMLU contains 6 categories, while MMLU-ProX contains 14 categories (see Appendix E). To ensure balanced representation, we sample 100 questions per category from each dataset, yielding 600 questions from Global-MMLU and 1,400 from MMLU-ProX, for a total of 2,000 unique questions. Since both datasets are available in all 22 languages, this results in 44,000 language-specific instances. We further verify that the MCQA labels are uniformly distributed in both datasets, ruling out positional bias (see App. J).

**Models:** Our evaluation spans nine models from both closed-source and open-source families. For closed-source models, we employ the Claude 4.5 Sonnet ([Anthropic, 2025](#)) - the state-of-the-art model at the time of writing. For open-source models, we consider instruction-tuned models including five Gemma3 models (0.27B, 1B, 4B, 12B, 27B) ([Team et al., 2025](#)) and three Qwen3 models (4B, 30B-A3B, 235B-A22B) ([Yang et al., 2025](#)), encompassing both dense and Mixture of Experts (MoE) architectures. We report a scaling study on

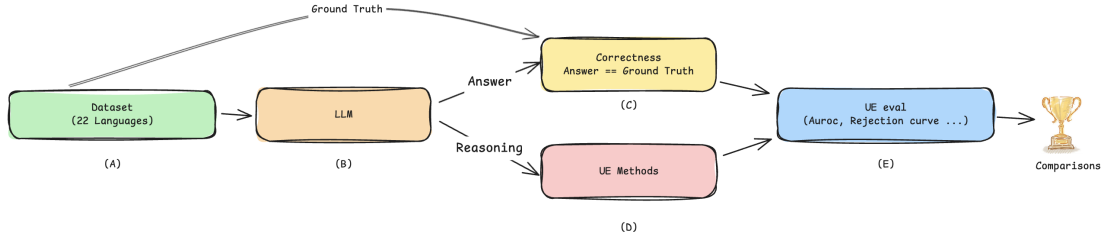


Figure 1: Evaluation pipeline. Unlike prior work that estimates uncertainty from predicted labels in English-only settings, our framework (1) extracts uncertainty from the model’s *reasoning trace* and (2) evaluates across 22 languages. The pipeline proceeds as: multilingual prompts (A) → LLM generation (B) → correctness evaluation against ground truth (C) → uncertainty estimation on reasoning (D) → method comparison via AUROC (E).

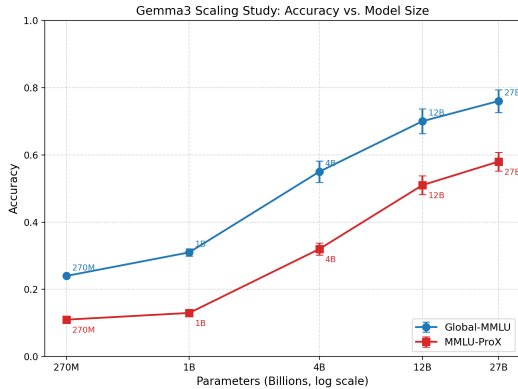


Figure 2: Scaling study on MCQA accuracy on both datasets (mean  $\pm$  95% CI) over 22 languages. Similar trends are observed in Qwen models.

the Q&A accuracy in Figure 2. In our experimentation we applied the quantization only to Qwen3 235B-A22B to 4bit due to hardware constrains (see our Hardware Infrastructure in Appendix H). In Appendix G, we investigate whether quantization impacts multilingual UE performance. We selected Qwen3-30B for this analysis as it is the largest model that fits in memory at full precision. Results show no statistically significant effect when quantizing to 8-bit and 4-bit.

**Uncertainty Estimation Methods:** Following the taxonomy in Section 2, we evaluate 9 methods, including both open-box methods (with access to model internal states) and closed-box methods. Open-box probability baselines include TokenEntropy, Maximum-Sequence-Probability (Max Prob), and Self Certainty (Fomicheva et al., 2020). Closed-box methods include Self-verbalized uncertainty (Tian et al., 2023a) and multi-sample disagreement based on lexical overlap as Lexical Similarity (ROUGE-L) (Fomicheva et al., 2020), semantic clustering such as Semantic Entropy (Kuhn et al., 2023), and graph-based dispersion

(EigValLaplacian\_Jaccard, DegMat\_Jaccard, Eccentricity\_NLI\_Jaccard) (Lin et al., 2024). More detailed method definitions are available in Appendix A. All the prompts used are listed in Appendix D. To run our experiments we leverage LM-Polygraph (Fadeeva et al., 2023), a framework that implements state-of-the-art UE methods. Generation parameters are those used in LM-Polygraph, following the original paper implementations.

## 4 Results And Discussion

**RQ1: Which UE methods are robust across models and languages?** To address this question, Figure 3 compares nine UE methods across 22 languages and two open-source models (Gemma3-27B and Qwen3-235B). We excluded closed-source models to ensure a fair comparison between open-box and closed-box approaches.

The best-performing method is “Self Verbalized”, which achieves an average AUROC of 0.72 and demonstrates consistent performance across all languages and models. The second-ranked method is the open-box “Token Entropy” with an AUROC of 0.71, followed by “Self Certainty” at 0.70. The open-box “Max Prob” performs reasonably well on average (0.68) but degrades substantially on low-resource languages, dropping to 0.56 on Yoruba and 0.57 on Nepali.

Sampling-based methods maintain reasonable performance on high-resource languages but fail on low-resource ones, with Eccentricity and Semantic Entropy dropping to near-random performance on Yoruba (0.49 and 0.50 AUROC). These methods generate multiple reasoning traces and use their diversity as an uncertainty signal—the intuition being that confident models produce consistent outputs while uncertain models vary. To understand why this fails, we analyze the diversity of reasoning traces across all 22 languages using pairwise Jac-

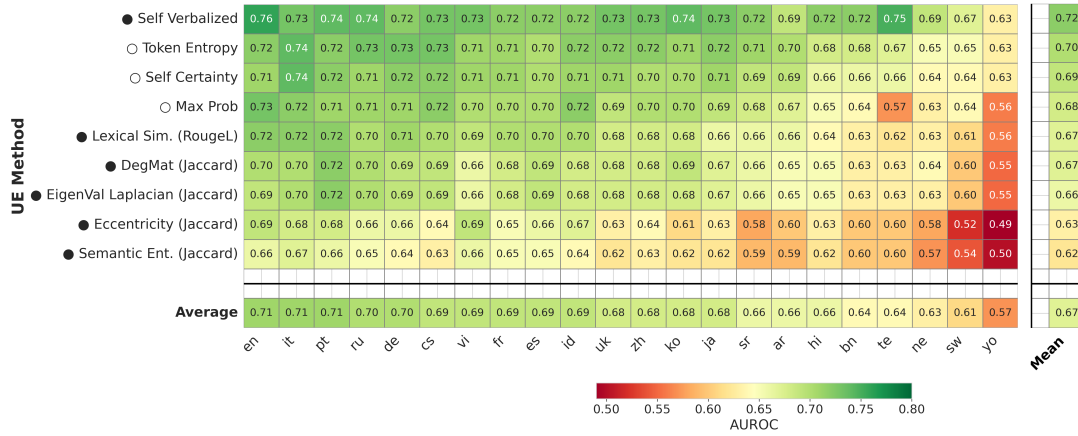


Figure 3: Per-language AUROC for nine UE methods across 22 languages on the dense Gemma3-27B and MoE Qwen3-235B-A22B models. Languages are sorted by mean AUROC (median sorting produces identical rankings). Closed-box UE methods are marked with a black circle while open-box are marked with a white circle. In the left block we average across models, in the middle block across languages, while in the right block we average across languages and models. Remarkably, the closed-box method *Self Verbalized* achieves top performance, rivaling open-box methods such as *Token Entropy* and *Self Certainty*—suggesting that effective UE is achievable without access to token-level probabilities.

card distance. For high-resource languages, the expected pattern holds: incorrect predictions produce more diverse reasoning than correct ones (English: 0.72 vs. 0.63,  $p < 10^{-6}$ ; German: 0.73 vs. 0.65; Spanish: 0.74 vs. 0.63; French: 0.69 vs. 0.64), with gaps of 0.08–0.11. For low-resource languages, this gap collapses to  $\leq 0.03$ : Yoruba (0.93 vs. 0.93), Swahili (0.93 vs. 0.92), Bengali (0.83 vs. 0.80). LLMs generate highly variable reasoning regardless of correctness, leaving no diversity signal to exploit. Self-verbalized approaches, by contrast, leverage meta-cognitive capabilities (Steyvers and Peters, 2025; Xu et al., 2025) that transfer more robustly across languages.

**RQ2: How does model scale affect multilingual UE method efficacy?** We investigate the effect of model parameters scale on UE performance, we evaluate the top-2 open-box (Token Entropy, Self Certainty) and top-2 closed-box (Self Verbalized, Lexical Sim.) methods (selected for figure clarity) that scale best across all languages (see RQ1); for all other methods, see App. I. In Figure 4, we computed the average across all the languages and averaging across datasets. We select Qwen3 as it offers the broadest parameter range among the open-source models in our study from 4B to 235B. We observe a clear divergence in scaling behavior. Token Entropy remains stable across all scales with overlapping confidence intervals indicating no statistically significant variation at the

change of the parameters scale. Self Verbalized, by contrast, improves substantially at increasing the parameter (from 0.65→0.66→0.77 AUROC); the difference between 4B and 30B is not statistically significant as the CIs are overlapping but when scaling to 235B, the difference becomes statistically significant. At 4B and 30B scale, Token Entropy significantly outperforms Self Verbalized, but this relationship reverses at 235B where Self Verbalized achieves the highest AUROC among all methods. These findings have practical implications: at smaller scales, open-box methods such as Token Entropy provide the most reliable uncertainty signal; at larger scales, Self Verbalized surpasses all alternatives by a significant margin, suggesting that meta-cognitive abilities required for accurate self-assessment emerge primarily at larger model scales, consistent with prior findings on self-knowledge in LLMs (Kadavath et al., 2022; Steyvers and Peters, 2025). However, that training pipelines are not uniform across model sizes within the same family: the Qwen3 Technical Report (Yang et al., 2025) discloses that smaller Qwen3 variants are distilled from larger teacher models (Qwen3-235B-A22B), so part of the observed gap may also reflect that self-verbalization capability are sensible to training procedures. Furthermore, in Figure 5 we show results across Gemma from 270M to 27B parameters, averaged across all nine UE methods. A clear scaling trend emerges again: models below 1B param-

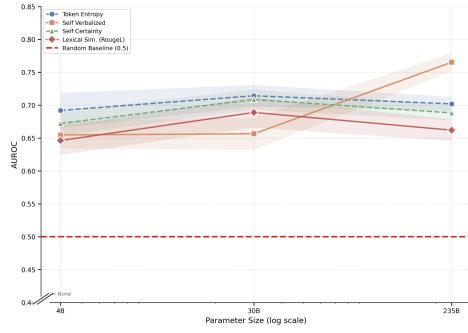


Figure 4: Effect of model scale on UE performance (Qwen3 family). Shaded regions denote 95% CIs. Dashed line: random baseline.

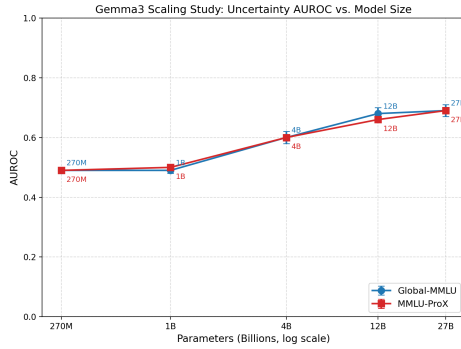


Figure 5: Uncertainty estimation performance (AUROC  $\pm$  CI 95%) across varying model scales. Results are averaged across 22 languages in both datasets. Increasing model size generally yields higher AUROC scores.

ters (Gemma3-0.27B, Gemma3-1B, Qwen3-0.6B) perform at random baseline (0.50 AUROC), indicating that UE methods fail to extract meaningful uncertainty signals from very small models. Performance increases substantially between 1B and 4B parameters, with achieving 0.60-0.65 AUROC. Beyond 12B parameters, performance plateaus around 0.67-0.70 AUROC across both model families. We report Expected Calibration Error in App. K, where larger models show a slight tendency toward overconfidence.

**RQ3: How does the language of reasoning affect UE quality across languages?** To investigate whether the language used for reasoning influences multilingual UE, we compare two configurations: *Multilingual*, where the model reasons in the target language (i.e., language of the question and answers), and *English Reasoning*, where reasoning is conducted in English regardless of the target language. Figure 6 reports AUROC with 95% confidence intervals across 22 languages, averaged across all datasets and the largest models

(Qwen 235B, Gemma3 27B, and Claude 4.5 Sonnet), grouped by WALS language genus<sup>1</sup>.

On average, English reasoning improves UE quality (0.72 vs. 0.68 AUROC). Crucially, this improvement is driven by low-resource languages: English reasoning effectively brings their UE performance to the level of high-resource languages. For Yoruba, AUROC improves from 0.58 to 0.68; for Swahili, from 0.58 to 0.64; for Nepali, from 0.66 to 0.71—all statistically significant gains (non-overlapping CIs). After applying English reasoning, these languages achieve AUROC scores comparable to Germanic and Romance languages (0.71–0.73), which show no significant change between conditions. These results reveal an asymmetry between natural language understanding and generation. In both conditions, the question remains in the target language—only the reasoning language changes. If comprehension were the bottleneck, changing the reasoning language would not help: the model would not know what to reason about. The fact that English reasoning closes the gap suggests that models comprehend low-resource questions adequately but struggle to generate coherent reasoning in those languages. By shifting generation to English, models leverage their stronger generative capabilities, producing stronger uncertainty signals.

**RQ4: How do UE methods perform under cross-lingual answer settings?** Multilingual applications often involve cross-lingual scenarios where different components of an input span multiple languages—a common situation in cross-lingual retrieval-augmented generation (RAG) systems where documents retrieved from multiple languages must be integrated. We evaluate whether UE methods remain robust when this cross-lingual complexity is introduced. We construct instances where the question and reasoning are in the target language, but each answer option is drawn from a different language. Leveraging the parallel nature of our multilingual datasets, we use aligned answer IDs across translations and sample each option from a different language variant. For instance, an English question asking “What day comes after Monday?” would present answer options such as “Mardi” (French, Tuesday), “Mercoledì” (Italian, Wednesday), “Donnerstag” (German, Thursday), and “Viernes” (Spanish, Friday). Crucially, unlike RQ3, the reasoning language always follows the

<sup>1</sup><https://wals.info/>

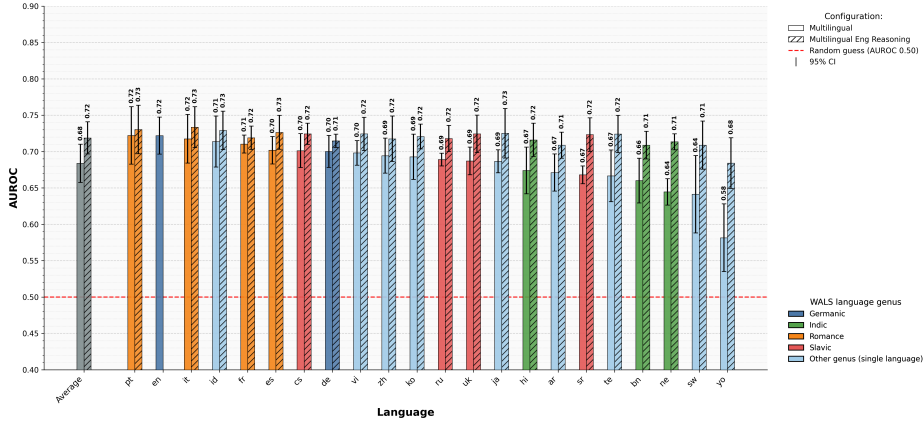


Figure 6: Effect of reasoning language on UE quality across 22 languages. Solid bars represent reasoning in target language condition; hatched bars represent the English-reasoning condition. AUROC scores are averaged across all UE methods, datasets, and the largest models (Qwen3-235B, Gemma3-27B, Claude Sonnet 4.5). Error bars denote 95% CI. English reasoning significantly narrows the performance gap for South Asian and African languages.

question language in both conditions. We evaluate using the largest models in our study: Gemma3-27B, Qwen3-235B-A22B, and Claude Sonnet 4.5. Figure 7 reports AUROC with 95% confidence intervals comparing the standard multilingual setting (solid bars) against the cross-lingual setting (hatched bars). This setting represents a more challenging evaluation scenario than standard multilingual UE. Despite this increased difficulty, the cross-lingual setting does not degrade UE performance. On average, AUROC remains stable (0.68 vs. 0.67) with overlapping confidence intervals across all language families. These results demonstrate that UE methods are robust to cross-lingual answer configurations, suggesting they can be reliably deployed in multilingual systems where language mixing is common.

#### RQ5: How can we find the best UE threshold across languages?

A key practical application of uncertainty estimation is selective prediction: using uncertainty scores to abstain from unreliable outputs (Madhusudhan et al., 2025), thereby improving system trustworthiness. The central challenge in multilingual settings is threshold calibration—selecting a filtering threshold on validation data for application at inference time. Given an uncertainty score  $u$  and threshold  $t$ , we accept the prediction if  $u < t$ , otherwise abstain (e.g., defer to human review). We select  $t$  by maximizing F1-score for error detection on validation data, comparing three calibration strategies: (1)  $t_{\text{EN}}$ , optimized on English data only and applied to all languages; (2)  $t_{\text{GLOBAL}}$ , optimized on pooled multilingual data;

and (3)  $t_{\text{LANG}}$ , using language-specific thresholds transferred to corresponding target languages. We use one dataset as validation set for threshold tuning and the other as test set with Claude Sonnet 4.5, comparing against Oracle upper bounds where thresholds are tuned directly on target test data.

| Strategy            | Acc $\uparrow$ | TP $\uparrow$ | FP $\downarrow$ | TN $\uparrow$ | FN $\downarrow$ |
|---------------------|----------------|---------------|-----------------|---------------|-----------------|
| No filtering        | 86.5           | 86.5          | 13.5            | 0.0           | 0.0             |
| $t_{\text{EN}}$     | 91.2           | 79.9          | 7.7             | 5.8           | 6.6             |
| $t_{\text{GLOBAL}}$ | 92.4           | 73.0          | 6.0             | 7.5           | 13.5            |
| $t_{\text{LANG}}$   | 93.0           | 72.1          | 5.4             | 8.1           | 14.4            |
| Oracle              | 97.2           | 72.2          | 5.1             | 8.4           | 14.4            |

Table 1: Impact of different threshold selection using Self-Verbalized (4.5 Sonnet). TP/TN: accepted predictions that are correct/wrong. FP/FN: filtered predictions that are wrong/correct. All values are percentages of total 44,000 samples averaged across the 2 datasets.

Table 1 reports the confusion matrix for each strategy. The selective prediction strategies reveal distinct trade-offs between error reduction, coverage, and deployment complexity. English-only calibration ( $t_{\text{EN}}$ ) offers the most practical entry point, requiring minimal data collection while achieving a 43% relative error reduction (from 13% to 7.7%) and maintaining high coverage with only 6.6% of correct predictions filtered out. Pooled multilingual optimization ( $t_{\text{GLOBAL}}$ ) strikes a middle ground, delivering 56% error reduction at the cost of moderate coverage loss (13.5% false negatives), making it suitable for applications where validation data spans multiple languages. Language-specific calibration ( $t_{\text{LANG}}$ ) achieves the strongest practical

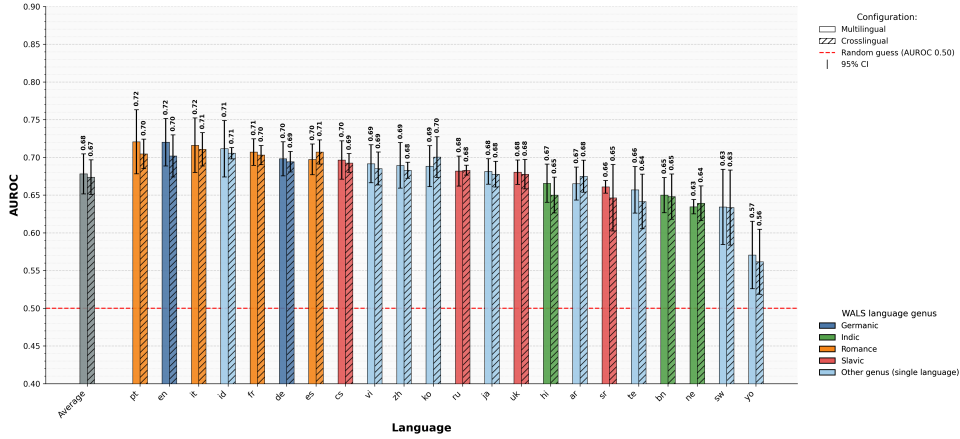


Figure 7: Effect of cross-lingual answer options on UE quality. Solid bars: multilingual setting; hatched bars: cross-lingual with each answer in a different language from the question language. Reasoning follows the question language in both conditions. AUROC scores are averaged across all UE methods, datasets, and the largest models (Qwen3-235B, Gemma3-27B, Sonnet 4.5). Error bars denote 95% CIs.

performance with 60% error reduction and 8.1% error detection, nearly matching Oracle thresholds tuned on target data itself, though this comes at the expense of filtering out 14.4% of correct predictions. The Oracle strategy establishes the performance ceiling at 62% error reduction but requires target-specific calibration data, limiting its practical applicability. These results demonstrate that the choice of strategy depends critically on application requirements: safety-critical systems benefit most from  $t_{\text{LANG}}$ 's superior error detection, resource-constrained deployments favor  $t_{\text{EN}}$ 's simplicity, and balanced applications achieve optimal accuracy-coverage trade-offs with  $t_{\text{GLOBAL}}$ .

## 5 Conclusion & Future Work

We present the largest multilingual evaluation of LLM uncertainty estimation to date, covering 9 UE methods, 9 models (5 Gemma3, 4 Qwen3, and Claude Sonnet 4.5) spanning 270M to 235B parameters, and 22 languages from high- to low-resource settings. Our experiments yield three actionable findings. First, generation language matters more than question language: prompting models to reason in English closes the UE performance gap for low-resource languages, indicating that the reliability bottleneck lies in generation rather than comprehension. Notably, Swahili, Yoruba, and Telugu—the lowest-performing languages in our benchmark—achieve performance on par with English when reasoning is elicited in English. Nevertheless, reasoning in the target language is important under certain scenarios (eg explanations

in target language for non English speakers), and closing the gap is left for future work.

Second, UE method selection should be scale-aware. All UE methods maintain stable performance across model scales, except for Self Verbalized, which improves at scale (0.65 AUROC at 4B to 0.77 at 235B), becoming the best-performing approach with larger LLMs. This trend is consistent with prior evidence that meta-cognitive capabilities emerge primarily in larger models (Tian et al., 2023c; Xiong et al., 2024; Steyvers and Peters, 2025), although training-pipeline differences across sizes (e.g., distillation) may also contribute.

Third, threshold calibration transfers across datasets but benefits from language awareness. English-only calibration halves error rates, while multilingual calibration reduces errors by up to 77%. We also observe that sampling-based consistency methods degrade on low-resource languages. Since these methods rely on lexical overlap measures, sensitivity to surface-level variation may explain this degradation, though further investigation is needed. We release our evaluation framework to support research on trustworthy multilingual AI. Future directions include extending to open-ended generation with reliable correctness assessment and developing UE methods explicitly designed for cross-linguistic robustness.

## 6 Limitations

We focus on two MCQA datasets, which, while providing reliable ground-truth labels without model-based proxies, may not fully represent the diversity

of real-world NLP applications. The landscape of uncertainty estimation methods is vast, and we considered in our analysis 9 state-of-the-art methods selected to cover the main paradigms (probability-based, verbalized, and consistency-based), we did not include any training-based methods.

## References

- Anthropic. 2025. Introducing Claude Sonnet 4.5. <https://www.anthropic.com/news/claude-sonnet-4-5>. Accessed: 2026-05-25.
- Rishi Bommasani. 2021. On the opportunities and risks of foundation models. *arXiv preprint arXiv:2108.07258*.
- Nicola Cecere, Andrea Bacciu, Ignacio Fernández-Tobías, and Amin Mantrach. 2025. [Monte Carlo temperature: a robust sampling strategy for LLM’s uncertainty quantification methods](#). In *Proceedings of the 5th Workshop on Trustworthy NLP (TrustNLP 2025)*, pages 305–320, Albuquerque, New Mexico. Association for Computational Linguistics.
- Ekaterina Fadeeva, Roman Vashurin, Akim Tsvigun, Artem Vazhentsev, Sergey Petrakov, Kirill Fedyanin, Daniil Vasilev, Elizaveta Goncharova, Alexander Panchenko, Maxim Panov, Timothy Baldwin, and Artem Shelmanov. 2023. [LM-polygraph: Uncertainty estimation for language models](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 446–461, Singapore. Association for Computational Linguistics.
- Sebastian Farquhar, Jannik Kossen, Lorenz Kuhn, and Yarin Gal. 2024. Detecting hallucinations in large language models using semantic entropy. *Nature*, 630(8017):625–630.
- Marina Fomicheva, Shuo Sun, Lisa Yankovskaya, Frédéric Blain, Francisco Guzmán, Mark Fishel, Nikolaos Aletras, Vishrav Chaudhary, and Lucia Specia. 2020. [Unsupervised quality estimation for neural machine translation](#). *Transactions of the Association for Computational Linguistics*, 8:539–555.
- Tobias Groot and Matias Valdenegro Toro. 2024. [Overconfidence is key: Verbalized uncertainty evaluation in large language and vision-language models](#). In *Proceedings of the 4th Workshop on Trustworthy Natural Language Processing (TrustNLP 2024)*, pages 145–171, Mexico City, Mexico. Association for Computational Linguistics.
- Mandar Joshi, Eunsol Choi, Daniel Weld, and Luke Zettlemoyer. 2017. [TriviaQA: A large scale distantly supervised challenge dataset for reading comprehension](#). In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1601–1611, Vancouver, Canada. Association for Computational Linguistics.
- Saurav Kadavath, Tom Conerly, Amanda Askell, Tom Henighan, Dawn Drain, Ethan Perez, Nicholas Schiefer, Zac Hatfield-Dodds, Nova DasSarma, Eli Tran-Johnson, and 1 others. 2022. Language models (mostly) know what they know. *arXiv preprint arXiv:2207.05221*.
- Jannik Kossen, Jiatong Han, Muhammed Razzak, Lisa Schut, Shreshth Malik, and Yarin Gal. 2024. Semantic entropy probes: Robust and cheap hallucination detection in LLMs. *arXiv preprint arXiv:2406.15927*.
- Lorenz Kuhn, Yarin Gal, and Sebastian Farquhar. 2023. Semantic uncertainty: Linguistic invariances for uncertainty estimation in natural language generation. *arXiv preprint arXiv:2302.09664*.
- Tom Kwiatkowski, Jennimaria Palomaki, Olivia Redfield, Michael Collins, Ankur Parikh, Chris Alberti, Danielle Epstein, Illia Polosukhin, Jacob Devlin, Kenton Lee, Kristina Toutanova, Llion Jones, Matthew Kelcey, Ming-Wei Chang, Andrew M. Dai, Jakob Uszkoreit, Quoc Le, and Slav Petrov. 2019. [Natural questions: A benchmark for question answering research](#). *Transactions of the Association for Computational Linguistics*, 7:452–466.
- Chin-Yew Lin. 2004. [ROUGE: A package for automatic evaluation of summaries](#). In *Text Summarization Branches Out*, pages 74–81, Barcelona, Spain. Association for Computational Linguistics.
- Zhen Lin, Shubhendu Trivedi, and Jimeng Sun. 2023. Generating with confidence: Uncertainty quantification for black-box large language models. *arXiv preprint arXiv:2305.19187*.
- Zhen Lin, Shubhendu Trivedi, and Jimeng Sun. 2024. Generating with confidence: Uncertainty quantification for black-box large language models. *Transactions on Machine Learning Research*. ArXiv preprint arXiv:2305.19187.
- Shayne Longpre, Yi Lu, and Joachim Daiber. 2021. [MKQA: A linguistically diverse benchmark for multilingual open domain question answering](#). *Transactions of the Association for Computational Linguistics*, 9:1389–1406.
- Nishanth Madhusudhan, Sathwik Tejaswi Madhusudhan, Vikas Yadav, and Masoud Hashemi. 2025. [Do LLMs know when to NOT answer? investigating abstention abilities of large language models](#). In *Proceedings of the 31st International Conference on Computational Linguistics*, pages 9329–9345, Abu Dhabi, UAE. Association for Computational Linguistics.
- Andrey Malinin and Mark Gales. 2021. [Uncertainty estimation in autoregressive structured prediction](#). In *International Conference on Learning Representations*.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang,

- Sandhini Agarwal, Katarina Slama, Alex Ray, and 1 others. 2022. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744.
- Jakub Podolak and Rajeev Verma. 2025. Read your own mind: Reasoning helps surface self-confidence signals in llms. *arXiv preprint arXiv:2505.23845*.
- Andrea Santilli, Adam Golinski, Michael Kirchhof, Federico Danieli, Arno Blaas, Miao Xiong, Luca Zappella, and Sinead Williamson. 2025. [Revisiting uncertainty quantification evaluation in language models: Spurious interactions with response length bias results](#). In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, page 743–759. Association for Computational Linguistics.
- Shivalika Singh, Angelika Romanou, Cl  mentine Fourrier, David I. Adelani, Jian Gang Ngui, Daniel Vila-Suero, Peerat Limkonchotiwat, Kelly Marchisio, Wei Qi Leong, Yosephine Susanto, Raymond Ng, Shayne Longpre, Wei-Yin Ko, Sebastian Ruder, Madeline Smith, Antoine Bosselut, Alice Oh, Andre F. T. Martins, Leshem Choshen, and 5 others. 2025. [Global mmlu: Understanding and addressing cultural and linguistic biases in multilingual evaluation](#). *Preprint*, arXiv:2412.03304.
- Mark Steyvers and Megan AK Peters. 2025. Metacognition and uncertainty communication in humans and large language models. *Current Directions in Psychological Science*, page 09637214251391158.
- Gemma Team, Aishwarya Kamath, Johan Ferret, Shreya Pathak, Nino Vieillard, Ramona Merhej, Sarah Perrin, Tatiana Matejovicova, Alexandre Ram  , Morgane Rivi  re, Louis Rouillard, Thomas Mesnard, Geoffrey Cideron, Jean bastien Grill, Sabela Ramos, Edouard Yvinec, Michelle Casbon, Etienne Pot, Ivo Penchev, and 197 others. 2025. [Gemma 3 technical report](#). *Preprint*, arXiv:2503.19786.
- Katherine Tian, Eric Mitchell, Hugh Zhao, Gintare Dziugaite, Amos Azaria, and Samuel R. Bowman. 2023a. [Just ask for calibration: Strategies for eliciting calibrated confidence scores from language models fine-tuned with human feedback](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 12302–12328, Singapore. Association for Computational Linguistics.
- Katherine Tian, Eric Mitchell, Allan Zhou, Archit Sharma, Rafael Rafailov, Huaxiu Yao, Chelsea Finn, and Christopher Manning. 2023b. Just ask for calibration: Strategies for eliciting calibrated confidence scores from language models fine-tuned with human feedback. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 5433–5442, Singapore. Association for Computational Linguistics.
- Katherine Tian, Eric Mitchell, Allan Zhou, Archit Sharma, Rafael Rafailov, Huaxiu Yao, Chelsea Finn, and Christopher D. Manning. 2023c. Just ask for calibration: Strategies for eliciting calibrated confidence scores from language models fine-tuned with human feedback. *arXiv preprint arXiv:2305.14975*.
- Roman Vashurin, Ekaterina Fadeeva, Artem Vazhentsev, Lyudmila Rvanova, Daniil Vasilev, Akim Tsvigun, Sergey Petrakov, Rui Xing, Abdelrahman Sadallah, Kirill Grishchenkov, Alexander Panchenko, Timothy Baldwin, Preslav Nakov, Maxim Panov, and Artem Shelmanov. 2025. [Benchmarking uncertainty quantification methods for large language models with LM-polygraph](#). *Transactions of the Association for Computational Linguistics*, 13:220–248.
- Artem Vazhentsev, Akim Tsvigun, Roman Vashurin, Sergey Petrakov, Daniil Vasilev, Maxim Panov, Alexander Panchenko, and Artem Shelmanov. 2023. Efficient out-of-domain detection for sequence to sequence models. In *Findings of the Association for Computational Linguistics: ACL 2023*, pages 1430–1454, Toronto, Canada. Association for Computational Linguistics.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Brian Ichter, Fei Xia, Ed H. Chi, Quoc V. Le, and Denny Zhou. 2022. Chain-of-thought prompting elicits reasoning in large language models. In *Proceedings of the 36th International Conference on Neural Information Processing Systems, NIPS ’22*, Red Hook, NY, USA. Curran Associates Inc.
- Miao Xiong, Zhiyuan Hu, Xinyang Lu, YIFEI LI, Jie Fu, Junxian He, and Bryan Hooi. 2024. Can llms express their uncertainty? an empirical evaluation of confidence elicitation in llms. In *The Twelfth International Conference on Learning Representations*.
- Zhengtao Xu, Tianqi Song, and Yi-Chieh Lee. 2025. Confronting verbalized uncertainty: Understanding how llm’s verbalized uncertainty influences users in ai-assisted decision-making. *International Journal of Human-Computer Studies*, 197:103455.
- Weihao Xuan, Rui Yang, Heli Qi, Qingcheng Zeng, Yunze Xiao, Aosong Feng, Dairui Liu, Yun Xing, Junjue Wang, Fan Gao, Jinghui Lu, Yuang Jiang, Huitao Li, Xin Li, Kunyu Yu, Ruihai Dong, Shangding Gu, Yuekang Li, Xiaofei Xie, and 13 others. 2025. [Mmlu-prox: A multilingual benchmark for advanced large language model evaluation](#). *Preprint*, arXiv:2503.10497.
- Boyang Xue, Hongru Wang, Rui Wang, Sheng Wang, Zezhong Wang, Yiming Du, Bin Liang, Wenxuan Zhang, and Kam-Fai Wong. 2025. [MlingConf: A comprehensive study of multilingual confidence estimation on large language models](#). In *Findings of the Association for Computational Linguistics: ACL 2025*, pages 2535–2556, Vienna, Austria. Association for Computational Linguistics.
- An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, Chujie Zheng, Dayiheng Liu, Fan Zhou, Fei Huang, Feng Hu, Hao

Ge, Haoran Wei, Huan Lin, Jialong Tang, and 41 others. 2025. [Qwen3 technical report](#). *Preprint*, arXiv:2505.09388.

## A Uncertainty Estimation Methods

We classify Uncertainty Estimation (UE) methods into two primary categories based on the level of access required to the underlying Large Language Model (LLM). **Open-box methods** derive uncertainty from the model’s internal probability distributions (logits), whereas **Closed-box methods** operate solely on the generated text, utilizing either self-reflection or the consistency between multiple stochastic samples.

Throughout, we denote the input prompt by  $x$ , the model’s parameters by  $\phi$ , and the generated sequence by  $\mathbf{y} = (y_1, \dots, y_T)$ . The model’s predictive distribution at step  $t$  is denoted as  $p_\phi(y_t \mid \mathbf{y}_{<t}, x)$ .

### A.1 Open-Box Methods

These methods require white-box access to token-level probabilities. They are computationally efficient as they typically operate on a single greedy generation, but they are inapplicable to API-based models where logits are unavailable.

**Maximum Sequence Probability** We estimate the confidence of a sequence  $\mathbf{y}$  using its joint probability. To normalize for varying sequence lengths, we compute the negative average log-likelihood (NLL) (Malinin and Gales, 2021):

$$u_{\text{MSP}}(x) = \text{Norm} \left( -\frac{1}{T} \sum_{t=1}^T \log p_\phi(y_t \mid \mathbf{y}_{<t}, x) \right). \quad (1)$$

Here,  $\text{Norm}(\cdot)$  denotes min-max normalization to the interval  $[0, 1]$  computed over the evaluation dataset.

**Mean Token Entropy** Rather than relying on the probability of the specific token chosen, this method measures the average distributional uncertainty (Shannon entropy) over the entire vocabulary  $\mathcal{V}$  at each generation step (Fomicheva et al., 2020):

$$u_{\text{MTE}}(x) = \text{Norm} \left( \frac{1}{T} \sum_{t=1}^T \sum_{v \in \mathcal{V}} -p_\phi(v) \log p_\phi(v) \right), \quad (2)$$

where  $p_\phi(v)$  is shorthand for  $p_\phi(v \mid \mathbf{y}_{<t}, x)$ .

### A.2 Closed-Box Methods

These methods treat the LLM as a black box. We further divide them into *Verbalized* methods (single generation, self-reflection) and *Consistency* methods (multiple generations, aggregation).

#### A.2.1 Verbalized Uncertainty

These approaches prompt the model to explicitly state its confidence, relying on the assumption that instruction-tuned models possess internal calibration capabilities (Kadavath et al., 2022; Tian et al., 2023c).

**Self-Certainty.** We append a query to the context asking the model to rate its confidence  $s \in [0, 1]$ . The uncertainty is defined as the complement of the extracted score:

$$u_{\text{SelfCert}}(x) = 1 - s_\phi(x). \quad (3)$$

If the model fails to output a valid numerical score, we assign maximum uncertainty ( $u = 1$ ).

**Qualitative Verbalization** We prompt the model to describe its confidence using natural language (e.g., “quite sure”, “very uncertain”). We map these linguistic markers to discrete scalar values  $u \in \{0.0, 0.25, 0.5, 0.75, 1.0\}$  based on a predefined lexical heuristic.

#### A.2.2 Consistency and Graph-based Uncertainty

These methods approximate uncertainty by measuring the semantic dispersion of  $k$  stochastic samples  $\mathcal{Y} = \{\mathbf{y}^{(1)}, \dots, \mathbf{y}^{(k)}\}$  generated at temperature  $T = 1$ .

**Semantic Entropy (Jaccard)** Following Kuhn et al. (2023), we cluster the generations in  $\mathcal{Y}$  based on Jaccard similarity. Let  $C_1, \dots, C_m$  be the resulting semantic clusters, and  $\hat{p}(C_j) = |C_j|/k$  be the empirical probability of cluster  $j$ . The uncertainty is the entropy of the cluster distribution:

$$u_{\text{SemEnt}}(x) = \text{Norm} \left( -\sum_{j=1}^m \hat{p}(C_j) \log \hat{p}(C_j) \right) \quad (4)$$

**Lexical Similarity.** We compute the average pairwise similarity using the ROUGE-L (Lin, 2004) metric as a proxy for consistency. Higher disagreement implies higher uncertainty:

$$u_{\text{LexSim}}(x) = \text{Norm} \left( 1 - \binom{k}{2}^{-1} \sum_{i < j} \text{RougeL}(\mathbf{y}^{(i)}, \mathbf{y}^{(j)}) \right) \quad (5)$$

**Graph-Theoretic Measures.** Inspired by Lin et al. (2023), we construct a semantic similarity graph  $G = (V, E)$  where nodes are generations and edges exist if  $\text{Jac}(\mathbf{y}^{(i)}, \mathbf{y}^{(j)}) \geq \tau$ . We extract three topological uncertainty features:

1. **Laplacian Eigenvalue:** We compute the spectral radius (largest eigenvalue) of the graph Laplacian  $L$ , denoted  $\lambda_{\max}(L)$ . A higher  $\lambda_{\max}$  indicates a partitioned graph structure (conflicting answers):

$$u_{\text{Eig}}(x) = \text{Norm}(\lambda_{\max}(L)). \quad (6)$$

2. **Degree Concentration:** We measure the dispersion of node degrees to detect if a dominant consensus cluster exists. We define the score using the mean  $\mu_d$  and variance  $\sigma_d^2$  of the degrees:

$$u_{\text{Deg}}(x) = \text{Norm}\left(\frac{1}{(\mu_d + \epsilon)(\sigma_d + \epsilon)}\right). \quad (7)$$

3. **Eccentricity:** We measure the geometric spread of the generations using a weighted graph combining Jaccard and NLI scores. Uncertainty is the mean eccentricity (maximum shortest-path distance from each node):

$$u_{\text{Ecc}}(x) = \text{Norm}\left(\frac{1}{k} \sum_{i=1}^k \max_j d_G(i, j)\right). \quad (8)$$

## B Why not open-ended QA?

An alternative to MCQA would be the use of open-ended QA datasets such as TriviaQA (Joshi et al., 2017), Natural Questions (Kwiatkowski et al., 2019), or their multilingual counterpart MKQA (Longpre et al., 2021). We do not adopt this setup, because correctness on these datasets must still be assessed on the open-ended string, which requires regex matching, an LLM-as-a-judge, or embedding similarity, reintroducing exactly the biases in the correctness signal that we want to avoid. Open-ended answers are also extremely short factoids (e.g. ‘‘Albert Einstein’’): the median is 1 word for TriviaQA and 2 words for Natural Questions, so they do not provide longer text for UE than MCQA labels do, and reasoning elicitation would be needed in either case.

MCQA with elicited reasoning is therefore the only configuration that simultaneously satisfies all

three requirements: (1) multilingual coverage including low-resource languages, (2) sufficiently long generations for UE, (3) and unbiased model-free correctness.

## C Languages

Our benchmark covers 22 languages given by the intersection of Global-MMLU and MMLU-ProX, so that every instance has a strictly parallel multiple-choice question and shared answer key across all languages. The languages and their ISO codes are: Arabic (ar), Bengali (bn), Czech (cs), German (de), English (en), Spanish (es), French (fr), Hindi (hi), Indonesian (id), Italian (it), Japanese (ja), Korean (ko), Nepali (ne), Portuguese (pt), Russian (ru), Serbian (sr), Swahili (sw), Telugu (te), Urdu (ur), Vietnamese (vi), Yoruba (yo), and Chinese (zh).

## D Prompts

This section details the prompts used in our experiments. All uncertainty estimation methods share the same base prompt for answer generation. The *self-verbalized* method additionally employs a follow-up prompt to elicit a confidence score from the model.

### D.1 Answer Generation Prompt

Provide your best guess for the following question about {subject} selecting one of the options. Before giving your answer, provide a step-by-step explanation of your thought process. Then on a new line give the guess with no other words or explanation.

For example:

Explanation: <one sentence step-by-step explanation of your thought process>  
Guess: <most likely guess, as short as possible; not a complete sentence, just the guess!>

Q: {question}  
A. {choice\_A}  
B. {choice\_B}  
C. {choice\_C}  
D. {choice\_D}

### D.2 Self-Verbalized Confidence Prompt

Provide the probability that your guess is correct. Give ONLY the probability, no other words or explanation.

For example:

Probability: <the probability between 0.0 and 1.0 that your guess is correct, without any extra commentary whatsoever; just the probability!>

To transform the probability into the uncertainty score we computed the uncertainty as  $1 - \text{probability score}$ .

## E Question Answering Categories

We report the list of categories in both datasets used. The categories of Global-MMLU are [business, humanities, medical, other, stem, social sciences]. The categories of MMLU-ProX are [biology, business, chemistry, computer science, economics, engineering, health, history, law, math, other, philosophy, physics, psychology]

## F Cross-lingual Q&A Accuracy comparison

This section provides task accuracy results for the cross-lingual evaluation discussed in RQ4 (Section 4). In the cross-lingual setting, the question and reasoning are in the target language, but each answer option is drawn from a different language.

We compared the task accuracy between the standard multilingual setting and the cross-lingual setting. As expected, the cross-lingual setting is more challenging: accuracy drops by 4–6 percentage points across all models across all the scales. This decrease reflects the additional difficulty of matching reasoning to answer options expressed in different languages.

Crucially, while task accuracy decreases, uncertainty estimation performance (AUROC) remains stable (see Figure 7). This indicates that UE methods successfully capture model confidence regardless of the cross-lingual configuration, and do not rely solely on lexical matching between reasoning and answer options.

Beyond uncertainty estimation, we also investigate whether English reasoning improves task accuracy. Table 2 reports a comparison between language-specific reasoning and English reasoning in terms of Q&A accuracy, showing the absolute difference ( $\Delta_{\text{abs}}$ ) and the relative improvement ( $\Delta_{\text{rel}}\%$ ) across the 22 languages considered in our study, averaged across the largest models of each family evaluated on Global-MMLU and MMLU-ProX.

The results indicate that English reasoning consistently improves task accuracy, with an average absolute gain of +0.037 (relative improvement of +5.4%). The improvement is most pronounced for lower-resource languages such as Yoruba (YO, +17.6%), Nepali (NE, +10.7%), and Swahili (SW, +10.5%), while higher-resource languages such as French (FR), Portuguese (PT), and German (DE) exhibit smaller gains. This suggests that English reasoning not only helps models better recognize their uncertainty but also enhances their ability to answer correctly, particularly in languages where the model’s native reasoning capabilities are weaker.

| Lang.      | Lang-Spec.        | Eng. Reason.      | $\Delta_{\text{abs}}$ | $\Delta_{\text{rel}}\%$ |
|------------|-------------------|-------------------|-----------------------|-------------------------|
| YO         | 0.581 $\pm$ 0.047 | 0.684 $\pm$ 0.035 | +0.103                | +17.6%                  |
| NE         | 0.644 $\pm$ 0.018 | 0.713 $\pm$ 0.011 | +0.069                | +10.7%                  |
| SW         | 0.641 $\pm$ 0.053 | 0.709 $\pm$ 0.033 | +0.068                | +10.5%                  |
| TE         | 0.666 $\pm$ 0.036 | 0.724 $\pm$ 0.026 | +0.057                | +8.6%                   |
| SR         | 0.668 $\pm$ 0.012 | 0.723 $\pm$ 0.023 | +0.055                | +8.3%                   |
| BN         | 0.660 $\pm$ 0.031 | 0.709 $\pm$ 0.019 | +0.049                | +7.4%                   |
| HI         | 0.674 $\pm$ 0.032 | 0.716 $\pm$ 0.023 | +0.042                | +6.2%                   |
| JA         | 0.686 $\pm$ 0.016 | 0.725 $\pm$ 0.034 | +0.039                | +5.6%                   |
| AR         | 0.671 $\pm$ 0.026 | 0.709 $\pm$ 0.018 | +0.038                | +5.6%                   |
| UK         | 0.687 $\pm$ 0.019 | 0.724 $\pm$ 0.026 | +0.037                | +5.4%                   |
| RU         | 0.689 $\pm$ 0.009 | 0.718 $\pm$ 0.018 | +0.029                | +4.2%                   |
| KO         | 0.693 $\pm$ 0.031 | 0.721 $\pm$ 0.017 | +0.028                | +4.0%                   |
| VI         | 0.698 $\pm$ 0.017 | 0.724 $\pm$ 0.023 | +0.026                | +3.8%                   |
| ES         | 0.702 $\pm$ 0.019 | 0.726 $\pm$ 0.023 | +0.025                | +3.5%                   |
| ZH         | 0.694 $\pm$ 0.024 | 0.718 $\pm$ 0.031 | +0.023                | +3.3%                   |
| CS         | 0.701 $\pm$ 0.023 | 0.724 $\pm$ 0.015 | +0.023                | +3.3%                   |
| IT         | 0.718 $\pm$ 0.034 | 0.734 $\pm$ 0.028 | +0.016                | +2.2%                   |
| ID         | 0.714 $\pm$ 0.035 | 0.729 $\pm$ 0.026 | +0.015                | +2.1%                   |
| DE         | 0.700 $\pm$ 0.022 | 0.715 $\pm$ 0.009 | +0.015                | +2.1%                   |
| FR         | 0.710 $\pm$ 0.012 | 0.719 $\pm$ 0.016 | +0.009                | +1.2%                   |
| PT         | 0.722 $\pm$ 0.039 | 0.730 $\pm$ 0.033 | +0.008                | +1.1%                   |
| <b>Avg</b> | <b>0.682</b>      | <b>0.719</b>      | <b>+0.037</b>         | <b>+5.4%</b>            |

Table 2: Q&A accuracy comparison between language-specific reasoning and English reasoning, averaged across the largest models of each family on Global-MMLU and MMLU-ProX.

## G Does the quantization impact multilingual UE performances?

In Figure 8, we compare the performance of full-precision Qwen3-30B against its 8-bit and 4-bit quantizations. We selected Qwen3-30B as it is the largest model in our study that fits in memory at full precision, making it the most relevant candidate for quantization. As can be noticed, there is no statistically significant impact from quantizing the model.

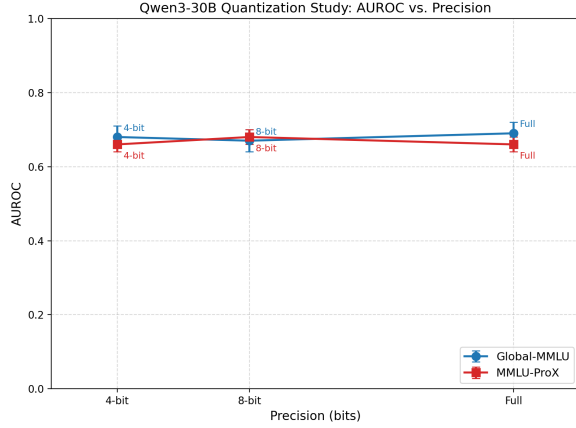


Figure 8: Uncertainty estimation performance for Qwen3-30B with quantization to 8-bit and 4-bit. Quantization has minimal impact on both accuracy and uncertainty estimation at this model scale. ( $\pm$  indicates 95% CI).

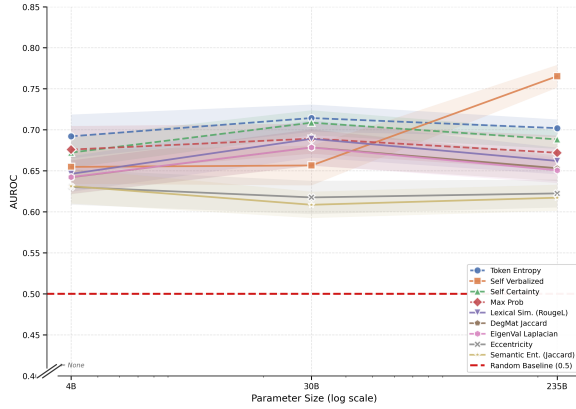


Figure 9: Effect of model scale on UE performance (Qwen3 family). Shaded regions denote 95% CIs. Self Verbalized improves with statistical significance at 235B. Dashed line: random baseline. This figure complements the Figure 4 presented in the main where for simplicity we show only 4 methods.

## H Hardware Infrastructure

All experiments were conducted on an Amazon EC2 g6e.48xlarge instance with an AMD EPYC 7R13 processor (192 vCPUs), 1.5 TiB of system memory, and 8 NVIDIA L40S GPUs (48GB VRAM each, 384GB total).

## I Complete Multilingual UE Method Scaling Effect

In Figure 9, we report the model size scaling factor with all the UE methods.

## J Label Distribution Analysis for Positional Bias

To rule out positional bias in our evaluation, we analyzed the distribution of correct answer labels across the two MCQA benchmarks used in our study. Tables 3 and 4 show that, for both Global MMLU (4 options, expected uniform = 25%) and MMLU-ProX (10 options, expected uniform = 10%), all labels fall within approximately  $\pm 2$  percentage points of the uniform expectation. We therefore conclude that label distributions are sufficiently balanced to exclude positional bias as a confounding factor in the results reported in the main paper.

| Label     | A     | B     | C     | D     |
|-----------|-------|-------|-------|-------|
| Freq. (%) | 22.96 | 24.65 | 25.50 | 26.90 |

Table 3: Label distribution for Global MMLU.

| Label | Freq. (%) |
|-------|-----------|
| A     | 11.57     |
| B     | 11.14     |
| C     | 10.93     |
| D     | 11.11     |
| E     | 9.58      |
| F     | 9.42      |
| G     | 9.78      |
| H     | 9.32      |
| I     | 9.22      |
| J     | 7.92      |

Table 4: Label distribution for MMLU-ProX.

## K Calibration Analysis

Table 9 extends the AUROC results from the main paper (Table 10) with Expected Calibration Error (ECE), which measures the alignment between predicted confidence and actual accuracy. Lower ECE indicates better calibration.

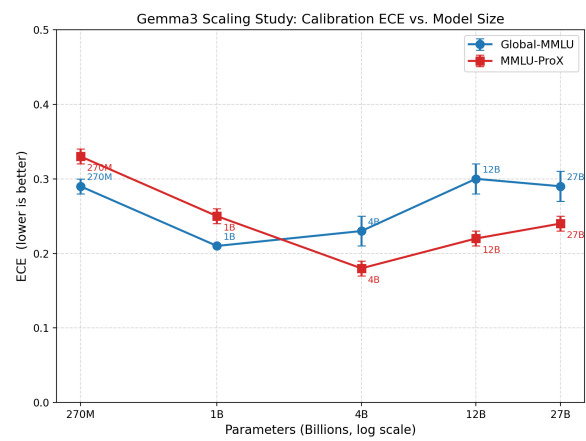


Figure 10: Effect of model scale on ECE. (Mean across langs and 95% CIs).