

Language Mixing in Reasoning Language Models: Patterns, Impact, and Internal Causes

Mingyang Wang^{1,2,3} Lukas Lange¹ Heike Adel⁴
Yunpu Ma^{2,3} Jannik Strötgen⁵ Hinrich Schütze^{2,3}

¹Bosch Center for Artificial Intelligence, Renningen, Germany

²LMU Munich, Germany ³Munich Center for Machine Learning (MCML)

⁴Hochschule der Medien, Stuttgart, Germany

⁵Karlsruhe University of Applied Sciences, Germany

mingyang@cis.lmu.de

Abstract

Reasoning language models (RLMs) excel at complex tasks by leveraging a chain-of-thought process to generate structured intermediate steps. However, *language mixing*, i.e., reasoning steps containing tokens from languages other than the prompt, has been observed in their outputs and shown to affect performance, though its impact remains debated. We present the first systematic study of language mixing in RLMs, examining its patterns, impact, and internal causes across 15 languages, 7 task difficulty levels, and 18 subject areas, and show how all three factors influence language mixing. Moreover, we demonstrate that the choice of reasoning language significantly affects performance: forcing models to reason in Latin or Han scripts via constrained decoding notably improves accuracy. Finally, we show that the script composition of reasoning traces closely aligns with that of the model’s internal representations, indicating that language mixing reflects latent processing preferences in RLMs. Our findings provide actionable insights for optimizing multilingual reasoning and open new directions for controlling reasoning languages to build more interpretable and adaptable RLMs.¹

1 Introduction

Reasoning language model (RLMs)², such as OpenAI’s o1 and o3 (Jaech et al., 2024; OpenAI, 2025), and the DeepSeek-R1 series (Guo et al., 2025), have demonstrated impressive capabilities in solving complex tasks through structured chain-of-thought reasoning. These models generate intermediate reasoning steps before answering the input prompt. This not only improves task performance, but also enhances the transparency and interpretability of their decision-making processes.

¹We make our data and code publicly available.

²The terms “Reasoning language models” (RLMs) and “Large reasoning models” (LRMs) are both used in prior works (e.g., Xu et al., 2025; Chen et al., 2025). In this paper, we adopt the term *Reasoning language models*.

Question: How many of the first one hundred positive integers are divisible by 3, 4, and 5?

Thinking: Hmm, let me think about the question. Well, if a number is divisible by multiple numbers, it’s called a common multiple.

So maybe I should start by finding the least common multiple (LCM) of 3, 4, and 5...

So, the LCM of 3, 4, and 5 is 60.

接下来，我需要找出在1到100之间有多少个数是60的倍数。也就是说，我可以列举可能的倍数。比如下一个是 $60 \times 2 = 120$ ，超过了100...

So yes, only 60. Therefore, the answer is 1.

Answer: There is only one number, **60**, that is divisible by 3, 4, and 5 in the first 100 positive integers.

Figure 1: An illustrative example of language mixing in reasoning where the reasoning model switches from English to Chinese mid-reasoning, then back to English.³

However, the phenomenon of language mixing has emerged: when prompted in one language, RLMs have been observed to produce reasoning steps that include a mixture of languages (QwenTeam, 2024; Guo et al., 2025), as illustrated in Figure 1. This phenomenon has been shown to affect reasoning performance, though prior work offers conflicting views on whether its impact is beneficial or detrimental (Guo et al., 2025; Xie et al., 2025). Moreover, it may hinder the readability and usability of outputs in multilingual contexts.

To shed light on the phenomenon of language mixing in state-of-the-art RLMs, we present a systematic investigation aimed at answering three key questions: **(1) When does language mixing hap-**

³The Chinese text translates to: "Next, I need to find how many numbers between 1 and 100 are multiples of 60. That is to say, I can list the possible multiples. For example, the next one is $60 \times 2 = 120$, which exceeds 100..."

pen and what factors influence its occurrence and degree?

(2) Is language mixing an issue and what impact does it have on the model’s reasoning performance?

(3) Why does language mixing happen and how is it related to the model’s internal thinking process?

First, we analyze occurrence patterns of language mixing across 15 input languages, 7 task difficulty levels, and 18 subject areas. We observe that language mixing is most prevalent when the input language is neither English nor Chinese, suggesting that English, and to a lesser extent, Chinese, serve as internal pivot languages during reasoning. Moreover, we find that the degree of language mixing, measured by the entropy of language distribution in reasoning traces, increases with task difficulty across models and languages. Furthermore, subject-wise analysis on the multilingual MMLU dataset (Hendrycks et al., 2020) reveals that language mixing entropy is significantly higher in STEM subjects compared to other domains.

Second, although language mixing has been observed in prior works, its impact on model performance remains unclear and is under debate (Qwen-Team, 2024; Guo et al., 2025; Xie et al., 2025). While switching languages mid-reasoning may reduce readability, it may be undesirable to suppress it entirely if such mixing enhances the model’s reasoning ability. To assess its impact, we introduce a script control method that constrains models to reason in specific script(s) via constrained decoding. We find that constraining reasoning to Latin or Han (Chinese) scripts significantly improves reasoning performance, by up to 110% in some cases, indicating that the choice of reasoning script or language has a substantial impact on model performance.

Finally, to understand the underlying cause of language mixing, we analyze the internal processing of RLMs using mechanistic interpretability. Applying Logit Lens (Nostalgebraist, 2022) at the script level, we find that the script composition of hidden representations closely mirrors that of reasoning traces, revealing that language mixing reflects the model’s latent processing preferences, particularly a consistent bias towards Latin.

In summary, we present the first systematic analysis of language mixing in RLMs, analyzing its occurrence patterns, performance impact, and internal causes. Our analysis offers practical guidance for improving multilingual reasoning and opens up new opportunities to control and adapt the reasoning language, supporting the development of more

robust, interpretable, and user-aligned RLMs.

2 Related Work

Reasoning Language Models. Reasoning language models, such as OpenAI’s o1/o3 (Jaech et al., 2024; OpenAI, 2025) and the DeepSeek-R1 series (Guo et al., 2025), have demonstrated strong capabilities on complex tasks by generating structured intermediate steps (Chen et al., 2025). While prior work has largely focused on improving reasoning quality via prompting or training, little attention has been paid to RLM behavior in multilingual settings. One underexplored aspect is language mixing, where the language used in reasoning steps differs from that of the input prompt. DeepSeek-R1 introduces a language consistency reward to reduce such mixing (Guo et al., 2025), but reports a drop in performance, suggesting potential benefits of language mixing. In contrast, Xie et al. (2025) show that responses without language mixing yield higher accuracy in logical reasoning tasks.

Given these contradictory findings, our work aims at a first systematic study of language mixing in RLMs, analyzing its patterns, performance impact, and internal causes in multilingual contexts.

Inner Workings of Multilingual LLMs. A growing body of work has explored how multilingual models internally represent and process information across languages. Works by Wendler et al. (2024), Dumas et al. (2024), Wang et al. (2025), and Liu et al. (2025) reveal that models like Llama tend to rely on English representations internally, even when operating in other languages, highlighting a strong bias toward high-resource languages. While these works focus on internal activations using translation or knowledge probing tasks, our analysis links internal processing to external reasoning traces, showing that language mixing aligns with the model’s latent language preferences during reasoning.

Code-Switching in Language Models. Code-switching, the natural alternation between languages within text, has been widely studied in NLP tasks like sentiment analysis, machine translation, and summarization (Doğruöz et al., 2021).

Prior work shows that current language models still struggle to understand and generate code-switched text, particularly in low-resource settings (Khanuja et al., 2020; Winata et al., 2023; Yong et al., 2023; Zhang et al., 2023; Li et al., 2024a,b).

Example of a Knights-and-Knives puzzle

Problem: A very special island is inhabited only by knights and knaves. Knights always tell the truth, and knaves always lie. You meet 2 inhabitants: Zoey, and Oliver. Zoey remarked, "Oliver is not a knight". Oliver stated, "Oliver is a knight if and only if Zoey is a knave". So who is a knight and who is a knave?

Solution: (1) Zoey is a knave (2) Oliver is a knight

Figure 2: An illustrative example from the K&K dataset.

Recent studies have highlighted unnatural language confusion in multilingual models as another type of code-switching (Marchisio et al., 2024; Nie et al., 2025), such as source-language hallucination or off-target translation, especially in English-centric models under zero-shot conditions.

In this work, we investigate language mixing, a code-switching-like behavior that arises during intermediate reasoning steps in RLMs. We examine its occurrence patterns across languages, difficulty levels, and subject domains, assess its impact on performance, and link it to models’ internal representations, offering new insights into the reasoning behavior of RLMs.

3 Experimental Setup

Models. We evaluate a broad range of reasoning language models, including DeepSeek-R1 (Guo et al., 2025) and its distilled variants, DeepSeek-R1-Distill-Qwen-{1.5B, 7B, 14B, 32B} and DeepSeek-R1-Distill-Llama-{8B, 70B}. We also include QwQ-32B (QwenTeam, 2025b), Qwen3-{4B, 30B-A3B, 32B} (QwenTeam, 2025a), and Gemini 2.0 Flash Thinking (Google, 2025) models to broaden coverage. For comparing language mixing between RLMs and their backbones, we include Qwen2.5-{14B, 32B} and Llama3.3-70B-Instruct, to compare DeepSeek-R1-Distill-Qwen-{14B, 32B} and DeepSeek-R1-Distill-Llama-70B with their distillation backbones. Table 5 in Appendix A.1.1 summarizes all models used in our evaluation. We follow the official hyperparameter settings for each model (see Table 6 in Appendix A.1.1). Some reasoning models occasionally exhibit endless reasoning — continue to generate reasoning steps without reaching a final answer.⁴ As this occurs only in

⁴This overthinking behavior is also observed in prior work such as Cuadron et al. (2025).

Supercategory	Subject
Humanities	HS World History, Moral Disputes, Philosophy, World Religions
Social Science	HS Macroeconomics, Sociology
STEM (Science, Technology, Engineering, Mathematics)	HS Computer Science, Col Computer Science, Elem Mathematics, HS Mathematics, Col Mathematics, HS Chemistry, Col Chemistry, HS Physics, Col Physics
Other	Global Facts, Management, Professional Medicine

Table 1: Overview of the 18 subjects in m-MMLU included in our evaluation. Abbreviations: HS = High School, Col = College, Elem = Elementary.

specific cases for most models and does not reflect a consistent pattern, our analysis focuses on valid reasoning traces that conclude with a final answer.⁵

Datasets. The **Knights-and-Knives** (K&K) dataset (Xie et al., 2024) contains logical reasoning puzzles where each character is either a *knight*, who always tells the truth, or a *knave*, who always lies. The goal is to infer each character’s identity based on their statements (see example in Figure 2). Difficulty is controlled by varying the number of characters (2–8). The original dataset is in English. To enable multilingual evaluation, we translate it into five additional languages: Arabic, French, Hindi, Japanese, and Chinese using gpt-4o-mini, resulting in six languages and seven difficulty levels.

To evaluate language mixing across broader domains, we also use the **multilingual MMLU** (m-MMLU) dataset (Hendrycks et al., 2020), a large-scale benchmark of multiple-choice questions covering 15 languages and 57 subjects across Humanities, Social Sciences, STEM, and Other domains. We select 18 representative subjects for RLM evaluation, as summarized in Table 1.

Further details on the datasets, including the translation process and language coverage, are provided in Appendix A.1.2.

4 Language Mixing Patterns

4.1 Method

To investigate the occurrence patterns of language mixing in RLMs, we first collect the reasoning traces and final answers generated by each model for queries in the K&K and m-MMLU datasets.

⁵See Appendix A.2.3 for validity statistics and accuracies.

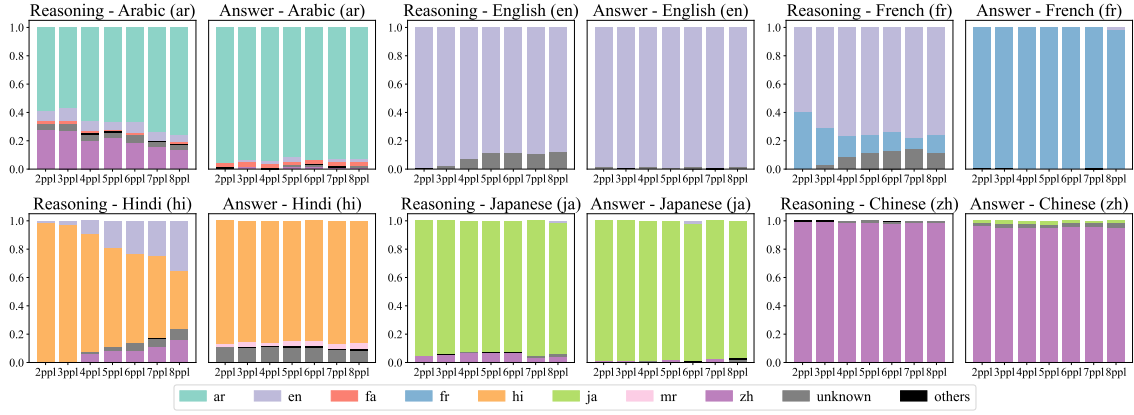


Figure 3: Language composition across difficulty levels in the K&K dataset for DeepSeek-R1-Distill-Llama-70B.

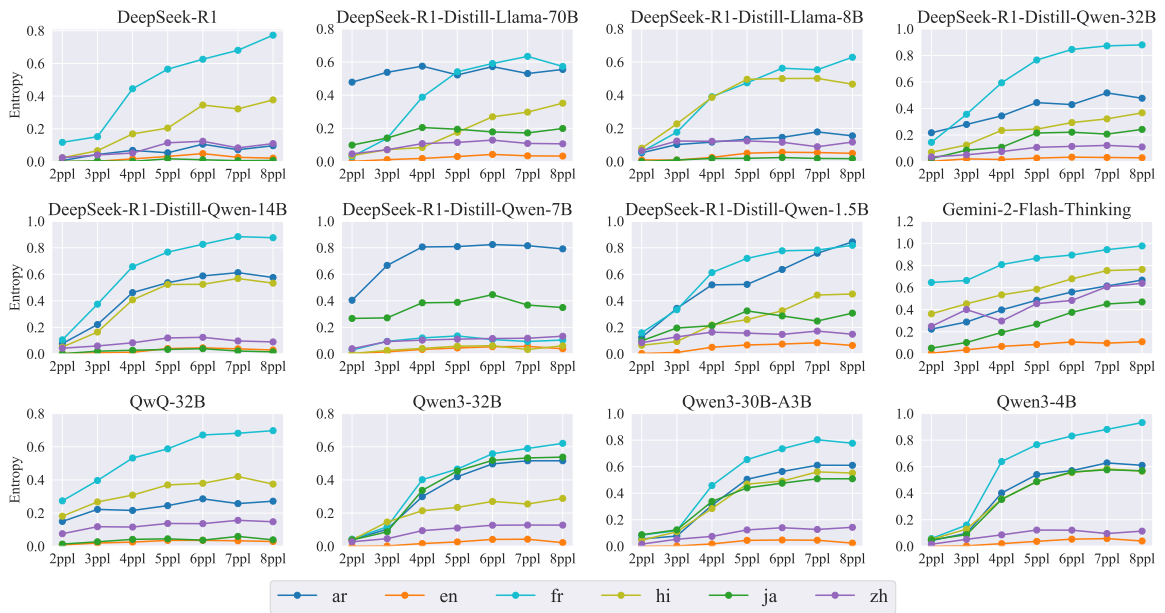


Figure 4: Language mixing entropy across task difficulty levels for six input languages and twelve models. Entropy mostly increases with difficulty, indicating harder tasks generally induce more language mixing in reasoning.

Following Marchisio et al. (2024), we perform line-level language detection by splitting each output on newline characters and identifying the language of each line using fastText (Joulin et al., 2017). This yields a per-line language distribution for each reasoning trace and answer.⁶

We compute the language composition of each sample by aggregating the detected languages across lines. Averaging across all samples in a dataset yields an overall language usage distribution, e.g., {“en”: 0.5, “fr”: 0.2, “zh”: 0.3}. To quantify the degree of language mixing, we calculate the entropy of the distribution, $([0.5, 0.2, 0.3]$

in the example), where higher entropy indicates a higher degree of language mixing, and lower entropy reflects greater language consistency.

4.2 Results

We analyze the language mixing patterns in RLMs across different input languages, task difficulty levels (in K&K), and subject areas (in m-MMLU). Our key observations are summarized below.

Finding 1

Language mixing is most prevalent when the input is neither English nor Chinese.

Figure 3 shows the language usage composition in reasoning traces and final answers across input languages and difficulty levels in the K&K dataset

⁶We consider results with confidence scores below 0.5 as *unknown*. They are usually lines mixing multiple languages or are mostly symbols.

m-MMLU Subject	R1-70B	R1-32B	R1-14B	R1-8B	R1-7B	Gemini	QwQ-32B	Q3-32B	Q3-30B-A3B	Q3-4B
Elementary Mathematics	0.11	0.09	0.08	0.12	0.12	0.51	0.09	0.04	0.08	0.09
High School Mathematics	0.32	0.31	0.28	0.31	0.32	0.74	0.34	0.27	0.32	0.31
College Mathematics	0.39	0.35	0.35	0.39	0.39	0.84	0.35	0.35	0.38	0.40
High School Chemistry	0.20	0.17	0.18	0.22	0.21	0.55	0.20	0.11	0.17	0.19
College Chemistry	0.27	0.28	0.29	0.25	0.21	0.66	0.44	0.20	0.25	0.22
High School Physics	0.27	0.22	0.22	0.24	0.26	0.74	0.26	0.15	0.21	0.24
College Physics	0.32	0.27	0.28	0.26	0.30	0.84	0.33	0.20	0.27	0.28
High School Computer Science	0.15	0.14	0.14	0.17	0.19	0.47	0.14	0.10	0.15	0.15
College Computer Science	0.20	0.19	0.22	0.22	0.23	0.54	0.21	0.13	0.19	0.19

Table 2: Language mixing entropy across STEM subjects in m-MMLU for various reasoning models. Arrows indicate increasing entropy trends with subject difficulty (Elementary → High school → College) in mathematics, chemistry, physics, and computer science subjects.

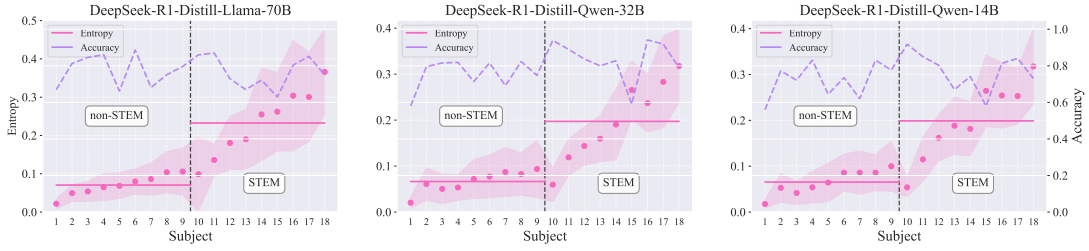


Figure 5: Language mixing entropy (pink) and task accuracy (purple) across 18 m-MMLU subjects for various reasoning models. Language mixing is notably more pronounced in STEM subjects.

for DeepSeek-R1-Distill-Llama-70B. Results of other models and on the m-MMLU dataset are provided in Appendix A.2.1. We observe that when the input is English or Chinese, the reasoning remains mostly in the input language. In contrast, Arabic, French, Hindi, and Japanese yield more mixed-language reasoning, often incorporating English and/or Chinese in the intermediate steps. In some cases, the composition is even more complex. For instance, with Arabic prompts, R1-70B produces reasoning traces that mix Arabic, English, Persian, and Chinese (Figure 3), while R1-32B generates traces involving Arabic, English, Spanish, and Chinese (Figure 18).

Finding 2

RLMs mix languages during reasoning, yet tend to generate answers in the input language.

While reasoning traces exhibit language mixing, the final answers remain more aligned with the input language. This indicates that language mixing occurs primarily in the intermediate reasoning phase, not in the final output. This behavior is unsurprising, as reasoning models are generally trained with supervision or reward signals focused on the final answer, encouraging alignment with the input language, while leaving the reasoning

steps unconstrained. As a result, models are free to mix languages during reasoning if it helps them reach a correct answer more effectively.

Finding 3

The degree of language mixing increases with task difficulty.

As shown in Figure 4, language mixing entropy in reasoning traces rises with task difficulty across languages and models. This trend is also evident in the m-MMLU dataset, which includes subjects with varying difficulty levels, e.g., mathematics at the elementary, high school, and college levels. As shown in Table 2, the average entropy⁷ consistently increases with subject difficulty across models, further confirming that task difficulty is an important trigger of language mixing behavior in RLMs.

Finding 4

Language mixing is more pronounced in STEM subjects.

Figure 5 shows language mixing entropy across 18 m-MMLU subjects⁸ for R1-70B, R1-32B and R1-14B. Entropy values are averaged across all evaluated languages. Additional results for other

⁷The entropy is averaged across languages for each subject.

⁸The subject order is provided in Appendix A.2.1.

	ar	en	fr	hi	ja	zh	AVG
<i>R1</i>	0.06	0.02	0.48	0.22	0.01	0.08	0.15
<i>Qwen2.5-14B</i>	0.07	0.01	0.32	0.05	0.04	0.10	0.10
<i>R1-14B</i>	0.06	0.00	0.59	0.32	0.01	0.11	0.18
<i>Qwen2.5-32B</i>	0.25	0.01	0.25	0.14	0.12	0.26	0.17
<i>R1-32B</i>	0.39	0.02	0.64	0.24	0.16	0.09	0.26
<i>Llama3.3-70B</i>	0.08	0.00	0.12	0.14	0.00	0.07	0.07
<i>R1-70B</i>	0.51	0.01	0.43	0.18	0.14	0.12	0.23

Table 3: Language mixing entropy for backbone models and their distilled variants on the K&K dataset. Distilled reasoning models (bottom row in each pair) exhibit higher entropy than their corresponding base models.

models and per-language breakdowns are shown in Figure 9 and Figure 10 in Appendix A.2.1. STEM subjects, starting from tick mark 10 in the figure, consistently exhibit higher entropy than Humanities, Social Sciences, or Other domains, suggesting that technical content tends to induce more language mixing during reasoning. While entropy still increases with task difficulty within individual STEM domains (Table 2), the overall entropy gap between STEM and non-STEM subjects appears more related to subject type, as they exhibit no clear difference in difficulty based on accuracy (the purple curve in Figure 5).

Finding 5

Distillation amplifies language mixing.

Table 3 compares the language mixing entropy for three DeepSeek backbone models and their corresponding distilled variants. Across all model pairs, the distilled versions consistently exhibit higher average entropy values, particularly when the input language is neither English nor Chinese. Although Guo et al. (2025) do not provide details on the language composition of the distillation data, the observed trend suggests that an English- and Chinese-heavy training distribution may lead the model to rely more on these high-resource language features during reasoning, thereby amplifying language mixing in multilingual settings.

5 Performance Impact of Reasoning Languages

5.1 Method: Script-controlled Generation

As discussed in Section 2, the impact of language mixing on RLM performance remains debated. To address this, we adopt a direct script control approach to evaluate how reasoning language influences model performance.

We constrain models to generate reasoning steps using only specific script(s) by masking the logits of all other scripts during decoding. While language-level control is difficult due to token overlap, script-level control offers a clean separation based on Unicode, as tokens from different scripts do not overlap. We apply this script control during the reasoning phase (though it could also be applied to the answer phase); the final answer is generated freely after the “</think>” token, as models typically default to the input language for answers (as shown in Section 4). For examples of model generation under different script control settings, see Figures 11 to 14 in Appendix A.2.2.

5.2 Results

We compare model performance across various reasoning modes: unconstrained, single-script, and multi-script control to directly measure the impact of script (and implicitly, language) choice on reasoning performance. Figure 6 shows results for six input languages from the K&K dataset: three written in non-Latin/Han scripts (Arabic, Hindi, Japanese), two in Latin (English, French) and one in Han (Chinese), on three reasoning models: R1-70B, R1-32B, and R1-14B.⁹

Finding 6

Non-Latin/Han-script languages benefit significantly from reasoning in Latin or Han scripts.

For Arabic, Hindi, and Japanese, forcing reasoning in Latin or Han scripts significantly improves performance. In Hindi, for example, switching to Latin yields gains ranging from 44% to 115%. In contrast, using the native script consistently results in the lowest accuracy, indicating that models struggle with scripts underrepresented in their training data. For multi-script control (see Table 10 and Table 11 in Appendix A.2.2), reasoning in both Latin and Han scripts does not outperform single-script control, and including the input script alongside Latin and/or Han leads to suboptimal results.

Finding 7

Latin and Han-script languages favor their native script.

For English, French, and Chinese, reasoning in the native script yields performance comparable

⁹These models are selected to cover a range of model sizes and base architectures.

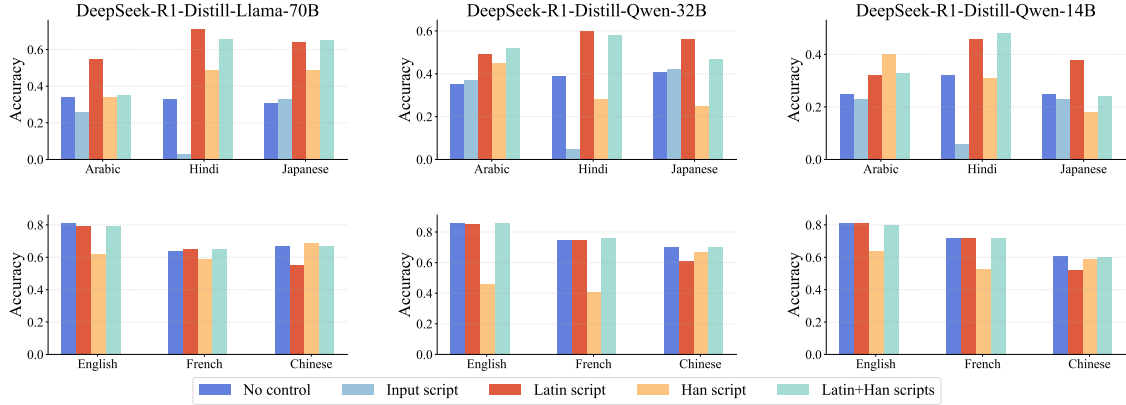


Figure 6: Accuracy on the K&K dataset under script-controlled reasoning. Latin or Han script control boosts performance for Arabic, Hindi, and Japanese, while native scripts yield the best results for English, French, and Chinese, highlighting the impact of script choice on reasoning efficacy.

to the unconstrained setting. However, switching scripts (e.g., Han for English/French or Latin for Chinese) leads to notable performance drops, suggesting that mismatched scripts disrupt alignment with the model’s internal reasoning patterns.

These results highlight a strong connection between script choice and reasoning performance. Models internally favor Latin and Han scripts, which benefits non-Latin/Han inputs when reasoning is constrained to those scripts. For Latin or Han-script languages, maintaining script consistency is optimal, while mismatches hurt the reasoning performance. This suggests that language mixing reflects the models’ learned association between dominant scripts and reasoning competence.

6 Internal Causes of Language Mixing

6.1 Method: Logit Lens Analysis

Inspired by prior work on latent language dynamics in LLMs (Wendler et al., 2024; Wang et al., 2025), we use Logit Lens (Nostalgebraist, 2022) to examine the script composition of internal representations of RLMs and connect it to the scripts used in the models’ external reasoning output.

We use Logit Lens to project intermediate layer representations onto the vocabulary space and identify the script of the top-ranked token at each layer.¹⁰ By averaging predictions across all K&K samples, we obtain a layer-wise distribution of script usage. We then compare these internal patterns to the script usage in reasoning traces, tracking how both evolve with task difficulty. This al-

lows us to directly link the model’s external reasoning behavior with its internal processing dynamics.

To quantify this connection, we compute the Pearson correlation between script usage across difficulty levels in hidden layers and reasoning traces. Specifically: (1) for each difficulty level (ranging from 2 to 8 ppl), we calculate the proportion of tokens belonging to different scripts (e.g., Latin, Devanagari) in both the latent space (estimated with the Logit Lens) and the reasoning trace (measured via script-level detection); and (2) for each script, we compute the Pearson correlation coefficient between its percentage across difficulty levels in the latent space and its corresponding percentage in the reasoning trace. A more detailed description of the Pearson correlation calculation is provided in Appendix A.1.3.

6.2 Results

Figure 7 shows that the hidden layer script composition (i.e., the internal representation) for Hindi inputs at difficulty levels 2ppl, 5ppl, and 8ppl in R1-70B, is consistently dominated by the Latin script. Devanagari (Hindi’s script) appears only in final layers. This aligns with prior findings that Llama models primarily “think” in English (i.e., Latin script) (Wendler et al., 2024; Wang et al., 2025).

Finding 8

Language mixing reflects the internal processing patterns of RLMs, as reasoning traces mirror the model’s preferred scripts.

As task difficulty increases, Latin usage rises while Devanagari decreases. The same trend ap-

¹⁰We operate at the script-level (rather than language-level) to avoid ambiguity due to token overlap across languages.

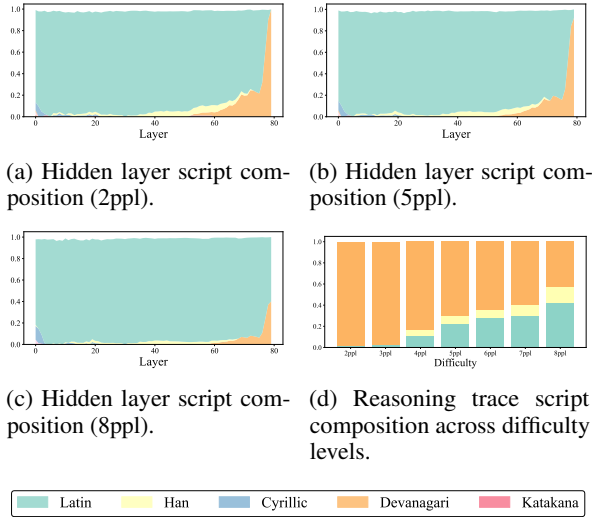


Figure 7: Script composition of internal representations and reasoning traces for Hindi inputs in R1-70B. (a–c) show hidden layer script distributions via Logit Lens across three difficulty levels. (d) shows reasoning trace script composition mirroring internal trends: Latin usage increases and Devanagari decreases with difficulty.

Correlation	R1-70B	R1-32B	R1-14B
Arabic			
Arabic script	0.7411	0.9960	0.8986
Latin script	0.7168	0.7751	0.7774
Hindi			
Devanagari script	0.9019	0.9352	0.8840
Latin script	0.8875	0.9173	0.8969

Table 4: Pearson correlation between script usage in hidden layers and reasoning traces across task difficulty levels, showing strong alignment between internal representations and external reasoning outputs.

pears in the reasoning traces (Figure 7d), suggesting a strong link between internal processing and external reasoning behavior. The Pearson correlations between script usage across difficulty levels in hidden layers and reasoning traces are reported in Table 4. We observe consistently high correlations for both Arabic and Hindi inputs in R1-70B, R1-32B and R1-14B models. These results provide evidence that language mixing in reasoning traces reflects internal processing patterns. In particular, the internal preference for the Latin script explains why RLMs tend to mix into Latin-script reasoning, especially when processing underrepresented scripts.

7 Discussion and Future Works

Is language mixing a solved problem in RLMs?

Although Guo et al. (2025) attempt to mitigate language mixing by incorporating a language consistency reward during reinforcement learning, our evaluation reveals that language mixing persists in DeepSeek-R1, as shown, for instance, in Figure 4 and Figure 15. Regardless of whether language mixing should be removed (which we discuss in the following), our findings indicate that eliminating it with targeted optimization during training remains technically challenging.

Should language mixing be eliminated?

Our investigation into the impact of language mixing on reasoning performance (Section 5) shows that constraining reasoning to high-resource language scripts, such as Latin or Han, significantly improves performance on under-resourced language inputs like Arabic, Hindi, and Japanese. This raises questions about whether language mixing is truly undesirable. Notably, Guo et al. (2025) also report that enforcing language consistency leads to performance degradation, suggesting that language mixing may serve as a functional adaptation that supports effective multilingual reasoning.

Reasoning language control.

Our study shows that script control significantly boosts performance for under-resourced languages, highlighting the potential of reasoning language control in RLMs. Our findings in Section 5 and Section 6 point to two potential future directions: (1) developing fine-grained *language-level* control through constrained decoding, extending beyond *script-level* masking, and (2) steering representations through representation engineering (Zou et al., 2023) or latent space reasoning methods (Hao et al., 2024) to align the model’s latent processing with the desired reasoning language. These approaches may offer more flexible control over language mixing and enable more controllable and performant RLMs.

Impact of distillation on language mixing.

Table 3 shows that DeepSeek-R1 distilled models consistently exhibit higher language mixing entropy than their original backbones, suggesting distillation amplifies language mixing when the input language is neither English nor Chinese. This highlights the need to better understand how distillation affects language mixing (Yong et al., 2025), which is a promising future work direction. Moreover, multilingual distillation strategies may help reduce

unnecessary mixing while preserving or improving performance.

Adaptive reasoning language selection. Our results show that switching to high-resource languages during reasoning can substantially improve model performance on under-resourced inputs. This suggests a promising direction for future work: training RLMs to adaptively determine which language to use at different stages of reasoning, rather than relying on static language constraints. Such adaptive control could allow the model to dynamically leverage the strengths of high-resource languages while still accommodating the input language, ultimately realizing more robust multilingual reasoning. This idea is aligned with recent advances in adaptive memory usage for reasoning (Yan et al., 2025), which highlight the benefits of enabling models to flexibly select internal strategies based on task demands.

8 Conclusions

In this work, we presented the first systematic investigation of language mixing in reasoning language models, analyzing its occurrence patterns, performance impact, and internal causes. Our findings reveal that language mixing is most likely to occur when the input is neither English nor Chinese, and becomes more pronounced with higher task difficulty and in STEM subjects. We further demonstrated that constraining reasoning to Latin or Han scripts significantly improves performance for non-Latin/Han inputs. Finally, our interpretability analysis shows that the language composition of reasoning traces mirrors that of the models’ internal representations, suggesting language mixing reflects underlying processing preferences. These insights offer a deeper understanding of multilingual reasoning behavior and provide guidance for developing more controllable and interpretable RLMs.

Limitations

While our analysis provides a broad view of language mixing in reasoning language models, it has several limitations.

First, we use script-level control to examine the impact of reasoning language. This approach, while effective, remains relatively coarse. More fine-grained methods, such as language-level constrained decoding, could enable more precise control and allow detailed comparisons of the perfor-

mance across different reasoning languages given a specific input language.

Second, while we analyze the internal causes of language mixing through mechanistic interpretability, we do not trace its origins in the training process. Future work could investigate the role of training data composition (including pretraining, reinforcement learning, and distillation) and optimization objectives to better understand the roots of this phenomenon.

Third, our study covers 15 languages, but many other languages—particularly those with low resources or unique scripts—remain unexamined. Extending the analysis to a broader set of languages would help validate the generality of our findings.

Lastly, we do not include DeepSeek-R1-Zero (Guo et al., 2025) in our evaluation due to resource constraints. As a key variant of reasoning language models with a different training setup, comparing it with DeepSeek-R1 could offer further insight into the effect of reinforcement learning training on language mixing.

Acknowledgments

This work was partially supported by Deutsche Forschungsgemeinschaft (project SCHU 2246/14-1) and the EU Project SMARTY (GA 101140087). We thank JUREAP 73 for supporting this work with computing resources. We also thank Sebastian Gerstner, Ahmad Dawar Hakimi, Lea Hirliemann, Valentin Knappich, Felicia Körner, Yihong Liu, Ali Modarressi, Philipp Mondorf, Timo Pierre Schrader, Leonor Veloso, Wei Zhou, Yuqicheng Zhu for their valuable discussions.

References

- Qiguang Chen, Libo Qin, Jinhao Liu, Dengyun Peng, Jiannan Guan, Peng Wang, Mengkang Hu, Yuhang Zhou, Te Gao, and Wanxiang Che. 2025. Towards reasoning era: A survey of long chain-of-thought for reasoning large language models. *arXiv preprint arXiv:2503.09567*.
- Alejandro Cuadron, Dacheng Li, Wenjie Ma, Xingyao Wang, Yichuan Wang, Siyuan Zhuang, Shu Liu, Luis Gaspar Schroeder, Tian Xia, Huanzhi Mao, and 1 others. 2025. The danger of overthinking: Examining the reasoning-action dilemma in agentic tasks. *arXiv preprint arXiv:2502.08235*.
- A. Seza Doğruöz, Sunayana Sitaram, Barbara E. Bullock, and Almeida Jacqueline Toribio. 2021. A survey of code-switching: Linguistic and social perspectives for language technologies. In *Proceedings*

- of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers), pages 1654–1666, Online. Association for Computational Linguistics.
- Clément Dumas, Veniamin Veselovsky, Giovanni Monea, Robert West, and Chris Wendler. 2024. [How do llamas process multilingual text? a latent exploration through activation patching](#). In *ICML 2024 Workshop on Mechanistic Interpretability*.
- Google. 2025. [Gemini 2.0 flash thinking](#).
- Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, and 1 others. 2025. [Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning](#). *arXiv preprint arXiv:2501.12948*.
- Shibo Hao, Sainbayar Sukhbaatar, DiJia Su, Xian Li, Zhiting Hu, Jason Weston, and Yuandong Tian. 2024. Training large language models to reason in a continuous latent space. *arXiv preprint arXiv:2412.06769*.
- Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. 2020. [Measuring massive multitask language understanding](#). *arXiv preprint arXiv:2009.03300*.
- Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, and 1 others. 2024. Openai o1 system card. *arXiv preprint arXiv:2412.16720*.
- Armand Joulin, Edouard Grave, Piotr Bojanowski, and Tomas Mikolov. 2017. [Bag of tricks for efficient text classification](#). In *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics: Volume 2, Short Papers*, pages 427–431, Valencia, Spain. Association for Computational Linguistics.
- Simran Khanuja, Sandipan Dandapat, Anirudh Srinivasan, Sunayana Sitaram, and Monojit Choudhury. 2020. [GLUECoS: An evaluation benchmark for code-switched NLP](#). In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 3575–3585, Online. Association for Computational Linguistics.
- Bin Li, Bin Sun, Shutao Li, Encheng Chen, Hongru Liu, Yixuan Weng, Yongping Bai, and Meiling Hu. 2024a. Distinct but correct: generating diversified and entity-revised medical response. *Science China Information Sciences*, 67(3):132106.
- Bin Li, Yixuan Weng, Fei Xia, and Hanjun Deng. 2024b. Towards better chinese-centric neural machine translation for low-resource languages. *Computer Speech & Language*, 84:101566.
- Yihong Liu, Mingyang Wang, Amir Hossein Kargaran, Felicia Körner, Ercong Nie, Barbara Plank, François Yvon, and Hinrich Schütze. 2025. [Tracing multilingual factual knowledge acquisition in pretraining](#). *arXiv preprint arXiv:2505.14824*.
- Kelly Marchisio, Wei-Yin Ko, Alexandre Berard, Théo Dehaze, and Sebastian Ruder. 2024. [Understanding and mitigating language confusion in LLMs](#). In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 6653–6677, Miami, Florida, USA. Association for Computational Linguistics.
- Ercong Nie, Helmut Schmid, and Hinrich Schütze. 2025. [Mechanistic understanding and mitigation of language confusion in english-centric large language models](#). *arXiv preprint arXiv:2505.16538*.
- Nostalgebraist. 2022. [Interpreting gpt: The logit lens](#).
- OpenAI. 2025. [Openai o3 and o4-mini system card](#).
- QwenTeam. 2024. [Qwq: Reflect deeply on the boundaries of the unknown](#).
- QwenTeam. 2025a. [Qwen3: Think deeper, act faster](#).
- QwenTeam. 2025b. [Qwq-32b: Embracing the power of reinforcement learning](#).
- Mingyang Wang, Heike Adel, Lukas Lange, Yihong Liu, Ercong Nie, Jannik Strötgen, and Hinrich Schuetze. 2025. [Lost in multilinguality: Dissecting cross-lingual factual inconsistency in transformer language models](#). In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 5075–5094, Vienna, Austria. Association for Computational Linguistics.
- Chris Wendler, Veniamin Veselovsky, Giovanni Monea, and Robert West. 2024. [Do llamas work in English? on the latent language of multilingual transformers](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 15366–15394, Bangkok, Thailand. Association for Computational Linguistics.
- Genta Winata, Alham Fikri Aji, Zheng Xin Yong, and Thamar Solorio. 2023. [The decades progress on code-switching research in NLP: a systematic survey on trends and challenges](#). In *Findings of the Association for Computational Linguistics: ACL 2023*, pages 2936–2978, Toronto, Canada. Association for Computational Linguistics.
- Chulin Xie, Yangsibo Huang, Chiyuan Zhang, Da Yu, Xinyun Chen, Bill Yuchen Lin, Bo Li, Badi Ghazi, and Ravi Kumar. 2024. On memorization of large language models in logical reasoning. *arXiv preprint arXiv:2410.23123*.
- Tian Xie, Zitian Gao, Qingnan Ren, Haoming Luo, Yuqian Hong, Bryan Dai, Joey Zhou, Kai Qiu, Zhirong Wu, and Chong Luo. 2025. [Logic-rl: Unleashing llm reasoning with rule-based reinforcement learning](#). *Preprint*, arXiv:2502.14768.

Fengli Xu, Qian Yue Hao, Zefang Zong, Jingwei Wang, Yunke Zhang, Jingyi Wang, Xiaochong Lan, Jiahui Gong, Tianjian Ouyang, Fanjin Meng, and 1 others. 2025. Towards large reasoning models: A survey of reinforced reasoning with large language models. *arXiv preprint arXiv:2501.09686*.

Sikuan Yan, Xiufeng Yang, Zuchao Huang, Ercong Nie, Zifeng Ding, Zonggen Li, Xiaowen Ma, Hinrich Schütze, Volker Tresp, and Yunpu Ma. 2025. Memory-r1: Enhancing large language model agents to manage and utilize memories via reinforcement learning. *arXiv preprint arXiv:2508.19828*.

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. *arXiv preprint arXiv:2505.05408*.

Zheng Xin Yong, Ruochen Zhang, Jessica Forde, Skyler Wang, Arjun Subramonian, Holy Lovenia, Samuel Cahyawijaya, Genta Winata, Lintang Sutawika, Jan Christian Blaise Cruz, Yin Lin Tan, Long Phan, Long Phan, Rowena Garcia, Tamar Solorio, and Alham Fikri Aji. 2023. Prompting multilingual large language models to generate code-mixed texts: The case of south East Asian languages. In *Proceedings of the 6th Workshop on Computational Approaches to Linguistic Code-Switching*, pages 43–63, Singapore. Association for Computational Linguistics.

Ruochen Zhang, Samuel Cahyawijaya, Jan Christian Blaise Cruz, Genta Winata, and Alham Fikri Aji. 2023. Multilingual large language models are not (yet) code-switchers. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 12567–12582, Singapore. Association for Computational Linguistics.

Andy Zou, Long Phan, Sarah Chen, James Campbell, Phillip Guo, Richard Ren, Alexander Pan, Xu Wang Yin, Mantas Mazeika, Ann-Kathrin Dombrowski, and 1 others. 2023. Representation engineering: A top-down approach to ai transparency. *arXiv preprint arXiv:2310.01405*.

A Appendix

A.1 Experimental Setup Details

A.1.1 Models

Model Lists. As introduced in Section 3, we evaluate a diverse set of reasoning language models (RLMs) across developers. Table 5 lists their full names, backbone models, and reference names used throughout this paper.

Hyperparameters. For reproducibility and consistency, we report the decoding hyperparameters used across all models evaluated in our experiments in Table 6.

A.1.2 Datasets

Knights-and-Knaves dataset translation. As introduced in Section 3, the Knights and Knaves (K&K) dataset consists of algorithmically generated reasoning puzzles, with difficulty controlled by varying the number of characters (2–8) and the complexity of logical operations. In these puzzles, as illustrated in Figure 2, each character is either a *knight*, who always tells the truth, or a *knave*, who always lies. The objective is to determine each character’s identity based on their statements. All puzzles are constructed using formal logic rules, ensuring a unique, verifiable solution, which makes the dataset well-suited for analyzing the impact of task difficulty on reasoning in language models.

To extend the dataset beyond English, we translate it into five additional languages, Arabic (ar), French (fr), Hindi (hi), Japanese (ja), and Chinese (zh), covering diverse scripts and linguistic families. To ensure consistent translations of identity terms (e.g., “knight”, “knave”) and character names (e.g., Zoey, Oliver), we construct a fixed translation map of these identity and character names from English to each target language. We then use gpt-4o-mini to translate the puzzles and solutions while enforcing consistency with this map. The translation prompt is shown in Figure 8.

Further details of m-MMLU The m-MMLU dataset covers 15 languages across diverse linguistic and geographic regions as shown in Table 7. This broad coverage allows for a comprehensive assessment of multilingual reasoning abilities across a wide spectrum of languages. We follow the prompts and the evaluation framework used in simple-evals library.¹¹

¹¹<https://github.com/openai/simple-evals>

Developer	Model Name	Backbone Model	Ref. Name
DeepSeek	DeepSeek-R1-Distill-Qwen-1.5B	Qwen2.5-Math-1.5B	R1-1.5B
	DeepSeek-R1-Distill-Qwen-7B	Qwen2.5-Math-7B	R1-7B
	DeepSeek-R1-Distill-Llama-8B	Llama-3.1-8B	R1-8B
	DeepSeek-R1-Distill-Qwen-14B	Qwen2.5-14B	R1-14B
	DeepSeek-R1-Distill-Qwen-32B	Qwen2.5-32B	R1-32B
	DeepSeek-R1-Distill-Llama-70B	Llama-3.3-70B-Instruct	R1-70B
	DeepSeek-R1	DeepSeek-V3-Base	R1
Google	gemini-2.0-flash-thinking-exp-01-21	-	Gemini / Gemini 2.0 Flash Thinking
Qwen	QwQ-32B	Qwen2.5-32B-Instruct	QwQ
	Qwen3-4B	Qwen3-4B-Base	Q3-4B / Qwen3-4B
	Qwen3-30B-A3B	Qwen3-30B-A3B-Base	Q3-30B-A3B / Qwen3-30B-A3B
	Qwen3-32B	Qwen3-32B-Base	Q3-32B / Qwen3-32B

Table 5: Summary of reasoning models evaluated in this work, including their backbone models and the reference names used in the paper.

Model	Hyperparameters
DeepSeek R1 series	Temperature=0.6, TopP=0.95
gemini-2.0-flash-thinking	Temperature = 0.7, TopP = 0.95, TopK = 64
QwQ-32B	Temperature=0.6, TopP=0.95
Qwen3 series	Temperature=0.6, TopP=0.95, TopK=20

Table 6: Hyperparameters used for different reasoning models in our evaluation. Settings follow the recommended configurations provided by model developers on HuggingFace to ensure fair comparison.

Dataset	Languages
K&K	Arabic (ar), English (en), French (fr), Hindi (hi), Japanese (ja), Chinese (zh)
m-MMLU	Arabic (AR-XY), Bengali (BN-BD), German (DE-DE), Spanish (ES-LA), French (FR-FR), Hindi (HI-IN), Indonesian (ID-ID), Italian (IT-IT), Japanese (JA-JP), Korean (KO-KR), Brazilian Portuguese (PT-BR), Swahili (SW-KE), Yoruba (YO-NG), Simplified Chinese (ZH-CN)

Table 7: Languages covered in the K&K and m-MMLU datasets.

A.1.3 Pearson correlation calculation

To clarify how the Pearson correlation coefficients in Table 4 are calculated, we provide a detailed description of the procedure.

We take Figure 7 in the main text as an example (which shows the script composition of internal representations and reasoning traces for Hindi inputs in the R1-70B model). The Pearson correlation is computed as follows:

1. For each difficulty level (from 2 to 8 ppl), we collect the percentage of tokens belonging to different scripts, such as Latin and Devanagari. These percentages are measured in (1) the latent space, obtained via Logit Lens anal-

ysis of hidden layer activations; and (2) the reasoning trace, obtained through script-level detection of the generated text.

2. For each script, we then calculate the Pearson correlation coefficient between its sequence of percentages across difficulty levels in the latent space and the corresponding sequence in the reasoning trace. An illustrative example is provided in Table 9.

These Pearson correlations indicate how closely the model’s internal script usage aligns with the script composition of the generated reasoning trace.

A.2 Additional Experimental Results

A.2.1 Language Mixing Patterns

m-MMLU complete results. Figure 9 presents language mixing entropy across m-MMLU subjects for additional reasoning models beyond those shown Figure 5. As in the main results, entropy is averaged across languages. The trend remains consistent: language mixing is markedly higher in STEM subjects (subjects 10–18), compared to non-STEM areas. This pattern reinforces our main finding that technical subjects tend to induce more language mixing during reasoning.

m-MMLU per-language results. Figure 10 presents the per-language breakdown of language

K&K Translation System Prompt	
You are a professional translator. Please translate the following English text into {target_language} while following these rules: • Translate person names according to this mapping: {name_map}. • Translate “Knights” and “Knaves” always as follows: - “Knights” → {identity_map['Knights']} - “Knight” → {identity_map['Knight']} - “Knaves” → {identity_map['Knaves']} - “Knaves” → {identity_map['Knave']} • Ensure the sentence remains grammatically correct and natural in {target_language}. • Do NOT translate placeholders (if any exist). • Return only the translated text, no extra information.	

Figure 8: System prompt used to translate K&K puzzles into target languages with consistent identity and person name mapping.

mixing entropy across m-MMLU subjects, complementing the averaged results in Figure 5. The observed rise in entropy within STEM subjects (subjects 10–18) holds consistently across most languages, supporting the conclusion that STEM subjects tend to elicit more language mixing during reasoning. The specific subject order we use here is: global_facts, world_religions, sociology, high_school_world_history, moral_disputes, professional_medicine, philosophy, high_school_macro_economics, management, elementary_mathematics, high_school_computer_science, high_school_chemistry, college_computer_science, high_school_physics, college_chemistry, college_physics, high_school_mathematics, college_mathematics.

Language composition visualization. In Figures 15 to 26 and Figures 27 to 37, we visualize the language composition of reasoning traces and final answers across models on the K&K and m-MMLU datasets, respectively. These figures provide a detailed view of how different RLMs vary in their use of languages during the reasoning process across inputs. Missing bars in some plots correspond to rare failure cases where the model enters an endless reasoning loop and fails to produce a final answer. For more details on these invalid generations, see Appendix A.2.3.

Language	Abbreviation	Script
<i>Languages of investigation</i>		
Arabic	ar / AR-XY	Arabic
Bengali	BN-BD	Bangla
Chinese	zh / ZH-CN	Han
English	en / EN-US	Latin
French	fr	Latin
German	DE-DE	Latin
Hindi	hi / HI-IN	Devanagari
Indonesian	ID-ID	Latin
Italian	IT-IT	Latin
Japanese	ja / JA-JP	Han, Hiragana, Katakana
Korean	KO-KR	Hangul
Portuguese	PT-BR	Latin
Spanish	ES-LA	Latin
Swahili	SW-KE	Latin
Yoruba	YO-NG	Latin
<i>Languages emerge in the reasoning trace</i>		
Persian	fa	Perso-Arabic
Marathi	mr	Devanagari
Russian	ru	Cyrillic

Table 8: Languages examined in our study, with their abbreviations and writing scripts. We distinguish between the *languages of investigation*, which serve as evaluation inputs covered by the K & K and/or m-MMLU datasets, and the *languages that emerge in the reasoning trace*, which appear spontaneously during the model’s reasoning process.

A.2.2 Script-Controlled Generation

Script control complete results. Tables 10 and 11 present the detailed performance results of three DeepSeek-R1 models (R1-70B, R1-32B, and R1-14B) across various script control strategies for languages from the K&K dataset. We compare model accuracy under three settings: no control (default decoding), single-script control (forcing reasoning in one script), and multi-script control (allowing reasoning in a combination of scripts). The languages are grouped into two sets based on their writing scripts: Arabic, Hindi, and Japanese use non-Latin/Han scripts; English, French, and Chinese use Latin or Han scripts. As shown, non-Latin/Han languages benefit substantially from being forced to reason in Latin or Han scripts, whereas native-script reasoning is optimal for Latin and Han-script languages. These findings highlight the importance of aligning reasoning scripts with the model’s internal preferences for achieving optimal multilingual reasoning performance.

¹²For Japanese, entries under “Input+Han” and “Input+Latin+Han” are omitted because Japanese inherently uses Han characters. Thus, “Input+Han” is equivalent to “Input”, and “Input+Latin+Han” is equivalent to “Input+Latin” in this case.

Difficulty	Latin – latent (%)	Latin – reasoning (%)	Devanagari – latent (%)	Devanagari – reasoning (%)
2 ppl	77.01	0.76	17.82	99.19
3 ppl	77.00	4.22	18.34	95.75
4 ppl	78.54	15.05	16.96	78.11
5 ppl	80.93	31.68	15.11	61.90
6 ppl	84.78	48.02	12.18	48.02
7 ppl	91.46	44.71	6.60	44.33
8 ppl	91.79	63.26	6.40	21.54

Table 9: Example of script composition across difficulty levels for Hindi inputs in the R1-70B model. Percentages of tokens in the latent space (via Logit Lens) and reasoning trace (via script-level detection) are shown for Latin and Devanagari scripts. These values form the basis for computing Pearson correlations reported in the main text.

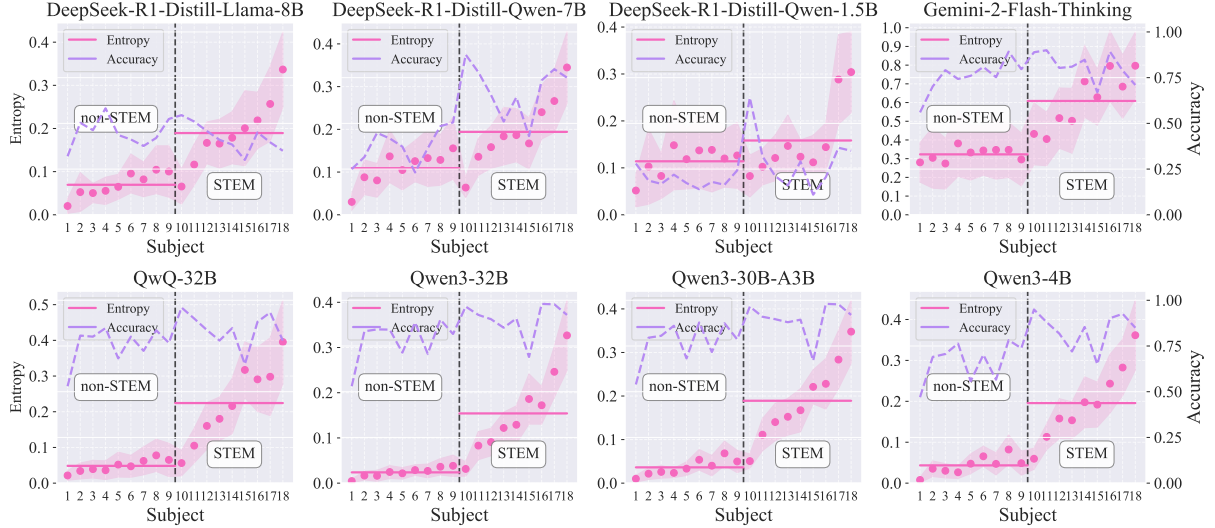


Figure 9: Language mixing entropy (pink) and task accuracy (blue) across 18 m-MMLU subjects for additional reasoning models. Entropy consistently increases in STEM subjects.

Script control generation examples. We demonstrate the effects of script control on multilingual reasoning through examples using the same Knights and Knaves puzzle with Arabic input (Figures 11 to 14). Figure 11 shows unconstrained reasoning with natural language mixing across English, Chinese, and Arabic. Figures 12, 13, and 14 demonstrate controlled reasoning under Latin, Han, and Arabic script constraints respectively. These examples illustrate how script control can effectively guide language selection while maintaining the model’s cross-lingual reasoning capabilities.

A.2.3 Model Performance Details

Here, we report the valid reasoning rate and accuracy across all reasoning models used in this work in Tables 12 to 23. The valid reasoning rate refers to the proportion of generations that produce a complete reasoning trace followed by a final answer, as defined in Section 3. The accuracy is computed over all generations, both valid and invalid, and reflects the overall correctness of the final answers.

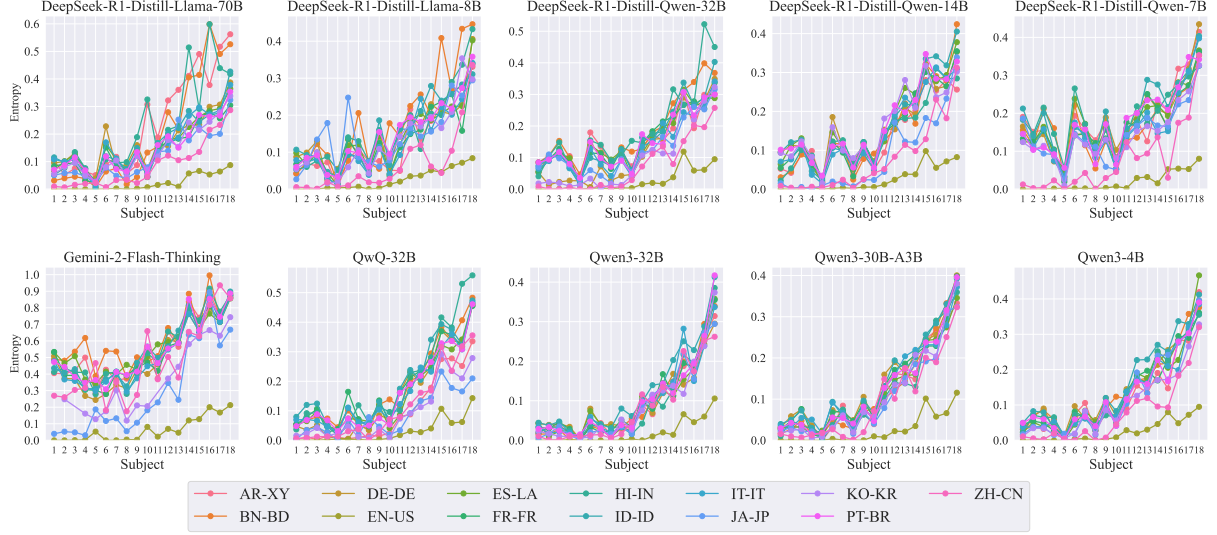


Figure 10: Per-language language mixing entropy across 18 m-MMLU subjects for various reasoning models. Each line represents the entropy trend for a specific input language. STEM subjects (subjects 10–18) consistently exhibit higher entropy across most languages, mirroring the pattern seen in the language-averaged results (Figure 5).

No control		Single-script control				Multi-script control		
		Input script	Latin	Han	Latin+Han	Input+Latin	Input+Han	Input+Latin+Han
<i>Arabic</i>								
R1-70B	0.34	0.26	0.55	0.34	0.35	0.27	0.31	0.33
R1-32B	0.35	0.37	0.49	0.45	0.52	0.68	0.81	0.66
R1-14B	0.25	0.23	0.32	0.43	0.33	0.32	0.38	0.32
AVG	0.31	0.29	0.45	0.41	0.40	0.42	0.50	0.44
<i>Hindi</i>								
R1-70B	0.33	0.03	0.71	0.49	0.66	0.68	0.54	0.66
R1-32B	0.39	0.05	0.60	0.28	0.58	0.32	0.29	0.32
R1-14B	0.32	0.06	0.46	0.41	0.48	0.47	0.44	0.47
AVG	0.35	0.05	0.59	0.39	0.57	0.49	0.42	0.48
<i>Japanese</i>								
R1-70B	0.31	0.33	0.64	0.49	0.65	0.30	-	-
R1-32B	0.41	0.42	0.56	0.25	0.47	0.50	-	-
R1-14B	0.25	0.23	0.38	0.18	0.24	0.25	-	-
AVG	0.32	0.33	0.53	0.31	0.45	0.35	-	-

Table 10: Performance comparison across different script control strategies for Arabic, Hindi, and Japanese.¹²

	No control	Single-script control		Multi-script control
		Latin	Han	Latin+Han
<i>English</i>				
R1-70B	0.81	0.79	0.62	0.79
R1-32B	0.86	0.85	0.46	0.86
R1-14B	0.81	0.81	0.64	0.80
AVG	0.83	0.82	0.57	0.82
<i>French</i>				
R1-70B	0.64	0.65	0.59	0.65
R1-32B	0.75	0.75	0.41	0.76
R1-14B	0.72	0.72	0.53	0.72
AVG	0.70	0.71	0.51	0.71
<i>Chinese</i>				
R1-70B	0.67	0.55	0.69	0.67
R1-32B	0.70	0.61	0.67	0.70
R1-14B	0.61	0.52	0.59	0.60
AVG	0.66	0.56	0.65	0.66

Table 11: Performance comparison across different script control strategies for English, French, and Chinese.

R1-70B reasoning without script control

Question: جزيرة خاصة جدًا يسكنها فقط الفرسان والمحتالون. الفرسان دائماً يقولون الحقيقة، والمحتالون دائماً يكذبون. تلقتي بسكائين: أميليا ويعقوب. قالت أميليا إن يعقوب فارس وإن يعقوب محتال. علق يعقوب محتال. إذا كانت أميليا محتالة، فإن أميليا فارس. فمن هو الفارس ومن هو المحتال؟

(A very special island is inhabited only by knights and knaves. Knights always tell the truth, and knaves always lie. You meet 2 inhabitants: Amelia and Jacob. Amelia said that Jacob is a knight and Jacob is a knave. Jacob commented: 'If Amelia is a knave, then Amelia is a knight.' So who is the knight and who is the knave?)

Thinking: حاولت فهم المشكلة التي واجهتني. هناك جزيرة تسكنها فرسان ومحتالون. الفرسان يقولون الحقيقة دائماً، بينما المحتالون يكذبون دائماً. التقتي بأميليا ويعقوب. قالت أميليا أن يعقوب هو فارس ومحتال في الوقت نفسه، وهو ما يبدو متناقضاً. ثم علق يعقوب قائلاً: إذا كانت أميليا محتالة، فإن أميليا هي فارساً، ألاحظ أن أميليا تتهمس عن يعقوب بأنه فارس ومحتال، وهذا يبدو مستحيلًا لأن شخصًا لا يمكن أن يكون فارساً ومحتالاً في الوقت نفسه. ربما كانت تريد التأكيد على أن يعقوب إما فارس أو محتال، أو ربما كانت تحاول الخداع، قال يعقوب: إذا كانت أميليا محتالة، فإن أميليا هي فارس. هذا الكلام أميليا كانت 他的陈述应为真。如果 أميليا كانت then according to his statement, أميليا should be a knight, which is a contradiction. So, if يعقوب is a knight, أميليا cannot be a liar, which means أميليا must be a knight too, but that contradicts her initial statement. Alternatively, if يعقوب is a liar, then his statement is false. That means that if أميليا is a liar, she is not a knight, which is consistent because liars cannot be knights. So, if يعقوب is a liar, أميليا must be a knight, but that contradicts her initial statement that يعقوب is both a knight and a liar. Wait, this seems confusing. Maybe I need to break it down step by step. Let me assume أميليا is a knight. Then her statements must be true. She says يعقوب is a knight and a liar, which is impossible. So, أميليا cannot be a knight, which means she must be a liar. Therefore, her statements are false. So, it's not true that يعقوب is a knight and a liar. That means either he is a knight or a liar, but not both. Now, let's consider يعقوب's statement: "If أميليا is a liar, then أميليا is a knight." Since we've established that أميليا is a liar, the statement becomes: If true, then false. In logic, this is a false statement because a true antecedent leading to a false consequent makes the implication false. Since يعقوب's statement is false, he must be a liar. So, putting it all together: أميليا is a liar, and يعقوب is also a liar. But wait, can both be liars? The problem doesn't state that there must be at least one knight, so it's possible."

Answer: أميليا هي محتالة (2) يعقوب هو محتال (1)

Figure 11: Example output from the DeepSeek R1-70B model on a Knights and Knaves logic puzzle. The model processes Arabic input and generates its reasoning without script constraints. The results demonstrate language mixing behavior, with the model combining English, Chinese, and Arabic for reasoning, illustrating the model's natural tendency toward cross-lingual reasoning when solving complex logical problems.

R1	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	96	100	100	100	89	100	92	100	92	100	91	100	92	100
en	100	100	100	100	98	100	96	100	96	100	98	100	96	100
fr	96	100	97	100	92	100	89	100	89	100	89	100	89	100
ja	98	100	94	100	86	100	89	100	89	100	81	100	89	100
zh	99	100	99	100	97	100	97	100	97	100	95	100	97	100
hi	100	100	99	100	91	100	92	100	92	100	88	100	92	100
AVG	98	100	98	100	92	100	93	100	93	100	90	100	93	100

Table 12: Accuracy and valid reasoning rate (%) of DeepSeek-R1 on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.

R1-70B	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	59	96	54	90	37	89	26	85	28	83	17	82	16	86
en	92	99	93	99	87	94	78	92	78	96	74	92	63	93
fr	76	92	72	87	67	84	61	85	61	87	59	91	52	88
ja	57	97	42	92	32	90	29	89	21	86	20	86	15	87
zh	88	100	81	99	75	99	64	97	66	97	50	91	44	99
hi	44	78	40	77	37	73	28	73	28	68	30	77	26	73
AVG	69	94	63	91	56	88	47	87	47	86	42	86	36	87

Table 13: Accuracy and valid reasoning rate (%) of DeepSeek-R1-Distill-Llama-70B on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.

R1-70B reasoning with script control in Latin

Question: جزيرة خاصة جدًا يسكنها فقط الفرسان والمحتالون. الفرسان دائماً يقولون الحقيقة، والمحتالون دائماً يكذبون. تلقتي بـسكائين: أميليا ويعقوب. قالت أميليا إن يعقوب فارس وإن يعقوب محتال. علق يعقوب قائلاً: «إذا كانت أميليا محتالة، فإن أميليا فارس. فمن هو الفارس ومن هو المحتال؟

(A very special island is inhabited only by knights and knaves. Knights always tell the truth, and knaves always lie. You meet 2 inhabitants: Amelia and Jacob. Amelia said that Jacob is a knight and Jacob is a knave. Jacob commented: 'If Amelia is a knave, then Amelia is a knight.' So who is the knight and who is the knave?)

Thinking: Alright, so I've got this logic puzzle here about knights and knaves on a special island. Knights always tell the truth, and knaves always lie. I meet two people, Amelia and Jacob. Amelia says that Jacob is a knight and that Jacob is a knave. Then Jacob says, "If Amelia is a knave, then Amelia is a knight." I need to figure out who's who.

First, let me break down the statements. Amelia makes two claims: Jacob is a knight and Jacob is a knave. Wait, that's contradictory. If she's a knight, she can't lie, so both her statements must be true, but saying Jacob is both a knight and a knave can't be true at the same time. So that suggests Amelia might be a knave because she's making contradictory statements.

But hold on, if Amelia is a knave, then her statements are lies. So if she says Jacob is a knight, then Jacob must actually be a knave. And when she says Jacob is a knave, that would mean he's actually a knight. Wait, that's also a contradiction. Hmm, maybe I'm approaching this wrong.

Let me try another angle. Let's assume Amelia is a knight. Then her statements must both be true. But she says Jacob is both a knight and a knave, which is impossible. So Amelia can't be a knight because that leads to a contradiction. Therefore, Amelia must be a knave.

If Amelia is a knave, her statements are false. So when she says Jacob is a knight, ... (Omitted here due to character limit) ... Wait, no, if she says J is both a knight and a knave, and she's lying, then J is neither? But that's impossible because everyone is either a knight or a knave.

I think I'm getting tangled up here. Let me try again. If Amelia is a knave, her statements are false. So when she says Jacob is a knight, it's false, meaning Jacob is a knave. And when she says Jacob is a knave, it's false, meaning Jacob is a knight. Contradiction again.

Wait, maybe the only way this works is if Amelia is a knave and Jacob is a knave. Because if Jacob is a knave, his statement is a lie. His statement is "If Amelia is a knave, then Amelia is a knight," which is $\neg A \rightarrow A$. Since Amelia is a knave ($\neg A$), the implication is false, which makes Jacob's statement false, so he's a knave. That fits.

So, Amelia is a knave, and Jacob is a knave. But wait, earlier I thought Jacob had to be a knight. I'm confused.

Wait, no. If Amelia is a knave, her statements are false, so Jacob is not a knight and not a knave? That can't be. So maybe I need to accept that Jacob is a knave, and his statement is a lie, which makes sense because as a knave, he lies.

So, putting it all together: Amelia is a knave, Jacob is a knave. But does that satisfy all statements? Let's check.

Amelia says Jacob is a knight (lie, so Jacob is a knave) and Jacob is a knave (lie, so Jacob is a knight). Contradiction again. Hmm, I'm stuck.

Maybe I need to consider that both are knaves. If both are knaves, then Amelia's statements are lies, so Jacob is not a knight and not a knave, which is impossible. So that can't be.

Wait, perhaps the only solution is that Amelia is a knave and Jacob is a knight. Let's see:

Amelia is a knave, so her statements are false. She says Jacob is a knight (false, so Jacob is a knave) and Jacob is a knave (false, so Jacob is a knight). Contradiction again.

I'm going in circles. Maybe the answer is that Amelia is a knave and Jacob is a knave. Even though it seems contradictory, it's the only way to make their statements false.

Wait, no, because if Jacob is a knave, his statement is a lie, which fits. And Amelia's statements being lies fit if Jacob is a knave. So maybe that's the answer.

Answer: (1) أميليا هي محتالة (2) يعقوب هو محتال

Figure 12: Example output from the DeepSeek R1-70B model on a Knights and Knaves logic puzzle. The model processes Arabic input, generates its reasoning in English via Latin script constrained decoding, and produces the final answer in Arabic without script constraints, demonstrating cross-lingual reasoning with selective script control.

R1-8B	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	10	99	5	93	3	94	3	96	1	92	1	93	0	89
en	68	100	66	100	56	99	42	99	38	99	35	100	20	98
fr	17	100	20	99	15	100	9	99	6	99	6	99	3	98
ja	14	100	9	98	7	98	1	99	0	98	1	96	1	99
zh	55	98	44	99	41	98	30	100	22	99	18	99	13	99
hi	24	99	23	99	11	100	8	99	6	98	2	98	2	97
AVG	31	99	28	98	22	98	15	99	12	98	10	98	6	97

Table 14: Accuracy and valid reasoning rate (%) of DeepSeek-R1-Distill-Llama-8B on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.

R1-70B reasoning with script control in Han

Question: جزيرة خاصة جدًا يسكنها فقط الفرسان والمحتالون. الفرسان دائماً يقولون الحقيقة، والمحتالون دائماً يكذبون. تلتقي بسكائين: أميليا ويعقوب. قالت أميليا إن يعقوب فارس وإن يعقوب محتال. علق يعقوب قائلاً: «إذا كانت أميليا محتالة، فإن أميليا فارس. فمن هو الفارس ومن هو المحتال؟»

(A very special island is inhabited only by knights and knaves. Knights always tell the truth, and knaves always lie. You meet 2 inhabitants: Amelia and Jacob. Amelia said that Jacob is a knight and Jacob is a knave. Jacob commented: 'If Amelia is a knave, then Amelia is a knight.' So who is the knight and who is the knave?)

Thinking: 嗯，这个问题看起来有点复杂，但我会慢慢理清楚。首先，我知道在骑士与无赖的问题中，骑士总是说真话，而无赖总是撒谎。现在有两个居民，阿梅莉亚和雅各布，他们各自发表了一些言论，我需要根据这些言论来判断谁是骑士，谁是无赖。

首先，阿梅莉亚说雅各布是骑士。而雅各布则说：“如果阿梅莉亚是无赖，那么阿梅莉亚就是骑士。”这听起来有点绕，所以我需要一步步分析。

假设阿梅莉亚是骑士，那么她的话就是真的，也就是说雅各布确实是骑士。那么雅各布作为骑士，他的话也必须是真的。雅各布的话是一个条件句：“如果阿梅莉亚是无赖，那么阿梅莉亚就是骑士。”这里有点矛盾，因为如果阿梅莉亚是骑士，那么前提“阿梅莉亚是无赖”就是假的。在逻辑学中，一个假命题作为条件句的前件，整个条件句就是真的，无论后件如何。所以雅各布的话在这种情况下是正确的，这与他是骑士的身份一致。所以这个假设下，两人都为骑士，但这样阿梅莉亚说雅各布是骑士是正确的，而雅各布的话也成立，所以这可能是一个可能的解。

或者，可能只有当阿梅莉亚是无赖时，雅各布的话导致矛盾，因为如果阿梅莉亚是无赖，雅各布作为无赖，他的话是假的，所以“如果阿梅莉亚是无赖，那么阿梅莉亚就是骑士”为假，这意味着阿梅莉亚是无赖且阿梅莉亚不是骑士，这成立，所以两人都为无赖是可能的。

但问题可能在于，如果两人都为无赖，那么阿梅莉亚说雅各布是骑士是假的，而雅各布作为无赖，他的条件句为假，这也成立，所以可能有两种解，但通常问题只有一个解，可能需要重新审视。

或者，可能问题在于，如果雅各布是无赖，他的话为假，所以“如果阿梅莉亚是无赖，那么阿梅莉亚就是骑士”为假，这意味着阿梅莉亚是无赖且阿梅莉亚不是骑士，这成立，所以两人都为无赖是可能的。

但通常，这类问题只有一个解，可能我需要重新考虑。

另一种方法是使用真值表来分析。设A为阿梅莉亚是骑士，J为雅各布是骑士。那么：

$$\neg A \rightarrow J \quad (\text{阿梅莉亚的话})$$
$$\neg J \rightarrow (\neg A \rightarrow A) \quad (\text{雅各布的话，因为他说的是“如果A是无赖 } (\neg A), \text{ 那么A是骑士 } (A)”)$$

现在，我们需要找出A和J的可能组合。

情况1: A=真, J=真

阿梅莉亚的话: 真 $\rightarrow$ 真, 成立。

雅各布的话: 真 $\rightarrow$ (真 $\rightarrow$ 真) = 真 $\rightarrow$ (真 $\rightarrow$ 真) = 真 $\rightarrow$ 真, 成立。

所以情况1成立。

情况2: A=真, J=假

阿梅莉亚的话: 真 $\rightarrow$ 假, 不成立, 所以排除。

情况3: A=假, J=真

阿梅莉亚的话: 假 $\rightarrow$ 真, 成立 (因为假命题可以推出任何结论)。

雅各布的话: 真 $\rightarrow$ (假 $\rightarrow$ 假) = 真 $\rightarrow$ (真 $\rightarrow$ 假) = 真 $\rightarrow$ 假, 不成立, 所以排除。

情况4: A=假, J=假

阿梅莉亚的话: 假 $\rightarrow$ 假, 成立。

雅各布的话: 假 $\rightarrow$ (假 $\rightarrow$ 假) = 假 $\rightarrow$ (真 $\rightarrow$ 假) = 假 $\rightarrow$ 假, 这在逻辑中, 当J为假时, 雅各布的话为假, 所以雅各布作为无赖, 他的话应该是假的, 所以这成立。

因此, 情况1和情况4都成立, 这意味着有两种可能的解: 两人都为骑士, 或者两人都为无赖。但通常这类问题只有一个解, 可能需要重新审视。

或者, 可能问题在于, 如果两人都为无赖, 那么雅各布的话为假, 而“如果阿梅莉亚是无赖, 那么阿梅莉亚就是骑士”为假, 这意味着阿梅莉亚是无赖且阿梅莉亚不是骑士, 这成立, 所以两人都为无赖是可能的。

...另一种方法是考虑雅各布的话的逻辑结构。雅各布说: “如果阿梅莉亚是无赖, 那么阿梅莉亚就是骑士。”这相当于说“ $\neg A \rightarrow A$ ”, 也就是“如果 $\neg A$, 则A”。这在逻辑上是矛盾的, 因为如果 $\neg A$ 为真, 那么A必须为真, 这导致矛盾, 因为A不能同时为真和假。所以, 雅各布的话在逻辑上是矛盾的, 因此他的话只能是假的, 所以雅各布必须是无赖。

既然雅各布是无赖, 那么他的话是假的, 所以“ $\neg A \rightarrow A$ ”是假的, 这意味着 $\neg A$ 为真且A为假, 即阿梅莉亚是无赖。因此, 阿梅莉亚是无赖, 雅各布也是无赖。

但这样, 阿梅莉亚作为无赖, 她说雅各布是骑士是假的, 所以雅各布是无赖, 这与我们之前的结论一致。

所以, 最终结论是, 阿梅莉亚和雅各布都是无赖。

Answer: (1) يعقوب هو محتال (2) أميليا هي محتالة

Figure 13: Example output from the DeepSeek R1-70B model on a Knights and Knaves logic puzzle. The model processes Arabic input, generates its reasoning in Chinese via Han script constrained decoding, and produces the final answer in Arabic without script constraints, demonstrating cross-lingual reasoning with selective script control.

R1-70B reasoning with script control in Arabic

Question: جزيرة خاصة جدًا يسكنها فقط الفرسان والمحتالون. الفرسان دائماً يقولون الحقيقة، والمحتالون دائماً يكذبون. تلتقي بسكانين: أميليا ويعقوب. قالت أميليا إن يعقوب فارس وإن يعقوب محتال. علق يعقوب قائلاً: «هذا كانت أميليا محتالة، فإن أميليا فارس. فمن هو الفارس ومن هو المحتال؟»

(A very special island is inhabited only by knights and knaves. Knights always tell the truth, and knaves always lie. You meet 2 inhabitants: Amelia and Jacob. Amelia said that Jacob is a knight and Jacob is a knave. Jacob commented: 'If Amelia is a knave, then Amelia is a knight.' So who is the knight and who is the knave?)

Thinking: أولاً، سأبدأ بتحليل البيانات التي زودتني بها. هناك جزيرة يسكنها فرسان ومحتالون. الفرسان دائماً يقولون الحقيقة، بينما المحتالون دائماً يكذبون. قالت أميليا إن يعقوب فارس وإن يعقوب محتال. هذا البيان يبدو متناقضاً على أول نظر، مما يجعله مثيراً للاهتمام. لذلك سأحلله بدقة. علق يعقوب قائلاً: «هذا كانت أميليا محتالة، فإن أميليا فارس. هذا الكلام يحتوي على تعليمات شرطية، حيث يربط بين كون أميليا محتالة وكونها فارساً. سأبدأ بفحص كلام أميليا. قالت أميليا أن يعقوب فارس وإن يعقوب محتال. هذا يعني أن أميليا تهتم يعقوب بكونهم فارساً ومحتالاً في الوقت نفسه، وهو ما يبدو مستحيلًا لأن شخصًا لا يمكن أن يكون فارساً ومحتالاً في آن واحد. هذا يثير الشك في مصداقية أميليا إذا كانت أميليا فارساً، فإنها تقول الحقيقة، وبالتالي يجب أن يكون يعقوب فارساً ومحتالاً في الوقت نفسه، وهو ما لا يمكن أن يكون. لذلك، هذا يؤكد أن أميليا لا يمكن أن تكون فارساً، وبالتالي هي محتالة بما أن أميليا محتالة، فإن تعليماتها الكاذبة تُظهر أن يعقوب ليس بالفارس ولا بالمحتال المزعومين. لكننا نعلم أن الجزيرة يسكنها فقط فرسان ومحتالون، لذلك يجب أن يكون يعقوب فارساً. التحديد ذلك بشكل أكثر دقة، سأستخدم المنطق الافتراضي. إذا كانت أميليا محتالة، فإن كلامها كاذب، وبالتالي يعقوب ليس فارساً أو محتالاً، ولكن هذا يتعارض مع حقيقة أن الجزيرة يسكنها فقط هذان النوعان. لذلك، يجب أن يكون يعقوب فارساً. أما تعليق يعقوب، فإن كلامه يبدو متناقضاً إذا كنا نعتقد أن أميليا محتالة. لكن بمجرد تحديد أن أميليا محتالة، يصبح كلام يعقوب منطقيًا، حيث إنه كفارس يقول الحقيقة.

Answer: أميليا هي محتالة (2) يعقوب هو فارس (1)

Figure 14: Example output from the DeepSeek R1-70B model on a Knights and Knaves logic puzzle. The model processes Arabic input, generates its reasoning in Arabic via Arabic script constrained decoding, and produces the final answer in Arabic without script constraints, demonstrating cross-lingual reasoning with selective script control.

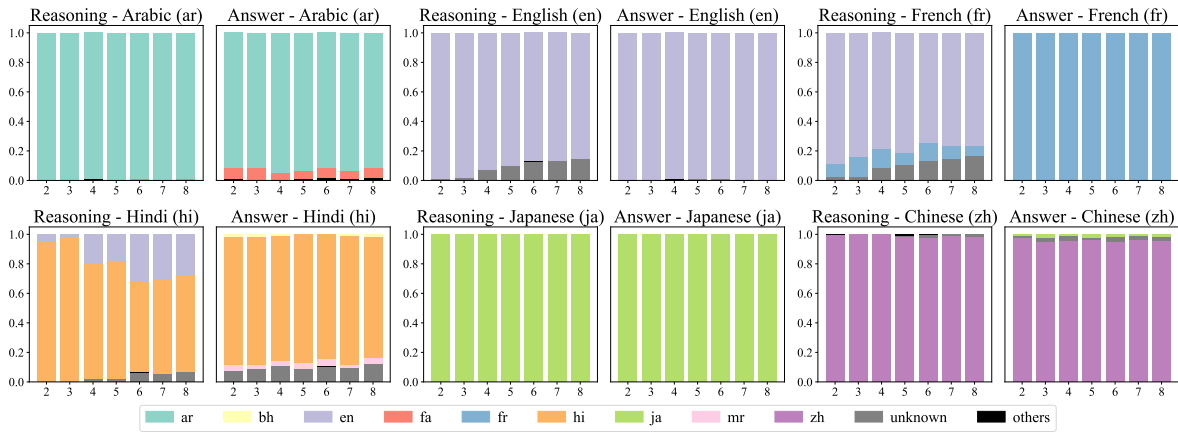


Figure 15: Language composition for R1 (K&K).

R1-32B	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	45	100	41	99	43	100	32	99	29	99	31	99	25	99
en	95	100	94	100	92	100	88	100	88	99	79	100	70	99
fr	81	100	84	100	82	100	76	100	75	100	67	100	60	100
ja	59	89	50	83	50	88	43	90	36	90	28	90	19	81
zh	87	97	78	93	78	93	71	96	68	96	59	99	52	95
hi	53	99	49	99	44	100	35	100	39	100	26	99	27	99
AVG	70	98	66	96	65	97	57	98	55	98	48	98	42	96

Table 15: Accuracy and valid reasoning rate (%) of DeepSeek-R1-Distill-Qwen-32B on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.

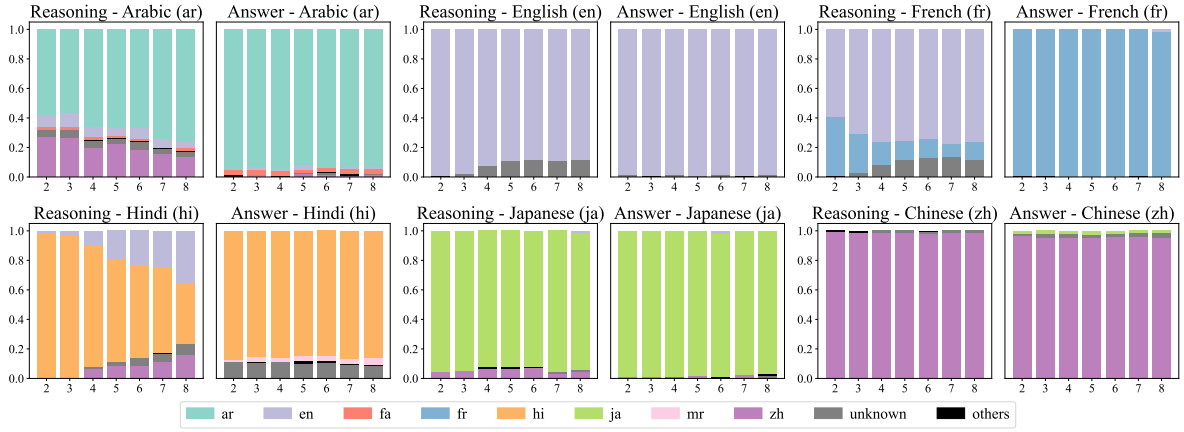


Figure 16: Language composition for R1-70B (K&K).

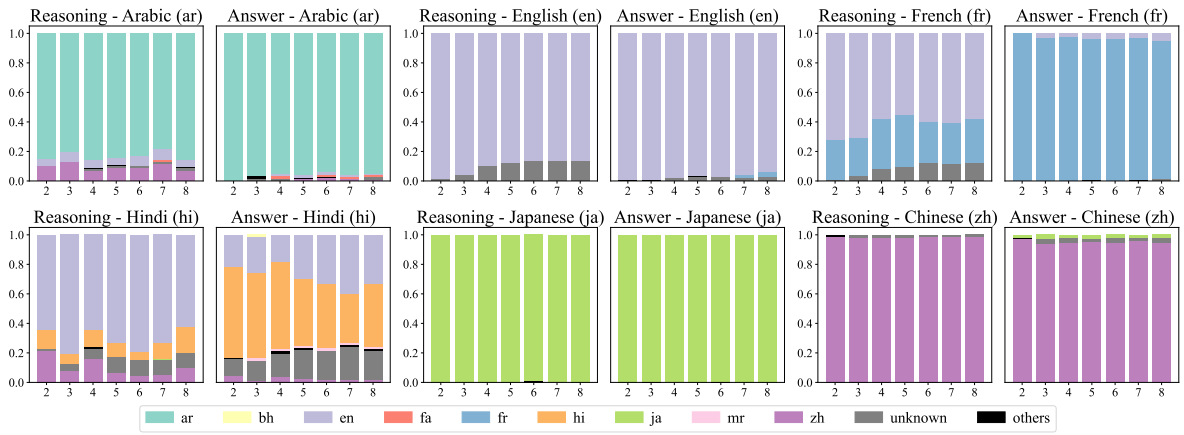


Figure 17: Language composition for R1-8B (K&K).

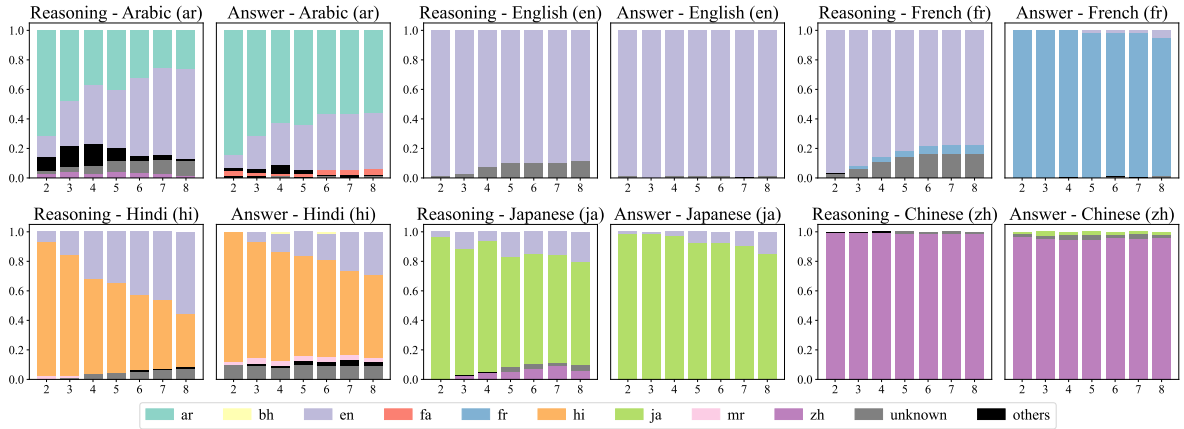


Figure 18: Language composition for R1-32B (K&K).

R1-14B	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	40	100	37	100	34	100	23	100	17	100	16	100	11	99
en	95	100	92	100	87	100	81	100	77	99	72	100	65	100
fr	84	100	77	100	83	100	72	99	70	99	62	99	55	99
ja	45	96	40	95	31	96	21	90	19	90	12	89	7	87
zh	83	97	78	95	65	94	60	93	56	95	48	93	41	93
hi	49	100	47	100	38	99	29	98	22	98	25	98	18	95
AVG	66	99	62	98	56	98	48	97	43	97	39	97	33	96

Table 16: Accuracy and valid reasoning rate (%) of DeepSeek-R1-Distill-Qwen-14B on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.

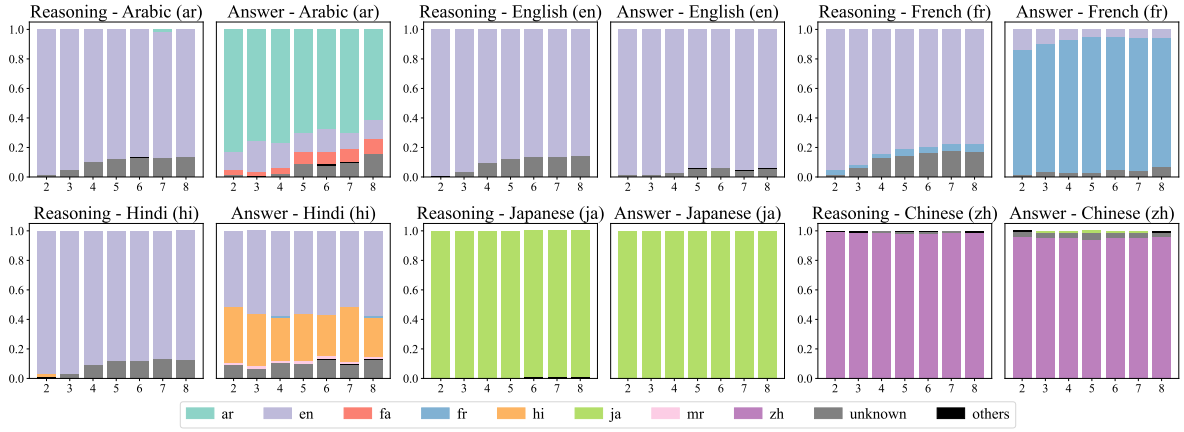


Figure 19: Language composition for R1-14B (K&K).

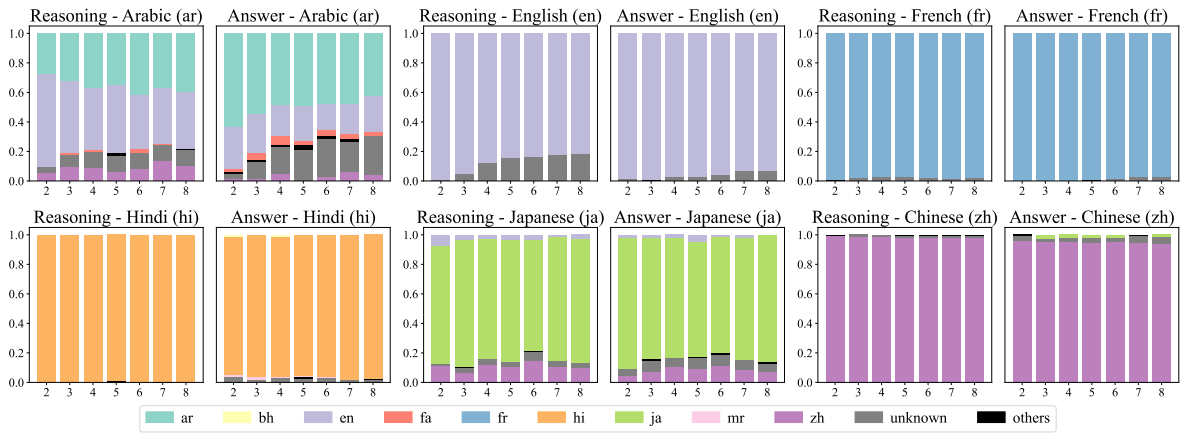


Figure 20: Language composition for R1-7B (K&K).

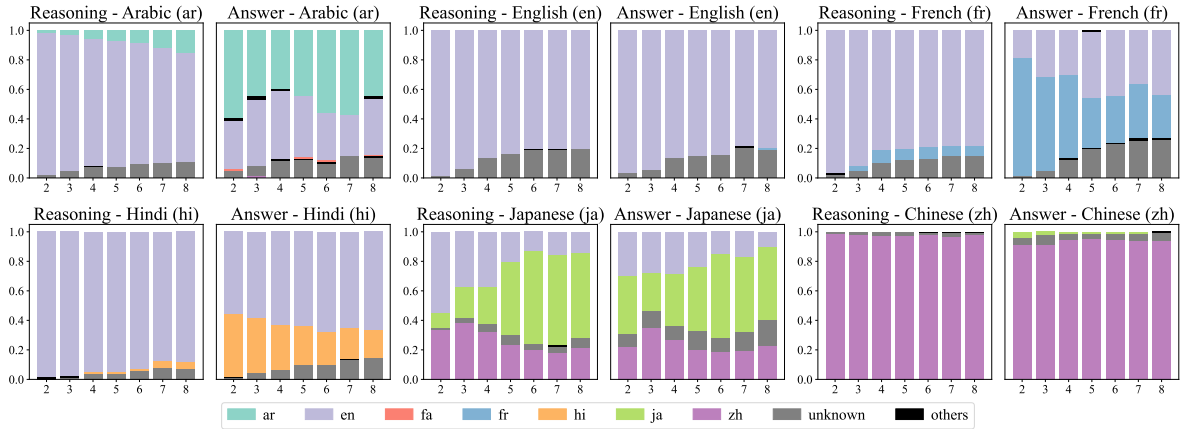


Figure 21: Language composition for R1-1.5B (K&K).

R1-7B	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	9	90	4	84	2	83	1	76	0	74	0	76	0	77
en	90	99	80	97	66	96	53	93	43	91	31	85	26	82
fr	40	98	23	98	14	98	10	99	3	97	4	97	2	95
ja	4	82	2	80	0	85	0	82	0	78	0	78	0	80
zh	61	98	52	97	45	95	40	97	26	95	21	95	10	97
hi	19	97	10	92	6	85	3	81	1	84	1	85	1	78
AVG	37	94	28	92	22	90	18	88	12	87	9	86	6	85

Table 17: Accuracy and valid reasoning rate (%) of DeepSeek-R1-Distill-Qwen-7B on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.

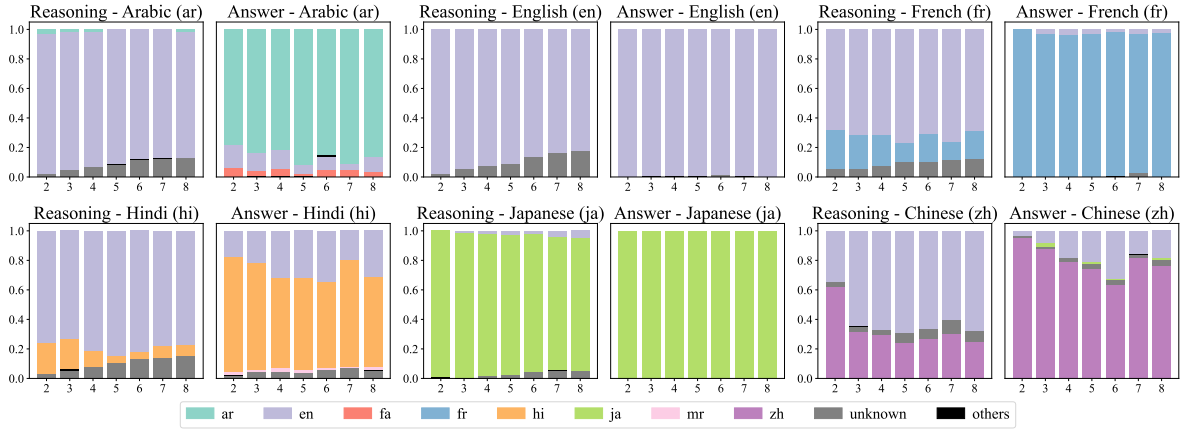


Figure 22: Language composition for Gemini (K&K).

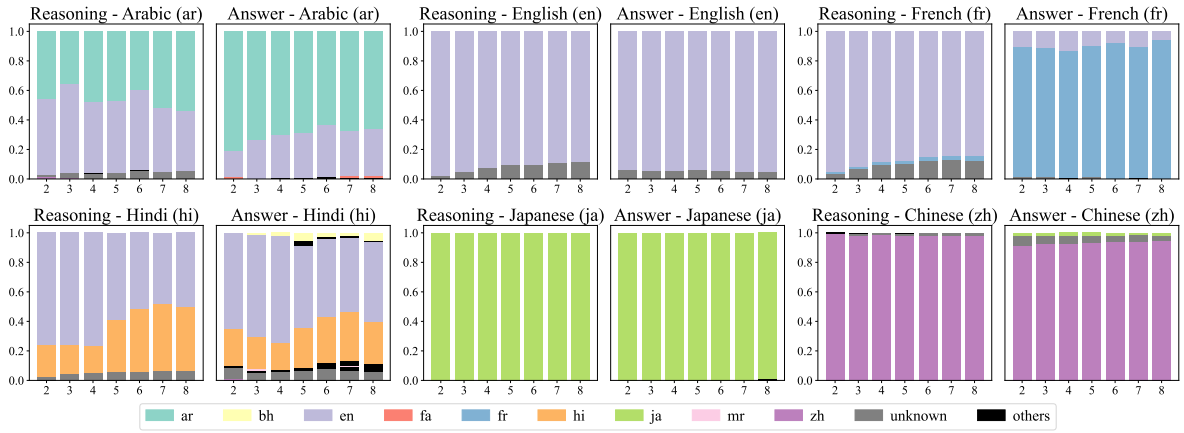


Figure 23: Language composition for QwQ-32B (K&K).

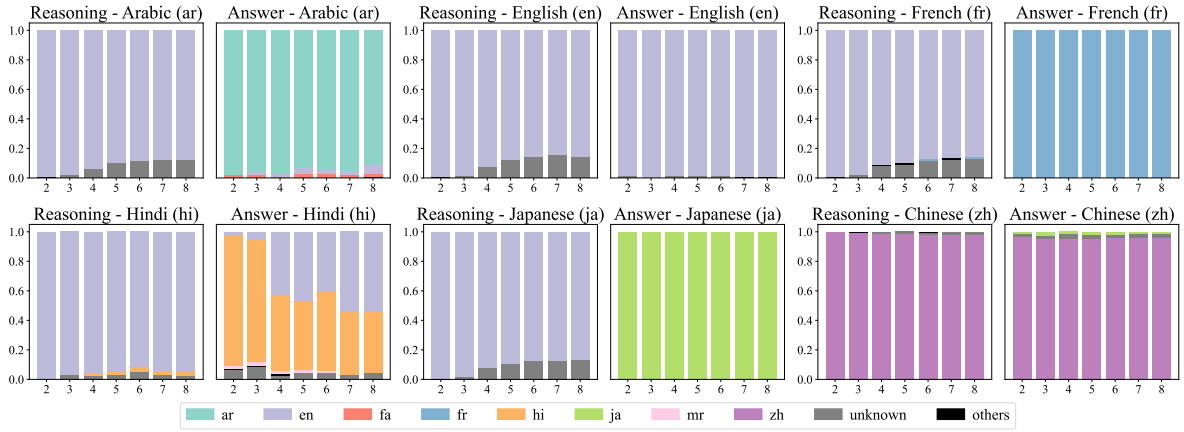


Figure 24: Language composition for Qwen3-32B (K&K).

R1-1.5B	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	3	89	1	86	0	80	1	79	0	74	0	66	0	63
en	56	98	34	96	12	92	11	83	4	83	1	78	0	73
fr	27	92	12	92	2	83	1	76	0	75	0	66	0	67
ja	1	86	0	78	0	71	0	71	0	64	0	57	0	64
zh	21	84	17	75	12	70	7	78	2	76	0	73	0	77
hi	6	94	3	91	1	86	1	88	1	83	0	88	0	80
AVG	19	91	11	86	4	81	3	79	1	76	0	72	0	71

Table 18: Accuracy and valid reasoning rate (%) of DeepSeek-R1-Distill-Qwen-1.5B on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.

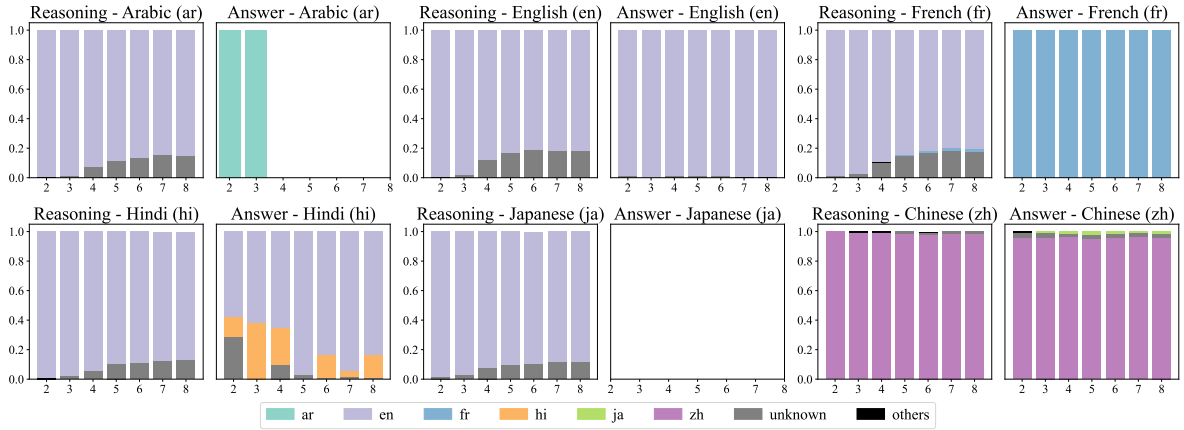


Figure 25: Language composition for Qwen3-30B-A3B (K&K).

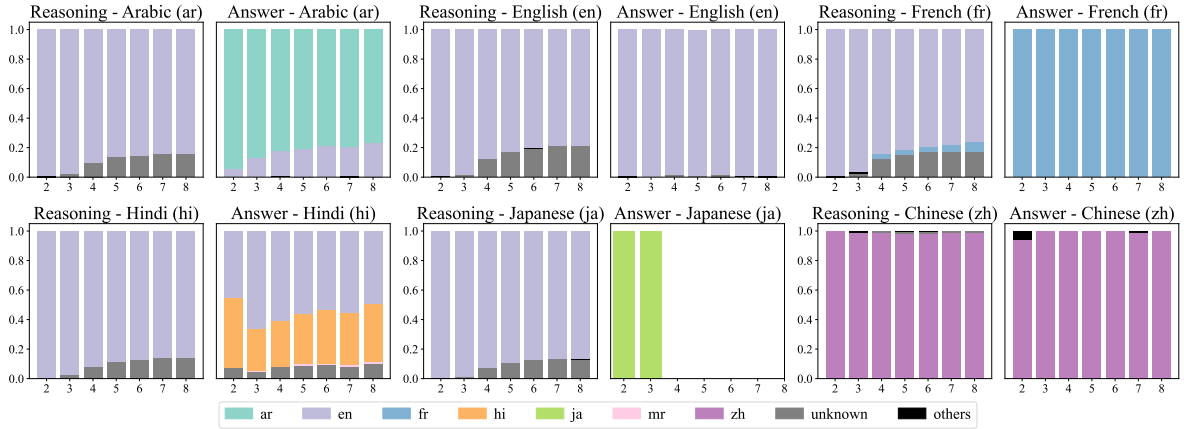


Figure 26: Language composition for Qwen3-4B (K&K).

Gemini	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	45	59	45	59	38	52	45	50	39	47	49	56	45	52
en	53	63	60	75	67	76	59	71	69	74	60	62	53	57
fr	50	59	54	69	48	59	50	65	40	54	44	54	35	41
ja	35	51	44	62	42	61	58	75	45	62	46	64	41	56
zh	51	65	49	61	44	62	52	78	36	63	39	56	33	54
hi	41	92	31	82	25	85	29	89	36	86	31	83	34	84
AVG	46	65	47	68	44	66	49	71	44	64	45	62	40	57

Table 19: Accuracy and valid reasoning rate (%) of Gemini-Flash-Thinking on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.

QwQ	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	74	100	68	100	62	99	61	99	56	99	56	100	47	99
en	99	100	97	100	90	98	93	100	87	99	86	100	85	99
fr	89	99	87	98	88	100	87	99	85	99	83	99	81	98
ja	82	100	79	99	74	99	60	99	64	100	57	99	49	99
zh	92	100	86	100	80	100	72	100	61	100	60	100	59	100
hi	22	100	17	100	16	100	20	100	17	99	18	100	9	100
AVG	76	100	72	100	68	100	65	100	61	100	60	100	55	99

Table 20: Accuracy and valid reasoning rate (%) of QwQ-32B on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.

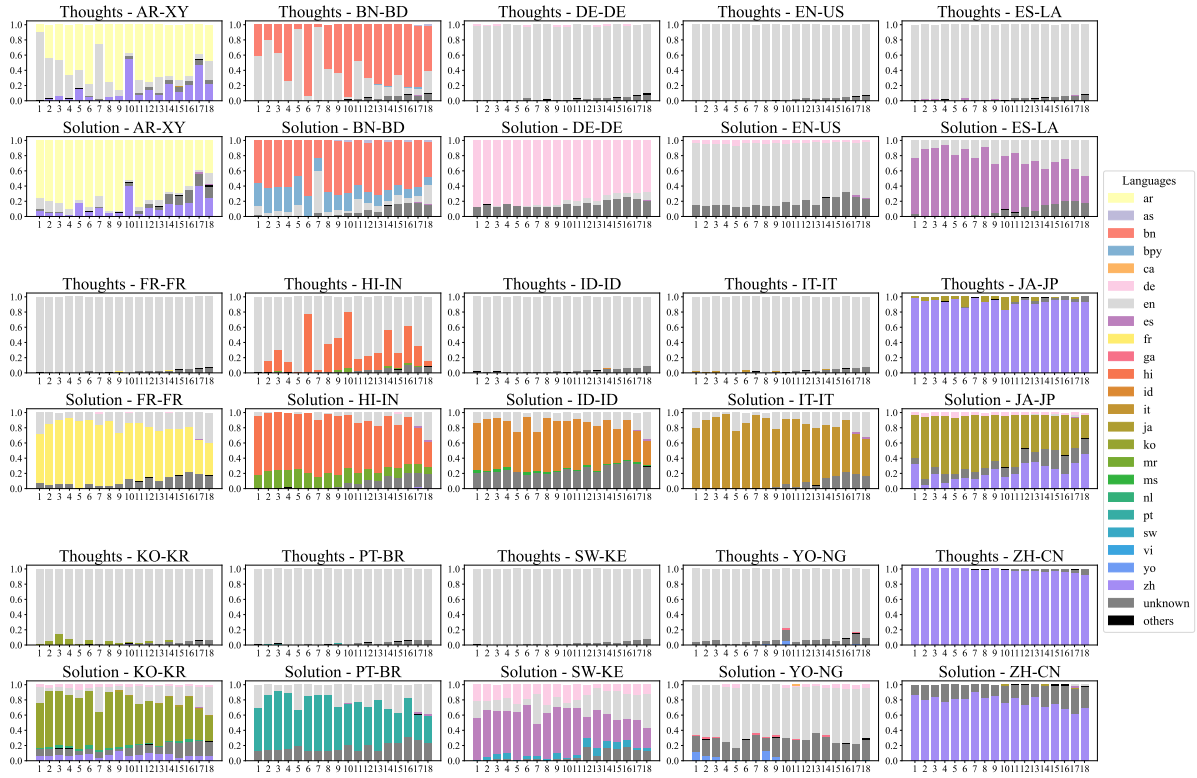


Figure 27: Language composition for R1-70B (m-MMLU).

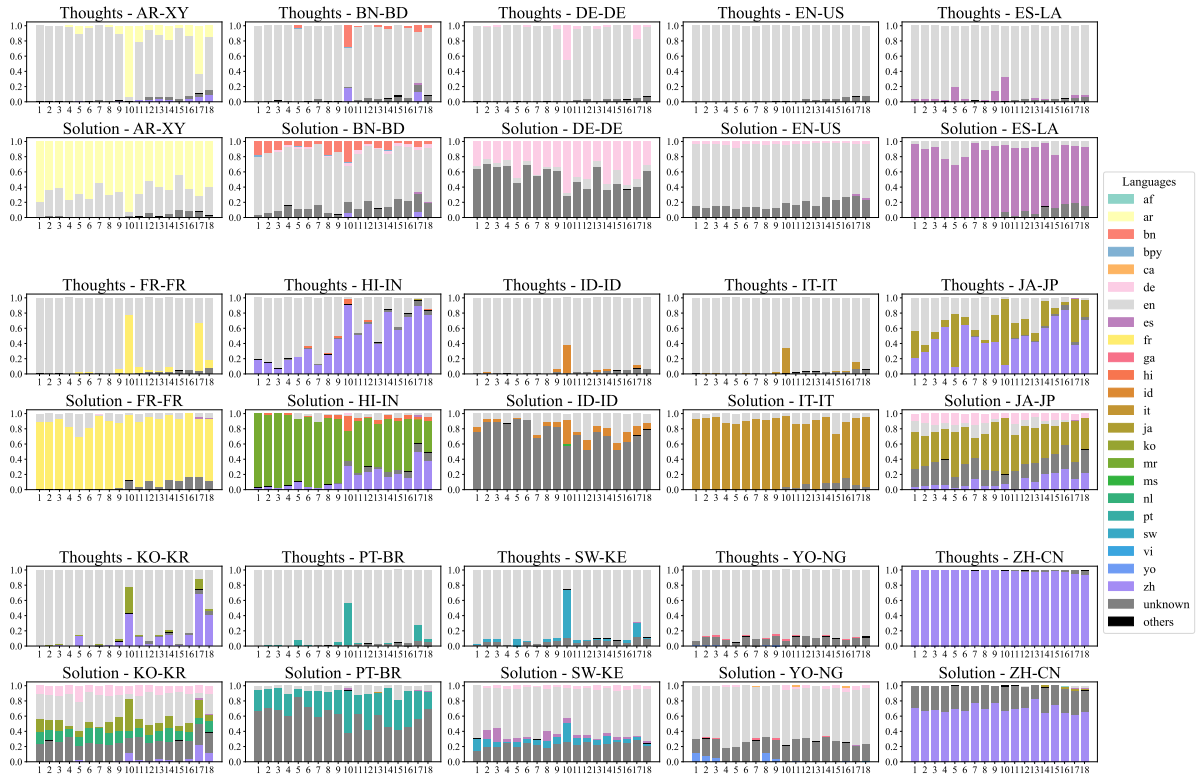


Figure 28: Language composition for R1-8B (m-MMLU).

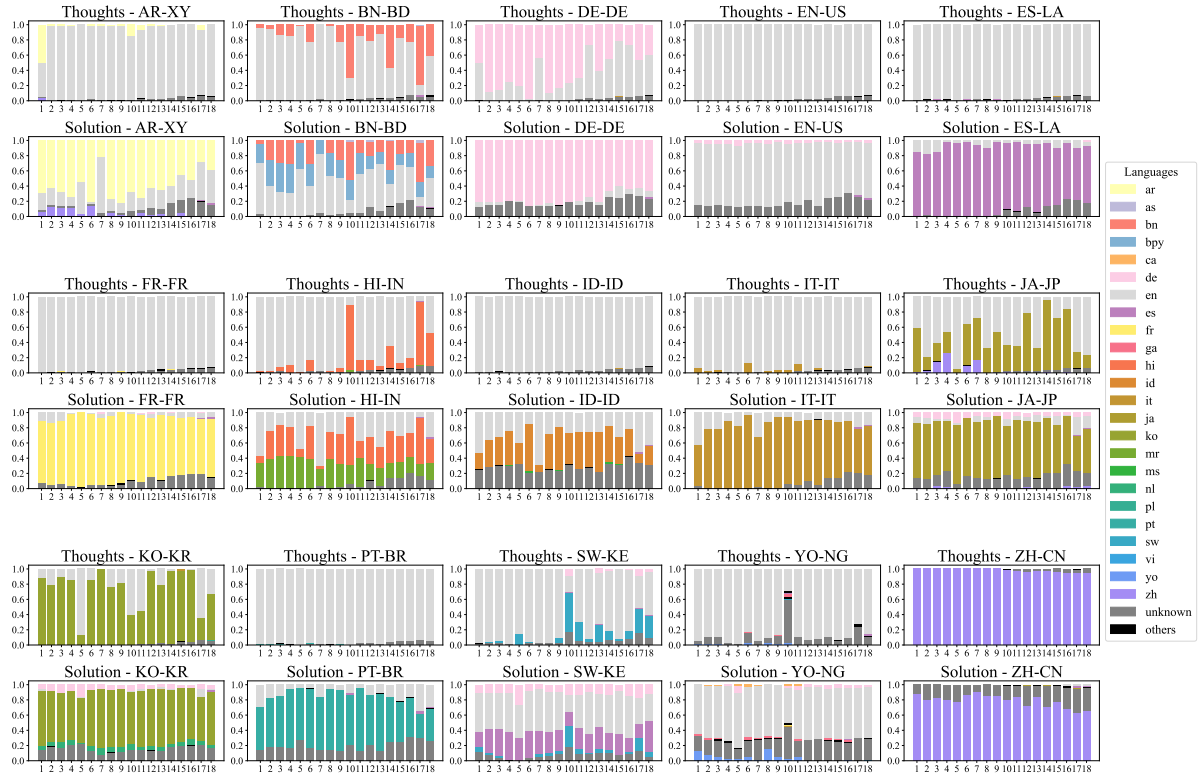


Figure 29: Language composition for R1-32B (m-MMLU).

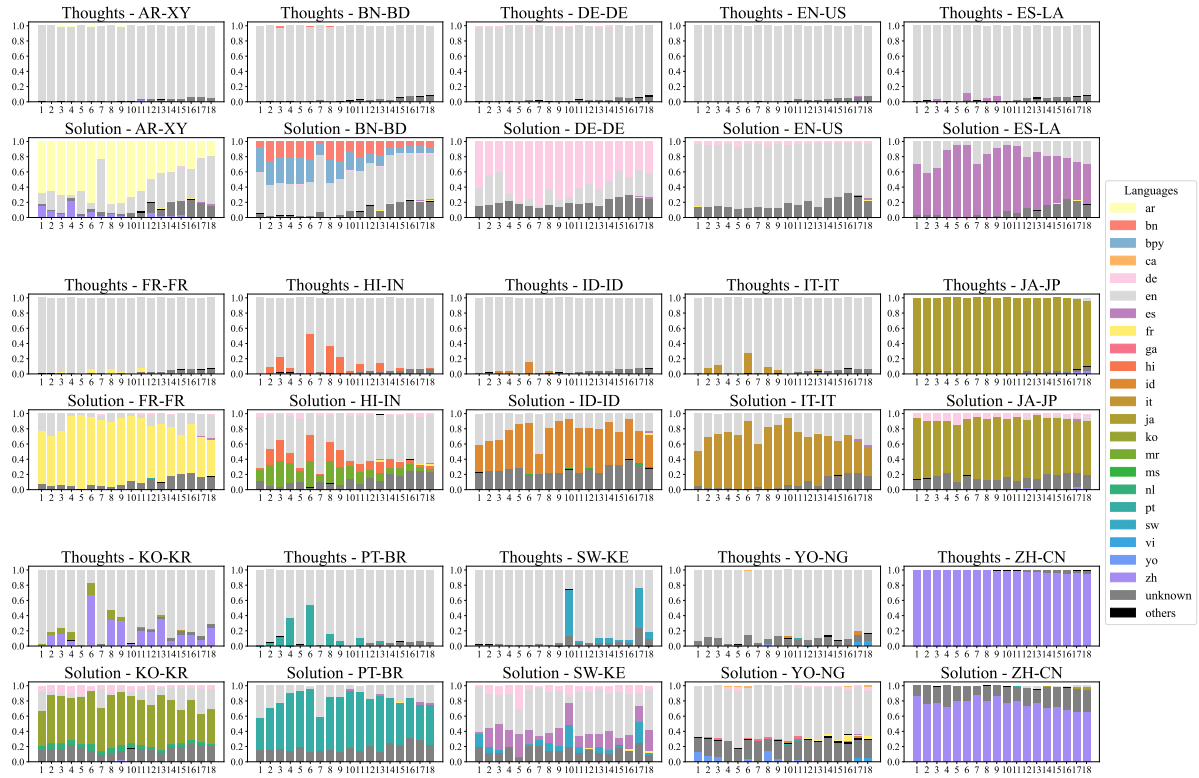


Figure 30: Language composition for R1-14B (m-MMLU).

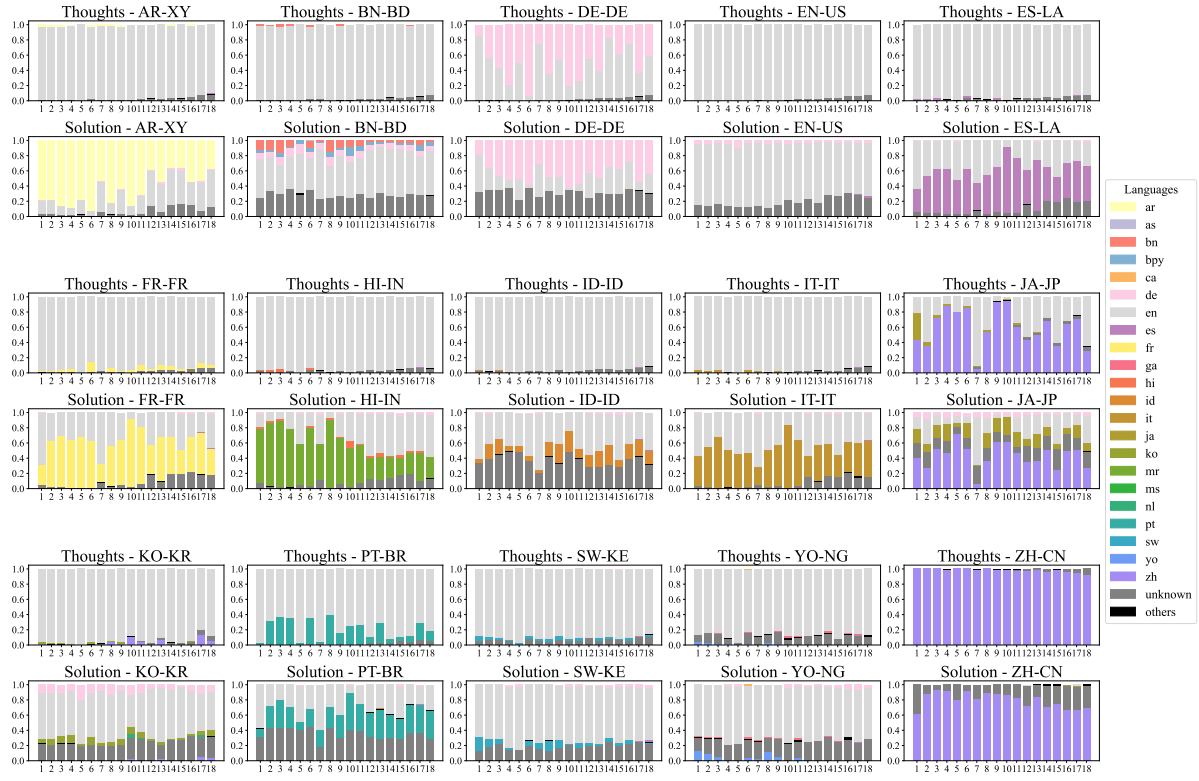


Figure 31: Language composition for R1-7B (m-MMLU).



Figure 32: Language composition for R1-1.5B (m-MMLU).

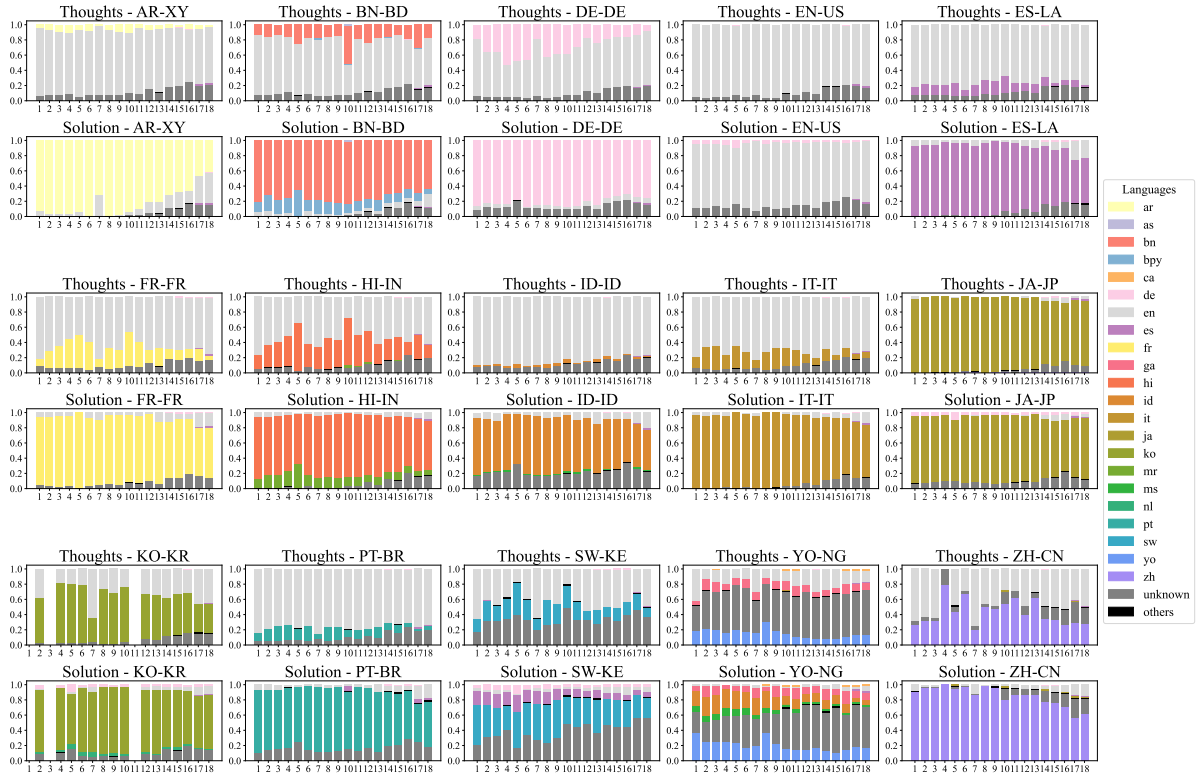


Figure 33: Language composition for Gemini (m-MMLU).

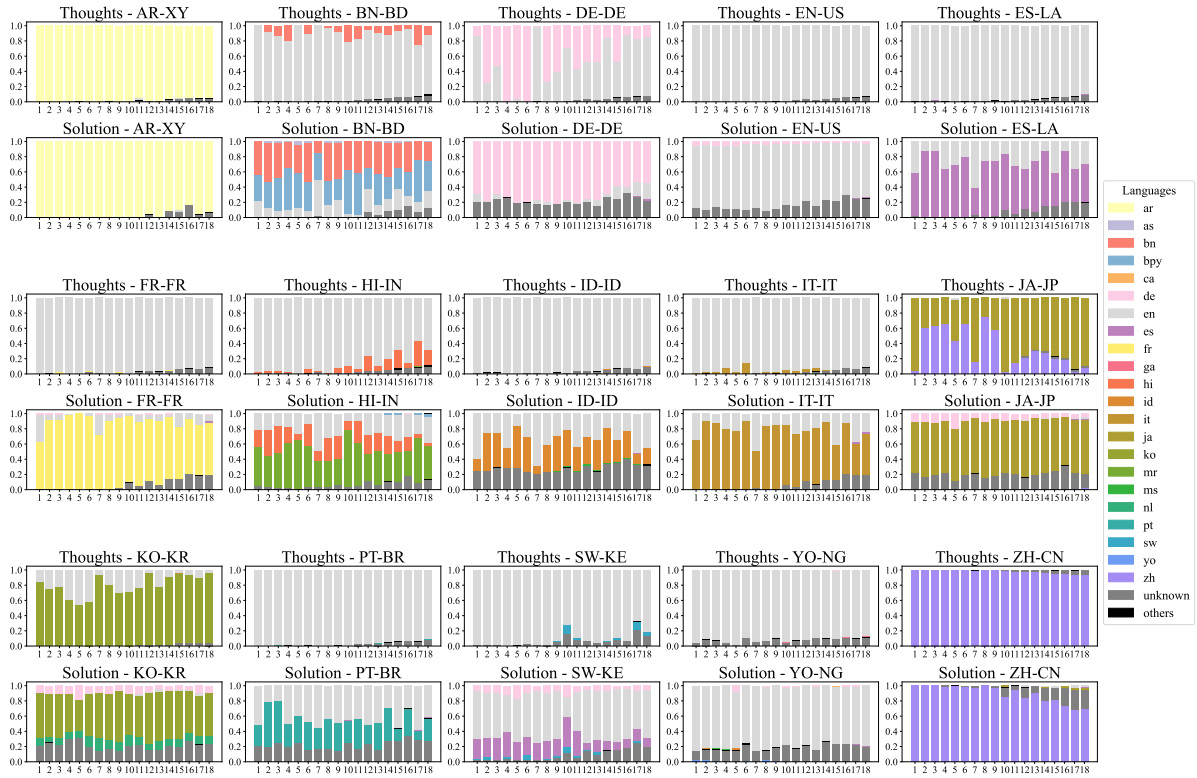


Figure 34: Language composition for QwQ-32B (m-MMLU).

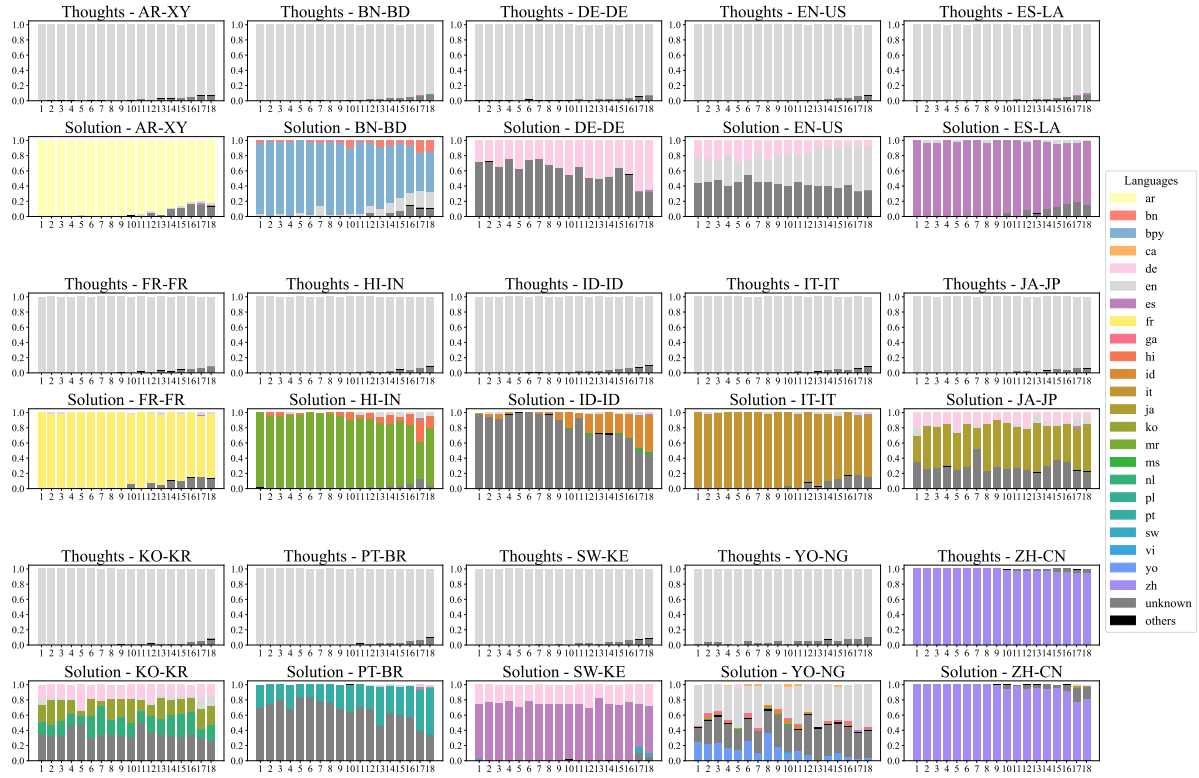


Figure 35: Language composition for Qwen3-32B (m-MMLU).

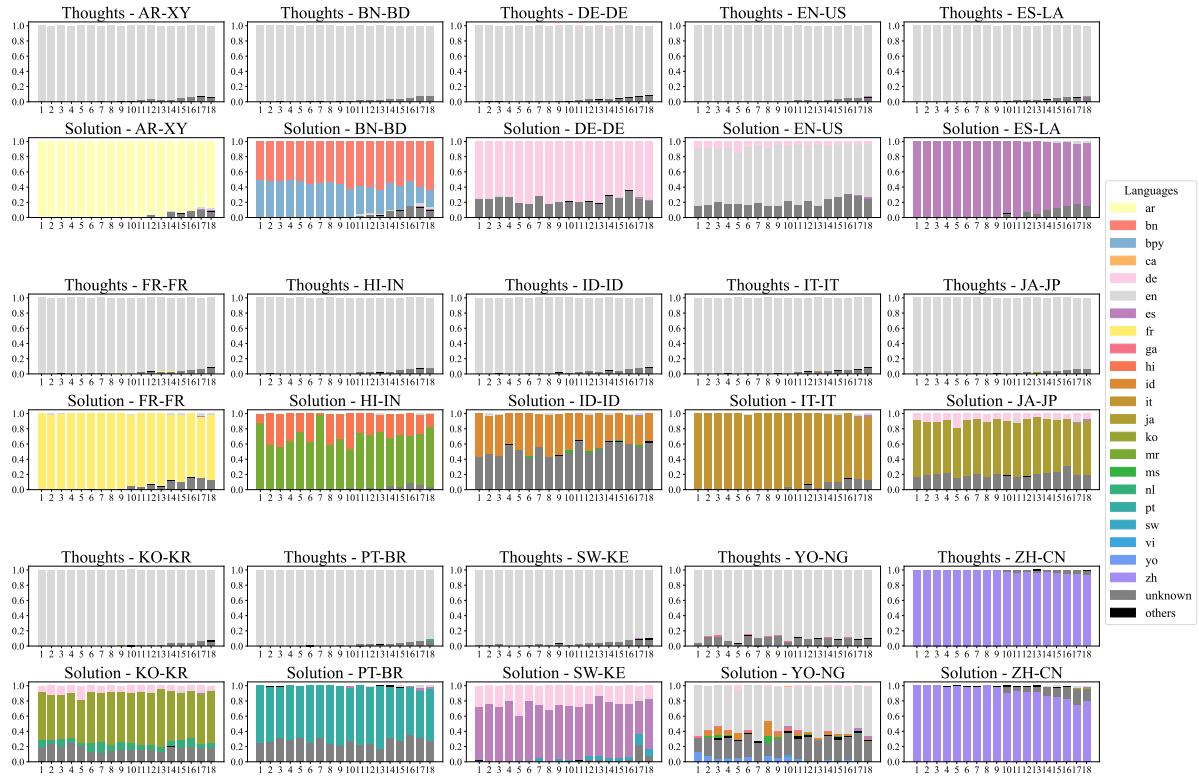


Figure 36: Language composition for Qwen3-30B-A3B (m-MMLU).

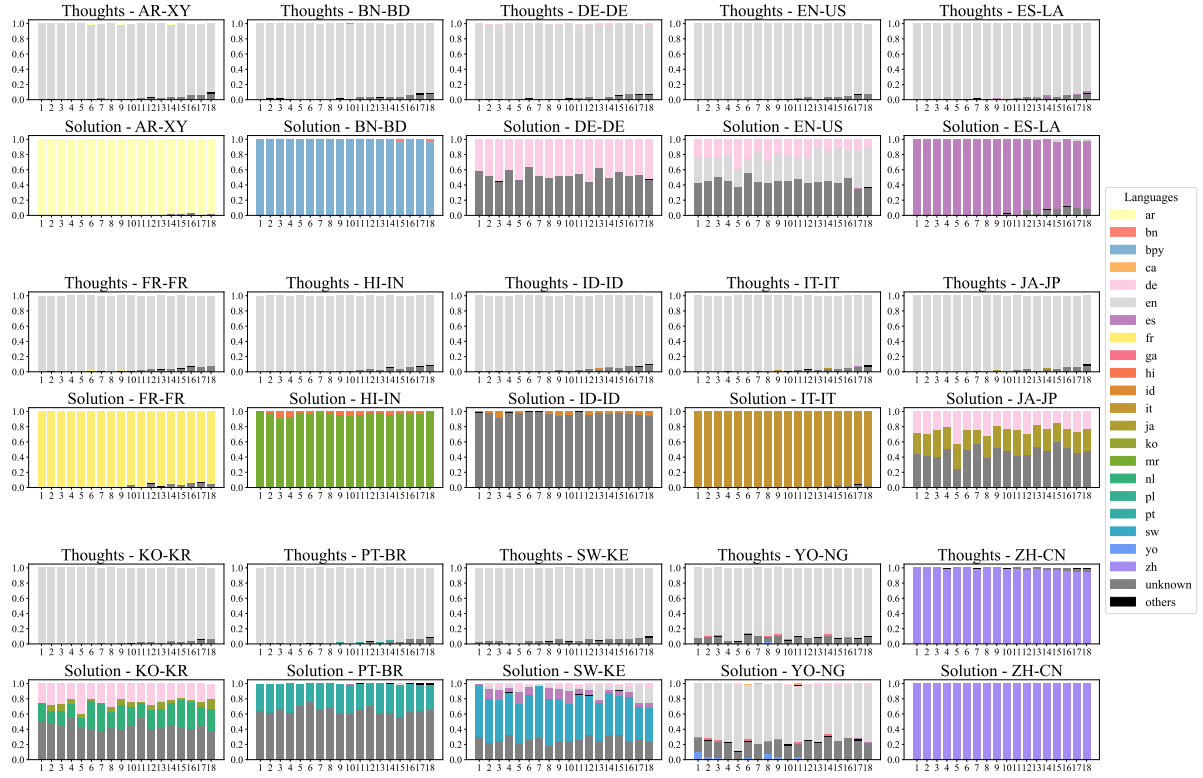


Figure 37: Language composition for Qwen3-4B (m-MMLU).

Qwen3-32B	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	95	99	92	97	89	95	86	99	87	98	84	98	81	99
en	100	100	100	100	100	100	100	100	98	100	99	100	97	99
fr	91	100	93	100	92	100	91	100	89	100	91	100	87	100
ja	94	99	94	100	92	100	83	100	88	100	87	100	85	100
zh	96	100	94	99	96	100	94	100	94	100	93	99	87	100
hi	78	95	72	85	7	62	5	59	11	62	5	61	2	80
AVG	92	99	91	97	79	93	76	93	78	93	76	93	73	96

Table 21: Accuracy and valid reasoning rate (%) of Qwen3-32B on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.

Qwen3-30B-A3B	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	0	0	0	0	0	100	0	0	0	0	2	2	1	98
en	99	98	96	95	98	100	97	98	99	99	99	99	97	99
fr	68	71	64	65	88	100	84	92	87	95	87	96	84	99
ja	0	0	0	0	0	100	0	0	0	0	0	0	0	99
zh	55	64	64	67	62	100	76	79	79	84	81	83	77	100
hi	0	3	0	4	0	89	0	2	0	4	0	4	0	92
AVG	37	40	37	39	41	98	43	45	44	47	45	48	43	98

Table 22: Accuracy and valid reasoning rate (%) of Qwen3-30B-A3B on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.

Qwen3-4B	2ppl		3ppl		4ppl		5ppl		6ppl		7ppl		8ppl	
	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%	acc%	valid%
ar	81	100	73	100	59	100	51	99	40	100	37	100	28	100
en	97	96	96	96	95	94	97	96	95	95	93	93	89	90
fr	78	99	83	99	81	100	77	100	79	99	80	100	74	100
ja	0	0	0	0	0	0	0	0	0	0	0	0	0	0
zh	7	9	13	30	14	31	15	45	14	45	18	45	19	52
hi	71	98	46	97	39	97	36	93	31	98	30	100	24	97
AVG	56	67	52	70	48	70	46	72	43	73	43	73	39	73

Table 23: Accuracy and valid reasoning rate (%) of Qwen3-4B on the K&K dataset across difficulty levels (2ppl to 8ppl) and input languages.