

From KMMLU-REDUX to PRO: A Professional Korean Benchmark Suite for LLM Evaluation

Anonymous ACL submission

Abstract

The development of Large Language Models (LLMs) requires robust benchmarks that encompass not only academic domains but also industrial fields to effectively evaluate their applicability in real-world scenarios. In this paper, we introduce two Korean expert-level benchmarks. **KMMLU-REDUX**, reconstructed from the existing KMMLU (Son et al., 2024a), consists of questions from the Korean National Technical Qualification exams, with critical errors removed to enhance reliability. **KMMLU-PRO** is based on Korean National Professional Licensure exams to reflect professional knowledge in Korea. Our experiments demonstrate that these benchmarks comprehensively represent industrial knowledge in Korea.

1 Introduction

As LLMs continue to achieve strong performance across a wide range of subjects (OpenAI et al., 2024b; Deepmind, 2024; DeepSeek-AI et al., 2025a; Research et al., 2025), the demand for comprehensive benchmarks has grown. MMLU (Hendrycks et al., 2021) is widely used for its broad coverage of general knowledge from elementary to college level. However, its publicly available online problems has raised concerns about reliability and potential data contamination (Gema et al., 2025; Vendrow et al., 2025; Zhao et al., 2024).

We identify similar issues in KMMLU (Son et al., 2024a), a widely used benchmark for evaluating Korean expert-level knowledge. The dataset was constructed by crawling websites that provide questions from various exams. We observe noisy samples, including problems that explicitly reveal the answer or non-existent reference, which can mislead performance evaluation. Additionally, we find evidence of contamination between the train and test splits, as well as with common web corpus such as FineWeb2 (Penedo et al., 2024).

Instead of collecting data from online sources directly, recent challenging benchmarks (Srivastava et al., 2023; Rein et al., 2024; Phan et al., 2025; Kazemi et al., 2025; Team et al., 2025b) have been constructed through problems and answers authored by human expert. Although this approach ensures high-quality, contamination-free benchmarks, it is costly to construct and maintain. The high construction cost hinders regular updates (White et al., 2025; Jain et al., 2025), leaving the benchmarks vulnerable to deprecation and potential contamination. Furthermore, existing benchmarks focus primarily on academic knowledge and evaluate models using a single overall score. They often overlook the practical applicability of models in industrial or professional contexts, specifically whether the models meet the standards for real-world professional certification.

We introduce two benchmarks to address the limitations above and reflect real-world professional knowledge. First, **KMMLU-REDUX**, a refined subset of KMMLU with 2,587 problems, is built through rigorous manual examination by authors to reduce errors and contamination within KMMLU. Furthermore, while the KMMLU contains problems from wide range of exams, even with high-school level, **KMMLU-REDUX** only selects Korean National Technical Qualification (KNTQ) exams as sources of the benchmark. The exams require applicants to have either a bachelor’s degree or at least nine years experience in industrial field, thus making the benchmark more challenging.

Second, we build **KMMLU-PRO**, a new challenging benchmark, which consists of 2,822 problems from acquisition exams for Korean National Professional Licensure (KNPL), representing highly specialized professions in Korea. We include 14 professions from diverse domains. Unlike KMMLU that crawls websites, we collect data directly from the official source of each license. After that, human annotators manually examine it to

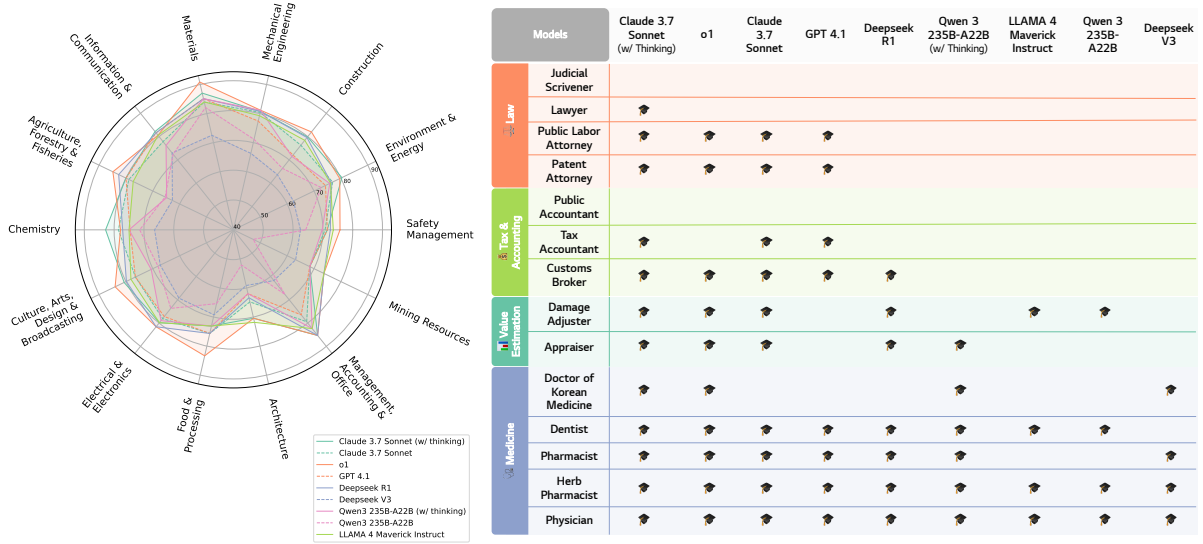


Figure 1: Model performance across industry domains on KMMLU-Redux (left) and profession acquisition status with 🎓 on KMMLU-Pro (right).

avoid noises. KMMLU-PRO only includes exams held in the most recent year and would be updated annually with the latest exam to maintain long-term reliability and prevent contamination.

We conduct extensive evaluations of various LLMs on the two benchmarks. Our benchmarks, based on real-world exams, enable aligned analysis with industrial and professional qualifications, effectively revealing the practical strengths of each model. As shown in Figure 1, KMMLU-PRO assesses which professional license a model is capable of obtaining, while KMMLU-REDUX evaluates the breadth of industrial knowledge. The results show that several state-of-the-art models perform strongly in the medicine domain, passing most licenses, yet nearly fail in law-related licenses.

Moreover, we observe the significant performance gaps between our datasets and merely translated datasets like MMMLU (OpenAI, 2024), especially in domains such as laws, where in-depth knowledge of specific countries is required (see Section 6.1). We argue that these findings underscore the practicality of our benchmarks for assessing the capabilities of models in professional fields within Korea.

To summarize, our contributions are as follows:

- We improve the previous benchmark, KMMLU, to construct the refined and compact version of the benchmark, **KMMLU-REDUX**, by correcting various errors.
- We introduce **KMMLU-PRO**, a new bench-

mark designed to evaluate high-level professional knowledge in Korea. By imitating real-world license acquisition systems, KMMLU-PRO assesses the industrial practicality across various professions in Korea.

- We comprehensively analyze the results of two benchmarks, highlighting the importance of benchmarks specialized in Korea-specific professional knowledge.

2 KMMLU-REDUX

We first revisit KMMLU (Son et al., 2024a) to examine its quality. Building on these insights, we construct KMMLU-REDUX, a cleaned and compact version of the KMMLU. We carefully denoise against the KMMLU and increase difficulty.

2.1 Revisiting KMMLU

KMMLU plays a significant role in the NLP community as a de facto standard for evaluating LLMs on Korea-specific expert knowledge (Yoo et al., 2024; Research et al., 2024; Team et al., 2025a). The dataset was constructed by crawling websites¹ where various exam questions are uploaded by people online, spanning from high school to professional qualification exams. Upon closer inspection, we identify several noises and limitations, which can be categorized into three types: 1) duplication issues, 2) dataset errors, and 3) contamination.

Duplication Issue As the KMMLU crawls problems from hundreds of exams in Korea, we observe

¹<https://www.kinz.kr/>

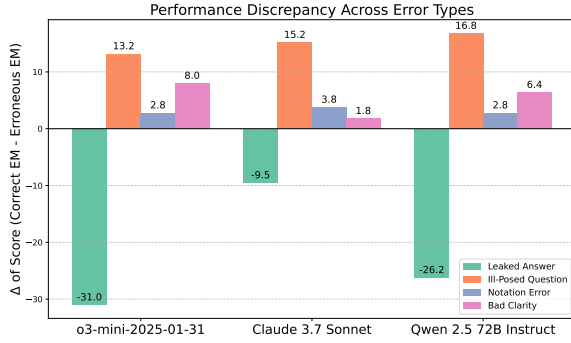


Figure 2: Performance differences in LLMs on the erroneous versus correct dataset. Leaked answer errors tend to overestimate model capabilities, while three other error types hinder LLMs to correctly predict true answers.

multiple duplicated questions across related exams. Notably, duplicated samples occur not only within the test set but also between the training and test sets. Using the Longest Common Sequence (LCS) algorithm to investigate overlaps, we could find 5.36% duplication within the test set and a 5.46% contamination between train and test set.

Dataset Errors Following Gema et al. (2025), we investigate the extent to which various error types appear in the KMMLU test set and their potential impact on LLM performance. We identify four representative error types: leaked answers, ill-posed questions, poor clarity, and notation errors. We then annotate the entire test set using GPT-4o² (OpenAI et al., 2024a). In total, we find that 7.66% of the data contains one of the errors mentioned above. We describe the details of the investigation of dataset errors in Appendix A.

In Figure 2, we observe significant discrepancies in the LLM’s performance between erroneous and correct samples. Notably, the performance of instances with leaked answers drops significantly when the LLMs are assessed on the clean dataset. Conversely, all models’ scores increase on ill-posed questions, underscoring their inability to identify the correct answer for poorly formulated questions.

Contamination The KMMLU was primarily sourced online, making them highly susceptible to contamination from web-crawled training corpora. When applying n-gram contamination detection (Lambert et al., 2025; Grattafiori et al., 2024) to the Korean subset of FineWeb2 (Penedo et al., 2024) and the KMMLU dataset, 1.88% of the data were flagged as contaminated.

²gpt-4o-2024-11-20

2.2 Dataset Construction

KMMLU consists of approximately 35k examples, making evaluation resource-intensive. To reduce this burden, we first restrict the scope to a subset of high-difficulty exams (Section 2.2.1). Furthermore, to ensure the reliability of the benchmark, we manually conduct a thorough examination of the dataset to eliminate errors (Section 2.2.2).

2.2.1 Filtering Non-Challenging Problems

Since KMMLU includes a variety of exams in Korea, spanning from high school to professional certification exams, we filter out the easier exams to build a more challenging benchmark. Specifically, we choose Korean National Technical Qualification (KNTQ) exams, which are primarily designed to assess practical technical competencies required in industrial field. The qualifications require applicants to have either a bachelor’s degree or at least nine years of professional experience. We adopt a collection of 100 KNTQ exams across the 14 domains in total. We only include the most recent exam for each qualification, thereby avoiding outdated knowledge being evaluated³.

To further discriminate simple problems, we leverage the performances of LLMs on the data (Wang et al., 2024; Zellers et al., 2019; Lee et al., 2023). By employing seven smaller LLMs⁴, we mark a data as *easy* if four or more models correctly predict its answer. Through this process, we remove 38.6% of the dataset.

2.2.2 Denoising

To remove noises, we follow the processes described in Section 2.1. Specifically, we first manually review the dataset to minimize errors. Next, we perform decontamination to prevent potential data leakage from pre-training corpora. Additionally, we detect inner duplication and against the training and test sets in KMMLU, thus finally remove all duplicates.

³We have compiled a list of all KNTQs exams alongside their most recent exam dates in Appendix B. Our aim is to make these available to LLM researchers and developers to help prevent data contamination.

⁴Llama 3.2 3B (Meta, 2024c), Qwen 2.5 3B (Qwen et al., 2025), Gemma 3 4B IT (Team, 2025a), Kanana Nano 2.1B Instruct (Team et al., 2025a), EXAONE 3.5 2.4B (Research et al., 2024), DeepSeek-R1-Distill-Qwen-1.5B (DeepSeek-AI et al., 2025a) EXAONE Deep 2.4B (Research et al., 2025), and Ko-R1-7B-v2.1 (OneLineAI).

Domain	Names of KNPLs	U.S. Equivalent	# of Instances
Law	Certified Judicial Scrivener	Paralegal, Legal Document Assistant, or Notary Public (no direct equivalent)	198
	Lawyer (Kim et al., 2024b)	Attorney-at-Law	150
	Certified Public Labor Attorney	Labor & Employment Lawyer (requires J.D. and bar admission; no separate certification in the U.S.)	239
	Certified Patent Attorney	Patent Attorney (JD + USPTO registration required)	109
Tax & Accounting	Certified Public Accountant (CPA)	Certified Public Accountant (CPA) – Exact Equivalent	208
	Certified Tax Accountant	Enrolled Agent (IRS) or CPA with Tax Specialization	238
	Certified Customs Broker	U.S. Customs Broker (licensed by U.S. Customs and Border Protection - CBP)	159
Value Estimation	Certified Damage Adjuster (CDA)	Claims Adjuster / Insurance Adjuster (state-licensed)	120
	Certified Appraiser	Certified Real Estate Appraiser (licensed at the state level)	196
Medicine	Doctor of Korean Medicine	Licensed Acupuncturist (L.Ac.) or Doctor of Acupuncture and Oriental Medicine (D.A.O.M.)	288
	Dentist (Kweon et al., 2024)	Doctor of Dental Surgery (D.D.S.) / Doctor of Dental Medicine (D.M.D.)	252
	Pharmacist (Kweon et al., 2024)	Doctor of Pharmacy (Pharm.D.)	271
	Herb Pharmacist	Herbalist (non-licensed or CAM-certified depending on state)	244
	Physician (Kweon et al., 2024)	Medical Doctor (M.D./D.O.)	150
Total			2822

Table 1: The list of National Professional Licenses (NPLs) used for KMMLU-PRO and their corresponding statistics. The names of NPLs are translated from those in Korea and we also report equivalent licences in U.S. We use KorMedMCQA (Kweon et al., 2024) for three licenses in the Medical category, and KBL (Kim et al., 2024b) for the bar exam of lawyer.

2.2.3 Final Statistics

For KMMLU-REDUX, we have collected 2,587 problems from 100 KNTQ exams. Among these, 596 problems are from exams that require over nine years of professional experience to acquire the qualification. To categorize the dataset into 14 domains, we follow the Korean Standard Industrial Classification (KSIC) published by Statistics Korea⁵ as the qualification system is primarily designed to align with industrial fields. Figure 6 in Appendix B illustrates the distribution of domain of KMMLU-REDUX.

3 KMMLU-PRO

A major challenge in building benchmarks from online sources is data contamination (Zhao et al., 2024; Jain et al., 2025; Roberts et al., 2024). While some studies address this by having experts manually create problems (Srivastava et al., 2023; Rein et al., 2024; Phan et al., 2025; Kazemi et al., 2025; Team et al., 2025b), this approach is costly and time-consuming. As an alternative, recent work explores periodic releases of fresh subsets. (White et al., 2025; Jain et al., 2025).

Motivated by these approaches, we focus on the Korean National Professional Licensure (KNPL) exams. These are high-stakes exams administered annually that pose a significant challenge. Unlike benchmarks crafted by a small set of experts (Rein et al., 2024; Phan et al., 2025), our approach leverages the well-established curricula of professional licensing systems, designed to assess real-world professional knowledge.

⁵<https://www.kostat.go.kr/>

3.1 Korean National Professional Licensures

We choose KNPL exams for main source of KMMLU-PRO. KNPL exams target high-level professionals, such as lawyers, accountants, or physicians, requiring advanced knowledge, critical reasoning, and ethical judgment. Among them, we select 14 KNPLs representing highly specialized and regulated professions in Korea (See Table 1 for the list of KNPLs and their equivalent licensure in U.S.). These licenses are legally mandated credentials required to practice in their respective domains. As such, they serve as institutionalized gateways to high-status occupations with significant entry barriers. Our evaluation simulates real-world assessment standards by incorporating official exam pass criteria, aligning model performance with human standards.

3.2 Dataset Collection and Annotation

In Korea, the government releases and manages the questions for KNPL acquisition exams. We directly download the PDF files from the government’s websites for each license and use GPT-4o (OpenAI et al., 2024a) for OCR parsing. As our dataset sources from the official PDFs, we can enhance the quality of the dataset, avoiding potential errors when collecting from online text (Team et al., 2025b). Since GPT-4o has difficulty handling tables and low-resolution PDFs, we employ human annotators to review parsed questions. When a problem contains an image, the annotators convert it into text that conveys the same meaning of the image, if possible⁶. Notably, our process remains relatively

⁶For further details, please see Appendix C

cost-efficient as it requires human annotators solely for error reviewing tasks, not full annotation.

We follow previous works, releasing a new set of questions periodically (White et al., 2025; Jain et al., 2025). Since the exams are conducted annually, we commit to collect and release questions from the exams held just before the current year.

3.3 Decontamination

We also adopt the same process outlined in Section 2.2.2 to ensure KMMLU-PRO is free from contamination. When conducting n-gram match between KMMLU-PRO, FineWeb2, and the training and validation sets of KMMLU, we did not find any contaminated examples. As a result, we can retain all 2,822 data points in the KMMLU-PRO, maintaining its contamination-free integrity.

4 Experiments

We select a diverse set of baseline models varying in size, multilingual capability, and reasoning ability. By default, we apply a zero-shot Chain-of-Thought (CoT) (Wei et al., 2022) prompt written in Korean. However, we observe that some models perform worse when prompted in Korean. For those models, we report results using the English prompt instead⁷. Our evaluation implementation is based on OpenAI’s simple-evals repository⁸. We use greedy decoding for non-reasoning models, while for reasoning models, we follow the settings in DeepSeek-AI et al. (2025a); Research et al. (2025), using a temperature of 0.6 and top-p (Holtzman et al., 2020) of 0.95. For more details, see Appendix D.

Metrics In addition to accuracy, our primary evaluation metric, we also report the number of licenses each LLM obtains in KMMLU-PRO. To better align with human evaluation standards, our procedure is designed to mirror the official license acquisition criteria. However, for licenses that rely heavily on image-based questions, full replication is not possible. See Appendix D.3 for details.

5 Results

5.1 Main Results

Table 2 presents the overall performance on KMMLU-REDUX and KMMLU-PRO. The o1 model achieves the highest average accuracy

(79.55), followed by Claude 3.7 with Thinking (78.49). Among open-weight models, DeepSeek’s R1 even outperforms many closed models. Models equipped with reasoning capabilities consistently perform better than their non-reasoning models, such as the Qwen3 series.

Beyond accuracy, we evaluate each model’s ability to obtain professional licenses under the official passing criteria for KMMLU-PRO. Specifically, in most licenses, a model must score at least 40% in each subject and achieve an overall average of 60% to pass. Claude 3.7 with Thinking obtains 12 out of 14 KNPL licenses, the highest among all evaluated models. In contrast, although the o1 model achieves higher accuracy, it qualifies for fewer licenses than Claude 3.7 with Thinking. This highlights the importance of balanced competence across subjects, as required in real-world certification exams.

5.2 Performance Across Industrial Domains in KMMLU-REDUX

Through KMMLU-REDUX, we assess the technical competencies of models across diverse industrial fields. As shown in Figure 1 (left), models with reasoning capabilities consistently outperform their non-reasoning counterparts across all domains. The improvement is particularly notable in Agriculture, Food & Processing, and Architecture. Despite these gains, all models continue to struggle in domains such as Mining Resources and Architecture, highlighting persistent challenges in underrepresented or highly specialized fields. The full results for all models across 14 domains are presented in Appendix E.3.

5.3 Professional License Acquisition Status on KMMLU-PRO

We provide a breakdown of the KMMLU-PRO results by analyzing which licenses are most frequently obtained by LLMs. Figure 1 (right) shows the pass rates of LLMs across 14 KNPLs. While many models obtain licenses in the medicine domain, most fail in Law and Tax & Accounting, with only DeepSeek R1 passing the Customs Broker exam among open-weight models. Notably, no models pass the Judicial Scrivener or Public Accountant exams. This trend becomes even more evident in the full results across a broader range of LLMs; the only licenses that relatively smaller models (<20B) are able to pass are in the medicine domain (see Table 8 in Appendix E.4).

Moreover, many LLMs fail to obtain licenses

⁷For further details, please see Section 6.3.

⁸<https://github.com/openai/simple-evals>

	KMMLU-Redux Acc	KMMLU-Pro Acc	# of passed KNPLs	Avg. Acc (micro)
Open-weight Models				
Aya Expanse 32B (Dang et al., 2024)	33.05	31.26	0/14	32.12
Gemma 3 12B IT (Team, 2025a)	46.70	45.82	2/14	46.24
Phi-4 (14B) (Abdin et al., 2024)	49.75	45.32	1/14	47.44
EXAONE 3.5 32B Instruct (Research et al., 2024)	49.40	46.71	2/14	48.00
Mistral Small 3.1 Instruct (24B) (Mistral, 2025)	52.92	49.49	3/14	51.13
Gemma 3 27B IT (Team, 2025a)	54.04	51.03	2/14	52.47
Llama 3.3 70B Instruct (Meta, 2024a)	56.17	53.24	3/14	54.64
Qwen3-14B (Yang et al., 2025)	57.25	53.02	3/14	55.04
EXAONE Deep 32B (Research et al., 2025)	58.33	52.33	1/14	55.20
Qwen3-30B-A3B (Yang et al., 2025)	58.41	52.33	3/14	55.24
C4AI Command A (111B) (Cohere, 2025)	62.93	57.48	3/14	60.07
Qwen3-32B (Yang et al., 2025)	64.98	58.86	3/14	61.79
Llama-4-Scout-17B-16E-Instruct (Meta, 2025)	67.49	58.14	4/14	62.61
Qwen3-30B-A3B (w/ thinking) (Yang et al., 2025)	65.25	60.52	3/14	62.78
Qwen3-14B (w/ thinking) (Yang et al., 2025)	65.71	60.18	2/14	62.82
DeepSeek V3 (671B) (DeepSeek-AI et al., 2025b)	65.64	60.77	4/14	63.10
Qwen3-32B (w/ thinking) (Yang et al., 2025)	68.77	61.14	3/14	64.79
QwQ 32B (Team, 2025b)	67.34	63.94	5/14	65.57
Qwen3-235B-A22B (Yang et al., 2025)	69.54	62.12	4/14	65.67
Qwen3-235B-A22B (w/ thinking) (Yang et al., 2025)	74.49	68.22	6/14	71.22
Llama-4-Maverick-17B-128E-Instruct (Meta, 2025)	77.58	68.10	4/14	72.63
DeepSeek R1 (671B) (DeepSeek-AI et al., 2025a)	78.51	71.33	7/14	74.76
Closed Models				
GPT-4.1 mini (2025-04-14) (OpenAI, 2025a)	67.03	62.18	4/14	64.50
o3-mini (2025-01-31) (OpenAI, 2025c)	67.84	62.05	3/14	64.82
Grok-3-mini-beta (xAI, 2025)	71.47	65.08	5/14	68.14
Grok-3-beta (xAI, 2025)	72.90	68.37	7/14	70.54
o4-mini (2025-04-16) (OpenAI, 2025b)	75.80	69.65	6/14	72.59
GPT-4.1 (2025-04-14) (OpenAI, 2025a)	75.86	72.99	10/14	74.36
Claude 3.7 Sonnet (Anthropic, 2025)	76.88	74.52	10/14	75.65
o3 (OpenAI, 2025b)	79.92	73.60	9/14	76.62
Claude 3.7 Sonnet (w/ thinking) (Anthropic, 2025)	79.36	77.70	12/14	78.49
o1 (OpenAI et al., 2024b)	81.14	78.09	10/14	79.55

Table 2: The main evaluation results of KMMLU-REDUX and KMMLU-PRO benchmarks on various LLMs. The gray-shaded models stand for reasoning models. The results of models with size < 10B are presented in Table 6 in Appendix E.2.

even when scoring above 60%; for example, o3-mini, Qwen3-235B-A22B, and Llama-4-Maverick score over 85% on the Pharmacist exam but still fail to qualify due to not meeting the threshold of law-related subject in the exam. These cases highlight the difficulty of acquiring region-specific domain knowledge, particularly in legal subjects governed by Korean law.

6 Analysis

6.1 The Importance of Locally Adapted Benchmarks

To highlight the importance of evaluation grounded in local context (Plaza et al., 2024; Singh et al., 2025), we compare category-level performance be-

tween benchmarks translated from English and KMMLU-PRO. Specifically, we focus on subjects related to law, accounting, and medicine,⁹ selecting relevant subjects from the Korean subset of MMMLU (OpenAI, 2024), as well as KMMLU-PRO.

As shown in Figure 3, the performance gap is relatively small in categories such as medicine, where domain knowledge is largely consistent across countries and cultures. In contrast, categories such as law, where substantial differences in content are expected, show a significantly larger gap. This sug-

⁹{*professional_law*, *jurisprudence*, *international_law*} for Law. {*professional_accounting*} for Accounting. {*professional_medicine*, *clinical_knowledge*, *college_medicine*, *medical_genetics*, *anatomy*} for Medicine.

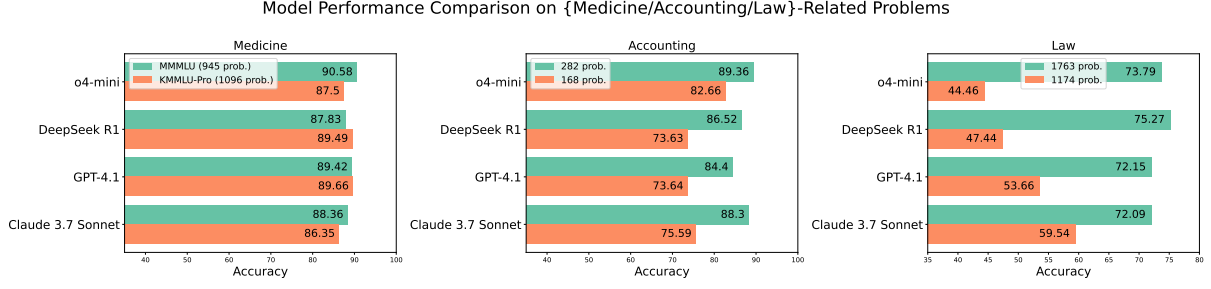


Figure 3: Performance of four LLMs on {Medical(left), Accounting(center), Law(right)}-relevant subsets from the MMMLU (Korean) (OpenAI, 2024) and KMMLU-PRO. While the discrepancies in scores are narrow in the medicine domain, they are wider in law-related problems, emphasizing the need for datasets that reflecting real professional knowledge in Korea.

gests that MMMLU, which relies on direct translation of law questions based on U.S. standards, cannot adequately represent knowledge of Korean law. These findings highlight the importance of our dataset, which reflects authentic professional knowledge specific to the Korean context.

6.2 Impact of Reasoning Budget

Recent studies have shown that increasing reasoning efforts can enhance model performance (Muenighoff et al., 2025; Anthropic, 2025; Qwen et al., 2025). To examine how reasoning *budget* affects performance on individual licenses in KMMLU-PRO, we conduct an experiment varying the number of tokens allocated to the reasoning path. We adopt Qwen3-32B (Yang et al., 2025) and Claude 3.7 Sonnet (Anthropic, 2025), for each budget $b \in B$, we generate $n = 4$ for Qwen3 and $n = 2$ for Claude.

As shown in Figure 4, we observe a positive correlation between the reasoning budget and the overall score of KMMLU-PRO for both models. However, this trend does not hold uniformly across all licenses. For example, we observe little to no improvement on the Judicial Scrivener and Herb Pharmacist licenses for both models, indicating that increased reasoning budget does not result in performance gains for certain licenses.

6.3 Impact of Prompt Language

Prompt language can significantly influence model behavior, raising concerns about consistency in multilingual settings (Wang et al., 2025; Zhang et al., 2025; Lai and Nissim, 2024). Since all exam questions in our datasets are written in Korean, using Korean prompts is a natural choice. However, we observe that some models perform worse when prompted in Korean. Table 3 presents the performance difference between English and Korean prompts. The Llama-4 model series exhibits

	KMMLU-REDUX			KMMLU-PRO		
	English	Korean	diff (%)	English	Korean	diff (%)
Qwen3-32B (w/ thinking)	68.77	69.08	+0.5%	61.14	60.66	-0.8%
o4-mini (2025-04-16)	75.80	76.17	+0.5%	69.65	69.10	-0.8%
Qwen3-235B-A22B (w/ thinking)	74.49	75.11	+0.8%	68.22	66.98	-1.8%
Grok-3-mini-beta	71.47	70.85	-0.9%	65.08	64.89	-0.3%
Qwen3-14B (w/ thinking)	65.71	65.40	-0.5%	60.18	59.48	-1.2%
EXAONE Deep 32B	58.33	56.17	-3.7%	52.33	52.19	-0.3%
DeepSeek R1 (671B)	78.51	75.38	-4.0%	71.33	70.62	-1.0%
QwQ 32B	67.34	62.66	-6.9%	63.94	59.95	-6.2%
EXAONE Deep 7.8B	44.99	40.82	-9.3%	41.53	38.98	-6.1%
Llama-4-Maverick-17B-128E-Instruct	77.58	72.52	-7.0%	68.10	57.15	-16.1%
Llama-4-Scout-17B-16E-Instruct	67.49	45.03	-33.3%	58.14	28.74	-50.6%

Table 3: Comparison results between English and Korean prompts of models whose main results are evaluated on English prompts. The *diff* values are the relative difference in scores between two prompts. The specific prompts used are detailed in Appendix D.1.

the most substantial drop in performance, while closed models such as Grok-3-mini-beta and o4-mini show minimal change.

7 Related Works

Reliability Issues of Benchmarks Recent works (Gema et al., 2025; Vendrow et al., 2025) have raised concerns about the reliability of LLM benchmarks due to dataset noise and contamination. MMLU-Redux (Gema et al., 2025) improved evaluation quality through systematic human reannotation, while GSM8K-Platinum (Vendrow et al., 2025) refined arithmetic benchmarks via automated and manual error detection. MMLU-CF (Zhao et al., 2024) prevents both unintentional and malicious contamination via sourcing diverse domains and question rewriting. LiveBench (White et al., 2025) and LiveCodeBench (Jain et al., 2025) adopted dynamic evaluation protocols with temporal cutoffs to prevent future leakage.

Professional Benchmark With the rapid advancement of LLMs, more challenging benchmarks have become essential. GPQA (Rein et al., 2024) and SuperGPQA (Team et al., 2025b) assess graduate-level knowledge; MMLU-Pro (Wang

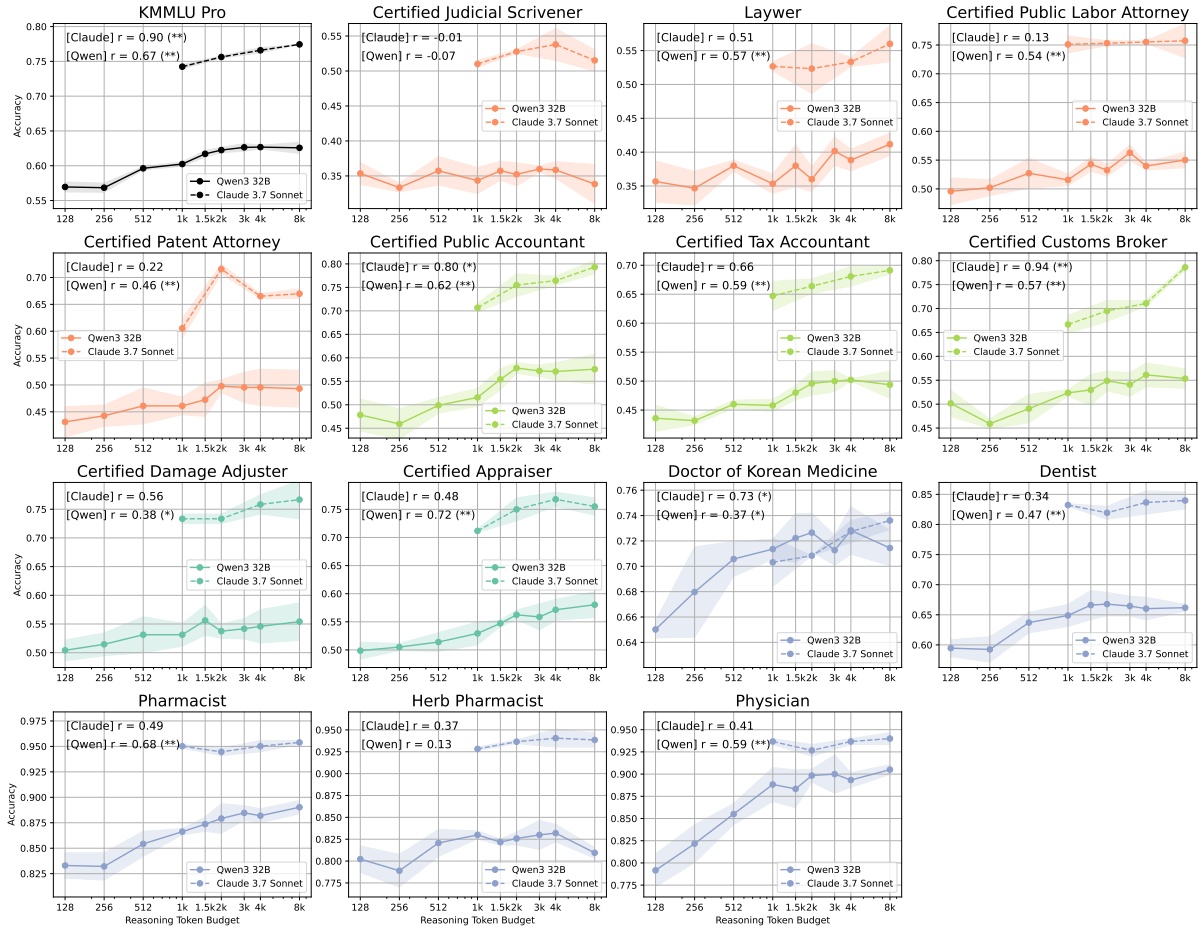


Figure 4: Reasoning budget results of Qwen3-32B and Claude 3.7 Sonnet on KMMLU-PRO. r indicates the Pearson Correlation coefficient value between the reasoning token budget and the accuracy. The responses are sampled multiple times for each thinking budget setting; $n = 4$ and $n = 2$, respectively, for Qwen and Claude.

et al., 2024) extends MMLU by increasing the share of college-level questions and expanding answer choices. Humanity’s Last Exam (Phan et al., 2025) introduces a frontier benchmark composed of manually authored, research-level questions.

Korean Benchmark While prior benchmarks focus primarily on English, recent efforts have produced Korean-specific evaluations (Son et al., 2024b; Kim et al., 2024a). Some rely on translated datasets (Park et al., 2024; Kim et al., 2025; OpenAI, 2024; Singh et al., 2024), often with human post-editing, but these lack regional context, institutional norms, and domain-specific fluency. In contrast, native Korean benchmarks such as KMMLU (Son et al., 2024a), KorMedMCQA (Kweon et al., 2024), and KBL (Kim et al., 2024b) address cultural and linguistic specificity. However, KMMLU (Son et al., 2024a) suffers from quality issues, including leaked answers, and KorMedMCQA (Kweon et al., 2024) and KBL (Kim

et al., 2024b) are limited to narrow domains.

In this work, we introduce KMMLU-REDUX and KMMLU-PRO, two contamination-free, industry grade benchmarks, providing a practical assessment of LLM capabilities in Korean industries.

8 Conclusion

We present two benchmarks constructed from real-world professional licensing exams, designed to reflect industrial domain knowledge and practical application standards. To ensure reliability, we identify and eliminate various sources of errors. Through extensive experiments, we evaluate the professional knowledge capabilities of LLMs across a wide range of domains. Our analysis further identifies key factors that influence performance, including region-specific knowledge, reasoning budget, and prompt language. We hope this work provides a foundation for more rigorous evaluation and continued advancement of real-world competence in language models.

Limitations

Our benchmarks are limited to text-only and multiple-choice questions for text-only LLMs. It restricts its coverage of real-world licensure exams. Many real-world professional qualification exams include non-textual modalities or require constructed responses such as essay. Our benchmark cannot fully assess all aspects of professional competence or reasoning required in such exams. Expanding to multimodal inputs and open-ended question formats is an important direction for future work.

Ethical Statements

All data used in our benchmarks are either publicly available or collected from official licensing materials released by government or professional institutions. For quality control, we hired human annotators to review parsed questions from PDF; they were compensated over the minimum wage in Korea. Our benchmarks would be released under CC-BY-NC-ND 4.0 license.

References

Marah Abidin, Jyoti Aneja, Harkirat Behl, Sébastien Bubeck, Ronen Eldan, Suriya Gunasekar, Michael Harrison, Russell J. Hewett, Mojan Javaheripi, Piero Kauffmann, James R. Lee, Yin Tat Lee, Yuanzhi Li, Weishung Liu, Caio C. T. Mendes, Anh Nguyen, Eric Price, Gustavo de Rosa, Olli Saarikivi, Adil Salim, Shital Shah, Xin Wang, Rachel Ward, Yue Wu, Dingli Yu, Cyril Zhang, and Yi Zhang. 2024. [Phi-4 technical report](#). *Preprint*, arXiv:2412.08905.

Anthropic. 2025. [Claude 3.7 sonnet and claude code](#).

Cohere. 2025. [Model card for c4ai command a](#).

John Dang, Shivalika Singh, Daniel D’souza, Arash Ahmadian, Alejandro Salamanca, Madeline Smith, Aidan Peppin, Sungjin Hong, Manoj Govindassamy, Terrence Zhao, Sandra Kublik, Meor Amer, Viraat Aryabumi, Jon Ander Campos, Yi-Chern Tan, Tom Kocmi, Florian Strub, Nathan Grinsztajn, Yannis Flet-Berliac, Acyr Locatelli, Hangyu Lin, Dwarak Talupuru, Bharat Venkitesh, David Cairuz, Bowen Yang, Tim Chung, Wei-Yin Ko, Sylvie Shang Shi, Amir Shukayev, Sammie Bae, Aleksandra Piktus, Roman Castagné, Felipe Cruz-Salinas, Eddie Kim, Lucas Crawhall-Stein, Adrien Morisot, Sudip Roy, Phil Blunsom, Ivan Zhang, Aidan Gomez, Nick Frosst, Marzieh Fadaee, Beyza Ermiş, Ahmet Üstün, and Sara Hooker. 2024. [Aya expand: Combining research breakthroughs for a new multilingual frontier](#). *Preprint*, arXiv:2412.04261.

Google Deepmind. 2024. [Introducing gemini 2.0: our new ai model for the agentic era](#).

DeepSeek-AI, Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, Xiaokang Zhang, Xingkai Yu, Yu Wu, Z. F. Wu, Zhibin Gou, Zhihong Shao, Zhuoshu Li, Ziyi Gao, Aixin Liu, Bing Xue, Bingxuan Wang, Bochao Wu, Bei Feng, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Qu, Hui Li, Jianzhong Guo, Jia Shi Li, Jiawei Wang, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, J. L. Cai, Jiaqi Ni, Jian Liang, Jin Chen, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Liang Zhao, Litong Wang, Liyue Zhang, Lei Xu, Leyi Xia, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Meng Li, Miaojun Wang, Mingming Li, Ning Tian, Panpan Huang, Peng Zhang, Qiancheng Wang, Qinyu Chen, Qiushi Du, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, R. J. Chen, R. L. Jin, Ruyi Chen, Shanghao Lu, Shangyan Zhou, Shanhuang Chen, Shengfeng Ye, Shiyu Wang, Shuiping Yu, Shunfeng Zhou, Shuting Pan, S. S. Li, Shuang Zhou, Shaoqing Wu, Shengfeng Ye, Tao Yun, Tian Pei, Tianyu Sun, T. Wang, Wangding Zeng, Wanbiao Zhao, Wen Liu, Wenfeng Liang, Wenjun Gao, Wenqin Yu, Wentao Zhang, W. L. Xiao, Wei An, Xiaodong Liu, Xiaohan Wang, Xiaokang Chen, Xiaotao Nie, Xin Cheng, Xin Liu, Xin Xie, Xingchao Liu, Xinyu Yang, Xinyuan Li, Xuecheng Su, Xuheng Lin, X. Q. Li, Xiangyue Jin, Xiaojin Shen, Xiaosha Chen, Xiaowen Sun, Xiaoxiang Wang, Xinnan Song, Xinyi Zhou, Xianzu Wang, Xinxia Shan, Y. K. Li, Y. Q. Wang, Y. X. Wei, Yang Zhang, Yanhong Xu, Yao Li, Yao Zhao, Yaofeng Sun, Yaohui Wang, Yi Yu, Yichao Zhang, Yifan Shi, Yiliang Xiong, Ying He, Yishi Piao, Yisong Wang, Yixuan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo, Yuan Ou, Yuduan Wang, Yue Gong, Yuheng Zou, Yujia He, Yunfan Xiong, Yuxiang Luo, Yuxiang You, Yuxuan Liu, Yuyang Zhou, Y. X. Zhu, Yanhong Xu, Yanping Huang, Yaohui Li, Yi Zheng, Yuchen Zhu, Yunxian Ma, Ying Tang, Yukun Zha, Yuting Yan, Z. Z. Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhean Xu, Zhenda Xie, Zhengyan Zhang, Zhewen Hao, Zhicheng Ma, Zhigang Yan, Zhiyu Wu, Zihui Gu, Zijia Zhu, Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song, Zizheng Pan, Zhen Huang, Zhipeng Xu, Zhongyu Zhang, and Zhen Zhang. 2025a. [Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning](#). *Preprint*, arXiv:2501.12948.

DeepSeek-AI, Aixin Liu, Bei Feng, Bing Xue, Bingxuan Wang, Bochao Wu, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Daya Guo, Dejian Yang, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Haowei

616	Zhang, Honghui Ding, Huajian Xin, Huazuo Gao,	Charlotte Caucheteux, Chaya Nayak, Chloe Bi,	678
617	Hui Li, Hui Qu, J. L. Cai, Jian Liang, Jianzhong Guo,	Chris Marra, Chris McConnell, Christian Keller,	679
618	Jiaqi Ni, Jiashi Li, Jiawei Wang, Jin Chen, Jingchang	Christophe Touret, Chunyang Wu, Corinne Wong,	680
619	Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, Junx-	Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Al-	681
620	iao Song, Kai Dong, Kai Hu, Kaige Gao, Kang	lonsius, Daniel Song, Danielle Pintz, Danny Livshits,	682
621	Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong	Danny Wyatt, David Esiobu, Dhruv Choudhary,	683
622	Zhang, Lei Xu, Leyi Xia, Liang Zhao, Litong Wang,	Dhruv Mahajan, Diego Garcia-Olano, Diego Perino,	684
623	Liyue Zhang, Meng Li, Miaojun Wang, Mingchuan	Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy,	685
624	Zhang, Minghua Zhang, Minghui Tang, Mingming	Elina Lobanova, Emily Dinan, Eric Michael Smith,	686
625	Li, Ning Tian, Panpan Huang, Peiyi Wang, Peng	Filip Radenovic, Francisco Guzmán, Frank Zhang,	687
626	Zhang, Qiancheng Wang, Qihao Zhu, Qinyu Chen,	Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis An-	688
627	Qiushi Du, R. J. Chen, R. L. Jin, Ruiqi Ge, Ruisong	derson, Govind Thattai, Graeme Nail, Gregoire Mi-	689
628	Zhang, Ruizhe Pan, Runji Wang, Runxin Xu, Ruoyu	alon, Guan Pang, Guillem Cucurell, Hailey Nguyen,	690
629	Zhang, Ruyi Chen, S. S. Li, Shanghao Lu, Shangyan	Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan	691
630	Zhou, Shanhuang Chen, Shaoqing Wu, Shengfeng	Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Is-	692
631	Ye, Shengfeng Ye, Shirong Ma, Shiyu Wang, Shuang	han Misra, Ivan Evtimov, Jack Zhang, Jade Copet,	693
632	Zhou, Shuiping Yu, Shunfeng Zhou, Shuting Pan,	Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park,	694
633	T. Wang, Tao Yun, Tian Pei, Tianyu Sun, W. L. Xiao,	Jay Mahadeokar, Jeet Shah, Jelmer van der Linde,	695
634	Wangding Zeng, Wanjia Zhao, Wei An, Wen Liu,	Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu,	696
635	Wenfeng Liang, Wenjun Guo, Wenqin Yu, Wentao	Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang,	697
636	Zhang, X. Q. Li, Xiangyue Jin, Xianzu Wang, Xiao	Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park,	698
637	Bi, Xiaodong Liu, Xiaohan Wang, Xiaojin Shen, Xi-	Joseph Rocca, Joshua Johnstun, Joshua Saxe, Jun-	699
638	aokang Chen, Xiaokang Zhang, Xiaosha Chen, Xiao-	teng Jia, Kalyan Vasuden Alwala, Karthik Prasad,	700
639	tao Nie, Xiaowen Sun, Xiaoxiang Wang, Xin Cheng,	Kartikkeya Upasani, Kate Plawiak, Ke Li, Kenneth	701
640	Xin Liu, Xin Xie, Xingchao Liu, Xingkai Yu, Xin-	Heafield, Kevin Stone, Khalid El-Arini, Krithika Iyer,	702
641	nan Song, Xinxia Shan, Xinyi Zhou, Xinyu Yang,	Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Kushal	703
642	Xinyuan Li, Xuecheng Su, Xuheng Lin, Y. K. Li,	Lakhotia, Lauren Rantala-Yearly, Laurens van der	704
643	Y. Q. Wang, Y. X. Wei, Y. X. Zhu, Yang Zhang, Yan-	Maaten, Lawrence Chen, Liang Tan, Liz Jenkins,	705
644	hong Xu, Yanhong Xu, Yanping Huang, Yao Li, Yao	Louis Martin, Lovish Madaan, Lubo Malo, Lukas	706
645	Zhao, Yaofeng Sun, Yaohui Li, Yaohui Wang, Yi Yu,	Blecher, Lukas Landzaat, Luke de Oliveira, Madeline	707
646	Yi Zheng, Yichao Zhang, Yifan Shi, Yiliang Xiong,	Muzzi, Mahesh Pasupuleti, Mannat Singh, Manohar	708
647	Ying He, Ying Tang, Yishi Piao, Yisong Wang, Yix-	Paluri, Marcin Kardas, Maria Tsimpoukelli, Mathew	709
648	uan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo,	Oldham, Mathieu Rita, Maya Pavlova, Melanie Kam-	710
649	Yu Wu, Yuan Ou, Yuchen Zhu, Yudian Wang, Yue	badur, Mike Lewis, Min Si, Mitesh Kumar Singh,	711
650	Gong, Yuheng Zou, Yujia He, Yukun Zha, Yunfan	Mona Hassan, Naman Goyal, Narjes Torabi, Niko-	712
651	Xiong, Yunxian Ma, Yuting Yan, Yuxiang Luo, Yuxi-	lay Bashlykov, Nikolay Bogoychev, Niladri Chatterji,	713
652	ang You, Yuxuan Liu, Yuyang Zhou, Z. F. Wu, Z. Z.	Ning Zhang, Olivier Duchenne, Onur Çelebi, Patrick	714
653	Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhean Xu,	Alrassy, Pengchuan Zhang, Pengwei Li, Petar Va-	715
654	Zhen Huang, Zhen Zhang, Zhenda Xie, Zhengyan	sic, Peter Weng, Prajjwal Bhargava, Pratik Dubal,	716
655	Zhang, Zhewen Hao, Zhibin Gou, Zhicheng Ma, Zhi-	Praveen Krishnan, Punit Singh Koura, Puxin Xu,	717
656	gang Yan, Zhihong Shao, Zhipeng Xu, Zhiyu Wu,	Qing He, Qingxiao Dong, Ragavan Srinivasan, Raj	718
657	Zhongyu Zhang, Zhuoshu Li, Zihui Gu, Zijia Zhu,	Ganapathy, Ramon Calderer, Ricardo Silveira Cabral,	719
658	Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song, Ziyi	Robert Stojnic, Roberta Raileanu, Rohan Maheswari,	720
659	Gao, and Zizheng Pan. 2025b. Deepseek-v3 techni-	Rohit Girdhar, Rohit Patel, Romain Sauvestre, Ron-	721
660	cal report . <i>Preprint</i> , arXiv:2412.19437.	nie Polidoro, Roshan Sumbaly, Ross Taylor, Ruan	722
661	Aryo Pradipta Gema, Joshua Ong Jun Leang, Gi-	Silva, Rui Hou, Rui Wang, Saghar Hosseini, Sa-	723
662	won Hong, Alessio Devoto, Alberto Carlo Maria	hana Chennabasappa, Sanjay Singh, Sean Bell, Seo-	724
663	Mancino, Rohit Saxena, Xuanli He, Yu Zhao, Xi-	hyun Sonia Kim, Sergey Edunov, Shaoliang Nie, Sha-	725
664	aotang Du, Mohammad Reza Ghasemi Madani,	ran Narang, Sharath Rapparthi, Sheng Shen, Shengye	726
665	Claire Barale, Robert McHardy, Joshua Harris, Jean	Wan, Shruti Bhosale, Shun Zhang, Simon Van-	727
666	Kaddour, Emile van Krieken, and Pasquale Min-	denhende, Soumya Batra, Spencer Whitman, Sten	728
667	ervini. 2025. Are we done with mmlu? <i>Preprint</i> ,	Sootla, Stephane Collot, Suchin Gururangan, Syd-	729
668	arXiv:2406.04127.	ney Borodinsky, Tamar Herman, Tara Fowler, Tarek	730
669	Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri,	Sheasha, Thomas Georgiou, Thomas Scialom, Tobias	731
670	Abhinav Pandey, Abhishek Kadian, Ahmad Al-	Speckbacher, Todor Mihaylov, Tong Xiao, Ujjwal	732
671	Dahle, Aiesha Letman, Akhil Mathur, Alan Schel-	Karn, Vedanuj Goswami, Vibhor Gupta, Vignesh	733
672	ten, Alex Vaughan, Amy Yang, Angela Fan, Anirudh	Ramanathan, Viktor Kerkez, Vincent Gouget, Vir-	734
673	Goyal, Anthony Hartshorn, Aobo Yang, Archi Mi-	ginie Do, Vish Vogeti, Vitor Albiero, Vladan Petro-	735
674	tra, Archie Sravankumar, Artem Korenev, Arthur	vic, Weiwei Chu, Wenhan Xiong, Wenxin Fu, Whit-	736
675	Hinsvark, Arun Rao, Aston Zhang, Aurelien Ro-	ney Meers, Xavier Martinet, Xiaodong Wang, Xi-	737
676	driguez, Austen Gregerson, Ava Spataru, Baptiste	aofang Wang, Xiaoqing Ellen Tan, Xide Xia, Xin-	738
677	Roziere, Bethany Biron, Binh Tang, Bobbie Chern,	feng Xie, Xuchao Jia, Xuwei Wang, Yaelle Gold-	739
		schlag, Yashesh Gaur, Yasmine Babaei, Yi Wen, Yi-	740
		wen Song, Yuchen Zhang, Yue Li, Yuning Mao,	741

742	Zacharie Delpierre Coudert, Zheng Yan, Zhengxing	806
743	Chen, Zoe Papakipos, Aaditya Singh, Aayushi Sri-	807
744	vastava, Abha Jain, Adam Kelsey, Adam Shajnfeld,	808
745	Adithya Gangidi, Adolfo Victoria, Ahuva Goldstand,	809
746	Ajay Menon, Ajay Sharma, Alex Boesenberg, Alexei	810
747	Baevski, Allie Feinstein, Amanda Kallet, Amit San-	811
748	gani, Amos Teo, Anam Yunus, Andrei Lupu, An-	812
749	dres Alvarado, Andrew Caples, Andrew Gu, Andrew	813
750	Ho, Andrew Poulton, Andrew Ryan, Ankit Ramchan-	814
751	dani, Annie Dong, Annie Franco, Anuj Goyal, Apar-	815
752	jita Saraf, Arkabandhu Chowdhury, Ashley Gabriel,	816
753	Ashwin Bharambe, Assaf Eisenman, Azadeh Yaz-	817
754	dan, Beau James, Ben Maurer, Benjamin Leonhardi,	818
755	Bernie Huang, Beth Loyd, Beto De Paola, Bhargavi	819
756	Paranjape, Bing Liu, Bo Wu, Boyu Ni, Braden Han-	820
757	cock, Bram Wasti, Brandon Spence, Brani Stojkovic,	821
758	Brian Gamido, Britt Montalvo, Carl Parker, Carly	822
759	Burton, Catalina Mejia, Ce Liu, Changan Wang,	823
760	Changkyu Kim, Chao Zhou, Chester Hu, Ching-	824
761	Hsiang Chu, Chris Cai, Chris Tindal, Christoph Fe-	825
762	ichtenhofer, Cynthia Gao, Damon Civin, Dana Beaty,	826
763	Daniel Kreymer, Daniel Li, David Adkins, David	827
764	Xu, Davide Testuggine, Delia David, Devi Parikh,	828
765	Diana Liskovich, Didem Foss, Dingkan Wang, Duc	829
766	Le, Dustin Holland, Edward Dowling, Eissa Jamil,	830
767	Elaine Montgomery, Eleonora Presani, Emily Hahn,	831
768	Emily Wood, Eric-Tuan Le, Erik Brinkman, Esteban	832
769	Arcaute, Evan Dunbar, Evan Smothers, Fei Sun,	833
770	Felix Kreuk, Feng Tian, Filippas Kokkinos, Firat Oz-	834
771	genel, Francesco Caggioni, Frank Kanayet, Frank	835
772	Seide, Gabriela Medina Florez, Gabriella Schwarz,	836
773	Gada Badeer, Georgia Sweet, Gil Halpern, Grant Her-	837
774	man, Grigory Sizov, Guangyi, Zhang, Guna Lak-	838
775	shminarayanan, Hakan Inan, Hamid Shojanazeri,	839
776	Han Zou, Hannah Wang, Hanwen Zha, Haroun	840
777	Habeeb, Harrison Rudolph, Helen Suk, Henry As-	841
778	pegren, Hunter Goldman, Hongyuan Zhan, Ibrahim	842
779	Damlaj, Igor Molybog, Igor Tufanov, Ilias Leon-	843
780	tiadis, Irina-Elena Veliche, Itai Gat, Jake Weiss-	
781	man, James Geboski, James Kohli, Janice Lam,	
782	Japhet Asher, Jean-Baptiste Gaya, Jeff Marcus, Jeff	
783	Tang, Jennifer Chan, Jenny Zhen, Jeremy Reizen-	
784	stein, Jeremy Teboul, Jessica Zhong, Jian Jin, Jingyi	
785	Yang, Joe Cummings, Jon Carvill, Jon Shepard,	
786	Jonathan McPhie, Jonathan Torres, Josh Ginsburg,	
787	Junjie Wang, Kai Wu, Kam Hou U, Karan Sax-	
788	ena, Kartikay Khandelwal, Katayoun Zand, Kathy	
789	Matosich, Kaushik Veeraraghavan, Kelly Michelena,	
790	Keqian Li, Kiran Jagadeesh, Kun Huang, Kunal	
791	Chawla, Kyle Huang, Lailin Chen, Lakshya Garg,	
792	Lavender A, Leandro Silva, Lee Bell, Lei Zhang,	
793	Liangpeng Guo, Licheng Yu, Liron Moshkovich,	
794	Luca Wehrstedt, Madian Khabza, Manav Avalani,	
795	Manish Bhatt, Martynas Mankus, Matan Hasson,	
796	Matthew Lennie, Matthias Reso, Maxim Groshev,	
797	Maxim Naumov, Maya Lathi, Meghan Keneally,	
798	Miao Liu, Michael L. Seltzer, Michal Valko, Michelle	
799	Restrepo, Mihir Patel, Mik Vyatskov, Mikayel	
800	Samvelyan, Mike Clark, Mike Macey, Mike Wang,	
801	Miquel Jubert Hermoso, Mo Metanat, Mohammad	
802	Rastegari, Munish Bansal, Nandhini Santhanam,	
803	Natascha Parks, Natasha White, Navyata Bawa,	
804	Nayan Singhal, Nick Egebo, Nicolas Usunier, Nikhil	
805	Mehta, Nikolay Pavlovich Laptev, Ning Dong, Nor-	
	man Cheng, Oleg Chernoguz, Olivia Hart, Omkar	
	Salpekar, Ozlem Kalinli, Parkin Kent, Parth Parekh,	
	Paul Saab, Pavan Balaji, Pedro Rittner, Philip Bon-	
	trager, Pierre Roux, Piotr Dollar, Polina Zvyag-	
	ina, Prashant Ratanchandani, Pritish Yuvraj, Qian	
	Liang, Rachad Alao, Rachel Rodriguez, Rafi Ayub,	
	Raghotham Murthy, Raghu Nayani, Rahul Mitra,	
	Rangaprabhu Parthasarathy, Raymond Li, Rebekkah	
	Hogan, Robin Battey, Rocky Wang, Russ Howes,	
	Ruty Rinott, Sachin Mehta, Sachin Siby, Sai Jayesh	
	Bondu, Samyak Datta, Sara Chugh, Sara Hunt, Sar-	
	gun Dhillon, Sasha Sidorov, Satadru Pan, Saurabh	
	Mahajan, Saurabh Verma, Seiji Yamamoto, Sharadh	
	Ramaswamy, Shaun Lindsay, Shaun Lindsay, Sheng	
	Feng, Shenghao Lin, Shengxin Cindy Zha, Shishir	
	Patil, Shiva Shankar, Shuqiang Zhang, Shuqiang	
	Zhang, Sinong Wang, Sneha Agarwal, Soji Sajuyigbe,	
	Soumith Chintala, Stephanie Max, Stephen Chen,	
	Steve Kehoe, Steve Satterfield, Sudarshan Govin-	
	daprasad, Sumit Gupta, Summer Deng, Sungmin	
	Cho, Sunny Virk, Suraj Subramanian, Sy Choud-	
	hury, Sydney Goldman, Tal Remez, Tamar Glaser,	
	Tamara Best, Thilo Koehler, Thomas Robinson,	
	Tianhe Li, Tianjun Zhang, Tim Matthews, Timo-	
	thy Chou, Tzook Shaked, Varun Vontimitta, Victoria	
	Ajayi, Victoria Montanez, Vijai Mohan, Vinay Satish	
	Kumar, Vishal Mangla, Vlad Ionescu, Vlad Poe-	
	nararu, Vlad Tiberiu Mihailescu, Vladimir Ivanov, Wei	
	Li, Wenchen Wang, Wenwen Jiang, Wes Bouaziz,	
	Will Constable, Xiaocheng Tang, Xiaojian Wu, Xi-	
	aolan Wang, Xilun Wu, Xinbo Gao, Yaniv Klein-	
	man, Yanjun Chen, Ye Hu, Ye Jia, Ye Qi, Yenda	
	Li, Yilin Zhang, Ying Zhang, Yossi Adi, Youngjin	
	Nam, Yu, Wang, Yu Zhao, Yuchen Hao, Yundi Qian,	
	Yunlu Li, Yuzi He, Zach Rait, Zachary DeVito, Zef	
	Rosnbrick, Zhaoduo Wen, Zhenyu Yang, Zhiwei	
	Zhao, and Zhiyu Ma. 2024. The llama 3 herd of	
	models . <i>Preprint</i> , arXiv:2407.21783.	
	Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou,	
	Mantas Mazeika, Dawn Song, and Jacob Steinhardt.	
	2021. Measuring massive multitask language under-	
	standing . In <i>International Conference on Learning</i>	
	<i>Representations</i> .	
	Ari Holtzman, Jan Buys, Li Du, Maxwell Forbes, and	
	Yejin Choi. 2020. The curious case of neural text de-	
	generation . In <i>International Conference on Learning</i>	
	<i>Representations</i> .	
	Naman Jain, King Han, Alex Gu, Wen-Ding Li, Fan-	
	jia Yan, Tianjun Zhang, Sida Wang, Armando Solar-	
	Lezama, Koushik Sen, and Ion Stoica. 2025. Live-	
	codebench: Holistic and contamination free evalua-	
	tion of large language models for code . In <i>The Thir-</i>	
	<i>teenth International Conference on Learning Repre-</i>	
	<i>sentations</i> .	
	Mehran Kazemi, Bahare Fatemi, Hritik Bansal, John	
	Palowitch, Chrysovalantis Anastasiou, Sanket Vaib-	
	hav Mehta, Lalit K. Jain, Virginia Aglietti, Disha	
	Jindal, Peter Chen, Nishanth Dikkala, Gladys Tyen,	
	Xin Liu, Uri Shalit, Silvia Chiappa, Kate Olszewska,	
	Yi Tay, Vinh Q. Tran, Quoc V. Le, and Orhan	

866	Firat. 2025. Big-bench extra hard . <i>Preprint</i> , arXiv:2502.19187.	923
867		924
868	Eunsu Kim, Juyoung Suk, Philhoon Oh, Haneul Yoo, James Thorne, and Alice Oh. 2024a. CLiCK: A benchmark dataset of cultural and linguistic intelligence in Korean . In <i>Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)</i> , pages 3335–3346, Torino, Italia. ELRA and ICCL.	925
869		926
870		927
871		928
872		929
873		
874		930
875		
876	Hyeonwoo Kim, Dahyun Kim, Jihoo Kim, Sukyung Lee, Yungi Kim, and Chanjun Park. 2025. Open ko-llm leaderboard2: Bridging foundational and practical evaluation for korean llms . <i>Preprint</i> , arXiv:2410.12445.	931
877		932
878		
879		933
880		934
881	Yeeun Kim, Youngrok Choi, Eunkyung Choi, JinHwan Choi, Hai Jin Park, and Wonseok Hwang. 2024b. Developing a pragmatic benchmark for assessing Korean legal language understanding in large language models . In <i>Findings of the Association for Computational Linguistics: EMNLP 2024</i> , pages 5573–5595, Miami, Florida, USA. Association for Computational Linguistics.	935
882		936
883		
884		937
885		
886		938
887		939
888		940
889	Sunjun Kweon, Byungjin Choi, Gyouk Chu, Junyeong Song, Daeun Hyeon, Sujin Gan, Jueon Kim, Minkyu Kim, Rae Woong Park, and Edward Choi. 2024. Kor-medmcqa: Multi-choice question answering benchmark for korean healthcare professional licensing examinations . <i>Preprint</i> , arXiv:2403.01469.	941
890		942
891		
892		943
893		944
894		
895	Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph E. Gonzalez, Hao Zhang, and Ion Stoica. 2023. Efficient memory management for large language model serving with pagedattention. In <i>Proceedings of the ACM SIGOPS 29th Symposium on Operating Systems Principles</i> .	945
896		946
897		947
898		
899		948
900		949
901		950
902	Huiyuan Lai and Malvina Nissim. 2024. mCoT: Multilingual instruction tuning for reasoning consistency in language models . In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 12012–12026, Bangkok, Thailand. Association for Computational Linguistics.	951
903		952
904		953
905		954
906		955
907		956
908		957
909	Nathan Lambert, Jacob Morrison, Valentina Pyatkin, Shengyi Huang, Hamish Ivison, Faeze Brahman, Lester James V. Miranda, Alisa Liu, Nouha Dziri, Shane Lyu, Yuling Gu, Saumya Malik, Victoria Graf, Jena D. Hwang, Jiangjiang Yang, Ronan Le Bras, Oyvind Tafjord, Chris Wilhelm, Luca Soldaini, Noah A. Smith, Yizhong Wang, Pradeep Dasigi, and Hannaneh Hajishirzi. 2025. Tulu 3: Pushing frontiers in open language model post-training . <i>Preprint</i> , arXiv:2411.15124.	958
910		959
911		960
912		961
913		962
914		963
915		964
916		965
917		966
918		967
919	Hwaran Lee, Seokhee Hong, Joonsuk Park, Takyoun Kim, Meeyoung Cha, Yejin Choi, Byoungpil Kim, Gunhee Kim, Eun-Ju Lee, Yong Lim, Alice Oh, Sangchul Park, and Jung-Woo Ha. 2023. SQuARE: A large-scale dataset of sensitive questions and acceptable responses created through human-machine collaboration . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 6692–6712, Toronto, Canada. Association for Computational Linguistics.	968
920		969
921		970
922		971
		972
		973
		974
		975
		976
		977
		978
		979

980	David Carr, David Farhi, David Mely, David Robin-	Peter Hoeschele, Peter Welinder, Phil Tillet, Philip	1044
981	son, David Sasaki, Denny Jin, Dev Valladares, Dim-	Pronin, Philippe Tillet, Prafulla Dhariwal, Qiming	1045
982	itris Tsipras, Doug Li, Duc Phong Nguyen, Duncan	Yuan, Rachel Dias, Rachel Lim, Rahul Arora, Ra-	1046
983	Findlay, Edele Oiwoh, Edmund Wong, Ehsan As-	jan Troll, Randall Lin, Rapha Gontijo Lopes, Raul	1047
984	dar, Elizabeth Proehl, Elizabeth Yang, Eric Antonow,	Puri, Reah Miyara, Reimar Leike, Renaud Gaubert,	1048
985	Eric Kramer, Eric Peterson, Eric Sigler, Eric Wal-	Reza Zamani, Ricky Wang, Rob Donnelly, Rob	1049
986	lace, Eugene Brevdo, Evan Mays, Farzad Khorasani,	Honsby, Rocky Smith, Rohan Sahai, Rohit Ramchan-	1050
987	Felipe Petroski Such, Filippo Raso, Francis Zhang,	dani, Romain Huet, Rory Carmichael, Rowan Zellers,	1051
988	Fred von Lohmann, Freddie Sulit, Gabriel Goh, Gene	Roy Chen, Ruby Chen, Ruslan Nigmatullin, Ryan	1052
989	Oden, Geoff Salmon, Giulio Starace, Greg Brockman,	Cheu, Saachi Jain, Sam Altman, Sam Schoenholz,	1053
990	Hadi Salman, Haiming Bao, Haitang Hu, Hannah	Sam Toizer, Samuel Miserendino, Sandhini Agar-	1054
991	Wong, Haoyu Wang, Heather Schmidt, Heather Whit-	wal, Sara Culver, Scott Ethersmith, Scott Gray, Sean	1055
992	ney, Heewoo Jun, Hendrik Kirchner, Henrique Ponde	Grove, Sean Metzger, Shamez Hermani, Shantanu	1056
993	de Oliveira Pinto, Hongyu Ren, Huiwen Chang,	Jain, Shengjia Zhao, Sherwin Wu, Shino Jomoto, Shi-	1057
994	Hyung Won Chung, Ian Kivlichan, Ian O’Connell,	rong Wu, Shuaiqi, Xia, Sonia Phene, Spencer Papay,	1058
995	Ian O’Connell, Ian Osband, Ian Silber, Ian Sohl,	Srinivas Narayanan, Steve Coffey, Steve Lee, Stew-	1059
996	Ibrahim Okuyucu, Ikai Lan, Ilya Kostrikov, Ilya	art Hall, Suchir Balaji, Tal Broda, Tal Stramer, Tao	1060
997	Sutskever, Ingmar Kanitscheider, Ishaan Gulrajani,	Xu, Tarun Gogineni, Taya Christianson, Ted Sanders,	1061
998	Jacob Coxon, Jacob Menick, Jakub Pachocki, James	Tejal Patwardhan, Thomas Cunningham, Thomas	1062
999	Aung, James Betker, James Crooks, James Lennon,	Degry, Thomas Dimson, Thomas Raoux, Thomas	1063
1000	Jamie Kiro, Jan Leike, Jane Park, Jason Kwon, Ja-	Shadwell, Tianhao Zheng, Todd Underwood, Todor	1064
1001	son Phang, Jason Teplitz, Jason Wei, Jason Wolfe,	Markov, Toki Sherbakov, Tom Rubin, Tom Stasi,	1065
1002	Jay Chen, Jeff Harris, Jenia Varavva, Jessica Gan	Tomer Kaftan, Tristan Heywood, Troy Peterson, Tyce	1066
1003	Lee, Jessica Shieh, Ji Lin, Jiahui Yu, Jiayi Weng,	Walters, Tyna Eloundou, Valerie Qi, Veit Moeller,	1067
1004	Jie Tang, Jieqi Yu, Joanne Jang, Joaquin Quinero	Vinnie Monaco, Vishal Kuo, Vlad Fomenko, Wayne	1068
1005	Candela, Joe Beutler, Joe Landers, Joel Parish, Jo-	Chang, Weiyei Zheng, Wenda Zhou, Wesam Manassra,	1069
1006	hannes Heidecke, John Schulman, Jonathan Lach-	Will Sheu, Wojciech Zaremba, Yash Patil, Yilei Qian,	1070
1007	man, Jonathan McKay, Jonathan Uesato, Jonathan	Yongjik Kim, Youlong Cheng, Yu Zhang, Yuchen	1071
1008	Ward, Jong Wook Kim, Joost Huizinga, Jordan	He, Yuchen Zhang, Yujia Jin, Yunxing Dai, and	1072
1009	Sitkin, Jos Kraaijeveld, Josh Gross, Josh Kaplan,	Yury Malkov. 2024a. Gpt-4o system card . <i>Preprint</i> ,	1073
1010	Josh Snyder, Joshua Achiam, Joy Jiao, Joyce Lee,	arXiv:2410.21276.	1074
1011	Juntang Zhuang, Justyn Harriman, Kai Fricke, Kai		
1012	Hayashi, Karan Singhal, Katy Shi, Kavin Karthik,	OpenAI, :, Aaron Jaech, Adam Kalai, Adam Lerer,	1075
1013	Kayla Wood, Kendra Rimbach, Kenny Hsu, Kenny	Adam Richardson, Ahmed El-Kishky, Aiden Low,	1076
1014	Nguyen, Keren Gu-Lemberg, Kevin Button, Kevin	Alec Helyar, Aleksander Madry, Alex Beutel, Alex	1077
1015	Liu, Kiel Howe, Krithika Muthukumar, Kyle Luther,	Carney, Alex Iftimie, Alex Karpenko, Alex Tachard	1078
1016	Lama Ahmad, Larry Kai, Lauren Itow, Lauren Work-	Passos, Alexander Neitz, Alexander Prokofiev,	1079
1017	man, Leher Pathak, Leo Chen, Li Jing, Lia Guy,	Alexander Wei, Allison Tam, Ally Bennett, Ananya	1080
1018	Liam Fedus, Liang Zhou, Lien Mamitsuka, Lil-	Kumar, Andre Saraiva, Andrea Vallone, Andrew Du-	1081
1019	lian Weng, Lindsay McCallum, Lindsey Held, Long	berstein, Andrew Kondrich, Andrey Mishchenko,	1082
1020	Ouyang, Louis Feuvrier, Lu Zhang, Lukas Kon-	Andy Applebaum, Angela Jiang, Ashvin Nair, Bar-	1083
1021	draciuk, Lukasz Kaiser, Luke Hewitt, Luke Metz,	ret Zoph, Behrooz Ghorbani, Ben Rossen, Benjamin	1084
1022	Lytic Doshi, Mada Aflak, Maddie Simens, Madelaine	Sokolowsky, Boaz Barak, Bob McGrew, Borys Mi-	1085
1023	Boyd, Madeleine Thompson, Marat Dukhan, Mark	naiev, Botao Hao, Bowen Baker, Brandon Houghton,	1086
1024	Chen, Mark Gray, Mark Hudnall, Marvin Zhang,	Brandon McKinzie, Brydon Eastman, Camillo Lu-	1087
1025	Marwan Aljube, Mateusz Litwin, Matthew Zeng,	garesi, Cary Bassin, Cary Hudson, Chak Ming Li,	1088
1026	Max Johnson, Maya Shetty, Mayank Gupta, Meghan	Charles de Bourcy, Chelsea Voss, Chen Shen, Chong	1089
1027	Shah, Mehmet Yatbaz, Meng Jia Yang, Mengchao	Zhang, Chris Koch, Chris Orsinger, Christopher	1090
1028	Zhong, Mia Glaese, Mianna Chen, Michael Jan-	Hesse, Claudia Fischer, Clive Chan, Dan Roberts,	1091
1029	ner, Michael Lampe, Michael Petrov, Michael Wu,	Daniel Kappler, Daniel Levy, Daniel Selsam, David	1092
1030	Michele Wang, Michelle Fradin, Michelle Pokrass,	Dohan, David Farhi, David Mely, David Robinson,	1093
1031	Miguel Castro, Miguel Oom Temudo de Castro,	Dimitris Tsipras, Doug Li, Dragos Oprica, Eben	1094
1032	Mikhail Pavlov, Miles Brundage, Miles Wang, Minal	Freeman, Eddie Zhang, Edmund Wong, Elizabeth	1095
1033	Khan, Mira Murati, Mo Bavarian, Molly Lin, Mur-	Proehl, Enoch Cheung, Eric Mitchell, Eric Wallace,	1096
1034	rat Yesildal, Nacho Soto, Natalia Gimelshein, Na-	Erik Ritter, Evan Mays, Fan Wang, Felipe Petroski	1097
1035	talie Cone, Natalie Staudacher, Natalie Summers,	Such, Filippo Raso, Florencia Leoni, Foivos Tsim-	1098
1036	Natan LaFontaine, Neil Chowdhury, Nick Ryder,	pourlas, Francis Song, Fred von Lohmann, Fred-	1099
1037	Nick Stathas, Nick Turley, Nik Tezak, Niko Felix,	die Sulit, Geoff Salmon, Giambattista Parascandolo,	1100
1038	Nithanth Kudige, Nitish Kesar, Noah Deutsch, Noel	Gildas Chabot, Grace Zhao, Greg Brockman, Guil-	1101
1039	Bundick, Nora Puckett, Ofir Nachum, Ola Okelola,	laume Leclerc, Hadi Salman, Haiming Bao, Hao	1102
1040	Oleg Boiko, Oleg Murk, Oliver Jaffe, Olivia Watkins,	Sheng, Hart Andrin, Hessam Bagherinezhad, Hongyu	1103
1041	Olivier Godement, Owen Campbell-Moore, Patrick	Ren, Hunter Lightman, Hyung Won Chung, Ian	1104
1042	Chao, Paul McMillan, Pavel Belov, Peng Su, Pe-	Kivlichan, Ian O’Connell, Ian Osband, Ignasi Clav-	1105
1043	ter Bak, Peter Bakkum, Peter Deng, Peter Dolan,	era Gilaberte, Ilge Akkaya, Ilya Kostrikov, Ilya	1106

1107	Sutskever, Irina Kofman, Jakub Pachocki, James	Bangkok, Thailand. Association for Computational	1167
1108	Lennon, Jason Wei, Jean Harb, Jerry Twore, Ji-	Linguistics.	1168
1109	acheng Feng, Jiahui Yu, Jiayi Weng, Jie Tang, Jieqi		
1110	Yu, Joaquin Quiñonero Candela, Joe Palermo, Joel	Guilherme Penedo, Hynek Kydlíček, Vinko Sabolčec,	1169
1111	Parish, Johannes Heidecke, John Hallman, John	Bettina Messmer, Negar Foroutan, Martin Jaggi, Le-	1170
1112	Rizzo, Jonathan Gordon, Jonathan Uesato, Jonathan	andro von Werra, and Thomas Wolf. 2024. Fineweb2:	1171
1113	Ward, Joost Huizinga, Julie Wang, Kai Chen, Kai	A sparkling update with 1000s of languages.	1172
1114	Xiao, Karan Singhal, Karina Nguyen, Karl Cobbe,		
1115	Katy Shi, Kayla Wood, Kendra Rimbach, Keren	Long Phan, Alice Gatti, Ziwen Han, Nathaniel Li,	1173
1116	Gu-Lemberg, Kevin Liu, Kevin Lu, Kevin Stone,	Josephina Hu, Hugh Zhang, Chen Bo Calvin Zhang,	1174
1117	Kevin Yu, Lama Ahmad, Lauren Yang, Leo Liu, Leon	Mohamed Shaaban, John Ling, Sean Shi, Michael	1175
1118	Maksin, Leyton Ho, Liam Fedus, Lilian Weng, Lin-	Choi, Anish Agrawal, Arnav Chopra, Adam Khoja,	1176
1119	den Li, Lindsay McCallum, Lindsey Held, Lorenz	Ryan Kim, Richard Ren, Jason Hausenloy, Oliver	1177
1120	Kuhn, Lukas Kondraciuk, Lukasz Kaiser, Luke Metz,	Zhang, Mantas Mazeika, Tung Nguyen, Daron Ander-	1178
1121	Madelaine Boyd, Maja Trebacz, Manas Joglekar,	son, Imad Ali Shah, Mikhail Doroshenko, Alun Cen-	1179
1122	Mark Chen, Marko Tintor, Mason Meyer, Matt Jones,	nyth Stokes, Mobeen Mahmood, Jaeho Lee, Olek-	1180
1123	Matt Kaufer, Max Schwarzer, Meghan Shah, Mehmet	sandr Pokutnyi, Oleg Iskra, Jessica P. Wang, Robert	1181
1124	Yatbaz, Melody Y. Guan, Mengyuan Xu, Mengyuan	Gerbicz, John-Clark Levin, Serguei Popov, Fiona	1182
1125	Yan, Mia Glaese, Mianna Chen, Michael Lampe,	Feng, Steven Y. Feng, Haoran Zhao, Michael Yu,	1183
1126	Michael Malek, Michele Wang, Michelle Fradin,	Varun Gangal, Chelsea Zou, Zihan Wang, Mstyslav	1184
1127	Mike McClay, Mikhail Pavlov, Miles Wang, Mingx-	Kazakov, Geoff Galgon, Johannes Schmitt, Alvaro	1185
1128	uan Wang, Mira Murati, Mo Bavarian, Mostafa Ro-	Sanchez, Yongki Lee, Will Yeadon, Scott Sauers,	1186
1129	haninejad, Nat McAleese, Neil Chowdhury, Neil	Marc Roth, Chidozie Agu, Søren Riis, Fabian Giska,	1187
1130	Chowdhury, Nick Ryder, Nikolas Tezak, Noam	Saiteja Utpala, Antrell Cheatom, Zachary Giboney,	1188
1131	Brown, Ofir Nachum, Oleg Boiko, Oleg Murk, Olivia	Gashaw M. Goshu, Sarah-Jane Crowson, Mohin-	1189
1132	Watkins, Patrick Chao, Paul Ashbourne, Pavel Iz-	der Maheshbhai Naiya, Noah Burns, Lennart Finke,	1190
1133	mailov, Peter Zhokhov, Rachel Dias, Rahul Arora,	Zerui Cheng, Hyunwoo Park, Francesco Fournier-	1191
1134	Randall Lin, Rapha Gontijo Lopes, Raz Gaon, Reah	Facio, Jennifer Zampese, John B. Wydallis, Ryan G.	1192
1135	Miyara, Reimar Leike, Renny Hwang, Rhythm	Hoerr, Mark Nandor, Tim Gehrunger, Jiaqi Cai,	1193
1136	Garg, Robin Brown, Roshan James, Rui Shu, Ryan	Ben McCarty, Jungbae Nam, Edwin Taylor, Jun	1194
1137	Cheu, Ryan Greene, Saachi Jain, Sam Altman, Sam	Jin, Gautier Abou Loume, Hangrui Cao, Alexis C	1195
1138	Toizer, Sam Toyer, Samuel Miserendino, Sandhini	Garretson, Damien Sileo, Qiuyu Ren, Doru Co-	1196
1139	Agarwal, Santiago Hernandez, Sasha Baker, Scott	joc, Pavel Arkhipov, Usman Qazi, Aras Bacho,	1197
1140	McKinney, Scottie Yan, Shengjia Zhao, Shengli	Lianghui Li, Sumeet Motwani, Christian Schroeder	1198
1141	Hu, Shibani Santurkar, Shraman Ray Chaudhuri,	de Witt, Alexei Kopylov, Johannes Veith, Eric Singer,	1199
1142	Shuyuan Zhang, Siyuan Fu, Spencer Papay, Steph	Paolo Rissone, Jaehyeok Jin, Jack Wei Lun Shi,	1200
1143	Lin, Suchir Balaji, Suvansh Sanjeev, Szymon Sidor,	Chris G. Willcocks, Ameya Prabhu, Longke Tang,	1201
1144	Tal Broda, Aidan Clark, Tao Wang, Taylor Gordon,	Kevin Zhou, Emily de Oliveira Santos, Andrey Pu-	1202
1145	Ted Sanders, Tejal Patwardhan, Thibault Sottiaux,	pasov Maksimov, Edward Vendrow, Kengo Zeni-	1203
1146	Thomas Degry, Thomas Dimson, Tianhao Zheng,	tani, Joshua Robinson, Aleksandar Mikov, Julien	1204
1147	Timur Garipov, Tom Stasi, Trapit Bansal, Trevor	Guillod, Yuqi Li, Ben Pageler, Joshua Vendrow,	1205
1148	Crech, Troy Peterson, Tyna Eloundou, Valerie Qi,	Vladyslav Kuchkin, Pierre Marion, Denis Efremov,	1206
1149	Vineet Kosaraju, Vinnie Monaco, Vitchyr Pong, Vlad	Jayson Lynch, Kaiqu Liang, Andrew Gritsevskiy,	1207
1150	Fomenko, Weiye Zheng, Wenda Zhou, Wes McCabe,	Dakotah Martinez, Nick Crispino, Dimitri Zvonk-	1208
1151	Wojciech Zaremba, Yann Dubois, Yinghai Lu, Yining	ine, Natanael Wildner Fraga, Saeed Soori, Ori	1209
1152	Chen, Young Cha, Yu Bai, Yuchen He, Yuchen Zhang,	Press, Henry Tang, Julian Salazar, Sean R. Green,	1210
1153	Yunyun Wang, Zheng Shao, and Zhuohan Li. 2024b.	Lina Brüssel, Moon Twayana, Aymeric Dieuleveut,	1211
1154	Openai o1 system card . <i>Preprint</i> , arXiv:2412.16720.	T. Ryan Rogers, Wenjin Zhang, Ross Finocchio,	1212
1155	OpenAI. 2024. Multilingual massive multitask lan-	Bikun Li, Jinzhou Yang, Arun Rao, Gabriel Loiseau,	1213
1156	guage understanding (mmmlu) .	Mikhail Kalinin, Marco Lukas, Ciprian Manolescu,	1214
1157	OpenAI. 2025a. Introducing gpt-4.1 in the api .	Nate Stambaugh, Subrata Mishra, Ariel Ghislain Ke-	1215
1158	OpenAI. 2025b. Introducing openai o3 and o4-mini .	mogne Kamdoun, Tad Hogg, Alvin Jin, Carlo Bosio,	1216
1159	OpenAI. 2025c. Openai o3-mini .	Gongbo Sun, Brian P Coppola, Haline Heidinger,	1217
1160	Chanjun Park, Hyeonwoo Kim, Dahyun Kim, SeongH-	Rafael Sayous, Stefan Ivanov, Joseph M Cavanagh,	1218
1161	wan Cho, Sanghoon Kim, Sukyung Lee, Yungi Kim,	Jiawei Shen, Joseph Marvin Imperial, Philippe	1219
1162	and Hwalsuk Lee. 2024. Open Ko-LLM leaderboard:	Schwaller, Shaipranesh Senthilkuma, Andres M Bran,	1220
1163	Evaluating large language models in Korean with Ko-	Andres Algaba, Brecht Verbeken, Kelsey Van den	1221
1164	h5 benchmark . In <i>Proceedings of the 62nd Annual</i>	Houte, Lynn Van Der Sypt, David Noever, Lisa Schut,	1222
1165	<i>Meeting of the Association for Computational Lin-</i>	Ilia Sucholutsky, Evgenii Zheltonozhskii, Qiaochu	1223
1166	<i>guistics (Volume 1: Long Papers)</i> , pages 3220–3234,	Yuan, Derek Lim, Richard Stanley, Shankar Sivara-	1224
		jan, Tong Yang, John Maar, Julian Wykowski, Martí	1225
		Oller, Jennifer Sandlin, Anmol Sahu, Cesare Giulio	1226
		Ardito, Yuzheng Hu, Felipe Meneguitti Dias, Tobias	1227
		Kreiman, Kaivalya Rawal, Tobias Garcia Vilchis,	1228

1229	Yuexuan Zu, Martin Lackner, James Koppel, Jeremy	Tordera, George Balabanian, Earth Anderson, Lynna	1293
1230	Nguyen, Daniil S. Antonenko, Steffi Chern, Bingchen	Kvistad, Alejandro José Moyano, Hsiaoyun Mill-	1294
1231	Zhao, Pierrot Arsene, Sergey Ivanov, Rafał Poświata,	iron, Ahmad Sakor, Murat Eron, Isaac C. McAl-	1295
1232	Chenguang Wang, Daofeng Li, Donato Crisostomi,	ister, Andrew Favre D. O., Shailesh Shah, Xiaox-	1296
1233	Ali Dehghan, Andrea Achilleos, John Arnold Am-	iang Zhou, Firuz Kamalov, Ronald Clark, Sher-	1297
1234	bay, Benjamin Myklebust, Archan Sen, David Per-	win Abdoli, Tim Santens, Harrison K Wang, Evan	1298
1235	rella, Nurdin Kaparov, Mark H Inlow, Allen Zang,	Chen, Alessandro Tomasiello, G. Bruno De Luca,	1299
1236	Kalyan Ramakrishnan, Daniil Orel, Vladislav Porit-	Shi-Zhuo Looi, Vinh-Kha Le, Noam Kolt, Niels	1300
1237	ski, Shalev Ben-David, Zachary Berger, Parker	Mündler, Avi Semler, Emma Rodman, Jacob Drori,	1301
1238	Whitfill, Michael Foster, Daniel Munro, Linh Ho,	Carl J Fossum, Luk Gloor, Milind Jagota, Ronak	1302
1239	Dan Bar Hava, Aleksey Kuchkin, Robert Lauff,	Pradeep, Honglu Fan, Tej Shah, Jonathan Eicher,	1303
1240	David Holmes, Frank Sommerhage, Anji Zhang,	Michael Chen, Kushal Thaman, William Merrill,	1304
1241	Richard Moat, Keith Schneider, Daniel Pyda, Zakayo	Moritz Firsching, Carter Harris, Stefan Ciobăcă, Ja-	1305
1242	Kazibwe, Mukhwinder Singh, Don Clarke, Dae Hyun	son Gross, Rohan Pandey, Ilya Gusev, Adam Jones,	1306
1243	Kim, Sara Fish, Veit Elser, Victor Efren Guadar-	Shashank Agnihotri, Pavel Zhelnov, Siranut Us-	1307
1244	rama Vilchis, Immo Klose, Christoph Demian,	awasutsakorn, Mohammadreza Mofayezi, Alexan-	1308
1245	Ujjwala Anantheswaran, Adam Zweiger, Guglielmo	der Piperski, Marc Carauleanu, David K. Zhang,	1309
1246	Albani, Jeffery Li, Nicolas Daans, Maksim Radionov,	Kostiantyn Dobarskyi, Dylan Ler, Roman Leven-	1310
1247	Václav Rozhoň, Vincent Ginis, Ziqiao Ma, Chris-	to, Ignat Soroko, Thorben Jansen, Scott Creighton,	1311
1248	tian Stump, Jacob Platnick, Volodymyr Nevirkovets,	Pascal Lauer, Joshua Duersch, Vage Taamazyan,	1312
1249	Luke Basler, Marco Piccardo, Niv Cohen, Viren-	Dario Bezzi, Wiktor Morak, Wenjie Ma, William	1313
1250	dra Singh, Josef Tkadlec, Paul Rosu, Alan Goldfarb,	Held, Tran Đức Huy, Ruicheng Xian, Armel Randy	1314
1251	Piotr Padlewski, Stanislaw Barzowski, Kyle Mont-	Zebaze, Mohanad Mohamed, Julian Noah Leser,	1315
1252	gomery, Aline Menezes, Arkil Patel, Zixuan Wang,	Michelle X Yuan, Laila Yacar, Johannes Lengler,	1316
1253	Jamie Tucker-Foltz, Jack Stade, Declan Grabb, Tom	Katarzyna Olszewska, Hossein Shahrtash, Edson	1317
1254	Goertzen, Fereshteh Kazemi, Jeremiah Milbauer, Ab-	Oliveira, Joseph W. Jackson, Daniel Espinosa Gon-	1318
1255	hishek Shukla, Hossam Elgnainy, Yan Carlos Leyva	zalez, Andy Zou, Muthu Chidambaram, Timothy	1319
1256	Labrador, Hao He, Ling Zhang, Alan Givré, Hew	Manik, Hector Haffenden, Dashiell Stander, Ali	1320
1257	Wolff, Gözdenur Demir, Muhammad Fayez Aziz,	Dasouqi, Alexander Shen, Emilien Duc, Bitá Gol-	1321
1258	Younesse Kaddar, Ivar Ångquist, Yanxu Chen, El-	shani, David Stap, Mikalai Uzhou, Alina Borisovna	1322
1259	liott Thornley, Robin Zhang, Jiayi Pan, Antonio Ter-	Zhidkovskaya, Lukas Lewark, Miguel Orbeagozo Ro-	1323
1260	pin, Niklas Muennighoff, Hailey Schoelkopf, Eric	driguez, Mátyás Vincze, Dustin Wehr, Colin Tang,	1324
1261	Zheng, Avishy Carmi, Jainam Shah, Ethan D. L.	Shaun Phillips, Fortuna Samuele, Jiang Muzhen,	1325
1262	Brown, Kelin Zhu, Max Bartolo, Richard Wheeler,	Fredrik Ekström, Angela Hammon, Oam Patel, Faraz	1326
1263	Andrew Ho, Shaul Barkan, Jiaqi Wang, Martin Ste-	Farhidi, George Medley, Forough Mohammadzadeh,	1327
1264	hberger, Egor Kretov, Peter Bradshaw, JP Heimo-	Madellene Peñaflor, Haile Kassahun, Alena Friedrich,	1328
1265	nen, Kaustubh Sridhar, Zaki Hossain, Ido Akov, Yury	Claire Sparrow, Rayner Hernandez Perez, Taom	1329
1266	Makarychev, Joanna Tam, Hieu Hoang, David M.	Sakal, Omkar Dhamane, Ali Khajegili Mirabadi,	1330
1267	Cunningham, Vladimir Goryachev, Demosthenes Pa-	Eric Hallman, Kenchi Okutsu, Mike Battaglia, Mo-	1331
1268	tramanis, Michael Krause, Andrew Redenti, David	hammad Maghsoudimehrabani, Alon Amit, Dave	1332
1269	Aldous, Jesyin Lai, Shannon Coleman, Jiangnan Xu,	Hulbert, Roberto Pereira, Simon Weber, Handoko,	1333
1270	Sangwon Lee, Ilias Magoulas, Sandy Zhao, Ning	Anton Peristyy, Stephen Malina, Samuel Albanie,	1334
1271	Tang, Michael K. Cohen, Micah Carroll, Orr Par-	Will Cai, Mustafa Mehkary, Rami Aly, Frank Rei-	1335
1272	adise, Jan Hendrik Kirchner, Stefan Steinerberger,	degeld, Anna-Katharina Dick, Cary Friday, Jasdeep	1336
1273	Maksym Ovchynnikov, Jason O. Matos, Adithya	Sidhu, Hassan Shapourian, Wanyoung Kim, Mariana	1337
1274	Shenoy, Michael Wang, Yuzhou Nie, Paolo Gior-	Costa, Hubeyb Gurdogan, Brian Weber, Harsh Ku-	1338
1275	dano, Philipp Petersen, Anna Szytber-Betley, Paolo	mar, Tong Jiang, Arunim Agarwal, Chiara Ceconello,	1339
1276	Faraboschi, Robin Riblet, Jonathan Crozier, Shiv Ha-	Warren S. Vaz, Chao Zhuang, Haon Park, Andrew R.	1340
1277	lasyamani, Antonella Pinto, Shreyas Verma, Prashant	Tawfeek, Daattavya Aggarwal, Michael Kirchhof,	1341
1278	Joshi, Eli Meril, Zheng-Xin Yong, Allison Tee,	Linjie Dai, Evan Kim, Johan Ferret, Yuzhou Wang,	1342
1279	Jérémy Andréoletti, Orion Weller, Raghav Singhal,	Minghao Yan, Krzysztof Burdzy, Lixin Zhang, An-	1343
1280	Gang Zhang, Alexander Ivanov, Seri Khoury, Nils	tonio Franca, Diana T. Pham, Kang Yong Loh,	1344
1281	Gustafsson, Hamid Mostaghimi, Kunvar Thaman,	Joshua Robinson, Abram Jackson, Shreen Gul, Gun-	1345
1282	Qijia Chen, Tran Quoc Khánh, Jacob Loader, Ste-	jan Chhablani, Zhehang Du, Adrian Cosma, Jesus	1346
1283	fano Cavalleri, Hannah Szlyk, Zachary Brown, Hi-	Colino, Colin White, Jacob Votava, Vladimir Vin-	1347
1284	manshu Narayan, Jonathan Roberts, William Alley,	nikov, Ethan Delaney, Petr Spelda, Vit Stritecky,	1348
1285	Kunyang Sun, Ryan Stendall, Max Lamparth, Anka	Syed M. Shahid, Jean-Christophe Mourrat, Lavr	1349
1286	Reuel, Ting Wang, Hanmeng Xu, Pablo Hernández-	Vetoshkin, Koen Sponselee, Renas Bacho, Floren-	1350
1287	Cámara, Freddie Martin, Thomas Preu, Tomek Kor-	cia de la Rosa, Xiuyu Li, Guillaume Malod, Leon	1351
1288	bak, Marcus Abramovitch, Dominic Williamson, Ida	Lang, Julien Laurendeau, Dmitry Kazakov, Fatimah	1352
1289	Bosio, Ziye Chen, Biró Bálint, Eve J. Y. Lo, Maria	Adesanya, Julien Portier, Lawrence Hollom, Victor	1353
1290	Inês S. Nunes, Yibo Jiang, M Saiful Bari, Peyman	Souza, Yuchen Anna Zhou, Julien Degorre, Yiğit	1354
1291	Kassani, Zihao Wang, Behzad Ansarinejad, Yewen	Yalın, Gbenga Daniel Obikoya, Luca Arnaboldi, Rai,	1355
1292	Sun, Stephane Durand, Guillaume Douville, Daniel	Filippo Bigi, M. C. Boscá, Oleg Shumar, Kani-	1356

1357	uar Bacho, Pierre Clavier, Gabriel Recchia, Mara	marks: is mmlu lost in translation? <i>Preprint</i> ,	1419
1358	Popescu, Nikita Shulga, Ngefor Mildred Tanwie, De-	arXiv:2406.17789.	1420
1359	nis Peskoff, Thomas C. H. Lux, Ben Rank, Colin		
1360	Ni, Matthew Brooks, Alesia Yakimchyk, Huanxu,	Qwen, :, An Yang, Baosong Yang, Beichen Zhang,	1421
1361	Liu, Olle Häggström, Emil Verkama, Hans Gund-	Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li,	1422
1362	lach, Leonor Brito-Santana, Brian Amaro, Vivek Va-	Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin,	1423
1363	jipey, Rynaa Grover, Yiyang Fan, Gabriel Poesia Reis	Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang,	1424
1364	e Silva, Linwei Xin, Yosi Kratish, Jakub Lucki, Wen-	Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang,	1425
1365	Ding Li, Sivakanth Gopi, Andrea Caciolai, Justin Xu,	Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li,	1426
1366	Kevin Joseph Scaria, Freddie Vargus, Farzad Habibi,	Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji	1427
1367	Long, Lian, Emanuele Rodolà, Jules Robins, Vin-	Lin, Tianhao Li, Tianyi Tang, Tingyu Xia, Xingzhang	1428
1368	cent Cheng, Tony Fruhauff, Brad Raynor, Hao Qi,	Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang	1429
1369	Xi Jiang, Ben Segev, Jingxuan Fan, Sarah Martin-	Zhang, Yu Wan, Yuqiong Liu, Zeyu Cui, Zhenru	1430
1370	son, Erik Y. Wang, Kaylie Hausknecht, Michael P.	Zhang, and Zihan Qiu. 2025. Qwen2.5 technical	1431
1371	Brenner, Mao Mao, Xinyu Zhang, David Avagian, Es-	report . <i>Preprint</i> , arXiv:2412.15115.	1432
1372	hawn Jessica Scipio, Alon Ragoler, Justin Tan, Blake		
1373	Sims, Rebeka Plecnik, Aaron Kirtland, Omer Faruk	David Rein, Betty Li Hou, Asa Cooper Stickland, Jack-	1433
1374	Bodur, D. P. Shinde, Zahra Adoul, Mohamed Zekry,	son Petty, Richard Yuanzhe Pang, Julien Dirani, Ju-	1434
1375	Ali Karakoc, Tania C. B. Santos, Samir Sham-	lian Michael, and Samuel R. Bowman. 2024. GPQA:	1435
1376	selddeen, Loukmane Karim, Anna Liakhovitskaia,	A graduate-level google-proof q&a benchmark . In	1436
1377	Nate Resman, Nicholas Farina, Juan Carlos Gonzalez,	<i>First Conference on Language Modeling</i> .	1437
1378	Gabe Maayan, Sarah Hoback, Rodrigo De Oliveira		
1379	Pena, Glen Sherman, Elizabeth Kelley, Hodjat Mar-	LG AI Research, Soyoung An, Kyunghoon Bae,	1438
1380	iji, Rasoul Pouriamanesh, Wentao Wu, Sandra Men-	Eunbi Choi, Kibong Choi, Stanley Jungkyu Choi,	1439
1381	doza, Ismail Alarab, Joshua Cole, Danyelle Fer-	Seokhee Hong, Junwon Hwang, Hyojin Jeon, Ger-	1440
1382	reira, Bryan Johnson, Mohammad Safdari, Liangti	rard Jeongwon Jo, Hyunjik Jo, Jiyeon Jung, YOUN-	1441
1383	Dai, Siriphan Arthornthurasuk, Alexey Pronin, Jing	tae Jung, Hyosang Kim, Joonkee Kim, Seonghwan	1442
1384	Fan, Angel Ramirez-Trinidad, Ashley Cartwright,	Kim, Soyeon Kim, Sunkyoung Kim, Yireun Kim,	1443
1385	Daphiny Pottmaier, Omid Taheri, David Outevsky,	Yongil Kim, Youchul Kim, Edward Hwayoung Lee,	1444
1386	Stanley Stepanic, Samuel Perry, Luke Askew, Raúl	Haeju Lee, Honglak Lee, Jinsik Lee, Kyungmin	1445
1387	Adrián Huerta Rodríguez, Ali M. R. Minissi, Sam Ali,	Lee, Woohyung Lim, Sangha Park, Sooyoun Park,	1446
1388	Ricardo Lorena, Krishnamurthy Iyer, Arshad Anil	Yongmin Park, Sihoon Yang, Heuiyeon Yeen, and	1447
1389	Fasiludeen, Sk Md Salauddin, Murat Islam, Juan	Hyeongu Yun. 2024. Exaone 3.5: Series of large	1448
1390	Gonzalez, Josh Ducey, Maja Somrak, Vasilios	language models for real-world use cases . <i>Preprint</i> ,	1449
1391	Mavroudis, Eric Vergo, Juehang Qin, Benjámín	arXiv:2412.04862.	1450
1392	Borbás, Eric Chu, Jack Lindsey, Anil Radhakrishnan,		
1393	Antoine Jallon, I. M. J. McInnis, Pawan Kumar, Lax-	LG AI Research, Kyunghoon Bae, Eunbi Choi, Ki-	1451
1394	man Prasad Goswami, Daniel Bugas, Nasser Heydari,	bong Choi, Stanley Jungkyu Choi, Yemuk Choi,	1452
1395	Ferenc Jeanplong, Archimedes Apronti, Abdallah	Seokhee Hong, Junwon Hwang, Hyojin Jeon, Ki-	1453
1396	Galal, Ng Ze-An, Ankit Singh, Joan of Arc Xavier,	jeong Jeon, Gerrard Jeongwon Jo, Hyunjik Jo, Jiyeon	1454
1397	Kanu Priya Agarwal, Mohammed Berkani, Bened-	Jung, Hyosang Kim, Joonkee Kim, Seonghwan	1455
1398	ito Alves de Oliveira Junior, Dmitry Malishev, Nico-	Kim, Soyeon Kim, Sunkyoung Kim, Yireun Kim,	1456
1399	las Remy, Taylor D. Hartman, Tim Tarver, Stephen	Yongil Kim, Youchul Kim, Edward Hwayoung Lee,	1457
1400	Mensah, Javier Gimenez, Roselynn Grace Monte-	Haeju Lee, Honglak Lee, Jinsik Lee, Kyungmin	1458
1401	cillo, Russell Campbell, Asankhaya Sharma, Khalida	Lee, Sangha Park, Yongmin Park, Sihoon Yang,	1459
1402	Meer, Xavier Alapont, Deepakkumar Patil, Rajat Ma-	Heuiyeon Yeen, Sihyuk Yi, and Hyeongu Yun. 2025.	1460
1403	heshwari, Abdelkader Dendane, Priti Shukla, Sergei	Exaone deep: Reasoning enhanced language models .	1461
1404	Bogdanov, Sören Möller, Muhammad Rehan Siddiqi,	<i>Preprint</i> , arXiv:2503.12524.	1462
1405	Prajvi Saxena, Himanshu Gupta, Innocent Enyekwe,		
1406	Ragavendran P V, Zienab EL-Wasif, Aleksandr Mak-	Manley Roberts, Himanshu Thakur, Christine Herlihy,	1463
1407	sapetyan, Vivien Rosssbach, Chris Harjadi, Mohsen	Colin White, and Samuel Dooley. 2024. To the cut-	1464
1408	Bahaloohoreh, Song Bian, John Lai, Justine Leon	off... and beyond? a longitudinal perspective on LLM	1465
1409	Uro, Greg Bateman, Mohamed Sayed, Ahmed Men-	data contamination . In <i>The Twelfth International</i>	1466
1410	shawy, Darling Duclosel, Yashaswini Jain, Ashley	<i>Conference on Learning Representations</i> .	1467
1411	Aaron, Murat Tiryakioğlu, Sheeshram Siddh, Keith		
1412	Krenek, Alex Hoover, Joseph McGowan, Tejal Pat-	Shivalika Singh, Angelika Romanou, Clémentine Four-	1468
1413	wardhan, Summer Yue, Alexandr Wang, and Dan	rier, David I. Adelani, Jian Gang Ngui, Daniel	1469
1414	Hendrycks. 2025. Humanity’s last exam . <i>Preprint</i> ,	Vila-Suero, Peerat Limkonchotiwat, Kelly Marchi-	1470
1415	arXiv:2501.14249.	sio, Wei Qi Leong, Yosephine Susanto, Raymond	1471
		Ng, Shayne Longpre, Wei-Yin Ko, Sebastian Ruder,	1472
		Madeline Smith, Antoine Bosselut, Alice Oh, Andre	1473
		F. T. Martins, Leshem Choshen, Daphne Ippolito,	1474
		Enzo Ferrante, Marzieh Fadaee, Beyza Ermiş, and	1475
		Sara Hooker. 2025. Global mmlu: Understanding	1476
1416	Irene Plaza, Nina Melero, Cristina del Pozo, Javier	and addressing cultural and linguistic biases in multi-	1477
1417	Conde, Pedro Reviriego, Marina Mayor-Rocher, and	lingual evaluation . <i>Preprint</i> , arXiv:2412.03304.	1478
1418	María Grandury. 2024. Spanish and llm bench-		

1479	Shivalika Singh, Angelika Romanou, Clémentine Fourrier, David I. Adelani, Jian Gang Ngui, Daniel Vila-Suero, Peerat Limkonchotiawat, Kelly Marchisio, Wei Qi Leong, Yosephine Susanto, Raymond Ng, Shayne Longpre, Wei-Yin Ko, Madeline Smith, Antoine Bosselut, Alice Oh, Andre F. T. Martins, Leshem Choshen, Daphne Ippolito, Enzo Ferrante, Marzieh Fadaee, Beyza Ermis, and Sara Hooker. 2024. Global mmlu: Understanding and addressing cultural and linguistic biases in multilingual evaluation . <i>Preprint</i> , arXiv:2412.03304.	
1480		
1481		
1482		
1483		
1484		
1485		
1486		
1487		
1488		
1489		
1490	Guijin Son, Hanwool Lee, Sungdong Kim, Seung-gone Kim, Niklas Muennighoff, Taekyoon Choi, Cheonbok Park, Kang Min Yoo, and Stella Biderman. 2024a. Kmmmlu: Measuring massive multi-task language understanding in korean . <i>Preprint</i> , arXiv:2402.11548.	
1491		
1492		
1493		
1494		
1495		
1496	Guijin Son, Hanwool Lee, Suwan Kim, Huiseo Kim, Jae cheol Lee, Je Won Yeom, Jihyu Jung, Jung woo Kim, and Songseong Kim. 2024b. HAE-RAE bench: Evaluation of Korean knowledge in language models . In <i>Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)</i> , pages 7993–8007, Torino, Italia. ELRA and ICCL.	
1497		
1498		
1499		
1500		
1501		
1502		
1503		
1504	Aarohi Srivastava, Abhinav Rastogi, Abhishek Rao, Abu Awal Md Shoeb, Abubakar Abid, Adam Fisch, Adam R. Brown, Adam Santoro, Aditya Gupta, Adrià Garriga-Alonso, Agnieszka Kluska, Aitor Lewkowycz, Akshat Agarwal, Alethea Power, Alex Ray, Alex Warstadt, Alexander W. Kocurek, Ali Safaya, Ali Tazarv, Alice Xiang, Alicia Parrish, Allen Nie, Aman Hussain, Amanda Askell, Amanda Dsouza, Ambrose Slone, Ameet Rahane, Anantharaman S. Iyer, Anders Johan Andreassen, Andrea Madotto, Andrea Santilli, Andreas Stuhlmüller, Andrew M. Dai, Andrew La, Andrew Kyle Lampinen, Andy Zou, Angela Jiang, Angelica Chen, Anh Vuong, Animesh Gupta, Anna Gottardi, Antonio Norelli, Anu Venkatesh, Arash Gholamidavoodi, Arfa Tabassum, Arul Menezes, Arun Kirubakaran, Asher Mullokandov, Ashish Sabharwal, Austin Herrick, Avia Efrat, Aykut Erdem, Ayla Karakaş, B. Ryan Roberts, Bao Sheng Loe, Barret Zoph, Bartłomiej Bojanowski, Batuhan Özyurt, Behnam Hedayatnia, Behnam Neyshabur, Benjamin Inden, Benno Stein, Berk Ekmekci, Bill Yuchen Lin, Blake Howald, Bryan Orinon, Cameron Diao, Cameron Dour, Catherine Stinson, Cedrick Argueta, Cesar Ferri, Chandan Singh, Charles Rathkopf, Chenlin Meng, Chitta Baral, Chiyu Wu, Chris Callison-Burch, Christopher Waites, Christian Voigt, Christopher D Manning, Christopher Potts, Cindy Ramirez, Clara E. Rivera, Clemencia Siro, Colin Raffel, Courtney Ashcraft, Cristina Garbacea, Damien Sileo, Dan Garrette, Dan Hendrycks, Dan Kilman, Dan Roth, C. Daniel Freeman, Daniel Khashabi, Daniel Levy, Daniel Moseguí González, Danielle Perszyk, Danny Hernandez, Danqi Chen, Daphne Ippolito, Dar Gilboa, David Dohan, David Drakard, David Jurgens, Debajyoti Datta, Deep Ganguli, Denis Emelin,	
1505		
1506		
1507		
1508		
1509		
1510		
1511		
1512		
1513		
1514		
1515		
1516		
1517		
1518		
1519		
1520		
1521		
1522		
1523		
1524		
1525		
1526		
1527		
1528		
1529		
1530		
1531		
1532		
1533		
1534		
1535		
1536		
1537		
1538		
1539		
	Denis Kleyko, Deniz Yuret, Derek Chen, Derek Tam, Dieuwke Hupkes, Diganta Misra, Dilyar Buzan, Dimitri Coelho Mollo, Diyi Yang, Dong-Ho Lee, Dylan Schrader, Ekaterina Shutova, Ekin Dogus Cubuk, Elad Segal, Eleanor Hagerman, Elizabeth Barnes, Elizabeth Donoway, Ellie Pavlick, Emanuele Rodolà, Emma Lam, Eric Chu, Eric Tang, Erkut Erdem, Ernie Chang, Ethan A Chi, Ethan Dyer, Ethan Jerzak, Ethan Kim, Eunice Engefu Manyasi, Evgenii Zheltonozhskii, Fanyue Xia, Fatemeh Siar, Fernando Martínez-Plumed, Francesca Happé, Francois Chollet, Frieda Rong, Gaurav Mishra, Genta Indra Winata, Gerard de Melo, Germán Kruszewski, Giambattista Parascandolo, Giorgio Mariani, Gloria Xinyue Wang, Gonzalo Jaimovitch-Lopez, Gregor Betz, Guy Gur-Ari, Hana Galijasevic, Hannah Kim, Hannah Rashkin, Hannaneh Hajishirzi, Harsh Mehta, Hayden Bogar, Henry Francis Anthony Shevlin, Hinrich Schuetze, Hiromu Yakura, Hongming Zhang, Hugh Mee Wong, Ian Ng, Isaac Noble, Jaap Jumelet, Jack Geissinger, Jackson Kernion, Jacob Hilton, Jaehoon Lee, Jaime Fernández Fisac, James B Simon, James Koppel, James Zheng, James Zou, Jan Kocon, Jana Thompson, Janelle Wingfield, Jared Kaplan, Jarema Radom, Jascha Sohl-Dickstein, Jason Phang, Jason Wei, Jason Yosinski, Jekaterina Novikova, Jelle Bosscher, Jennifer Marsh, Jeremy Kim, Jeroen Taal, Jesse Engel, Jesujoba Alabi, Jiacheng Xu, Jiaming Song, Jillian Tang, Joan Waweru, John Burden, John Miller, John U. Balis, Jonathan Batchelder, Jonathan Berant, Jörg Froberg, Jos Rozen, Jose Hernandez-Orallo, Joseph Boudeman, Joseph Guerr, Joseph Jones, Joshua B. Tenenbaum, Joshua S. Rule, Joyce Chua, Kamil Kanclerz, Karen Livescu, Karl Krauth, Karthik Gopalakrishnan, Katerina Ignatyeva, Katja Markert, Kaustubh Dhole, Kevin Gimpel, Kevin Omondi, Kory Wallace Mathewson, Kristen Chiafullo, Ksenia Shkaruta, Kumar Shridhar, Kyle McDonell, Kyle Richardson, Laria Reynolds, Leo Gao, Li Zhang, Liam Dugan, Lianhui Qin, Lidia Contreras-Ochando, Louis-Philippe Morency, Luca Moschella, Lucas Lam, Lucy Noble, Ludwig Schmidt, Luheng He, Luis Oliveros-Colón, Luke Metz, Lütfi Kerem Senel, Maarten Bosma, Maarten Sap, Maartje Ter Hoeve, Maheen Farooqi, Manaal Faruqui, Mantas Mazeika, Marco Baturan, Marco Marelli, Marco Maru, Maria Jose Ramirez-Quintana, Marie Tolkiehn, Mario Giulianelli, Martha Lewis, Martin Potthast, Matthew L Leavitt, Matthias Hagen, Mátyás Schubert, Medina Orduna Baitemirova, Melody Arnaud, Melvin McElrath, Michael Andrew Yee, Michael Cohen, Michael Gu, Michael Ivanitskiy, Michael Starritt, Michael Strube, Michał Śwędrowski, Michele Bevilacqua, Michihiro Yasunaga, Mihir Kale, Mike Cain, Mimeo Xu, Mirac Suzgun, Mitch Walker, Mo Tiwari, Mohit Bansal, Moin Aminnaseri, Mor Geva, Mozhdeh Gheini, Mukund Varma T, Nanyun Peng, Nathan Andrew Chi, Nayeon Lee, Neta Gur-Ari Krakover, Nicholas Cameron, Nicholas Roberts, Nick Doiron, Nicole Martinez, Nikita Nangia, Niklas Deckers, Niklas Muennighoff, Nitish Shirish Keskar, Niveditha S. Iyer, Noah Constant, Noah Fiedel, Nuan Wen, Oliver Zhang, Omar Agha, Omar Elbaghdadi, Omer Levy, Owain Evans, Pablo Anto-	1540 1541 1542 1543 1544 1545 1546 1547 1548 1549 1550 1551 1552 1553 1554 1555 1556 1557 1558 1559 1560 1561 1562 1563 1564 1565 1566 1567 1568 1569 1570 1571 1572 1573 1574 1575 1576 1577 1578 1579 1580 1581 1582 1583 1584 1585 1586 1587 1588 1589 1590 1591 1592 1593 1594 1595 1596 1597 1598 1599 1600 1601 1602 1603

1604	nio Moreno Casares, Parth Doshi, Pascale Fung,	Gaeun Seo. 2025a. Kanana: Compute-efficient bilin-	1665
1605	Paul Pu Liang, Paul Vicol, Pegah Alipoormolabashi,	gual language models . <i>Preprint</i> , arXiv:2502.18934.	1666
1606	Peiyuan Liao, Percy Liang, Peter W Chang, Pe-		
1607	ter Eckersley, Phu Mon Htut, Pinyu Hwang, Piotr	M-A-P Team, Xinrun Du, Yifan Yao, Kaijing Ma, Bingli	1667
1608	Milkowski, Piyush Patil, Pouya Pezeshkpour, Priti	Wang, Tianyu Zheng, Kang Zhu, Minghao Liu, Yim-	1668
1609	Oli, Qiaozhu Mei, Qing Lyu, Qinlang Chen, Rabin	ing Liang, Xiaolong Jin, Zhenlin Wei, Chujie Zheng,	1669
1610	Banjade, Rachel Etta Rudolph, Raefer Gabriel, Rahel	Kaixing Deng, Shuyue Guo, Shian Jia, Sichao Jiang,	1670
1611	Habacker, Ramon Risco, Raphaël Millière, Rhythm	Yiyan Liao, Rui Li, Qinrui Li, Sirun Li, Yizhi Li,	1671
1612	Garg, Richard Barnes, Rif A. Saurous, Riku Arakawa,	Yunwen Li, Dehua Ma, Yuansheng Ni, Haoran Que,	1672
1613	Robbe Raymaekers, Robert Frank, Rohan Sikand, Ro-	Qiyao Wang, Zhoufutu Wen, Siwei Wu, Tianshun	1673
1614	man Novak, Roman Sitelew, Ronan Le Bras, Rosanne	Xing, Ming Xu, Zhenzhu Yang, Zekun Moore Wang,	1674
1615	Liu, Rowan Jacobs, Rui Zhang, Russ Salakhutdinov,	Junting Zhou, Yuelin Bai, Xingyuan Bu, Chenglin	1675
1616	Ryan Andrew Chi, Seungjae Ryan Lee, Ryan Stovall,	Cai, Liang Chen, Yifan Chen, Chengtuo Cheng,	1676
1617	Ryan Teehan, Rylan Yang, Sahib Singh, Saif M. Mo-	Tianhao Cheng, Keyi Ding, Siming Huang, Yun	1677
1618	hammad, Sajant Anand, Sam Dillavou, Sam Shleifer,	Huang, Yaoru Li, Yizhe Li, Zhaoqun Li, Tianhao	1678
1619	Sam Wiseman, Samuel Gruetter, Samuel R. Bow-	Liang, Chengdong Lin, Hongquan Lin, Yinghao Ma,	1679
1620	man, Samuel Stern Schoenholz, Sanghyun Han, San-	Zhongyuan Peng, Zifan Peng, Qige Qi, Shi Qiu, Xing-	1680
1621	jeev Kwatra, Sarah A. Rous, Sarik Ghazarian, Sayan	wei Qu, Yizhou Tan, Zili Wang, Chenqing Wang,	1681
1622	Ghosh, Sean Casey, Sebastian Bischoff, Sebastian	Hao Wang, Yiya Wang, Yubo Wang, Jiajun Xu,	1682
1623	Gehrmann, Sebastian Schuster, Sepideh Sadeghi,	Kexin Yang, Ruibin Yuan, Yuanhao Yue, Tianyang	1683
1624	Shadi Hamdan, Sharon Zhou, Shashank Srivastava,	Zhan, Chun Zhang, Jingyang Zhang, Xiyue Zhang,	1684
1625	Sherry Shi, Shikhar Singh, Shima Asaadi, Shixi-	Xingjian Zhang, Yue Zhang, Yongchi Zhao, Xiangyu	1685
1626	ang Shane Gu, Shubh Pachchigar, Shubham Tosh-	Zheng, Chenghua Zhong, Yang Gao, Zhoujun Li,	1686
1627	niwal, Shyam Upadhyay, Shyamolima Shammie	Dayiheng Liu, Qian Liu, Tianyu Liu, Shiwen Ni, Jun-	1687
1628	Debnath, Siamak Shakeri, Simon Thormeyer, Si-	ran Peng, Yujia Qin, Wenbo Su, Guoyin Wang, Shi	1688
1629	mone Melzi, Siva Reddy, Sneha Priscilla Makini,	Wang, Jian Yang, Min Yang, Meng Cao, Xiang Yue,	1689
1630	Soo-Hwan Lee, Spencer Torene, Sriharsha Hatwar,	Zhaoxiang Zhang, Wangchunshu Zhou, Jiaheng Liu,	1690
1631	Stanislas Dehaene, Stefan Divic, Stefano Ermon,	Qunshu Lin, Wenhao Huang, and Ge Zhang. 2025b.	1691
1632	Stella Biderman, Stephanie Lin, Stephen Prasad,	Supergpqa: Scaling llm evaluation across 285 gradu-	1692
1633	Steven Piantadosi, Stuart Shieber, Summer Mish-	ate disciplines . <i>Preprint</i> , arXiv:2502.14739.	1693
1634	erghi, Svetlana Kiritchenko, Swaroop Mishra, Tal		
1635	Linzen, Tal Schuster, Tao Li, Tao Yu, Tariq Ali,	Qwen Team. 2025b. Qwq-32b: Embracing the power	1694
1636	Tatsunori Hashimoto, Te-Lin Wu, Théo Desbor-	of reinforcement learning .	1695
1637	des, Theodore Rothschild, Thomas Phan, Tianle		
1638	Wang, Tiberius Nkinyili, Timo Schick, Timofei Ko-	Joshua Vendrow, Edward Vendrow, Sara Beery, and	1696
1639	rnev, Titus Tunduny, Tobias Gerstenberg, Trenton	Aleksander Madry. 2025. Do large language	1697
1640	Chang, Trishala Neeraj, Tushar Khot, Tyler Shultz,	model benchmarks test reliability? <i>Preprint</i> ,	1698
1641	Uri Shaham, Vedant Misra, Vera Demberg, Victo-	arXiv:2502.03461.	1699
1642	ria Nyamai, Vikas Raunak, Vinay Venkatesh Ra-		
1643	masesh, vinay uday prabhu, Vishakh Padmakumar,	Yiming Wang, Pei Zhang, Jialong Tang, Haoran Wei,	1700
1644	Vivek Srikumar, William Fedus, William Saunders,	Baosong Yang, Rui Wang, Chenshu Sun, Feitong Sun,	1701
1645	William Zhang, Wout Vossen, Xiang Ren, Xiaoyu	Jiran Zhang, Junxuan Wu, Qiqian Cang, Yichang	1702
1646	Tong, Xinran Zhao, Xinyi Wu, Xudong Shen, Yadol-	Zhang, Fei Huang, Junyang Lin, Fei Huang, and Jin-	1703
1647	lah Yaghoobzadeh, Yair Lakretz, Yangqiu Song,	gren Zhou. 2025. Polymath: Evaluating mathemat-	1704
1648	Yasaman Bahri, Yejin Choi, Yichi Yang, Sophie	ical reasoning in multilingual contexts . <i>Preprint</i> ,	1705
1649	Hao, Yifu Chen, Yonatan Belinkov, Yu Hou, Yufang	arXiv:2504.18428.	1706
1650	Hou, Yuntao Bai, Zachary Seid, Zhuoye Zhao, Zijian		
1651	Wang, Zijie J. Wang, Zirui Wang, and Ziyi Wu. 2023.	Yubo Wang, Xueguang Ma, Ge Zhang, Yuansheng Ni,	1707
1652	Beyond the imitation game: Quantifying and extrap-	Abhranil Chandra, Shiguang Guo, Weiming Ren,	1708
1653	olating the capabilities of language models . <i>Trans-</i>	Aaran Arulraj, Xuan He, Ziyang Jiang, Tianle Li, Max	1709
1654	<i>actions on Machine Learning Research</i> . Featured	Ku, Kai Wang, Alex Zhuang, Rongqi Fan, Xiang	1710
1655	Certification.	Yue, and Wenhui Chen. 2024. MMLU-pro: A more	1711
		robust and challenging multi-task language under-	1712
		standing benchmark . In <i>The Thirty-eight Conference</i>	1713
1656	Google Team. 2025a. Gemma 3 technical report .	<i>on Neural Information Processing Systems Datasets</i>	1714
		<i>and Benchmarks Track</i> .	1715
1657	Kanana LLM Team, Yunju Bak, Hojin Lee, Minho Ryu,		
1658	Jiyeon Ham, Seungjae Jung, Daniel Wontae Nam,	Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten	1716
1659	Taegyeong Eo, Donghun Lee, Doohae Jung, Boseop	Bosma, brian ichter, Fei Xia, Ed H. Chi, Quoc V Le,	1717
1660	Kim, Nayeon Kim, Jaesun Park, Hyunho Kim, Hyun-	and Denny Zhou. 2022. Chain of thought prompt-	1718
1661	woong Ko, Changmin Lee, Kyoung-Woon On, Seu-	ing elicits reasoning in large language models . In	1719
1662	lye Baeg, Junrae Cho, Sunghee Jung, Jieun Kang,	<i>Advances in Neural Information Processing Systems</i> .	1720
1663	EungGyun Kim, Eunhwa Kim, Byeongil Ko, Daniel		
1664	Lee, Minchul Lee, Miok Lee, Shinbok Lee, and	Colin White, Samuel Dooley, Manley Roberts, Arka Pal,	1721
		Benjamin Feuer, Siddhartha Jain, Ravid Shwartz-Ziv,	1722

1723	Neel Jain, Khalid Saifullah, Sreemanti Dey, Shubh-	Sangho Choi, Seongjae Choi, Wooyong Choi, Se-	1784
1724	Agrawal, Sandeep Singh Sandha, Siddhartha Venkat	whan Chun, Dong Young Go, Chiheon Ham, Danbi	1785
1725	Naidu, Chinmay Hegde, Yann LeCun, Tom Gold-	Han, Jaemin Han, Moonyoung Hong, Sung Bum	1786
1726	stein, Willie Neiswanger, and Micah Goldblum. 2025.	Hong, Dong-Hyun Hwang, Seongchan Hwang, Jin-	1787
1727	Livebench: A challenging, contamination-limited	bae Im, Hyuk Jin Jang, Jaehyung Jang, Jaeni Jang,	1788
1728	LLM benchmark . In <i>The Thirteenth International</i>	Sihyeon Jang, Sungwon Jang, Joonha Jeon, Daun	1789
1729	<i>Conference on Learning Representations</i> .	Jeong, Joonhyun Jeong, Kyeongseok Jeong, Mini	1790
		Jeong, Sol Jin, Hanbyeol Jo, Hanju Jo, Minjung	1791
1730	xAI. 2025. Grok 3 beta — the age of reasoning agents .	Jo, Chaeyoon Jung, Hyungsik Jung, Jaeuk Jung,	1792
		Ju Hwan Jung, Kwangsun Jung, Seungjae Jung, Soon-	1793
1731	An Yang, Anfeng Li, Baosong Yang, Beichen Zhang,	won Ka, Donghan Kang, Soyoung Kang, Taeho	1794
1732	Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao,	Kil, Areum Kim, Beomyoung Kim, Byeongwook	1795
1733	Chengen Huang, Chenxu Lv, Chujie Zheng, Dayi-	Kim, Daehee Kim, Dong-Gyun Kim, Donggook	1796
1734	heng Liu, Fan Zhou, Fei Huang, Feng Hu, Hao Ge,	Kim, Donghyun Kim, Euna Kim, Eunuchul Kim, Gee-	1797
1735	Haoran Wei, Huan Lin, Jialong Tang, Jian Yang, Jian-	wook Kim, Gyu Ri Kim, Hanbyul Kim, Heesu Kim,	1798
1736	hong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang,	Isaac Kim, Jeonghoon Kim, Jihye Kim, Joonghoon	1799
1737	Jing Zhou, Jingren Zhou, Junyang Lin, Kai Dang,	Kim, Minjae Kim, Minsub Kim, Pil Hwan Kim,	1800
1738	Keqin Bao, Kexin Yang, Le Yu, Lianghao Deng,	Sammy Kim, Seokhun Kim, Seonghyeon Kim, Soo-	1801
1739	Mei Li, Mingfeng Xue, Mingze Li, Pei Zhang, Peng	jin Kim, Soong Kim, Soyeon Kim, Sunyoung Kim,	1802
1740	Wang, Qin Zhu, Rui Men, Ruize Gao, Shixuan Liu,	Taeho Kim, Wonho Kim, Yoonsik Kim, You Jin Kim,	1803
1741	Shuang Luo, Tianhao Li, Tianyi Tang, Wenbiao Yin,	Yuri Kim, Beomseok Kwon, Ohsung Kwon, Yoo-	1804
1742	Xingzhang Ren, Xinyu Wang, Xinyu Zhang, Xu-	Hwan Kwon, Anna Lee, Byungwook Lee, Changho	1805
1743	ancheng Ren, Yang Fan, Yang Su, Yichang Zhang,	Lee, Daun Lee, Dongjae Lee, Ha-Ram Lee, Hodong	1806
1744	Yinger Zhang, Yu Wan, Yuqiong Liu, Zekun Wang,	Lee, Hwiyeong Lee, Hyunmi Lee, Injae Lee, Jae-	1807
1745	Zeyu Cui, Zhenru Zhang, Zhipeng Zhou, and Zi-	ung Lee, Jeongsang Lee, Jisoo Lee, Jongsoo Lee,	1808
1746	han Qiu. 2025. Qwen3 technical report . <i>Preprint</i> ,	Joongjae Lee, Juhan Lee, Jung Hyun Lee, Junghoon	1809
1747	arXiv:2505.09388.	Lee, Junwoo Lee, Se Yun Lee, Sujin Lee, Sung-	1810
		jae Lee, Sungwoo Lee, Wonjae Lee, Zoo Hyun	1811
1748	Kang Min Yoo, Jaegun Han, Sookyo In, Hee-	Lee, Jong Kun Lim, Kun Lim, Taemin Lim, Nuri	1812
1749	won Jeon, Jisu Jeong, Jaewook Kang, Hyunwook	Na, Jeongyeon Nam, Kyeong-Min Nam, Yeonseog	1813
1750	Kim, Kyung-Min Kim, Munhyong Kim, Sungju	Noh, Biro Oh, Jung-Sik Oh, Solgil Oh, Yeontaek	1814
1751	Kim, Donghyun Kwak, Hanock Kwak, Se Jung	Oh, Boyoun Park, Cheonbok Park, Dongju Park,	1815
1752	Kwon, Bado Lee, Dongsoo Lee, Gichang Lee,	Hyeonjin Park, Hyun Tae Park, Hyunjung Park, Ji-	1816
1753	Jooho Lee, Baeseong Park, Seongjin Shin, Joon-	hye Park, Jooseok Park, Junghwan Park, Jungsoo	1817
1754	sang Yu, Seolki Baek, Sumin Byeon, Eungsup	Park, Miru Park, Sang Hee Park, Seunghyun Park,	1818
1755	Cho, Dooseok Choe, Jeesung Han, Youngkyun	Soyoung Park, Taerim Park, Wonkyeong Park, Hyun-	1819
1756	Jin, Hyein Jun, Jaeseung Jung, Chanwoong Kim,	joon Ryu, Jeonghun Ryu, Nahyeon Ryu, Soonshin	1820
1757	Jinhong Kim, Jinuk Kim, Dokyeong Lee, Dong-	Seo, Suk Min Seo, Yoonjeong Shim, Kyuyong Shin,	1821
1758	wook Park, Jeong Min Sohn, Sujung Han, Jiae	Wonkwang Shin, Hyun Sim, Woongseob Sim, Hyejin	1822
1759	Heo, Sungju Hong, Mina Jeon, Hyunhoon Jung,	Soh, Bokyeong Son, Hyunjun Son, Seulah Son, Chi-	1823
1760	Jungeun Jung, Wangkyo Jung, Chungjoon Kim, Hy-	Yun Song, Chiyoung Song, Ka Yeon Song, Minchul	1824
1761	eri Kim, Jonghyun Kim, Min Young Kim, Soeun	Song, Seungmin Song, Jisung Wang, Yonggoo Yeo,	1825
1762	Lee, Joonhee Park, Jieun Shin, Sojin Yang, Jung-	Myeong Yeon Yi, Moon Bin Yim, Taehwan Yoo,	1826
1763	soon Yoon, Hwaran Lee, Sanghwan Bae, Jeehwan	Yoonjoon Yoo, Sungmin Yoon, Young Jin Yoon,	1827
1764	Cha, Karl Gylleus, Donghoon Ham, Mihak Hong,	Hangyeol Yu, Ui Seon Yu, Xingdong Zuo, Jeon-	1828
1765	Youngki Hong, Yunki Hong, Dahyun Jang, Hyo-	gin Bae, Joungeun Bae, Hyunsoo Cho, Seonghyun	1829
1766	jun Jeon, Yujin Jeon, Yeji Jeong, Myunggeun Ji,	Cho, Yongjin Cho, Taekyoon Choi, Yera Choi, Ji-	1830
1767	Yeguk Jin, Chansong Jo, Shinyoung Joo, Seungh-	wan Chung, Zhenghui Han, Byeongho Heo, Euisuk	1831
1768	wan Jung, Adrian Jungmyung Kim, Byoung Hoon	Hong, Taebaek Hwang, Seonyeol Im, Sumin Jegal,	1832
1769	Kim, Hyomin Kim, Jungwhan Kim, Minkyoung	Sumin Jeon, Yelim Jeong, Yonghyun Jeong, Can	1833
1770	Kim, Minseung Kim, Sungdong Kim, Yonghee Kim,	Jiang, Juyong Jiang, Jiho Jin, Ara Jo, Younhyun	1834
1771	Youngjun Kim, Youngkwan Kim, Donghyeon Ko,	Jo, Hoyoun Jung, Juyoung Jung, Seunghyeong Kang,	1835
1772	Dugyun Lee, Ha Young Lee, Jaehong Lee, Jieun	Dae Hee Kim, Ginam Kim, Hangyeol Kim, Heeseung	1836
1773	Lee, Jonghyun Lee, Jongjin Lee, Min Young Lee,	Kim, Hyojin Kim, Hyojun Kim, Hyun-Ah Kim, Jee-	1837
1774	Yehbin Lee, Taehong Min, Yuri Min, Kiyoon Moon,	hye Kim, Jin-Hwa Kim, Jiseon Kim, Jonghak Kim,	1838
1775	Hyangnam Oh, Jaesun Park, Kyuyon Park, Younghun	Jung Yoon Kim, Rak Yeong Kim, Seongjin Kim,	1839
1776	Park, Hanbae Seo, Seunghyun Seo, Mihyun Sim,	Seoyoon Kim, Sewon Kim, Sooyoung Kim, Suky-	1840
1777	Gyubin Son, Matt Yeo, Kyung Hoon Yeom, Won-	oung Kim, Taeyong Kim, Naeun Ko, Bonseung Koo,	1841
1778	joon Yoo, Myungin You, Doheon Ahn, Homin Ahn,	Heeyoung Kwak, Haena Kwon, Youngjin Kwon, Bo-	1842
1779	Joohee Ahn, Seongmin Ahn, Chanwoo An, Hy-	ram Lee, Bruce W. Lee, Dageyeong Lee, Erin Lee,	1843
1780	eryun An, Junho An, Sang-Min An, Boram Byun,	Euijin Lee, Ha Gyeong Lee, Hyojin Lee, Hyun-	1844
1781	Eunbin Byun, Jongho Cha, Minji Chang, Seung-	jeong Lee, Jeeyoon Lee, Jeonghyun Lee, Jongheok	1845
1782	gyu Chang, Haesong Cho, Youngdo Cho, Dalnim	Lee, Joonhyung Lee, Junhyuk Lee, Mingu Lee,	1846
1783	Choi, Daseul Choi, Hyoseok Choi, Minseong Choi,	Nayeon Lee, Sangkyu Lee, Se Young Lee, Seulgi	1847

Lee, Seung Jin Lee, Suhyeon Lee, Yeonjae Lee, Yesol Lee, Youngbeom Lee, Yujin Lee, Shaocong Li, Tianyu Liu, Seong-Eun Moon, Taehong Moon, Max-Lasse Nihlenramstroem, Wonseok Oh, Yuri Oh, Hongbeen Park, Hyekyung Park, Jaeho Park, Nohil Park, Sangjin Park, Jiwon Ryu, Miru Ryu, Simo Ryu, Ahreum Seo, Hee Seo, Kangdeok Seo, Jamin Shin, Seungyoun Shin, Heetae Sin, Jiangping Wang, Lei Wang, Ning Xiang, Longxiang Xiao, Jing Xu, Seonyeong Yi, Haanju Yoo, Haneul Yoo, Hwanhee Yoo, Liang Yu, Youngjae Yu, Weijie Yuan, Bo Zeng, Qian Zhou, Kyunghyun Cho, Jung-Woo Ha, Joonsuk Park, Jihyun Hwang, Hyoung Jo Kwon, Soonyong Kwon, Jungyeon Lee, Seungho Lee, Seonghyeon Lim, Hyunkyung Noh, Seungho Choi, Sang-Woo Lee, Jung Hwa Lim, and Nako Sung. 2024. [Hyperclova x technical report](#). *Preprint*, arXiv:2404.01954.

Rowan Zellers, Ari Holtzman, Yonatan Bisk, Ali Farhadi, and Yejin Choi. 2019. [HellaSwag: Can a machine really finish your sentence?](#) In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 4791–4800, Florence, Italy. Association for Computational Linguistics.

Yidan Zhang, Yu Wan, Boyi Deng, Baosong Yang, Hao-ran Wei, Fei Huang, Bowen Yu, Junyang Lin, Fei Huang, and Jingren Zhou. 2025. [P-mmeval: A parallel multilingual multitask benchmark for consistent evaluation of llms](#). *Preprint*, arXiv:2411.09116.

Qihao Zhao, Yangyu Huang, Tengchao Lv, Lei Cui, Qinzhen Sun, Shaoguang Mao, Xin Zhang, Ying Xin, Qiufeng Yin, Scarlett Li, and Furu Wei. 2024. [Mmlu-cf: A contamination-free multi-task language understanding benchmark](#). *Preprint*, arXiv:2412.15194.

Lianmin Zheng, Liangsheng Yin, Zhiqiang Xie, Chuyue Sun, Jeff Huang, Cody Hao Yu, Shiyi Cao, Christos Kozyrakis, Ion Stoica, Joseph E. Gonzalez, Clark Barrett, and Ying Sheng. 2024. [SGLang: Efficient execution of structured language model programs](#). In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*.

A Details of KMMLU-REDUX Errors

A.1 Error Statistics

We define a set of error types based on recurring issues observed during our analysis of the KMMLU. Then, we find out the number of data instances per each error type as shown in Table 4. To identify leaked answer cases, we apply rule-based filtering using string-overlap heuristics. For other error types, including notation errors, bad clarity, and ill-posed questions, we leverage GPT-4o to assist in annotation. Also, we provide the prompt used for annotation in Figure 5.

- **Ill-posed Question** : The question lacks critical references or contextual information.
- **Leaked Answer** : The ground truth is explicitly stated within the question itself.
- **Notation Error** : Errors in mathematical expressions or chemical equations.
- **Bad Clarity** : The data itself is unclear and contains grammatical errors.

Error Type	# of Questions (Ratio)
Ill-posed Question	512 (1.46 %)
Leaked Answer	42 (0.11%)
Notation Error	846 (2.42%)
Bad Clarity	1284 (3.67%)

Table 4: Statistics of the error types

A.2 Example of Error Types

Table 10 provides examples of the error types in 2.2 identified in KMMLU. Each example illustrates a specific issue that affects benchmark reliability.

B Korean National Technical Qualification List of KMMLU-REDUX

Table 9 presents the list of Korean National Technical Qualifications (KNTQs) included in our benchmark along with their corresponding official exam dates. To facilitate analysis, we categorize the 100 NTQs into Korean Standard Industrial Classification (KSIC) where mapping to the exam (see Figure 6). This categorization enables structured evaluation across diverse domains and better reflects the real-world industrial fields.

GPT-4o Error Annotation Prompt

You are a helpful assistant that annotates error types in questions.

- Ill-posed question stands for which question miss the critical reference information to solve the question (such as table, formular, image, etc.).
- Notation correct stands for which question has correct math and chemical notation. Criteria of correctness whether the notation is impact to solve the question. (e.g. m2 is incorrect, but m^2 is correct. 센티미터 is correct.)
- Grammar correct stands for which question has correct grammar and spelling. Criteria of correctness whether the notation is impact to understand the question. (e.g. 상담기법보다는 상담자의 인간적 자질과 진솔한 태도를 중시한다. is grammatically incorrect due to spacing error.)

Please check the following question whether it has error types.

Check ill-posed question True/False.

Check notation correct True/False.

Check grammar correct True/False.

Then, if there is any error, please explain the reason.

Question

{{question}}

Figure 5: The prompt is used for error type annotation. Each sample is annotated as an error if the respective field returns True. The 'Grammar Correct' field is used to detect 'Bad Clarity' cases.

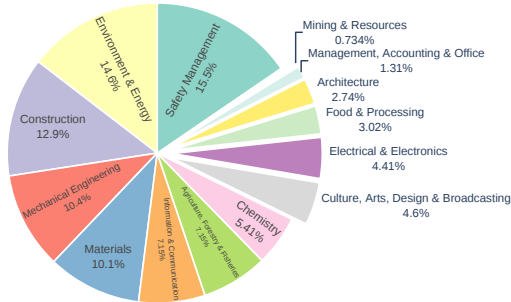


Figure 6: Domain distribution of problems in KMMLU-REDUX. The total size of the dataset is 2,587.

C Details of KMMLU-PRO Annotation

C.1 Annotation Pipeline

We first conduct OCR parsing with GPT-4o on PDF files of KNPL acquisition exams. With the parsed data, the main tasks for human annotators are: 1) reviewing parsing errors, 2) converting tables into latex format, and 3) converting images into text which conveys same meaning. If it is impossible to convert an image into text, we remove it. For the cases where multiple answers are allowed, commonly due to the ambiguity of question itself, we discard them.

As we leverage the official PDF files managed and controlled by the government, we can guarantee the correctness of answer label. This help us to save costs because we do not need experts for annotations nor the answer relabeling to avoid risk of data from online (Gema et al., 2025; Team et al., 2025b). Before annotation, we explained the context of benchmark construction about the Korean professional license exams to human annotators. We present annotation instructions in Fig 7.

KMMLU-Pro Annotation Guideline

작업 개괄

 요약: pdb 로 부터 추출된 json 이 잘 작성 되었는지 검수 및 수정하는 테스트입니다

- [illegible]

KMMLU-Pro Annotation Guideline

KMMLU-Pro Annotation Guideline

표/수식 검수 방법

2. 생성된 프로젝트 패에가져서 code editor 부분에 기본 템플릿을 붙여넣습니다.
 - 기본 템플릿

```
(documentclass{article})
\usepackage{graphicx} % Required for inserting images
\usepackage{notes}
\usepackage{booktabs, multirow, multicol, amsmath}

\begin{document}

\end{document}
```

4. `\begin{document}` `\end{document}` 사이에서 latex 문법으로 적힌 부분들을 적어넣습니다.
- 만약, **java** 파일에 있는 **표/수식들**을 **overleaf**에 옮겨서 **간수할** 경우, 아래를 잘 지켜주세요 합니다.
 - 1. 복사해줘서 모두 두 이렇게 두개의 세지기를 짓습니다. 이를 모두, 1차분 바리우셔야 합니다
 - 2. **인**를 넣어주셔야 합니다
- 이 사항들은 **갑수할때만** 바뀝니다. **갑수한** 확인된, **넣을**할 뿐, **원래** 값으로 **넣을**주세요.
5. '다시 컴파일하기' 버튼을 눌러 오른쪽 화면에 생성된 **표/수식**을 원본 **pdf**의 **표/수식**과 비교합니다.

표 데이터 제작 방법

1. <https://www.tahagenerator.com/> 에 들어갑니다.
2. 원본과 동일한 포문 제작입니다. (이 때, 테두리, bold, 밑줄, 열 병합, 글자 정렬 등 호프 최대한 재현
수요, 원본

	종류	반응 시간(min)			
		1	2	3	4
순환 반응 속도 (분당량)	1회	66	105		x
	2회	115		4	

생식전 다이어트



3. **이름, 수식적 명사, label로 표현하고, Noun으로 읽기**
3. **Generative 버전은 골짜기와 Result 하얀색 위안제 코드를 보자하면 볼까?**
4. **의미적으로,**
- a. 이렇게 생긴 **label** 텍스트는 모두 이렇게 한개에 보자를 갖는다. **날짜와 월, 이름 모두 두개에 하위주어 하나나** (그렇지 않거나, Jean 에만)
 - b. **출처한 원 문장을 하나의 문장으로 만들어주나** (그렇지 않거나, Jean 에만)
 - c. **일모두 같은 문장을 지워주세요**, 일모두는 문장을 삭제해 갑니다.
- a. % 로 시작하는 주석 line
 - b. \caption으로 시작하는 line
 - c. \label로 시작하는 line

we use `r"정답[^A-E]*:[^A-E]*([A-E])"`. The regex expressions for the flexible parsing are `r"Answer[^A-E]*([A-E])|([A-E])\\"` and `r"정답[^A-E]*([A-E])|([A-E])\\"`, respectively.

D.2 Inference Engines

Excluding closed models, the main inference engine we use is SGLang (Zheng et al., 2024). However, for the Gemma 3 series and Mistral Small 3.1 Instruct models, we adopt vLLM (Kwon et al., 2023) due to their incompatibility with SGLang at the time of the paper writing. Evaluating open-sourced models was conducted over four days using sixteen NVIDIA A100 GPUs. For closed models, we accessed them via API calls, which incurred a total cost of approximately 4,000 USD.

D.3 License Acquisition Criteria for KMMLU-PRO

We follow the official scoring criteria used for each license examination. All licensing exams in KMMLU-PRO are composed of multiple subjects. Candidates are typically required to score at least 40% in every subject and achieve an average score of at least 60%, except for the cases of the Certified

Figure 7: Excerpt from the annotation guidelines for converting PDF documents into structured text. We carefully instruct LaTeX table formatting.

Category	Demographics	Counts
Gender	Female	23
	Male	-
Age	20s	13
	30s	8
	40s	2
Academic background	Bachelor's degrees	20
	Master's degrees	3

Table 5: Demographics of the human annotators for KMMLU-PRO.

C.2 Annotators Demographics

The detailed demographics of annotators are presented in Table 5. The average hourly wage is 8.83 U.S. dollars, which is higher than the legal minimum wage at the time of hiring in South Korea.

D Evaluation Setup

D.1 Evaluation Prompts

The figure 8 and figure 9 present the prompt for the evaluation written in English and Korean, respectively. For the English prompt, we use the regex expression of `r"(?i)Answer[^\-E]*:[^\-E]*([^\-E])"`. For the Korean prompt,

English Prompt

Answer the following multiple choice question. The last line of your response should be of the following format: 'Answer: \$LETTER' (without quotes) where LETTER is one of ABCD. Think step by step before answering.

{{question}}

- A) {{option A}}
- B) {{option B}}
- C) {{option C}}
- D) {{option D}}

Figure 8: The English prompt used for evaluating LLMs on our KMMLU-REDUX and KMMLU-PRO. This prompt is exactly same with the prompt used for Multiple-Choices Question Answering (MCQA) in OpenAI’s simple-evals repository. The number of options is adjusted for each problems.

Korean Prompt

다음 문제에 대해 정답을 고르세요. 당신의 최종 정답은 ABCD 중 하나이고, "정답:" 뒤에 와야 합니다. 정답을 고르기 전에 차근차근 생각하고 추론하세요.

{{question}}

- A) {{option A}}
- B) {{option B}}
- C) {{option C}}
- D) {{option D}}

Figure 9: The Korean prompt used for evaluating LLMs on our KMMLU-REDUX and KMMLU-PRO. This prompt is translated version of the English prompt. The number of options is adjusted for each problems.

Judicial Scrivener and the Lawyer. For the Certified Judicial Scrivener exam, candidates need only score at least 40% in each subject, with no requirement regarding the average. The Lawyer license exam uses a relative grading system where only a certain proportion of top-scoring candidates pass. The usual cut-off point is approximately 54.22 (900 out of 1660), which we use as the passing threshold.

It is also important to note that our evaluation benchmark is text-based, not multi-modal. Therefore, we exclude questions that include images. In addition, candidates often need to go through multiple exam stages to obtain a license, with the later stages, such as the second or third, containing descriptive questions. However, since we only collect multiple-choice questions, the descriptive problems

are excluded. Lastly, we exclude questions with multiple answers introduced by the ambiguity of the question.

E Detailed Results

E.1 KMMLU vs KMMLU-REDUX

Figure 10 illustrates the performance of LLMs on KMMLU and KMMLU-REDUX. The LLMs’ performances on KMMLU-REDUX is lower than on KMMLU, due to our filtration process which aims to retain only challenging problems from KMMLU (see Section 2.2.1). Despite this decrease, there is a near-perfect monotonic association between the results, with a Spearman’s rank correlation coefficient (ρ) of 0.995, suggesting they are highly correlated.

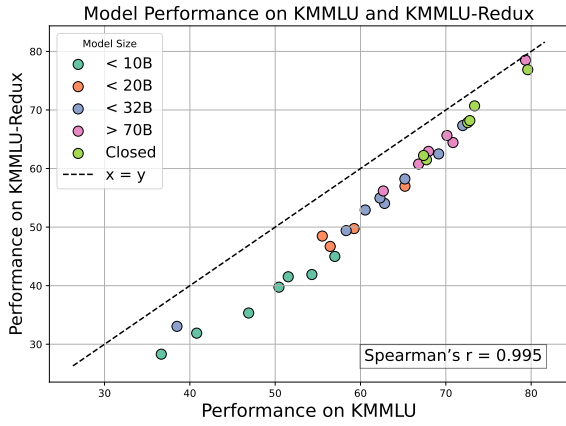


Figure 10: Performance of LLMs on KMMLU and KMMLU-REDUX. A high ρ value indicates a strong correlation between the results of the two benchmarks.

E.2 The Results for Smaller ($<10B$) Models

The Table 6 presents the results of KMMLU-PRO and KMMLU-REDUX for smaller ($<10B$) models. Since many of the *tiny* models in this table were used to construct KMMLU-REDUX through adversarial filtration, their KMMLU-REDUX scores are biased. Nevertheless, as shown by the results for larger models, models equipped with dense reasoning capabilities usually outperform their counterparts without reasoning (e.g., Qwen3-8B with and without “thinking”).

E.3 Breakdown of Results of KMMLU-REDUX

The Table 7 presents the breakdown results for 14 categories in KMMLU-REDUX across various LLMs.

E.4 Breakdown of Results of KMMLU-PRO

The Table 8 shows the breakdown results for all NPLs in KMMLU-PRO across various LLMs. While the models relatively easily pass licensing in the medicine domain, they struggle in the Law and Tax&Accounting domains.

	KMMLU-Redux Acc	KMMLU-Pro Acc	# of passed KNPLs	Avg. Acc (micro)
DeepSeek-R1-Distill-Qwen-1.5B (DeepSeek-AI et al., 2025a)	21.30*	20.55	0/14	20.91
Llama 3.2 3B Instruct (Meta, 2024c)	17.59*	25.53	0/14	21.73
Gemma 3 4B IT (Team, 2025a)	25.09*	32.86	0/14	29.14
Qwen 2.5 3B Instruct (Qwen et al., 2025)	24.74*	33.27	0/14	29.19
Qwen3-1.7B (Yang et al., 2025)	28.99	30.42	0/14	29.74
Kanana Nano 2.1B Instruct (Team et al., 2025a)	27.25*	32.60	0/14	30.04
Aya Expanse 8B (Dang et al., 2024)	28.30	31.65	0/14	30.05
EXAONE Deep 2.4B (Research et al., 2025)	28.84*	32.08	0/14	30.53
EXAONE 3.5 2.4B Instruct (Research et al., 2024)	27.72*	34.49	0/14	31.25
HyperCLOVAX-SEED-Text-Instruct-1.5B (naver hyperclovax, 2025)	33.94	30.13	0/14	31.95
Llama 3.1 8B Instruct (Meta, 2024b)	31.89	33.81	0/14	32.89
Qwen3-1.7B (w/ thinking) (Yang et al., 2025)	37.80	38.27	1/14	38.05
Ko-R1-7B-v2.1 (OneLineAI)	41.94	38.70	1/14	40.25
EXAONE 3.5 7.8B Instruct (Research et al., 2024)	41.90	41.71	0/14	41.80
EXAONE Deep 7.8B (Research et al., 2025)	44.99	41.53	0/14	43.18
Qwen3-8B (Yang et al., 2025)	49.25	46.92	1/14	48.03
Qwen3-8B (w/ thinking) (Yang et al., 2025)	58.79	55.27	3/14	56.95

Table 6: The main evaluation results of KMMLU-REDUX and KMMLU-PRO benchmarks on smaller (< 10B) LLMs. The gray-shaded models are the dense-reasoning models. The KMMLU-REDUX scores with * are biased as these models are used for the dataset filtration (Section 2.2.1).

Domain		Safety Management	Environment & Energy	Construction	Mechanical Engineering	Materials	Information & Communication	Agriculture, Forestry & Fisheries	Chemistry	Culture, Arts, Design & Broadcasting	Electrical & Electronics	Food & Processing	Architecture	Management, Accounting & Office	Mining & Resources	Avg.
Opensourced Models																
< 5B	DeepSeek-R1-Distill-Qwen-1.5B	20.50	17.77	24.32	20.00	24.43	26.49	17.84	20.71	19.33	20.18	15.38	18.31	20.59	31.58	21.30*
	Llama 3.2 3B Instruct	13.00	12.20	13.81	14.81	11.07	15.68	16.76	11.43	17.65	7.89	16.67	11.27	5.88	15.79	17.59*
	Qwen 2.5 3B Instruct	16.50	17.24	15.32	17.41	15.27	20.54	18.38	11.43	26.05	16.67	11.54	15.49	14.71	10.53	24.74*
	Gemma 3 4B IT	23.00	23.87	22.52	25.56	20.23	28.65	32.43	25.71	21.01	27.19	21.79	21.13	14.71	10.53	25.09*
	Kanana Nano 2.1B Instruct	24.75	27.32	25.83	26.67	29.77	22.16	26.49	23.57	25.21	28.07	28.21	28.17	26.47	10.53	27.25*
	EXAONE 3.5 2.4B Instruct	26.75	26.79	22.22	28.15	29.39	28.65	29.19	32.86	31.93	28.95	23.08	33.80	26.47	31.58	27.72*
	EXAONE Deep 2.4B	27.50	30.77	23.72	28.52	32.06	30.27	22.16	40.00	25.21	34.21	25.64	18.31	17.65	15.79	28.84*
	Qwen3-1.7B	30.50	27.06	27.93	36.67	32.82	29.73	23.24	30.00	23.53	28.95	25.64	23.94	20.59	15.79	28.99
	HyperCLOVAX-SEED-Text-Instruct-1.5B	28.00	25.99	30.93	32.59	33.21	29.19	32.97	22.86	35.29	28.07	32.05	32.39	17.65	26.32	33.94
	Qwen3-1.7B (w/ thinking)	37.25	35.54	33.03	39.63	40.08	41.62	36.76	40.71	35.29	46.49	42.31	29.58	41.18	42.11	37.80
< 10B	Aya Expanse 8B	25.00	20.42	23.72	21.11	24.05	29.19	29.19	21.43	23.53	22.81	20.51	30.99	26.47	21.05	28.30
	Llama 3.1 8B Instruct	30.50	25.99	23.12	28.15	30.53	33.51	32.97	22.14	35.29	24.56	29.49	26.76	29.41	31.58	31.89
	EXAONE 3.5 7.8B Instruct	41.75	37.40	41.14	44.81	42.37	48.11	41.62	40.00	47.06	42.11	38.46	46.48	38.24	21.05	41.90
	Ko-R1-7B-v2.10	42.50	37.93	39.94	41.48	38.55	50.81	37.30	52.14	49.58	50.00	33.33	33.80	50.00	36.84	41.94
	EXAONE Deep 7.8B	40.75	45.09	40.54	49.26	46.56	51.35	31.89	57.14	42.02	55.26	47.44	39.44	44.12	26.32	44.99
	Qwen3-8B	47.00	48.01	43.84	49.63	56.11	56.22	41.62	58.57	47.06	60.53	46.15	43.66	50.00	31.58	49.25
	Qwen3-8B (w/ thinking)	55.75	57.82	53.15	67.78	64.89	62.16	45.95	65.71	57.14	66.67	56.41	53.52	58.82	63.16	58.79
< 20B	Gemma 3 12B IT	44.50	42.18	46.25	44.81	46.95	50.81	51.35	50.71	57.98	48.25	46.15	39.44	50.00	42.11	46.70
	Phi-4 (14B)	48.50	47.75	48.05	52.59	49.62	52.43	44.86	58.57	54.62	56.14	50.00	36.62	52.94	36.84	49.75
	Qwen3-14B	52.75	55.70	51.95	63.33	61.45	65.95	47.57	61.43	57.98	67.54	52.56	59.15	61.76	47.37	57.25
	Qwen3-14B (w/ thinking)	64.50	66.58	60.36	70.00	72.90	68.65	54.05	73.57	63.87	75.44	51.28	57.75	70.59	68.42	65.71
< 32B	Aya Expanse 32B	30.50	32.10	27.03	38.15	35.50	39.46	37.84	28.57	33.61	35.09	32.05	35.21	26.47	21.05	33.05
	EXAONE 3.5 32B Instruct	46.75	45.89	46.25	46.30	55.34	54.05	55.14	47.14	57.98	50.00	50.00	50.70	55.88	31.58	49.40
	Mistral Small 3.1 Instruct (24B)	45.00	50.40	51.95	55.93	59.54	58.38	51.35	58.57	56.30	58.77	47.44	52.11	58.82	31.58	52.92
	Gemma 3 27B IT	49.50	51.19	47.45	57.41	58.02	62.70	49.73	60.71	59.66	62.28	56.41	47.89	67.65	31.58	54.04
	EXAONE Deep 32B	53.25	62.07	53.75	66.30	59.54	67.03	53.51	62.86	57.14	60.53	56.41	45.07	58.82	21.05	58.33
	Qwen3-30B-A3B	54.25	57.56	54.65	58.52	64.50	61.08	51.89	65.00	62.18	68.42	58.97	56.34	58.82	52.63	58.41
	Qwen3-32B	59.50	64.19	60.66	68.52	72.52	69.19	58.38	68.57	74.79	72.81	58.97	53.52	70.59	63.16	64.98
	Qwen3-30B-A3B (w/ thinking)	63.00	64.19	62.46	69.63	70.23	69.19	56.76	69.29	68.91	72.81	53.85	56.34	70.59	68.42	65.25
	QwQ 32B	61.25	67.64	68.17	71.11	74.05	72.97	58.38	71.43	72.27	71.05	53.85	56.34	79.41	52.63	67.34
	Qwen3-32B (w/ thinking)	64.50	68.97	66.97	74.07	75.57	70.81	62.70	74.29	68.91	72.81	64.10	53.52	73.53	57.89	68.77
> 70B	Llama 3.3 70B Instruct	52.25	54.11	52.85	62.59	59.92	64.32	52.43	50.71	60.50	60.53	55.13	53.52	64.71	36.84	56.17
	C4AI Command A (111B)	56.75	66.58	63.06	63.33	64.89	67.57	62.16	60.71	68.91	64.04	62.82	59.15	55.88	47.37	62.93
	DeepSeek V3 (671B)	62.50	62.33	63.66	66.30	72.52	72.97	62.70	66.43	67.23	69.30	69.23	59.15	61.76	63.16	65.64
	Llama-4-Scout-17B-16E-Instruct	64.00	67.64	65.77	69.63	77.10	71.35	61.08	69.29	66.39	71.93	62.82	57.75	64.71	57.89	67.49
	Qwen3-235B-A22B	64.25	72.68	66.37	71.11	82.06	72.43	65.41	71.43	68.07	73.68	65.38	52.11	67.65	47.37	69.54
	Qwen3-235B-A22B (w/ thinking)	69.75	75.86	71.47	81.11	85.11	76.22	64.86	75.00	70.59	78.95	73.08	61.97	82.35	68.42	74.49
	Llama-4-Maverick-17B-128E-Instruct	73.50	76.13	78.38	79.26	83.97	80.54	77.30	75.00	78.15	79.82	73.08	71.83	82.35	73.68	77.58
	DeepSeek R1 (671B)	72.50	76.39	79.58	80.74	85.11	81.62	82.70	77.14	79.83	81.58	75.64	63.38	85.29	73.68	78.51
Closed Models																
	GPT-4.1 mini (2024-04-14)	61.00	64.99	63.66	69.26	74.81	72.43	63.78	73.57	70.59	66.67	75.64	56.34	73.53	57.89	67.03
	o3-mini (2025-01-31)	63.75	66.05	63.96	72.96	73.66	75.14	63.78	69.29	72.27	75.44	66.67	46.48	76.47	57.89	67.84
	Grok-3-mini-beta	69.75	69.76	69.07	73.33	83.97	77.30	65.95	70.71	68.07	73.68	61.54	59.15	76.47	73.68	71.47
	Grok-3-beta	70.00	70.56	68.47	76.30	83.97	76.76	69.19	75.71	74.79	69.30	70.51	71.83	73.53	57.89	72.90
	o4-mini (2025-04-16)	67.75	73.74	76.28	78.52	82.06	81.08	76.76	78.57	80.67	79.82	71.79	61.97	79.41	78.95	75.80
	GPT-4.1	71.50	74.27	72.07	77.41	85.11	81.62	80.00	74.65	76.47	77.19	75.64	61.97	76.47	68.42	75.86
	Claude 3.7 Sonnet	70.75	76.92	76.28	80.37	83.97	77.84	78.92	78.57	76.47	78.07	75.64	64.79	79.41	68.42	76.88
	Claude 3.7 Sonnet (w/ thinking)	71.50	80.37	79.88	80.74	87.02	82.16	80.00	82.86	80.67	79.82	73.08	70.42	85.29	68.42	79.36
	o3	77.75	78.25	77.78	80.37	85.50	78.92	83.24	78.17	84.03	85.09	82.05	66.20	85.29	78.95	79.92
	o1	75.75	79.58	81.98	81.11	90.84	80.54	84.86	77.86	84.03	81.58	83.33	70.42	85.29	73.68	81.14

Table 7: The break down results for 14 categories in KMMLU-REDUX. The gray-shaded models are the dense-reasoning models. The scores with * are biased as these models are used for the dataset filtration (Section 2.2.1).

Domain		Law				Tax & Accounting			Value Estimation		Medical						
Names of KNPLs		Judicial Scrivener	Lawyer	Public Labor Attorney	Patent Attorney	Public Accountant (CPA)	Tax Accountant	Customs Broker	Damage Adjuster (CDA)	Appraiser	Doctor of Korean Medicine	Dentist	Pharmacist	Herb Pharmacist	Physician	Avg.	# of passed NPLs
Opensourced Models																	
< 5B	DeepSeek-R1-Distill-Qwen-1.5B	15.66	13.33	24.27	25.69	23.56	18.07	18.24	20.00	21.94	20.49	18.49	23.25	23.77	18.67	20.55	0/14
	Llama 3.2 3B Instruct	27.27	24.67	25.94	20.18	16.35	23.95	23.9	23.33	19.39	23.26	26.45	36.53	30.74	28.67	25.53	0/14
	HyperCLOVAX-SEED-Text-Instruct-1.5B	24.24	24.67	30.13	26.61	25.48	29.83	30.19	41.67	25.00	30.90	27.31	36.16	35.66	33.33	30.13	0/14
	Qwen3-1.7B	19.19	20.67	30.96	32.11	21.63	21.01	23.9	30.0	29.08	30.56	30.32	46.49	46.72	33.33	30.42	0/14
	EXAONE Deep 2.4B	16.67	21.33	31.38	36.70	25.00	23.53	32.08	37.50	30.61	30.21	29.25	52.77	44.67	32.00	32.08	0/14
	Kanana Nano 2.1B Instruct	19.19	24.67	30.13	22.02	23.08	26.47	26.42	28.33	26.53	42.01	31.61	52.40	44.67	38.67	32.60	0/14
	Gemma 3 4B IT	23.23	26.67	32.22	26.61	27.4	26.47	27.67	29.17	23.47	37.15	36.56	49.08	43.03	36.00	32.86	0/14
	Qwen 2.5 3B Instruct	20.20	25.33	35.15	24.77	21.63	23.11	30.19	35.00	27.04	37.85	35.05	47.97	52.46	34.67	33.27	0/14
	EXAONE 3.5 2.4B Instruct	28.28	24.67	32.22	28.44	33.17	21.43	27.67	35.00	31.63	36.81	35.91	53.14	43.44	38.67	34.49	0/14
	Qwen3-1.7B (w/ thinking)	25.25	25.33	31.8	34.86	25.48	25.63	29.56	34.17	36.22	44.79	38.06	61.99	59.43	44.67	38.27	1/14
< 10B	Aya Expanse 8B	20.20	20.00	30.54	22.02	23.56	18.49	37.11	33.33	22.45	36.11	29.25	50.18	46.31	42.00	31.65	0/14
	Llama 3.1 8B Instruct	28.28	21.33	30.96	23.85	21.15	23.53	33.33	33.33	26.02	35.42	33.76	54.61	50.82	42.00	33.81	0/14
	Ko-R1-7B-v2.1	40.59	28.85	38.33	39.8	35.22	33.94	26.47	29.86	42.21	29.29	48.67	46.24	64.21	30.67	38.70	1/14
	EXAONE Deep 7.8B	26.26	28.00	38.91	42.20	35.10	31.09	38.99	39.17	40.31	42.36	44.73	64.94	53.69	52.67	41.53	0/14
	EXAONE 3.5 7.8B Instruct	28.79	28.00	38.49	35.78	32.21	31.09	37.74	40.00	33.16	44.10	43.44	67.53	56.56	50.67	41.71	0/14
	Qwen3-8B	29.29	34.0	46.03	35.78	33.17	35.71	34.59	43.33	40.82	50.69	46.02	71.96	73.36	59.33	46.92	1/14
	Qwen3-14B (w/ thinking)	27.78	36.0	49.37	41.28	48.08	44.12	49.69	50.83	53.06	61.81	54.19	83.03	80.74	75.33	55.27	3/14
< 20B	Phi-4 (14B)	33.33	34.00	41.00	36.70	37.50	37.82	42.77	44.17	43.37	45.49	41.29	72.69	57.38	51.33	45.32	1/14
	Gemma 3 12B IT	28.79	23.33	40.59	39.45	34.62	34.87	42.14	39.17	35.71	52.08	49.46	71.59	62.30	68.00	45.82	2/14
	Qwen3-14B	32.83	30.67	39.75	47.71	48.56	38.24	42.14	49.17	50.00	57.64	55.91	81.18	75.0	75.33	53.02	3/14
	Qwen3-14B (w/ thinking)	30.81	36.67	48.95	47.71	58.65	51.26	52.2	47.50	54.59	65.28	61.94	87.82	80.74	82.67	59.48	3/14
< 32B	Aya Expanse 32B	24.75	18.67	24.27	32.11	20.67	21.01	27.67	34.17	23.98	32.64	31.83	53.14	45.49	38.67	31.26	0/14
	EXAONE 3.5 32B Instruct	33.33	26.67	43.10	38.53	34.62	32.35	42.77	52.50	41.33	48.96	49.89	70.11	61.89	66.00	46.71	2/14
	Mistral Small 3.1 Instruct (24B)	27.27	33.33	43.51	39.45	39.90	38.66	42.77	47.50	41.33	52.08	49.89	79.34	68.03	72.00	49.49	3/14
	Gemma 3 27B IT	29.80	31.33	44.35	33.94	38.46	43.28	40.88	43.33	44.39	52.08	58.49	81.18	73.36	72.67	51.03	2/14
	Qwen3-30B-A3B	31.31	26.00	47.28	35.78	45.19	35.71	44.03	49.17	47.96	56.60	53.98	83.03	77.87	72.00	52.33	3/14
	EXAONE Deep 32B	30.30	36.00	48.54	46.79	45.67	39.92	46.54	55.00	48.98	53.47	52.26	78.97	65.98	72.67	52.33	1/14
	Qwen3-32B	33.33	35.33	53.14	43.12	57.69	45.80	53.46	53.33	48.98	61.46	60.86	85.24	84.84	84.00	58.86	3/14
	Qwen3-30B-A3B (w/ thinking)	35.35	37.33	50.21	45.87	57.21	50.0	48.43	54.17	54.08	70.49	59.78	88.19	84.43	84.67	60.52	3/14
	Qwen3-32B (w/ thinking)	33.33	34.67	55.23	43.12	55.29	47.06	54.09	57.50	55.61	70.14	66.88	87.08	81.56	88.67	61.14	3/14
	QwQ 32B	35.35	39.33	55.65	44.04	60.58	55.04	57.23	56.67	64.29	71.53	66.24	88.93	84.43	88.67	63.94	5/14
> 70B	Llama 3.3 70B Instruct	32.83	41.33	46.44	42.20	42.31	38.66	49.69	54.17	43.37	51.39	56.77	85.98	71.72	74.00	53.24	3/14
	C4AI Command A (111B)	41.92	34.00	51.46	49.54	45.67	45.80	46.54	54.17	52.04	61.46	57.63	83.39	79.10	80.67	57.48	3/14
	Llama-4-Scout-17B-16E-Instruct	35.86	34.00	53.14	50.46	47.12	43.70	47.80	54.17	47.45	61.46	68.17	81.92	85.25	82.67	58.14	4/14
	DeepSeek V3 (671B)	38.89	36.00	56.49	48.62	50.00	42.44	47.80	52.50	53.06	71.53	64.95	87.82	87.70	84.67	60.77	4/14
	Qwen3-235B-A22B	38.38	39.33	54.81	45.87	55.77	49.58	52.83	60.0	57.65	64.24	70.97	86.72	86.07	84.67	62.12	4/14
	Llama-4-Maverick-17B-128E-Instruct	38.38	41.33	62.34	55.96	61.06	56.72	56.60	69.17	66.84	75.00	76.34	89.67	90.98	90.67	68.10	4/14
	Qwen3-235B-A22B (w/ thinking)	43.43	42.67	56.90	57.80	69.23	61.76	61.64	59.17	62.76	71.88	77.42	89.67	89.34	88.00	68.22	6/14
	DeepSeek R1 (671B)	45.45	38.00	61.51	57.80	66.83	60.50	69.18	68.33	64.80	79.51	81.29	92.99	94.67	92.67	71.33	7/14
Closed Models																	
	GPT-4.1 mini (2024-04-14)	38.38	32.00	56.90	56.88	54.81	48.32	50.31	50.00	55.10	67.36	72.47	90.77	81.97	90.00	62.18	4/14
	o3-mini (2025-01-31)	38.38	38.67	51.46	47.71	60.10	47.06	50.31	52.50	57.65	64.24	76.56	88.93	80.33	91.33	62.05	3/14
	Grok-3-mini-beta	37.37	34.67	57.74	48.62	67.79	49.58	65.41	54.17	64.80	66.32	73.98	91.14	84.84	90.0	65.08	5/14
	Grok-3-beta	42.42	41.33	59.00	55.96	63.94	57.98	62.89	61.67	66.84	70.49	77.85	89.30	91.80	94.67	68.37	7/14
	o4-mini (2025-04-16)	37.37	46.0	62.34	49.54	69.23	55.04	66.67	59.17	67.35	76.74	82.15	93.36	89.75	92.0	69.65	6/14
	GPT-4.1	47.47	50.00	66.95	63.30	66.83	59.24	67.30	66.67	68.37	80.21	82.80	94.10	92.62	94.67	72.99	10/14
	o3	45.96	41.33	71.13	57.80	73.56	61.76	70.44	65.83	71.94	77.43	87.96	92.99	91.39	94.67	73.60	9/14
	Claude 3.7 Sonnet	55.56	53.33	73.22	63.30	68.27	64.29	67.92	77.50	73.47	70.83	81.51	94.46	92.21	93.33	74.52	10/14
	Claude 3.7 Sonnet (w/ thinking)	50.00	56.00	77.41	74.31	78.85	70.17	75.47	70.00	81.12	76.39	84.09	93.36	93.03	92.67	77.70	12/14
	o1	54.55	49.33	71.55	65.14	75.48	67.23	76.73	78.33	78.06	83.33	88.39	94.10	95.49	96.67	78.09	10/14

Table 8: The break down results for all KNPLs in KMMLU-PRO. The gray-shaded models are the dense-reasoning models. The blue scores indicate that the LLM obtain the license. The details for the pass criteria of each license are described in Appendix D.3.

Year	Tests
2005	Master Craftsman Construction Equipment Maintenance, Master Craftsman Building General Work, Master Craftsman Precious Metal Processing, Master Craftsman Confectionary Making, Master Craftsman Casting, Master Craftsman Sheet-Metal & Boiler Making
2008	Master Craftsman Architectural Carpentering, Master Craftsman Surface Treatment
2010	Master Craftsman Railway Vehicles Maintenance
2014	Master Craftsman Welding
2016	Master Craftsman Steel Making
2017	Master Craftsman Metal Material, Master Craftman Metal Mould, Master Craftsman Cook
2018	Master Craftsman Gas, Master Craftsman Machinery Maintenance, Master Craftsman Plumbing, Master Craftsman Rolling, Master Craftsman Energy Management, Master Craftsman Hazardous Material, Master Craftsman Motor Vehicles Maintenance, Master Craftsman Electricity, Master Craftman Electronics, Master Craftsman Iron Making
2020	Engineer Radio Electronic Communication, Engineer Floral Design
2021	Engineer Construction Equipment, Engineer Machinery Design, Engineer Agricultural Health and Safety, Engineer Leak Nondestructive Testing, Engineer Radiation Nondestructive Testing, Engineer Biology Classification—Animal, Engineer Aquaculture, Engineer Visual Communication Design, Engineer Eddy Current Nondestructive Testing, Engineer Welding, Engineer Biomedical, Engineer Magnetic Nondestructive Testing, Engineer Electric Railway, Engineer Computer, Engineer Concrete, Engineer Explosives Handling
2022	Engineer Gas, Engineer Construction Safety, Engineer Construction Material Testing, Engineer Architecture, Engineer Building Facilities, Engineer Air-Conditioning Refrigerating Machinery, Engineer Transportation, Engineer Metal, Engineer Meteorology, Engineer Air Pollution Environmental, Engineer Urban Planning, Engineer Bioprocess, Engineer Forest, Engineer Industrial Safety, Engineer Industrial Hygiene Management, Engineer Plant Maintenance, Engineer Fire Protection System—Mechanical, Engineer Fire Protection System—Electrical, Engineer Noise & Vibration, Engineer Water Pollution Environmental, Engineer Elevator, Engineer Plant Protection, Engineer Food Processing Safety, New and Renewable Energy Equipment (Photovoltaic) Engineer, Engineer Interior Architecture, Engineer Energy Management, Engineer Greenhouse Gas Management, Engineer Organic Agriculture, Engineer Ergonomics, Engineer General Machinery, Engineer Motor Vehicles Maintenance, Engineer in Nature Environment and Ecological Restoration, Engineer Electric Work, Engineer Electricity, Engineer Computer System Application, Engineer Electronics, Engineer Information Processing, Engineer Landscape Architecture, Engineer Seeds, Engineer Cadastral Surveying, Engineer Railroad Signal Apparatus, Engineer Ultrasonic Nondestructive Testing, Engineer Livestock, Engineer Surveying Geo-Spatial Information, Engineer Penetrante Nondestructive Testing, Engineer Colorist, Engineer Civil Engineering, Engineer Soil Environment, Master Craftsman Telecommunication Apparatus, Engineer Wastes Treatment, Engineer Quality Management, Engineer Ocean Environment, Engineer Chemical Industry, Fire Investigation & Evaluation Engineer, Engineer Chemical Analysis
2023	Engineer Radio Telecommunication Equipment, Engineer Broadcasting Communication, Engineer Information Communication

Table 9: Redux Years and National Qualification Test Additions

Error Type	Examples
Ill-posed Question	Category: Political science and sociology A, B에 대한 설명으로 옳은 것만을 <보기>에서 고르면? <i>What is the correct explanation about A, B in <Reference>?</i>
Leaked Answer	Category: Ecology 산복수로에서 쌓기공작물의 높이가 3m이고, 수로깊이가 1m일 때 수로받이의 근사적 길이는? (문제 오류로 현재 복원중입니다. 보기 내용을 아시는 분들에게서는 오류 신고를 통하여 보기 작성 부탁 드립니다. 정답은 3번입니다.) <i>What is the approximate length of the culvert if the pile is 3 meters high and the channel is 1 meter deep? (This is currently being restored due to a question error. If you know the referebce, please report the error. The correct answer is 3.)</i>
Notation Error	Category: Math 다항식 $x^{2017}-1$을 x^2-x로 나누었을 때의 나머지를 $R(x)$라 할 때, $R(2017)$의 값은? <i>If the remainder of the polynomial $x^{2017}-1$ divided by x^2-x is called $R(x)$, what is the value of $R(2017)$?</i>
Bad Clarity	Category: Education 정신분석 상담과 행동주의 상담의 공통점에 해당하는 것은? A. 상담과정에서 과거 경험보다 미래 경험을 중시한 다 . B. 상담 기 법 보 다는 상담 자의 인 간 적 자 질 과 진 솔 한 태 도 를 중 시 한 다 . C. 인간의 행동을 인과적 관계로 해석하는 결정론적 관점을 가진 다 . D. 비합리적 신념을 인식하고 수정하는 논박 과정을 중시한 다 . <i>Which is a common feature of psychoanalytic counseling and behavioral counseling?</i> A. Emphasizes future experiences over past experiences in the counseling process. B. Prioritizes the counselor's human qualities and sincerity over counseling techniques. C. Interprets human behavior through a deterministic perspective based on causal relationships. D. Focuses on the disputation process to recognize and modify irrational beliefs.

Table 10: Examples of error types in KMMLU. Each example demonstrates a specific issue that impacts the reliability of the benchmark. Gray text represents translation of the examples in English