

EXAONE Deep: Reasoning Enhanced Language Models

LG AI Research*

Abstract

We present EXAONE Deep series, which exhibits superior capabilities in various reasoning tasks, including math and coding benchmarks. We train our models mainly on the reasoning-specialized dataset that incorporates long streams of thought processes. Evaluation results show that our smaller models, EXAONE Deep 2.4B and 7.8B, outperform other models of comparable size, while the largest model, EXAONE Deep 32B, demonstrates competitive performance against leading open-weight models. All EXAONE Deep models are openly available for research purposes and can be downloaded from <https://huggingface.co/LGAI-EXAONE>.



Figure 1: Overall performance comparison. EXAONE Deep 32B model demonstrates competitive performance compared to leading open-weight reasoning models such as QwQ-32B and DeepSeek-R1. It also outperforms both DeepSeek-R1-Distill-Qwen-32B and DeepSeek-R1-Distill-Llama-70B. The lightly colored regions in AIME 2024 and 2025 show the performance of majority vote (consensus).

*The complete list of authors who contributed to this work can be found in Appendix A.

1 Introduction

Recently, there has been a growing trend in research to enhance reasoning performance by adjusting computing resources during the testing phase [16]. In response to this trend, LG AI Research is introducing a new model lineup called EXAONE Deep 2.4B, 7.8B, and 32B. These models are fine-tuned versions of the EXAONE 3.5 series [7], specifically optimized for reasoning tasks. We have trained these models using three prominent techniques widely employed for fine-tuning: Supervised Fine-Tuning (SFT), Direct Preference Optimization (DPO), and Online Reinforcement Learning (Online RL).

The performance evaluation of the models indicates that, for the 2.4B variant, it demonstrates superior performance compared to DeepSeek-R1-Distill-Qwen-1.5B [3]. For the 7.8B variant, it outperforms not only open-weight models of comparable scale such as DeepSeek-R1-Distill-Qwen-7B and DeepSeek-R1-Distill-Llama-8B [3] but also a proprietary reasoning model OpenAI o1-mini [13]. In the case of the 32B model, its performance is competitive to that of the leading open-weight reasoning models such as QwQ-32B [14] and DeepSeek-R1 [3], and is superior to that of DeepSeek-R1-Distill-Qwen-32B and DeepSeek-R1-Distill-Llama-70B [3] as illustrated in Figure 1.

Thoughts become words, words become deeds, deeds become habits, habits become character, and character becomes destiny. Therefore watch the <thought>s of your mind with the sleepless eye of your mind.

- TRYON EDWARDS

2 Modeling

2.1 Data

To enhance the reasoning capabilities of language models, we have utilized 1.6M instances for SFT, 20K instances of preference data for DPO, and an additional 10K instances for Online RL. The SFT dataset contains approximately 12B tokens, with its length distribution illustrated in Figure 2. The dataset is designed to guide models in performing reasoning through an extended chain-of-thought (CoT) process, as illustrated in Figure 3.

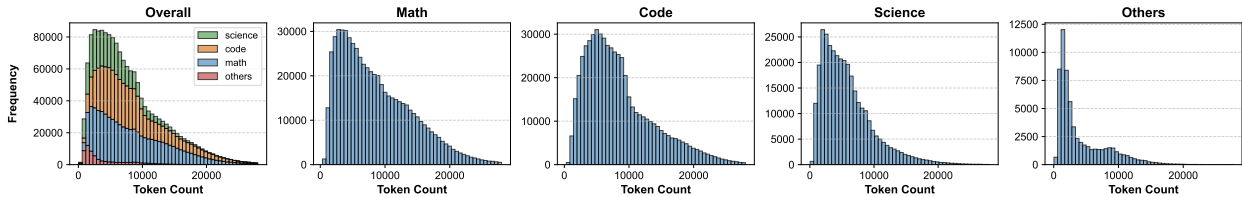


Figure 2: Distribution of token counts of the SFT dataset. Data points in the Code domain are notably longer on average, whereas those in the Others domain tend to be shorter.

2.2 Training

The base models of EXAONE Deep are EXAONE 3.5 Instruct models [8, 9, 10], which are instruction-tuned models possessing instruction-following capabilities. To enhance the reasoning abilities of EXAONE Deep, we structured SFT and DPO data in a templated format as illustrated in Figure 3. Each training instance consists of a structured thought process followed by a final answer that synthesizes the reasoning steps into a coherent and precise response. Specifically, the EXAONE 3.5 Instruct models are trained to engage in reasoning within the <thought> and </thought> tags, performing step-by-step logical progression along with reflection, self-checking, and correction. The final answer, generated after reasoning, is designed to be self-contained, summarizing the key insights derived from the thought process in a clear and concise manner. This structured approach enables EXAONE Deep to engage in robust reasoning and deliver well-founded answers to a given query. We utilize SIMPER [19] as the training algorithm for DPO and our designed GRPO [2] variant for Online RL.

In terms of training compute, the EXAONE Deep models are trained using NVIDIA H100 GPU clusters provided by Google Cloud Platform and NVIDIA NeMo Framework. The amount of computation used for pretraining of base models and fine-tuning of enhancing reasoning is presented in Table 1.

SFT Dataset Example

Query: Given the complex numbers u , v , and w such that $uv = 12 + 16i$, $vw = 20 - 15i$, and $wu = -30 + 40i$, find real numbers p and q such that $u + v + w = p + qi$. Determine the value of $p^2 + q^2$.

Response:

<thought>

Okay, so I need to find real numbers p and q such that $u + v + w = p + qi$, and then compute $p^2 + q^2$. The given information is the products of the complex numbers: $uv = 12 + 16i$, $vw = 20 - 15i$, and $wu = -30 + 40i$. Hmm, complex numbers can sometimes be tricky because of their real and imaginary parts, but maybe I can find expressions for u , v , w individually?

Let me think. If I have three equations involving the products of pairs of these complex numbers, maybe I can multiply all three equations together? Let's see: $uv \times vw \times wu = (uvw)^2$. So $(uvw)^2 = (12 + 16i)(20 - 15i)(-30 + 40i)$. If I compute that product, I can then take the square root to find uvw . But taking the square root of a complex number might give multiple values, but since we're dealing with complex numbers u , v , w , maybe we can choose a principal root or something. Wait, but maybe there's a better way.

Alternatively, if I can express each variable in terms of another. For example, from $uv = 12 + 16i$, maybe $u = (12 + 16i)/v$. Similarly, from $vw = 20 - 15i$, $w = (20 - 15i)/v$. Then from $wu = -30 + 40i$, substituting the expressions for u and w in terms of v , we get $[(20 - 15i)/v] \times [(12 + 16i)/v] = -30 + 40i$. That simplifies to $(20 - 15i)(12 + 16i)/(v^2) = -30 + 40i$. Then solving for v^2 , we get $v^2 = (20 - 15i)(12 + 16i)/(-30 + 40i)$. Then once we find v^2 , we can take square roots to find v , and subsequently find u and w .

... (thought process continues)

</thought>

For each v , we compute u and w :

Case 1: $v = \frac{9-13i}{5}$

- $u = -2 + 6i$
- $w = \frac{15+5i}{2}$

Case 2: $v = \frac{-9+13i}{5}$

- $u = 2 - 6i$
- $w = \frac{-15-5i}{2}$

Summing u , v , and w in both cases:

Case 1: $u + v + w = \frac{73}{10} + \frac{59}{10}i$

Case 2: $u + v + w = -\frac{73}{10} - \frac{59}{10}i$

In both cases, $p^2 + q^2$ is:

$$\left(\frac{73}{10}\right)^2 + \left(\frac{59}{10}\right)^2 = \frac{73^2 + 59^2}{100} = \frac{5329 + 3481}{100} = \frac{8810}{100} = \frac{881}{10}$$

Thus, the value of $p^2 + q^2$ is:

$\frac{881}{10}$

Figure 3: An example of SFT dataset. The dataset is specifically designed to facilitate models in conducting reasoning tasks through an extended chain-of-thought methodology.

Model size	Pretraining	Fine-tuning	Total
32B	1.25×10^{24}	7.04×10^{21}	1.26×10^{24}
7.8B	4.21×10^{23}	1.71×10^{21}	4.23×10^{23}
2.4B	9.36×10^{22}	5.27×10^{20}	9.41×10^{22}

Table 1: Amount of computation (FLOPs) for model training.

3 Evaluation

3.1 Benchmarks

We evaluate models on MATH-500 [11], American Invitational Mathematics Examination (AIME) 2024 and 2025 [12], the mathematics sections of South Korea’s College Scholastic Ability Test (CSAT) 2025 [6], GPQA Diamond [15], LiveCodeBench (24.08-25.02) [5], MMLU [4], and MMLU-Pro [18].

In the case of CSAT, which includes textual problems and supplementary graphical information, the graphical content is excluded when evaluating the performance. Since most of this graphical information is also available in textual descriptions, its impact on results is considered minimal. In addition, students can choose from three elective subjects — calculus, statistics, and geometry — in the mathematics section. The final score is calculated as the average of the scores obtained in these three elective subjects.

3.2 Baselines

We have conducted comprehensive evaluations against strong baseline such as DeepSeek-R1 as well as comparable-scale baselines such as QwQ-32B, DeepSeek-R1-Distill-Qwen-32B, 7B, 1.5B, DeepSeek-R1-Distill-Llama-70B, 8B, and OpenAI o1-mini (2024-09-12).

3.3 Evaluation Setup

Following the setup described in the DeepSeek-R1 technical report [3], the maximum length of model generations is set to 32K tokens. Additionally, since the lengthy outputs of the reasoning models can vary greatly, as noted in the report, we adopt the $\text{pass}@k$ metric [1] to ensure the reliability of model performances. Specifically, we generate k responses for each test case using a sampling temperature of 0.6 and a top- p value of 0.95¹, and calculate $\text{pass}@1$ as followed:

$$\text{pass}@1 = \frac{1}{k} \sum_{i=1}^k p_i, \quad (1)$$

where p_i denotes the correctness of the i -th response. We also report $\text{cons}@k$ [3] where the model generates k answers and the final answer is the one generated the most frequently [17].

The evaluation prompts used for assessing EXAONE Deep models are shown in Figure 4 and 5. For the CSAT benchmark, both prompts are employed, as it includes short-answer and multiple-choice questions. For baseline models, we report scores from official reports if available; otherwise, we measure performance ourselves. Specifically, for multiple-choice question answering types, we adopt prompts from OpenAI simple-evals framework². When evaluating short-answer questions, we utilize the MATH-500 prompt from simple-evals for OpenAI o-series models, but use the same prompt with ours (Figure 4) for others, as it aligns with their recommended prompt.

3.4 Experimental Results

The performance comparison between EXAONE Deep and baseline models is conducted across four categories: mathematics, science, coding, and general knowledge. The evaluation results for the mathematics category are presented in Table 2, while those for the other categories are given in Table 3.

¹For QwQ-32B, we additionally applied the recommended top- $k = 40$ during evaluation. For OpenAI o1-mini, we were unable to adjust sampling parameters to our desired settings; thus, we used temperature = 1 and top- $p = 1$.

²<https://github.com/openai/simple-evals>

Prompt for Short-Answer Questions

{{question}}

Please reason step by step, and put your final answer within \boxed{ }.

Figure 4: Prompt for evaluating models on short-answer questions. We apply the prompt to the MATH-500, AIME 2024/2025, and CSAT 2025 benchmarks.

Prompt for Multiple-Choice Questions

Question : {{question}}

Options :

A) {{option A}}

B) {{option B}}

...

Please reason step by step, and you should write the correct option alphabet within \boxed{ }.

Figure 5: Prompt used for evaluating EXAONE Deep models on multiple-choice questions. We apply the prompt to the CSAT 2025, GPQA Diamond, MMLU, and MMLU-Pro benchmarks. The number of options is adjusted for each test case.

We find that the EXAONE Deep 32B model exhibits competitive performance, rivaling leading open-weight reasoning models such as DeepSeek-R1 and QwQ-32B. Notably, it outperforms the distilled versions of DeepSeek-R1, including DeepSeek-R1-Distill-Qwen-32B and DeepSeek-R1-Distill-Llama-70B.

Furthermore, the EXAONE Deep 7.8B model demonstrates superior performance compared to models of similar scale such as DeepSeek-R1-Distill-Qwen-7B and DeepSeek-R1-Distill-Llama-8B. It also outperforms the proprietary reasoning model: OpenAI o1-mini.

Regarding the EXAONE Deep 2.4B, it outperforms the DeepSeek-R1-Distill-Qwen-1.5B. Our experimental results highlight the EXAONE Deep models demonstrate enhanced reasoning capabilities across different model sizes.

Model	MATH-500 (pass@1)	AIME 2024 (pass@1)	AIME 2024 (cons@64)	AIME 2025 (pass@1)	AIME 2025 (cons@64)	CSAT 2025 (pass@1)
EXAONE Deep 32B	<u>95.7</u>	72.1	90.0	65.8	80.0	94.5
DeepSeek-R1-Distill-Qwen-32B	94.3*	72.6*	83.3*	55.2	73.3	84.1
QwQ-32B	95.5	<u>79.5*</u>	<u>86.7</u>	67.1	<u>76.7</u>	<u>94.4</u>
DeepSeek-R1-Distill-Llama-70B	94.5*	70.0*	<u>86.7*</u>	53.9	66.7	88.8
DeepSeek-R1 (671B)	97.3*	79.8*	<u>86.7</u>	<u>66.8</u>	80.0	89.9
EXAONE Deep 7.8B	94.8	70.0	83.3	59.6	76.7	89.9
DeepSeek-R1-Distill-Qwen-7B	<u>92.8*</u>	55.5*	83.3*	38.5	56.7	79.7
DeepSeek-R1-Distill-Llama-8B	89.1*	50.4*	<u>80.0*</u>	33.6	53.3	74.1
OpenAI o1-mini	90.0*	<u>63.6*</u>	<u>80.0*</u>	<u>54.8</u>	<u>66.7</u>	<u>84.4</u>
EXAONE Deep 2.4B	92.3	52.5	76.7	47.9	73.3	79.2
DeepSeek-R1-Distill-Qwen-1.5B	<u>83.9*</u>	<u>28.9*</u>	<u>52.7*</u>	<u>23.9</u>	<u>36.7</u>	<u>65.6</u>

Table 2: Comparison between EXAONE Deep and baseline models in the Mathematics category. The asterisk (*) indicates figures that have been officially reported. When calculating pass@1 score, we set k in Equation 1 as 8 for MATH-500, 16 for CSAT 2025, and 64 for AIME 2024 and 2025. For AIME 2024 and 2025, we also report cons@64. For CSAT 2025, the individual scores for the elective subjects can be found in Table 4 within Appendix C. **Bold** scores indicate the best performance, and underlined scores mean the second best.

Model	GPQA Diamond (pass@1)	LiveCodeBench (pass@1)	MMLU (accuracy)	MMLU-Pro (accuracy)
EXAONE Deep 32B	<u>66.1</u>	59.5	83.0	74.0
DeepSeek-R1-Distill-Qwen-32B	62.1*	57.2*	83.6	77.2
QwQ-32B	63.3	<u>63.4</u> *	87.4	79.1
DeepSeek-R1-Distill-Llama-70B	65.2*	57.5*	<u>89.1</u>	79.4
DeepSeek-R1 (671B)	71.5*	65.9*	90.8*	84.0
EXAONE Deep 7.8B	62.6	55.2	<u>75.0</u>	<u>65.7</u>
DeepSeek-R1-Distill-Qwen-7B	49.1*	37.6*	57.7	59.3
DeepSeek-R1-Distill-Llama-8B	49.0*	39.6*	74.4	60.6
OpenAI o1-mini	<u>60.0</u> *	<u>53.8</u> *	85.2*	80.3
EXAONE Deep 2.4B	54.3	46.6	65.8	56.4
DeepSeek-R1-Distill-Qwen-1.5B	<u>33.8</u> *	<u>16.9</u> *	<u>33.8</u>	<u>39.3</u>

Table 3: Comparison between EXAONE Deep and baseline models in the Science (GPQA Diamond), Coding (LiveCodeBench), and General (MMLU and MMLU-Pro) categories. The asterisk (*) indicates figures that have been officially reported. When calculating pass@1 score, we set k in Equation 1 as 6 for LiveCodeBench and 16 for GPQA Diamond. **Bold** scores indicate the best performance, and underlined scores mean the second best.

4 Limitations

The EXAONE Deep models introduced in this document are specifically fine-tuned to excel at reasoning tasks. Though their base models are instruction-fine-tuned and generally capable of following instructions, for addressing a wider range of real-world use cases, we strongly recommend utilizing the EXAONE 3.5 Instruct models [8, 9, 10], which are optimized for practical application scenarios.

5 Deployment

Section B in the Appendix provides license information for using the EXAONE Deep models. Understanding the license information is essential for the legal utilization of the language model.

6 Conclusion

In this document, we presented three specialized reasoning models: EXAONE Deep 2.4B, 7.8B, and 32B. Despite the emergence of various methodologies aimed at improving reasoning capabilities, we have chosen to rely on well-established approaches such as SFT, DPO, and Online RL, achieving superior or competitive performance relative to models of comparable scale. Our results highlight the effectiveness and practicality of these proven techniques in advancing reasoning performance. At present, these models are primarily engaged in problem-solving in domains where clear answers are available, such as mathematics, science, and coding. Looking ahead, we aim to extend their capabilities into areas where answers are less clear or yet to be discovered, thereby broadening their impact and utility.

Our models are available to everyone for research purposes, and we welcome your feedback to help us improve the models. If you have any feedback or are interested in exploring commercial opportunities with our models, please reach out to contact_us@lgresearch.ai.

A Contributors

All authors are listed in alphabetical order by last name.

Core Contributors Eunbi Choi, Kibong Choi, Seokhee Hong, Junwon Hwang, Hyojin Jeon, Hyunjik Jo, Joonkee Kim, Seonghwan Kim, Soyeon Kim, Sunkyoung Kim, Yireun Kim, Yongil Kim, Haeju Lee, Jinsik Lee, Kyungmin Lee, Sangha Park, Heuiyeen Yeen, Hyeongu Yun

Contributors Kyunghoon Bae, Stanley Jungkyu Choi, Yemuk Choi, Kijeong Jeon, Gerrard Jeongwon Jo, Jiyeon Jung, Hyosang Kim, Youchul Kim, Edward Hwayoung Lee, Honglak Lee, Yongmin Park, Sihoon Yang, Sihyuk Yi

B Model License

EXAONE AI Model License Agreement 1.1 - NC

This License Agreement (“Agreement”) is entered into between you (“Licensee”) and LG Management Development Institute Co., Ltd. (“Licensor”), governing the use of the EXAONE AI Model (“Model”). By downloading, installing, copying, or using the Model, you agree to comply with and be bound by the terms of this Agreement. If you do not agree to all the terms, you must not download, install, copy, or use the Model. This Agreement constitutes a binding legal agreement between the Licensee and Licensor.

1. Definitions

1.1 Model: The artificial intelligence model provided by Licensor, which includes any software, algorithms, machine learning models, or related components supplied by Licensor. This definition extends to encompass all updates, enhancements, improvements, bug fixes, patches, or other modifications that may be provided by Licensor from time to time, whether automatically or manually implemented.

1.2 Derivatives: Any modifications, alterations, enhancements, improvements, adaptations, or derivative works of the Model created by Licensee or any third party. This includes changes made to the Model’s architecture, parameters, data processing methods, or any other aspect of the Model that results in a modification of its functionality or output.

1.3 Output: Any data, results, content, predictions, analyses, insights, or other materials generated by the Model or Derivatives, regardless of whether they are in their original form or have been further processed or modified by the Licensee. This includes, but is not limited to, textual or numerical produced directly or indirectly through the use of the Model.

1.4 Licensor: LG Management Development Institute Co., Ltd., the owner, developer, and provider of the EXAONE AI Model. The Licensor holds all rights, title, and interest in the Model and is responsible for granting licenses to use the Model under the terms specified in this Agreement.

1.5 Licensee: The individual, organization, corporation, academic institution, government agency, or other entity using or intending to use the Model under the terms and conditions of this Agreement. The Licensee is responsible for ensuring compliance with the Agreement by all authorized users who access or utilize the Model on behalf of the Licensee.

2. License Grant

2.1 Grant of License: Subject to the terms and conditions outlined in this Agreement, the Licensor hereby grants the Licensee a limited, non-exclusive, non-transferable, worldwide, and revocable license to:

- a. Access, download, install, and use the Model solely for research purposes. This includes evaluation, testing, academic research, experimentation, and participation in competitions, provided that such participation is in a non-commercial context. Notwithstanding Section 3.1, the Licensee may only provide the Model or Derivatives for a competition if no commercial license is granted to the competition organizer or any third party.
- b. Publicly disclose research results and findings derived from the use of the Model or Derivatives, including publishing papers or presentations.
- c. Modify the Model and create Derivatives based on the Model, provided that such modifications and Derivatives are used exclusively for research purposes. The Licensee may conduct experiments, perform analyses, and apply custom modifications to the Model to explore its capabilities and performance under various scenarios. If the Model is modified, the modified Model must include “EXAONE” at the beginning of its name.
- d. Distribute the Model and Derivatives in each case with a copy of this Agreement.

2.2 Scope of License: The license granted herein does not authorize the Licensee to use the Model for any purpose not explicitly permitted under this Agreement. Any use beyond the scope of this license, including any commercial application or external distribution, is strictly prohibited unless explicitly agreed upon in writing by the Licensor.

3. Restrictions

3.1 Commercial Use: The Licensee is expressly prohibited from using the Model, Derivatives, or Output for any commercial purposes, including but not limited to, developing or deploying products, services, or applications that generate revenue, whether directly or indirectly. Any commercial exploitation of the Model or its derivatives requires a separate commercial license agreement with the Licensor. Furthermore, the Licensee shall not use the Model, Derivatives or Output to develop or improve other models.

3.2 Reverse Engineering: The Licensee shall not decompile, disassemble, reverse engineer, or attempt to derive the source code, underlying ideas, algorithms, or structure of the Model, except to the extent that such activities are expressly permitted by applicable law. Any attempt to bypass or circumvent technological protection measures applied to the Model is strictly prohibited.

3.3 Unlawful Use: The Licensee shall not use the Model and Derivatives for any illegal, fraudulent, or unauthorized activities, nor for any purpose that violates applicable laws or regulations. This includes but is not limited to the creation, distribution, or dissemination of malicious, deceptive, or unlawful content.

3.4 Ethical Use: The Licensee shall ensure that the Model or Derivatives is used in an ethical and responsible manner, adhering to the following guidelines:

a. The Model and Derivatives shall not be used to generate, propagate, or amplify false, misleading, or harmful information, including fake news, misinformation, or disinformation.

b. The Model and Derivatives shall not be employed to create, distribute, or promote content that is discriminatory, harassing, defamatory, abusive, or otherwise offensive to individuals or groups based on race, gender, sexual orientation, religion, nationality, or other protected characteristics.

c. The Model and Derivatives shall not infringe on the rights of others, including intellectual property rights, privacy rights, or any other rights recognized by law. The Licensee shall obtain all necessary permissions and consents before using the Model and Derivatives in a manner that may impact the rights of third parties.

d. The Model and Derivatives shall not be used in a way that causes harm, whether physical, mental, emotional, or financial, to individuals, organizations, or communities. The Licensee shall take all reasonable measures to prevent misuse or abuse of the Model and Derivatives that could result in harm or injury.

4. Ownership

4.1 Intellectual Property: All rights, title, and interest in and to the Model, including any modifications, Derivatives, and associated documentation, are and shall remain the exclusive property of the Licensor. The Licensee acknowledges that this Agreement does not transfer any ownership rights to the Licensee. All trademarks, service marks, and logos associated with the Model are the property of the Licensor.

4.2 Output: All rights, title, and interest in and to the Output generated by the Model and Derivatives whether in its original form or modified, are and shall remain the exclusive property of the Licensor. Licensee may use, modify, and distribute the Output and its derivatives for research purpose. The Licensee shall not claim ownership of the Output except as expressly provided in this Agreement. The Licensee may use the Output solely for the purposes permitted under this Agreement and shall not exploit the Output for unauthorized or commercial purposes.

4.3 Attribution: In any publication or presentation of results obtained using the Model, the Licensee shall provide appropriate attribution to the Licensor, citing the Model's name and version, along with any relevant documentation or references specified by the Licensor.

5. No Warranty

5.1 “As-Is” Basis: The Model, Derivatives, and Output are provided on an “as-is” and “as-available” basis, without any warranties or representations of any kind, whether express, implied, or statutory. The Licensor disclaims all warranties, including but not limited to, implied warranties of merchantability, fitness for a particular purpose, accuracy, reliability, non-infringement, or any warranty arising from the course of dealing or usage of trade.

5.2 Performance and Reliability: The Licensor does not warrant or guarantee that the Model, Derivatives or Output will meet the Licensee’s requirements, that the operation of the Model, Derivatives or Output will be uninterrupted or error-free, or that defects in the Model will be corrected. The Licensee acknowledges that the use of the Model, Derivatives or Output is at its own risk and that the Model, Derivatives or Output may contain bugs, errors, or other limitations.

5.3 No Endorsement: The Licensor does not endorse, approve, or certify any results, conclusions, or recommendations derived from the use of the Model. The Licensee is solely responsible for evaluating the accuracy, reliability, and suitability of the Model for its intended purposes.

6. Limitation of Liability

6.1 No Liability for Damages: To the fullest extent permitted by applicable law, in no event shall the Licensor be liable for any special, incidental, indirect, consequential, exemplary, or punitive damages, including but not limited to, damages for loss of business profits, business interruption, loss of business information, loss of data, or any other pecuniary or non-pecuniary loss arising out of or in connection with the use or inability to use the Model, Derivatives or any Output, even if the Licensor has been advised of the possibility of such damages.

6.2 Indemnification: The Licensee agrees to indemnify, defend, and hold harmless the Licensor, its affiliates, officers, directors, employees, and agents from and against any claims, liabilities, damages, losses, costs, or expenses (including reasonable attorneys’ fees) arising out of or related to the Licensee’s use of the Model, any Derivatives, or any Output, including any violation of this Agreement or applicable laws.

7. Termination

7.1 Termination by Licensor: The Licensor reserves the right to terminate this Agreement and revoke the Licensee’s rights to use the Model at any time, with or without cause, and without prior notice if the Licensee breaches any of the terms or conditions of this Agreement. Termination shall be effective immediately upon notice.

7.2 Effect of Termination: Upon termination of this Agreement, the Licensee must immediately cease all use of the Model, Derivatives, and Output and destroy all copies of the Model, Derivatives, and Output in its possession or control, including any backup or archival copies. The Licensee shall certify in writing to the Licensor that such destruction has been completed.

7.3 Survival: The provisions of this Agreement that by their nature should survive termination, including but not limited to, Sections 4 (Ownership), 5 (No Warranty), 6 (Limitation of Liability), and this Section 7 (Termination), shall continue to apply after termination.

8. Governing Law

8.1 Governing Law: This Agreement shall be governed by and construed in accordance with the laws of the Republic of Korea, without regard to its conflict of laws principles.

8.2 Arbitration: Any disputes, controversies, or claims arising out of or relating to this Agreement, including its existence, validity, interpretation, performance, breach, or termination, shall be referred to and finally resolved by arbitration administered by the Korean Commercial Arbitration Board (KCAB) in accordance with the International Arbitration Rules of the Korean Commercial Arbitration Board in force at the time of the commencement of the arbitration. The seat of arbitration shall be Seoul, Republic of Korea. The tribunal shall consist of one arbitrator. The language of the arbitration shall be English.

9. Alterations

9.1 Modifications: The Licensor reserves the right to modify or amend this Agreement at any time, in its sole discretion. Any modifications will be effective upon posting the updated Agreement on the Licensor's website or through other means of communication. The Licensee is responsible for reviewing the Agreement periodically for changes. Continued use of the Model after any modifications have been made constitutes acceptance of the revised Agreement.

9.2 Entire Agreement: This Agreement constitutes the entire agreement between the Licensee and Licensor concerning the subject matter hereof and supersedes all prior or contemporaneous oral or written agreements, representations, or understandings. Any terms or conditions of any purchase order or other document submitted by the Licensee in connection with the Model that are in addition to, different from, or inconsistent with the terms and conditions of this Agreement are not binding on the Licensor and are void.

By downloading, installing, or using the EXAONE AI Model, the Licensee acknowledges that it has read, understood, and agrees to be bound by the terms and conditions of this Agreement.

C Evaluation Details

Model	CSAT 2025 Calculus (pass@1)	CSAT 2025 Statistics (pass@1)	CSAT 2025 Geometry (pass@1)
EXAONE Deep 32B	95.1	<u>95.0</u>	93.5
QwQ-32B	<u>94.5</u>	95.5	<u>93.3</u>
DeepSeek-R1 (671B)	<u>89.4</u>	90.8	<u>89.6</u>
DeepSeek-R1-Distill-Llama-70B	88.1	90.3	88.1
DeepSeek-R1-Distill-Qwen-32B	84.3	84.4	83.5
EXAONE Deep 7.8B	89.5	91.8	88.5
OpenAI o1-mini	<u>82.8</u>	<u>85.9</u>	<u>84.6</u>
DeepSeek-R1-Distill-Qwen-7B	75.6	83.8	79.6
DeepSeek-R1-Distill-Llama-8B	70.6	79.0	72.6
EXAONE Deep 2.4B	74.4	82.4	80.8
DeepSeek-R1-Distill-Qwen-1.5B	<u>64.4</u>	<u>67.6</u>	<u>64.6</u>

Table 4: Comparison of EXAONE Deep with baseline models in the individual scores for the elective subjects at the CSAT 2025 benchmark. **Bold** scores indicate the best performance, and underlined scores mean the second best.

References

- [1] Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, Alex Ray, Raul Puri, Gretchen Krueger, Michael Petrov, Heidy Khlaaf, Girish Sastry, Pamela Mishkin, Brooke Chan, Scott Gray, Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz Kaiser, Mohammad Bavarian, Clemens Winter, Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebgen Guss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Josh Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. Evaluating Large Language Models Trained on Code. <https://arxiv.org/abs/2107.03374>, 2021.
- [2] DeepSeek-AI. DeepSeek-V2: A Strong, Economical, and Efficient Mixture-of-Experts Language Model. <https://arxiv.org/abs/2405.04434>, 2024.
- [3] DeepSeek-AI. DeepSeek-R1: Incentivizing Reasoning Capability in LLMs via Reinforcement Learning. <https://arxiv.org/abs/2501.12948>, 2025.
- [4] Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. Measuring Massive Multitask Language Understanding. <https://arxiv.org/abs/2009.03300>, 2021.
- [5] Naman Jain, King Han, Alex Gu, Wen-Ding Li, Fanjia Yan, Tianjun Zhang, Sida Wang, Armando Solar-Lezama, Koushik Sen, and Ion Stoica. LiveCodeBench: Holistic and Contamination Free Evaluation of Large Language Models for Code. <https://arxiv.org/abs/2403.07974>, 2024.
- [6] Korea Institute for Curriculum and Evaluation. College Scholastic Ability Test. <https://www.suneung.re.kr>, 2025.
- [7] LG AI Research. EXAONE 3.5: Series of Large Language Models for Real-world Use Cases. <https://arxiv.org/abs/2412.04862>, 2024.
- [8] LG AI Research. EXAONE-3.5-2.4B-Instruct. <https://huggingface.co/LGAI-EXAONE/EXAONE-3.5-2.4B-Instruct>, 2025.
- [9] LG AI Research. EXAONE-3.5-32B-Instruct. <https://huggingface.co/LGAI-EXAONE/EXAONE-3.5-32B-Instruct>, 2025.
- [10] LG AI Research. EXAONE-3.5-7.8B-Instruct. <https://huggingface.co/LGAI-EXAONE/EXAONE-3.5-7.8B-Instruct>, 2025.
- [11] Hunter Lightman, Vineet Kosaraju, Yura Burda, Harri Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s Verify Step by Step. <https://arxiv.org/abs/2305.20050>, 2023.
- [12] Mathematical Association of America. American Invitational Mathematics Examination. <https://maa.org>, 2025.
- [13] OpenAI. OpenAI o1 System Card. <https://cdn.openai.com/o1-system-card-20241205.pdf>, 2024.
- [14] Qwen Team. QwQ-32B: Embracing the Power of Reinforcement Learning. <https://qwenlm.github.io/blog/qwq-32b>, 2025.
- [15] David Rein, Betty Li Hou, Asa Cooper Stickland, Jackson Petty, Richard Yuanzhe Pang, Julien Dirani, Julian Michael, and Samuel R. Bowman. GPQA: A Graduate-Level Google-Proof Q&A Benchmark. <https://arxiv.org/abs/2311.12022>, 2023.
- [16] Charlie Snell, Jaehoon Lee, Kelvin Xu, and Aviral Kumar. Scaling LLM Test-Time Compute Optimally can be More Effective than Scaling Model Parameters. <https://arxiv.org/abs/2408.03314>, 2024.
- [17] Xuezhi Wang, Jason Wei, Dale Schuurmans, Quoc Le, Ed Chi, Sharan Narang, Aakanksha Chowdhery, and Denny Zhou. Self-Consistency Improves Chain of Thought Reasoning in Language Models. <https://arxiv.org/abs/2203.11171>, 2023.

- [18] Yubo Wang, Xueguang Ma, Ge Zhang, Yuansheng Ni, Abhranil Chandra, Shiguang Guo, Weiming Ren, Aaran Arulraj, Xuan He, Ziyang Jiang, Tianle Li, Max Ku, Kai Wang, Alex Zhuang, Rongqi Fan, Xiang Yue, and Wenhui Chen. MMLU-Pro: A More Robust and Challenging Multi-Task Language Understanding Benchmark. <https://arxiv.org/abs/2406.01574>, 2024.
- [19] Teng Xiao, Yige Yuan, Zhengyu Chen, Mingxiao Li, Shangsong Liang, Zhaochun Ren, and Vasant G Honavar. SimPER: A Minimalist Approach to Preference Alignment without Hyperparameters. <https://arxiv.org/abs/2502.00883>, 2025.