
Self-Correction Bench: Revealing the Self-Correction Blind Spot in LLMs

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 Although large language models (LLMs) have become transformative, they still
2 make mistakes and can explore unproductive reasoning paths. Self-correction is
3 an important capability for a trustworthy LLM. While LLMs can identify error in
4 user input, they exhibit a systematic 'Self-Correction Blind Spot' - failing to cor-
5 rect identical error in their own outputs. To systematically study this phenomenon,
6 we introduce **Self-Correction Bench**, a systematic framework to measure this
7 phenomenon through controlled error injection at three complexity levels. Testing
8 14 open-source non-reasoning models, we find an average 64.5% blind spot rate.
9 Remarkably, simply appending "*Wait*" reduces blind spots by 89.3%, suggesting
10 that the capability exists but requires activation. Our work highlights a critical lim-
11 itation in current LLMs and offers potential avenues for improving their reliability
12 and trustworthiness.

1 Introduction

14 Large Language Models (LLMs) have rapidly advanced natural language processing, achieving
15 state-of-the-art results on a diverse range of tasks [OpenAI et al., 2024, Anthropic, 2024, Gem-
16 ini Team, 2025, Yang et al., 2025, Meta, 2025, DeepSeek-AI et al., 2025a]. However, despite their
17 impressive capabilities, LLMs are known to exhibit unpredictable failures and generate inaccurate
18 information [Huang et al., 2025, Bang et al., 2023, Shi et al., 2023, Nezhurina et al., 2025].

19 Observing LLM self-correction behavior in natural settings is challenging due to their inherent ac-
20 curacy; the rarity of naturally occurring errors makes systematic evaluation difficult. We therefore
21 ask 1) Do LLMs exhibit a systematic failure to self-correct when their own output contains an error?
22 2) If so, can a simple prompt intervention recover this capability without retraining?

23 To answer these questions, we construct Self-Correction Bench by systematically injecting error
24 into the LLM reasoning traces. This allows us to create controlled environments for testing self-
25 correction and reliably quantifying performance using dedicated benchmarks.

2 Background and related Works

27 Recent work has shown that intrinsic self-correction improves performance by prompting the same
28 LLM to generate feedback on their own responses [Kim et al., 2023, Shinn et al., 2023, Madaan
29 et al., 2023, Kamoi et al., 2024] in multi-step prompting. Lanham et al. [2023] injects mistakes into
30 reasoning chains to evaluate the faithfulness of chain-of-thought reasoning, but not self-correction.
31 Simple test-time scaling [Muennighoff et al., 2025] demonstrates that appending "*Wait*" can improve
32 reasoning significantly. We provide a behavior explanation for why it works by conducting an
33 intervention experiment on unfinetuned models. Koo et al. [2024], Echterhoff et al. [2024], Jones and

Steinhardt [2022] show that LLMs exhibit different cognitive biases given their training in human data. We study how the bias blind spot [Pronin et al., 2002] affects self-correction. Our work contributes a systematic methodology for testing self-correction and demonstrates that test-time interventions can activate self-correction mechanism that otherwise was inactive.

3 Methodology

3.1 Self-Correction Blind Spot

Pronin et al. [2002] introduces bias blind spot, which is the tendency of individuals to recognize bias while not recognizing their own. Inspired by the study, we further differentiate error comes from either external or internal.

Let e denote an error injected into a response fragment. We compare the probability that a model M produces a correct final answer given an internal error (in its own output, r_m) versus an external error (in the user prompt, r_u).

Self-Correction Blind Spot is defined as:

$$\text{Self-Correction Blind Spot} = \begin{cases} 1 - \frac{P_M(r_{\text{correct}}|r_m, e)}{P_M(r_{\text{correct}}|r_u, e)} & \text{if } P_M(r_{\text{correct}}|r_u, e) > 0 \\ 0 & \text{if } P_M(r_{\text{correct}}|r_u, e) = 0 \end{cases} \quad (1)$$

If the Blind Spot equals 1, it means that an LLM can self-correct the external error but not the internal error. By design, the Self-Correction Blind Spot isolates confounding factors, including internal model knowledge.

3.2 Self-Correction Bench

We introduce three datasets for measuring Self-Correction blind spot in LLMs across varying complexities. To understand self-correction failures, we must first ask: can models self-correct the simplest possible error? If not, a more complex error correction is unlikely.

We begin with artificially simple errors (Self-correct Like I am 5, SCLI5) to isolate the self-correction mechanism from confounding factors such as knowledge limitations or complexity of reasoning. GSM8K-SC introduces controlled errors in multi-step reasoning tasks, while PRM800K-SC provides realistic erroneous completions from actual LLM outputs. This progressive evaluation allows us to systematically map where self-correction breaks down while controlling for reasoning complexity and error realism. Our benchmark can be summarized in Table 1. See construction details in Appendix A.

For each dataset, we systematically inject an identical error into both model response (r_m) and user prompt (r_u), allowing us to empirically estimate both $P_M(r_{\text{correct}}|r_m, e_{\text{controlled}})$ and $P_M(r_{\text{correct}}|r_u, e_{\text{controlled}})$ under identical error conditions.

We are aware that the Self-Correction Bench introduces various distribution mismatch for various models because controlled error injection may not perfectly mirror naturally occurring error patterns. However, this methodology enables systematic isolation of self-correction capabilities from confounding factors.

Table 1: Dataset comparison

Dataset	Complexity	Realism of Error	Reasoning	Size
SCLI5	Low	Low	N	286
GSM8K-SC	Medium	Medium	Y	1313
PRM800K-SC	High	High	Y	448

4 Results

We test a wide range of open-source state-of-the-art models with a temperature of 0.0 and fixed token budget of 1024. We also report the result under different settings in Appendix B, which does not change our conclusion. We do not include close-source model because only open-source models support fine-grained control of injecting a prefix into the model response, limiting generalizability to closed-source LLMs. We use the DeepInfra completion API.

We use ‘gemini-2.5-flash-preview-05-20’ to compare LLMs’ completion against the ground-truth answer. Due to the objectivity of the task and the provision of ground truth in the prompt, we do not believe there is significant bias.

We report the Blind Spot and its standard error for each model. The standard error is estimated using the formula $\sigma_M = \frac{\sigma}{\sqrt{N}}$, where N is the population size and σ is the population standard deviation.

4.1 Non-reasoning models

We observe statistically significant Self-Correction Blind Spot for most models across datasets - LLMs fail to self-correct error when injected in model completion (internal error); while they are able to do so when the error is in the user prompt (external error) in Figure 1. The Blind Spot, on average, 64.5%, exists across models of different model sizes.

In Figure 3, we also observe a moderate correlation between the datasets, suggesting that there is a blind spot between tasks of different complexity, and from artificial (SCLI5) to realistic (PRM800K-SC) errors, demonstrating the generality of phenomenon across error types.

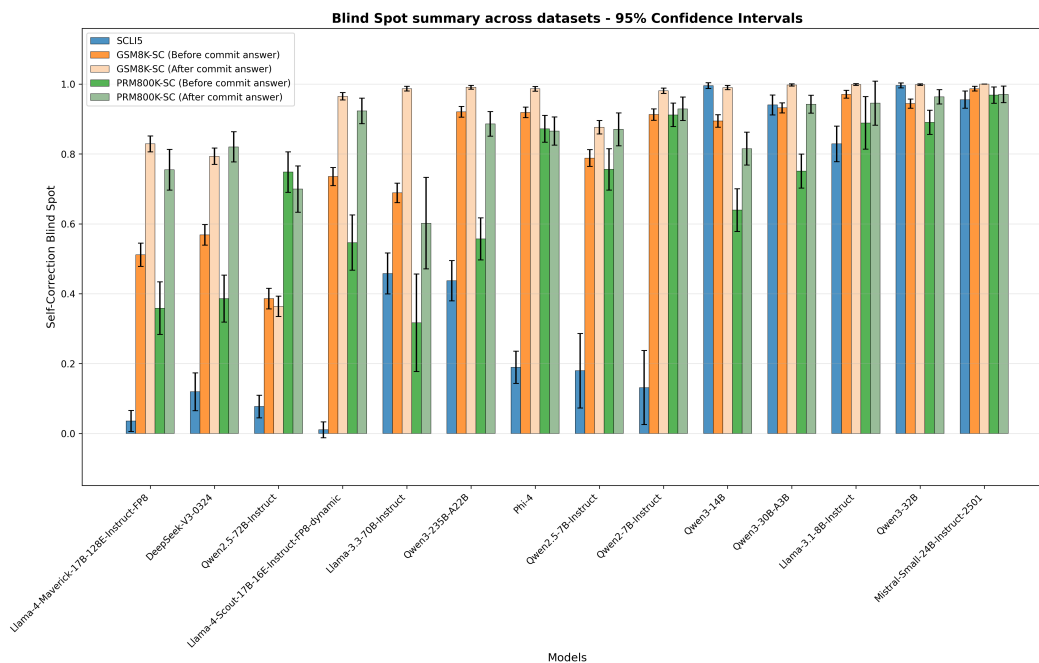


Figure 1: Self-Correction Blind Spot across models

The Blind Spot motivates us to further investigate the qualitative difference of how LLMs respond when error is injected into the model response vs. the user response. We observe that when error is external, correction markers are generated 179.5% and 73.6% more frequently in GSM8K-SC and PRM800K-SC. This motivates our intervention experiment, which tests whether latent capabilities can be activated without retraining.

After appending “Wait”, Figure 2 shows significant reductions in Blind Spot, in some cases, a negative Blind Spot even without any finetuning. Averaging across models and datasets, the reduction amounts to 89.3%. This evidence leads us to believe that “Wait” and similar correction markers

95 serve as a strong conditioning token that shifts the model’s conditional probability distribution. This
 96 new distribution favors sequences associated with re-evaluation and correction.

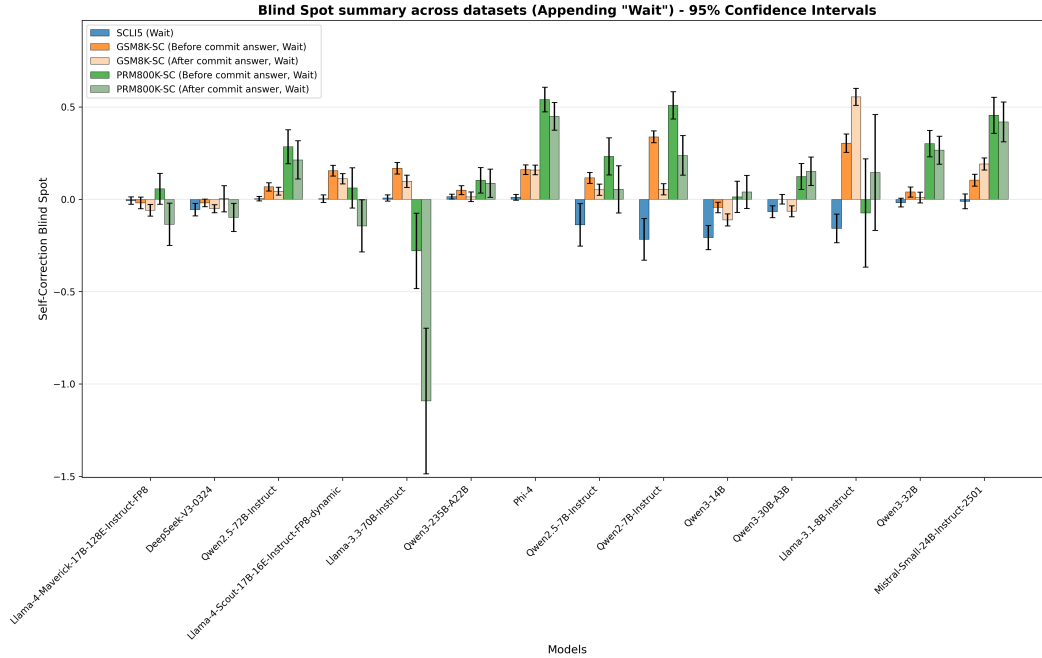


Figure 2: Self-Correction Blind Spot across models after appending “Wait”

97 4.2 Reasoning models

98 Unlike non reasoning model, we observe smaller, even negative, self-correction blind spots in rea-
 99 soning models in Figure 4.

100 5 Conclusion

101 We have uncovered Self-Correction Blind Spot in contemporary open-source non-reasoning LLMs
 102 and quantified it with the Self-Correction Bench. We do not notice significant Blind Spot in rea-
 103 soning models. A simple prompt intervention by appending “Wait” almost eliminates this blind
 104 spot in non-reasoning models, revealing a latent self-correction capability that is otherwise dormant.
 105 Future work could extend Self-Correction Bench to cover programming, logic, and common sense
 106 reasoning by injecting realistic error.

107 6 Reproducibility statement

108 Our experiments utilize various open source models, close sourced models, and datasets. Our codes
 109 for running the experiment are released in github. Self-Correction Bench is available at Hugging
 110 Face.

111 7 Acknowledgment and disclosure of funding

112 Information regarding funding and acknowledgments will be provided upon acceptance to maintain
 113 anonymity during the review process.

114 References

115 OpenAI, Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Floren-
 116 cia Leoni Aleman, Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, Red

117 Avila, Igor Babuschkin, Suchir Balaji, Valerie Balcom, Paul Baltescu, Haiming Bao, Moham-
 118 mad Bavarian, Jeff Belgum, Irwan Bello, Jake Berdine, Gabriel Bernadett-Shapiro, Christopher
 119 Berner, Lenny Bogdonoff, Oleg Boiko, Madelaine Boyd, Anna-Luisa Brakman, Greg Brock-
 120 man, Tim Brooks, Miles Brundage, Kevin Button, Trevor Cai, Rosie Campbell, Andrew Cann,
 121 Brittany Carey, Chelsea Carlson, Rory Carmichael, Brooke Chan, Che Chang, Fotis Chantzis,
 122 Derek Chen, Sully Chen, Ruby Chen, Jason Chen, Mark Chen, Ben Chess, Chester Cho, Casey
 123 Chu, Hyung Won Chung, Dave Cummings, Jeremiah Currier, Yunxing Dai, Cory Decareaux,
 124 Thomas Degry, Noah Deutsch, Damien Deville, Arka Dhar, David Dohan, Steve Dowling, Sheila
 125 Dunning, Adrien Ecoffet, Atty Eleti, Tyna Eloundou, David Farhi, Liam Fedus, Niko Felix,
 126 Simón Posada Fishman, Juston Forte, Isabella Fulford, Leo Gao, Elie Georges, Christian Gib-
 127 son, Vik Goel, Tarun Gogineni, Gabriel Goh, Rapha Gontijo-Lopes, Jonathan Gordon, Morgan
 128 Grafstein, Scott Gray, Ryan Greene, Joshua Gross, Shixiang Shane Gu, Yufei Guo, Chris Hal-
 129 lacy, Jesse Han, Jeff Harris, Yuchen He, Mike Heaton, Johannes Heidecke, Chris Hesse, Alan
 130 Hickey, Wade Hickey, Peter Hoeschele, Brandon Houghton, Kenny Hsu, Shengli Hu, Xin Hu,
 131 Joost Huizinga, Shantanu Jain, Shawn Jain, Joanne Jang, Angela Jiang, Roger Jiang, Haozhun
 132 Jin, Denny Jin, Shino Jomoto, Billie Jonn, Heewoo Jun, Tomer Kaftan, Łukasz Kaiser, Ali Ka-
 133 mali, Ingmar Kanitscheider, Nitish Shirish Keskar, Tabarak Khan, Logan Kilpatrick, Jong Wook
 134 Kim, Christina Kim, Yongjik Kim, Jan Hendrik Kirchner, Jamie Kiros, Matt Knight, Daniel
 135 Kokotajlo, Łukasz Kondraciuk, Andrew Kondrich, Aris Konstantinidis, Kyle Kopic, Gretchen
 136 Krueger, Vishal Kuo, Michael Lampe, Ikai Lan, Teddy Lee, Jan Leike, Jade Leung, Daniel
 137 Levy, Chak Ming Li, Rachel Lim, Molly Lin, Stephanie Lin, Mateusz Litwin, Theresa Lopez,
 138 Ryan Lowe, Patricia Lue, Anna Makanju, Kim Malfacini, Sam Manning, Todor Markov, Yaniv
 139 Markovski, Bianca Martin, Katie Mayer, Andrew Mayne, Bob McGrew, Scott Mayer McKinney,
 140 Christine McLeavey, Paul McMillan, Jake McNeil, David Medina, Aalok Mehta, Jacob Menick,
 141 Luke Metz, Andrey Mishchenko, Pamela Mishkin, Vinnie Monaco, Evan Morikawa, Daniel
 142 Mossing, Tong Mu, Mira Murati, Oleg Murk, David Mély, Ashvin Nair, Reiichiro Nakano, Ra-
 143 jeev Nayak, Arvind Neelakantan, Richard Ngo, Hyeonwoo Noh, Long Ouyang, Cullen O’Keefe,
 144 Jakub Pachocki, Alex Paino, Joe Palermo, Ashley Pantuliano, Giambattista Parascandolo, Joel
 145 Parish, Emy Parparita, Alex Passos, Mikhail Pavlov, Andrew Peng, Adam Perelman, Filipe
 146 de Avila Belbute Peres, Michael Petrov, Henrique Ponde de Oliveira Pinto, Michael, Pokorny,
 147 Michelle Pokrass, Vitchyr H. Pong, Tolly Powell, Alethea Power, Boris Power, Elizabeth Proehl,
 148 Raul Puri, Alec Radford, Jack Rae, Aditya Ramesh, Cameron Raymond, Francis Real, Kendra
 149 Rimbach, Carl Ross, Bob Rotsted, Henri Roussez, Nick Ryder, Mario Saltarelli, Ted Sanders,
 150 Shibani Santurkar, Girish Sastry, Heather Schmidt, David Schnurr, John Schulman, Daniel Sel-
 151 sam, Kyla Sheppard, Toki Sherbakov, Jessica Shieh, Sarah Shoker, Pranav Shyam, Szymon Sidor,
 152 Eric Sigler, Maddie Simens, Jordan Sitkin, Katarina Slama, Ian Sohl, Benjamin Sokolowsky,
 153 Yang Song, Natalie Staudacher, Felipe Petroski Such, Natalie Summers, Ilya Sutskever, Jie Tang,
 154 Nikolas Tezak, Madeleine B. Thompson, Phil Tillet, Amin Tootoonchian, Elizabeth Tseng, Pre-
 155 ston Tuggle, Nick Turley, Jerry Tworek, Juan Felipe Cerón Uribe, Andrea Vallone, Arun Vi-
 156 jayvergiya, Chelsea Voss, Carroll Wainwright, Justin Jay Wang, Alvin Wang, Ben Wang, Jonathan
 157 Ward, Jason Wei, CJ Weinmann, Akila Welihinda, Peter Welinder, Jiayi Weng, Lilian Weng,
 158 Matt Wiethoff, Dave Willner, Clemens Winter, Samuel Wolrich, Hannah Wong, Lauren Work-
 159 man, Sherwin Wu, Jeff Wu, Michael Wu, Kai Xiao, Tao Xu, Sarah Yoo, Kevin Yu, Qiming
 160 Yuan, Wojciech Zaremba, Rowan Zellers, Chong Zhang, Marvin Zhang, Shengjia Zhao, Tianhao
 161 Zheng, Juntang Zhuang, William Zhuk, and Barret Zoph. Gpt-4 technical report, 2024. URL
 162 <https://arxiv.org/abs/2303.08774>.

163 Anthropic. The claude 3 model family: Opus, sonnet, haiku, Mar 2024. URL [https://](https://www-cdn.anthropic.com/de8ba9b01c9ab7cbabf5c33b80b7bbc618857627/Model_Card_Claude_3.pdf)
 164 [www-cdn.anthropic.com/de8ba9b01c9ab7cbabf5c33b80b7bbc618857627/](https://www-cdn.anthropic.com/de8ba9b01c9ab7cbabf5c33b80b7bbc618857627/Model_Card_Claude_3.pdf)
 165 [Model_Card_Claude_3.pdf](https://www-cdn.anthropic.com/de8ba9b01c9ab7cbabf5c33b80b7bbc618857627/Model_Card_Claude_3.pdf).

166 Google Gemini Team. Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality,
 167 long context, and next generation agentic capabilities., June 2025. URL [https://storage.](https://storage.googleapis.com/deepmind-media/gemini/gemini_v2_5_report.pdf)
 168 [googleapis.com/deepmind-media/gemini/gemini_v2_5_report.pdf](https://storage.googleapis.com/deepmind-media/gemini/gemini_v2_5_report.pdf).

169 An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang
 170 Gao, Chengen Huang, Chenxu Lv, Chujie Zheng, Dayiheng Liu, Fan Zhou, Fei Huang, Feng Hu,
 171 Hao Ge, Haoran Wei, Huan Lin, Jialong Tang, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin
 172 Yang, Jiaxi Yang, Jing Zhou, Jingren Zhou, Junyang Lin, Kai Dang, Keqin Bao, Kexin Yang,
 173 Le Yu, Lianghao Deng, Mei Li, Mingfeng Xue, Mingze Li, Pei Zhang, Peng Wang, Qin Zhu, Rui

174 Men, Ruize Gao, Shixuan Liu, Shuang Luo, Tianhao Li, Tianyi Tang, Wenbiao Yin, Xingzhang
175 Ren, Xinyu Wang, Xinyu Zhang, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yinger
176 Zhang, Yu Wan, Yuqiong Liu, Zekun Wang, Zeyu Cui, Zhenru Zhang, Zhipeng Zhou, and Zihan
177 Qiu. Qwen3 technical report, 2025. URL <https://arxiv.org/abs/2505.09388>.

178 Meta. The llama 4 herd: The beginning of a new era of natively multimodal ai innovation, Apr 2025.
179 URL <https://ai.meta.com/blog/llama-4-multimodal-intelligence/>.

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

209 Lei Huang, Weijiang Yu, Weitao Ma, Weihong Zhong, Zhangyin Feng, Haotian Wang, Qianglong
210 Chen, Weihua Peng, Xiaocheng Feng, Bing Qin, and Ting Liu. A survey on hallucination in large
211 language models: Principles, taxonomy, challenges, and open questions. *ACM Transactions on*
212 *Information Systems*, 43(2):1–55, January 2025. ISSN 1558-2868. doi: 10.1145/3703155. URL
213 <http://dx.doi.org/10.1145/3703155>.

214 Yejin Bang, Samuel Cahyawijaya, Nayeon Lee, Wenliang Dai, Dan Su, Bryan Wilie, Holy Lovenia,
215 Ziwei Ji, Tiezheng Yu, Willy Chung, Quyet V. Do, Yan Xu, and Pascale Fung. A multitask, mul-
216 tilingual, multimodal evaluation of chatgpt on reasoning, hallucination, and interactivity, 2023.
217 URL <https://arxiv.org/abs/2302.04023>.

218 Freda Shi, Xinyun Chen, Kanishka Misra, Nathan Scales, David Dohan, Ed Chi, Nathanael Schärli,
219 and Denny Zhou. Large language models can be easily distracted by irrelevant context, 2023.
220 URL <https://arxiv.org/abs/2302.00093>.

221 Marianna Nezhurina, Lucia Cipolina-Kun, Mehdi Cherti, and Jenia Jitsev. Alice in wonderland:
222 Simple tasks showing complete reasoning breakdown in state-of-the-art large language models,
223 2025. URL <https://arxiv.org/abs/2406.02061>.

224 Geunwoo Kim, Pierre Baldi, and Stephen McAleer. Language models can solve computer tasks,
225 2023. URL <https://arxiv.org/abs/2303.17491>.

- 226 Noah Shinn, Federico Cassano, Edward Berman, Ashwin Gopinath, Karthik Narasimhan, and
227 Shunyu Yao. Reflexion: Language agents with verbal reinforcement learning, 2023. URL
228 <https://arxiv.org/abs/2303.11366>.
- 229 Aman Madaan, Niket Tandon, Prakhar Gupta, Skyler Hallinan, Luyu Gao, Sarah Wiegrefe, Uri
230 Alon, Nouha Dziri, Shrimai Prabhumoye, Yiming Yang, Shashank Gupta, Bodhisattwa Prasad
231 Majumder, Katherine Hermann, Sean Welleck, Amir Yazdanbakhsh, and Peter Clark. Self-
232 refine: Iterative refinement with self-feedback, 2023. URL [https://arxiv.org/abs/](https://arxiv.org/abs/2303.17651)
233 2303.17651.
- 234 Ryo Kamoi, Yusen Zhang, Nan Zhang, Jiawei Han, and Rui Zhang. When can LLMs actually correct
235 their own mistakes? a critical survey of self-correction of LLMs. *Transactions of the Association*
236 *for Computational Linguistics*, 12:1417–1440, 2024. doi: 10.1162/tacl.a.00713. URL [https:](https://aclanthology.org/2024.tacl-1.78/)
237 [//aclanthology.org/2024.tacl-1.78/](https://aclanthology.org/2024.tacl-1.78/).
- 238 Tamera Lanham, Anna Chen, Ansh Radhakrishnan, Benoit Steiner, Carson Denison, Danny Her-
239 nandez, Dustin Li, Esin Durmus, Evan Hubinger, Jackson Kernion, Kamilė Lukošūtė, Karina
240 Nguyen, Newton Cheng, Nicholas Joseph, Nicholas Schiefer, Oliver Rausch, Robin Larson,
241 Sam McCandlish, Sandipan Kundu, Saurav Kadavath, Shannon Yang, Thomas Henighan, Tim-
242 othy Maxwell, Timothy Telleen-Lawton, Tristan Hume, Zac Hatfield-Dodds, Jared Kaplan, Jan
243 Brauner, Samuel R. Bowman, and Ethan Perez. Measuring faithfulness in chain-of-thought rea-
244 soning, 2023. URL <https://arxiv.org/abs/2307.13702>.
- 245 Niklas Muennighoff, Zitong Yang, Weijia Shi, Xiang Lisa Li, Li Fei-Fei, Hannaneh Hajishirzi, Luke
246 Zettlemoyer, Percy Liang, Emmanuel Candès, and Tatsunori Hashimoto. s1: Simple test-time
247 scaling, 2025. URL <https://arxiv.org/abs/2501.19393>.
- 248 Ryan Koo, Minhwa Lee, Vipul Raheja, Jong Inn Park, Zae Myung Kim, and Dongyeop Kang.
249 Benchmarking cognitive biases in large language models as evaluators. In Lun-Wei Ku, Andre
250 Martins, and Vivek Srikumar, editors, *Findings of the Association for Computational Linguistics:*
251 *ACL 2024*, pages 517–545, Bangkok, Thailand, August 2024. Association for Computational
252 Linguistics. doi: 10.18653/v1/2024.findings-acl.29. URL [https://aclanthology.org/](https://aclanthology.org/2024.findings-acl.29/)
253 2024.findings-acl.29/.
- 254 Jessica Echterhoff, Yao Liu, Abeer Alessa, Julian McAuley, and Zexue He. Cognitive bias in
255 decision-making with llms, 2024. URL <https://arxiv.org/abs/2403.00811>.
- 256 Erik Jones and Jacob Steinhardt. Capturing failures of large language models via human cognitive
257 biases, 2022. URL <https://arxiv.org/abs/2202.12299>.
- 258 Emily Pronin, Daniel Y. Lin, and Lee Ross. The bias blind spot: Perceptions of bias in
259 self versus others. *Personality and Social Psychology Bulletin*, 28(3):369–381, 2002. doi:
260 10.1177/0146167202286008. URL <https://doi.org/10.1177/0146167202286008>.
- 261 Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser,
262 Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John
263 Schulman. Training verifiers to solve math word problems, 2021. URL [https://arxiv.](https://arxiv.org/abs/2110.14168)
264 [org/abs/2110.14168](https://arxiv.org/abs/2110.14168).
- 265 OpenAI. Introducing gpt-4.1 in the api, Apr 2025. URL [https://openai.com/index/](https://openai.com/index/gpt-4-1/)
266 [gpt-4-1/](https://openai.com/index/gpt-4-1/).
- 267 Hunter Lightman, Vineet Kosaraju, Yuri Burda, Harrison Edwards, Bowen Baker, Teddy Lee, Jan
268 Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s verify step by step. In *The Twelfth*
269 *International Conference on Learning Representations*, 2024. URL [https://openreview.](https://openreview.net/forum?id=v8L0pN6EOi)
270 [net/forum?id=v8L0pN6EOi](https://openreview.net/forum?id=v8L0pN6EOi).
- 271 Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul Arora, Steven Basart, Eric Tang, Dawn
272 Song, and Jacob Steinhardt. Measuring mathematical problem solving with the MATH dataset.
273 In *Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks*
274 *Track (Round 2)*, 2021. URL <https://openreview.net/forum?id=7Bywt2mQsCe>.

DeepSeek-AI, Aixin Liu, Bei Feng, Bing Xue, Bingxuan Wang, Bochao Wu, Chengda Lu, Cheng-gang 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 Zhang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Li, Hui Qu, J. L. Cai, Jian Liang, Jianzhong Guo, Jiaqi Ni, Jiashi Li, Jiawei Wang, Jin Chen, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, Junxiao Song, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Lei Xu, Leyi Xia, Liang Zhao, Litong Wang, Liyue Zhang, Meng Li, Miaojun Wang, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Mingming Li, Ning Tian, Panpan Huang, Peiyi Wang, Peng Zhang, Qiancheng Wang, Qihao Zhu, Qinyu Chen, Qiushi Du, R. J. Chen, R. L. Jin, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, Runxin Xu, Ruoyu Zhang, Ruyi Chen, S. S. Li, Shanghao Lu, Shangyan Zhou, Shanhuang Chen, Shaoqing Wu, Shengfeng Ye, Shengfeng Ye, Shirong Ma, Shiyu Wang, Shuang Zhou, Shuiping Yu, Shunfeng Zhou, Shut-ing Pan, T. Wang, Tao Yun, Tian Pei, Tianyu Sun, W. L. Xiao, Wangding Zeng, Wanjia Zhao, Wei An, Wen Liu, Wenfeng Liang, Wenjun Gao, Wenqin Yu, Wentao Zhang, X. Q. Li, Xiangyue Jin, Xianzu Wang, Xiao Bi, Xiaodong Liu, Xiaohan Wang, Xiaojin Shen, Xiaokang Chen, Xi-aokang Zhang, Xiaosha Chen, Xiaotao Nie, Xiaowen Sun, Xiaoxiang Wang, Xin Cheng, Xin Liu, Xin Xie, Xingchao Liu, Xingkai Yu, Xinnan Song, Xinxia Shan, Xinyi Zhou, Xinyu Yang, Xinyuan Li, Xuecheng Su, Xuheng Lin, Y. K. Li, Y. Q. Wang, Y. X. Wei, Y. X. Zhu, Yang Zhang, Yanhong Xu, Yanhong Xu, Yanping Huang, Yao Li, Yao Zhao, Yaofeng Sun, Yaohui Li, Yaohui Wang, Yi Yu, Yi Zheng, Yichao Zhang, Yifan Shi, Yiliang Xiong, Ying He, Ying Tang, Yishi Piao, Yisong Wang, Yixuan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo, Yu Wu, Yuan Ou, Yuchen Zhu, Yudian Wang, Yue Gong, Yuheng Zou, Yujia He, Yukun Zha, Yunfan Xiong, Yunxian Ma, Yuting Yan, Yuxiang Luo, Yuxiang You, Yuxuan Liu, Yuyang Zhou, Z. F. Wu, Z. Z. Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhean Xu, Zhen Huang, Zhen Zhang, Zhenda Xie, Zhengyan Zhang, Zhewen Hao, Zhibin Gou, Zhicheng Ma, Zhigang Yan, Zhihong Shao, Zhipeng Xu, Zhiyu Wu, Zhongyu Zhang, Zhuoshu Li, Zihui Gu, Zijia Zhu, Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song, Ziyi Gao, and Zizheng Pan. Deepseek-v3 technical report, 2025b. URL <https://arxiv.org/abs/2412.19437>.

Qwen, :, An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li, Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tianhao Li, Tianyi Tang, Tingyu Xia, Xingzhang Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan, Yuqiong Liu, Zeyu Cui, Zhenru Zhang, and Zihan Qiu. Qwen2.5 technical report, 2025. URL <https://arxiv.org/abs/2412.15115>.

Meta. Llama 3.3, Dec 2024. URL https://www.llama.com/docs/model-cards-and-prompt-formats/llama3_3/.

Marah Abdin, 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. Phi-4 technical report, 2024. URL <https://arxiv.org/abs/2412.08905>.

An Yang, Baosong Yang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Zhou, Chengpeng Li, Chengyuan Li, Dayiheng Liu, Fei Huang, Guanting Dong, Haoran Wei, Huan Lin, Jialong Tang, Jialin Wang, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Ma, Jianxin Yang, Jin Xu, Jin-gren Zhou, Jinze Bai, Jinzheng He, Junyang Lin, Kai Dang, Keming Lu, Keqin Chen, Kexin Yang, Mei Li, Mingfeng Xue, Na Ni, Pei Zhang, Peng Wang, Ru Peng, Rui Men, Ruize Gao, Runji Lin, Shijie Wang, Shuai Bai, Sinan Tan, Tianhang Zhu, Tianhao Li, Tianyu Liu, Wen-bin Ge, Xiaodong Deng, Xiaohuan Zhou, Xingzhang Ren, Xinyu Zhang, Xipin Wei, Xuancheng Ren, Xuejing Liu, Yang Fan, Yang Yao, Yichang Zhang, Yu Wan, Yunfei Chu, Yuqiong Liu, Zeyu Cui, Zhenru Zhang, Zhifang Guo, and Zhihao Fan. Qwen2 technical report, 2024. URL <https://arxiv.org/abs/2407.10671>.

Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Alex Vaughan, Amy Yang, Angela Fan,

331 Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Ko-
 332 renev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurelien Rodriguez, Austen Gregerson, Ava
 333 Spataru, Baptiste Roziere, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux,
 334 Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret,
 335 Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Allonsius,
 336 Daniel Song, Danielle Pintz, Danny Livshits, Danny Wyatt, David Esiobu, Dhruv Choudhary,
 337 Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab
 338 AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Francisco
 339 Guzmán, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Govind That-
 340 tai, Graeme Nail, Gregoire Mialon, Guan Pang, Guillem Cucurell, Hailey Nguyen, Hannah Kore-
 341 vaar, Hu Xu, Hugo Touvron, Iliyan Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Ishan Misra,
 342 Ivan Evtimov, Jack Zhang, Jade Copet, Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay Ma-
 343 hadeokar, Jeet Shah, Jelmer van der Linde, Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu,
 344 Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu, Joanna Bitton, Joe Spisak, Jong-
 345 soo Park, Joseph Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia, Kalyan Vasuden Alwala,
 346 Karthik Prasad, Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth Heafield, Kevin Stone, Khalid
 347 El-Arini, Krithika Iyer, Kshitiz Malik, Kuenley Chiu, Kunal Bhalla, Kushal Lakhotia, Lauren
 348 Rantala-Yearly, Laurens van der Maaten, Lawrence Chen, Liang Tan, Liz Jenkins, Louis Martin,
 349 Lovish Madaan, Lubo Malo, Lukas Blecher, Lukas Landzaat, Luke de Oliveira, Madeline Muzzi,
 350 Mahesh Pasupuleti, Mannat Singh, Manohar Paluri, Marcin Kardas, Maria Tsimpoukelli, Mathew
 351 Oldham, Mathieu Rita, Maya Pavlova, Melanie Kambadur, Mike Lewis, Min Si, Mitesh Ku-
 352 mar Singh, Mona Hassan, Naman Goyal, Narjes Torabi, Nikolay Bashlykov, Nikolay Bogoy-
 353 chev, Niladri Chatterji, Ning Zhang, Olivier Duchenne, Onur Çelebi, Patrick Alrassy, Pengchuan
 354 Zhang, Pengwei Li, Petar Vasic, Peter Weng, Prajjwal Bhargava, Pratik Dubal, Praveen Krishnan,
 355 Punit Singh Koura, Puxin Xu, Qing He, Qingxiao Dong, Ragavan Srinivasan, Raj Ganapathy, Ra-
 356 mon Calderer, Ricardo Silveira Cabral, Robert Stojnic, Roberta Raileanu, Rohan Maheswari, Ro-
 357 hit Girdhar, Rohit Patel, Romain Sauvestre, Ronnie Polidoro, Roshan Sumbaly, Ross Taylor, Ruan
 358 Silva, Rui Hou, Rui Wang, Saghar Hosseini, Sahana Chennabasappa, Sanjay Singh, Sean Bell,
 359 Seohyun Sonia Kim, Sergey Edunov, Shaoliang Nie, Sharan Narang, Sharath Rapparthi, Sheng
 360 Shen, Shengye Wan, Shruti Bhosale, Shun Zhang, Simon Vandenhende, Soumya Batra, Spencer
 361 Whitman, Sten Sootla, Stephane Collet, Suchin Gururangan, Sydney Borodinsky, Tamar Herman,
 362 Tara Fowler, Tarek Sheasha, Thomas Georgiou, Thomas Scialom, Tobias Speckbacher, Todor Mi-
 363 haylov, Tong Xiao, Ujjwal Karn, Vedanuj Goswami, Vibhor Gupta, Vignesh Ramanathan, Viktor
 364 Kerkez, Vincent Gonguet, Virginie Do, Vish Vogeti, Vitor Albiero, Vladan Petrovic, Weiwei
 365 Chu, Wenhan Xiong, Wenyin Fu, Whitney Meers, Xavier Martinet, Xiaodong Wang, Xiaofang
 366 Wang, Xiaoqing Ellen Tan, Xide Xia, Xinfeng Xie, Xuchao Jia, Xuwei Wang, Yaelle Gold-
 367 schlag, Yashesh Gaur, Yasmine Babaei, Yi Wen, Yiwen Song, Yuchen Zhang, Yue Li, Yuning
 368 Mao, Zacharie DelPierre Coudert, Zheng Yan, Zhengxing Chen, Zoe Papakipos, Aaditya Singh,
 369 Aayushi Srivastava, Abha Jain, Adam Kelsey, Adam Shajnfeld, Adithya Gangidi, Adolfo Victoria,
 370 Ahuva Goldstand, Ajay Menon, Ajay Sharma, Alex Boesenberg, Alexei Baevski, Allie Feinstein,
 371 Amanda Kallet, Amit Sangani, Amos Teo, Anam Yunus, Andrei Lupu, Andres Alvarado, An-
 372 drew Caples, Andrew Gu, Andrew Ho, Andrew Poulton, Andrew Ryan, Ankit Ramchandani, An-
 373 nie Dong, Annie Franco, Anuj Goyal, Aparajita Saraf, Arkabandhu Chowdhury, Ashley Gabriel,
 374 Ashwin Bharambe, Assaf Eisenman, Azadeh Yazdan, Beau James, Ben Maurer, Benjamin Leon-
 375 hardi, Bernie Huang, Beth Loyd, Beto De Paola, Bhargavi Paranjape, Bing Liu, Bo Wu, Boyu
 376 Ni, Braden Hancock, Bram Wasti, Brandon Spence, Brani Stojkovic, Brian Gamido, Britt Mon-
 377 talvo, Carl Parker, Carly Burton, Catalina Mejia, Ce Liu, Changan Wang, Changkyu Kim, Chao
 378 Zhou, Chester Hu, Ching-Hsiang Chu, Chris Cai, Chris Tindal, Christoph Feichtenhofer, Cynthia
 379 Gao, Damon Civin, Dana Beaty, Daniel Kreymer, Daniel Li, David Adkins, David Xu, Davide
 380 Testuggine, Delia David, Devi Parikh, Diana Liskovich, Didem Foss, Dingkan Wang, Duc Le,
 381 Dustin Holland, Edward Dowling, Eissa Jamil, Elaine Montgomery, Eleonora Presani, Emily
 382 Hahn, Emily Wood, Eric-Tuan Le, Erik Brinkman, Esteban Arcaute, Evan Dunbar, Evan Smoth-
 383 ers, Fei Sun, Felix Kreuk, Feng Tian, Filippos Kokkinos, Firat Ozgenel, Francesco Caggioni,
 384 Frank Kanayet, Frank Seide, Gabriela Medina Florez, Gabriella Schwarz, Gada Badeer, Georgia
 385 Swee, Gil Halpern, Grant Herman, Grigory Sizov, Guangyi, Zhang, Guna Lakshminarayanan,
 386 Hakan Inan, Hamid Shojanazeri, Han Zou, Hannah Wang, Hanwen Zha, Haroun Habeeb, Harri-
 387 son Rudolph, Helen Suk, Henry Aspegren, Hunter Goldman, Hongyuan Zhan, Ibrahim Damlaj,
 388 Igor Molybog, Igor Tufanov, Ilias Leontiadis, Irina-Elena Veliche, Itai Gat, Jake Weissman, James
 389 Geboski, James Kohli, Janice Lam, Japhet Asher, Jean-Baptiste Gaya, Jeff Marcus, Jeff Tang, Jen-

390 nifer Chan, Jenny Zhen, Jeremy Reizenstein, Jeremy Teboul, Jessica Zhong, Jian Jin, Jingyi Yang,
 391 Joe Cummings, Jon Carvill, Jon Shepard, Jonathan McPhie, Jonathan Torres, Josh Ginsburg, Jun-
 392 jie Wang, Kai Wu, Kam Hou U, Karan Saxena, Kartikay Khandelwal, Katayoun Zand, Kathy
 393 Matosich, Kaushik Veeraraghavan, Kelly Michelena, Keqian Li, Kiran Jagadeesh, Kun Huang,
 394 Kunal Chawla, Kyle Huang, Lailin Chen, Lakshya Garg, Lavender A, Leandro Silva, Lee Bell,
 395 Lei Zhang, Liangpeng Guo, Licheng Yu, Liron Moshkovich, Luca Wehrstedt, Madian Khabsa,
 396 Manav Avalani, Manish Bhatt, Martynas Mankus, Matan Hasson, Matthew Lennie, Matthias
 397 Reso, Maxim Groshev, Maxim Naumov, Maya Lathi, Meghan Keneally, Miao Liu, Michael L.
 398 Seltzer, Michal Valko, Michelle Restrepo, Mihir Patel, Mik Vyatskov, Mikayel Samvelyan, Mike
 399 Clark, Mike Macey, Mike Wang, Miquel Jubert Hermoso, Mo Metanat, Mohammad Rastegari,
 400 Munish Bansal, Nandhini Santhanam, Natascha Parks, Natasha White, Navyata Bawa, Nayan
 401 Singhal, Nick Egebo, Nicolas Usunier, Nikhil Mehta, Nikolay Pavlovich Laptev, Ning Dong,
 402 Norman Cheng, Oleg Chernoguz, Olivia Hart, Omkar Salpekar, Ozlem Kalinli, Parkin Kent,
 403 Parth Parekh, Paul Saab, Pavan Balaji, Pedro Rittner, Philip Bontrager, Pierre Roux, Piotr Dollar,
 404 Polina Zvyagina, Prashant Ratanchandani, Pritish Yuvraj, Qian Liang, Rachad Alao, Rachel Ro-
 405 driguez, Rafi Ayub, Raghotham Murthy, Raghu Nayani, Rahul Mitra, Rangaprabhu Parthasarathy,
 406 Raymond Li, Rebekkah Hogan, Robin Battey, Rocky Wang, Russ Howes, Ruty Rinott, Sachin
 407 Mehta, Sachin Siby, Sai Jayesh Bondu, Samyak Datta, Sara Chugh, Sara Hunt, Sargun Dhillon,
 408 Sasha Sidorov, Satadru Pan, Saurabh Mahajan, Saurabh Verma, Seiji Yamamoto, Sharadh Ra-
 409 maswamy, Shaun Lindsay, Shaun Lindsay, Sheng Feng, Shenghao Lin, Shengxin Cindy Zha,
 410 Shishir Patil, Shiva Shankar, Shuqiang Zhang, Shuqiang Zhang, Sinong Wang, Sneha Agarwal,
 411 Soji Sajuyigbe, Soumith Chintala, Stephanie Max, Stephen Chen, Steve Kehoe, Steve Satter-
 412 field, Sudarshan Govindaprasad, Sumit Gupta, Summer Deng, Sungmin Cho, Sunny Virk, Suraj
 413 Subramanian, Sy Choudhury, Sydney Goldman, Tal Remez, Tamar Glaser, Tamara Best, Thilo
 414 Koehler, Thomas Robinson, Tianhe Li, Tianjun Zhang, Tim Matthews, Timothy Chou, Tzook
 415 Shaked, Varun Vontimitta, Victoria Ajayi, Victoria Montanez, Vijai Mohan, Vinay Satish Ku-
 416 mar, Vishal Mangla, Vlad Ionescu, Vlad Poenaru, Vlad Tiberiu Mihailescu, Vladimir Ivanov,
 417 Wei Li, Wenchen Wang, Wenwen Jiang, Wes Bouaziz, Will Constable, Xiaocheng Tang, Xiao-
 418 jian Wu, Xiaolan Wang, Xilun Wu, Xinbo Gao, Yaniv Kleinman, Yanjun Chen, Ye Hu, Ye Jia,
 419 Ye Qi, Yenda Li, Yilin Zhang, Ying Zhang, Yossi Adi, Youngjin Nam, Yu, Wang, Yu Zhao,
 420 Yuchen Hao, Yundi Qian, Yunlu Li, Yuzi He, Zach Rait, Zachary DeVito, Zef Rosnbrick, Zhao-
 421 duo Wen, Zhenyu Yang, Zhiwei Zhao, and Zhiyu Ma. The llama 3 herd of models, 2024. URL
 422 <https://arxiv.org/abs/2407.21783>.

423 Mistral AI Team. Mistral small 3, Jan 2025. URL [https://mistral.ai/news/](https://mistral.ai/news/mistral-small-3)
 424 [mistral-small-3](https://mistral.ai/news/mistral-small-3).

A Dataset construction

A.1 SCLI5

SCLI5 is constructed by introducing error to answer of simple tasks, that can be corrected simply with a recall of answer, or limited reasoning. Starting with basic task can isolate any confounding error such as model internal knowledge and reasoning capability. Most error types are off-by-one error or flip error. The composition of the task is shown in Table 2. Due to the simplicity of the task, we do not add any reasoning before answering.

After designing the composition and error, we generated the dataset programmatically.

Although LLMs are unlikely to make these simple mistakes, self-correction requires error detection. If it cannot detect an obvious error, it cannot detect a subtle one.

Table 2: Task composition of SCLI5

Task	Count	Error Type	Question and Answer
Add one	20	Off-by-one	Q: What is the answer of $1 + 1$? A: The answer is 3.
Subtract one	20	Off-by-one	Q: What is the answer of $3 - 1$? A: The answer is 1.
Next character	52	Off-by-one	Q: What letter comes after A? A: The answer is C.
Previous character	52	Off-by-one	Q: What letter comes before C? A: The answer is A.
Larger number	71	Flip	Q: Which one is smaller, 1 or 2? A: The answer is 2.
Smaller number	71	Flip	Q: Which one is larger, 2 or 5? A: The answer is 2.

A.2 GSM8K-SC

Unlike SCLI5, solving GSM8K (Cobbe et al. [2021]) requires multiple steps reasoning. As such, we injected an error in reasoning. The answer is also incorrect in a way that follows the incorrect reasoning. We introduce various types of error as stated in Table 3.

Using the test set, we generated wrong reasoning using ‘gpt-4.1-2025-04-14’[OpenAI, 2025] with a prompt that contains the question, reasoning steps, correct answer, and error type, and order of reasoning step. The prompt can be found in the appendix.

To ensure consistency of incorrect reasoning steps and incorrect answer, we prompted ‘gemini-2.5-flash-preview-05-20’[Gemini Team, 2025] to follow the given incorrect reasoning to arrive at the incorrect answer. Of 1,319 questions, 6 cannot arrive at the answer, and therefore, we dropped them.

A.3 PRM800K-SC

PRM800K[Lightman et al., 2024], derived from a subset of MATH[Hendrycks et al., 2021], provides step-by-step annotations of multistep reasoning. The solution is pre-generated by finetuned GPT4[OpenAI et al., 2024], reflecting the realism of error.

We excluded quality control and initial screening questions and those identified as a bad problem and the annotators had given up. To increase complexity, we only selected samples whose pre-generated answer does not match golden answer. Finally, we randomly sampled 1 completion for each question.

Table 3: Error composition of GSM8K-SC

Category	Description
Problem Representation Errors	These errors arise when the solver misunderstands or misinterprets the problem’s requirements or given information. This can involve misreading the problem statement, confusing the relationships between quantities, or failing to grasp what is being asked.
Planning Errors	These occur when the solver devises an incorrect or incomplete strategy to tackle the problem. This might include choosing the wrong operations, setting up flawed equations, or overlooking key components of the problem.
Execution Errors	These are mistakes made while carrying out the planned steps, such as errors in calculations, misapplication of mathematical rules, or procedural slip-ups, even if the plan itself is sound.

B Sensitivity analysis

B.1 Result of different temperature

Mean accuracy measures accuracy given an error in model generation. Result under a different temperature does not change our conclusion.

Table 4: Mean accuracy of models at temperature 0.0

Model	SCLi5	GSM8K-SC	PRM800K-SC
Llama-4-Maverick-17B-128E-Instruct-FP8 [Meta, 2025]	0.948	0.416	0.455
DeepSeek-V3-0324 [DeepSeek-AI et al., 2025b]	0.825	0.399	0.475
Qwen2.5-72B-Instruct [Qwen et al., 2025]	0.92	0.58	0.154
Llama-4-Scout-17B-16E-Instruct [Meta, 2025]	0.976	0.24	0.263
Llama-3.3-70B-Instruct [Meta, 2024]	0.538	0.275	0.246
Qwen3-235B-A22B [Yang et al., 2025]	0.563	0.073	0.348
phi-4 [Abdin et al., 2024]	0.808	0.076	0.092
Qwen2.5-7B-Instruct [Qwen et al., 2025]	0.559	0.19	0.141
Qwen2-7B-Instruct [Yang et al., 2024]	0.601	0.078	0.058
Qwen3-14B [Yang et al., 2025]	0.004	0.092	0.254
Qwen3-30B-A3B [Yang et al., 2025]	0.056	0.061	0.194
Llama-3.1-8B-Instruct [Grattafiori et al., 2024]	0.136	0.019	0.02
Qwen3-32B [Yang et al., 2025]	0.004	0.05	0.083
Mistral-Small-24B-Instruct-2501 [Team, 2025]	0.042	0.011	0.016

Table 5: Mean accuracy of models at temperature 0.6

Model	SCLI5	GSM8K-SC	PRM800K-SC
Llama-4-Maverick-17B-128E-Instruct-FP8	0.955	0.423	0.469
DeepSeek-V3-0324	0.874	0.42	0.504
Qwen2.5-72B-Instruct	0.902	0.573	0.165
Llama-4-Scout-17B-16E-Instruct	0.976	0.248	0.272
Llama-3.3-70B-Instruct	0.497	0.273	0.243
Qwen3-235B-A22B	0.57	0.091	0.4
phi-4	0.794	0.093	0.116
Qwen2.5-7B-Instruct	0.563	0.183	0.127
Qwen2-7B-Instruct	0.601	0.071	0.065
Qwen3-14B	0.007	0.101	0.27
Qwen3-30B-A3B	0.108	0.07	0.232
Qwen3-32B	0.038	0.068	0.105
Meta-Llama-3.1-8B-Instruct	0.182	0.025	0.022
Mistral-Small-24B-Instruct-2501	0.122	0.02	0.038

C Figures

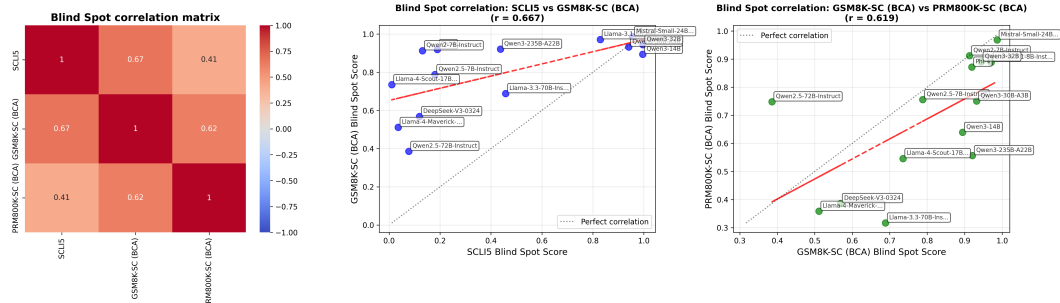


Figure 3: *left*: Blind spot correlation matrix *middle*: Scatter plot between SCL15 vs GSM8K-SC *right*: Scatter plot between GSM8K-SC vs PRM800K-SC
BCA: Before commit an answer

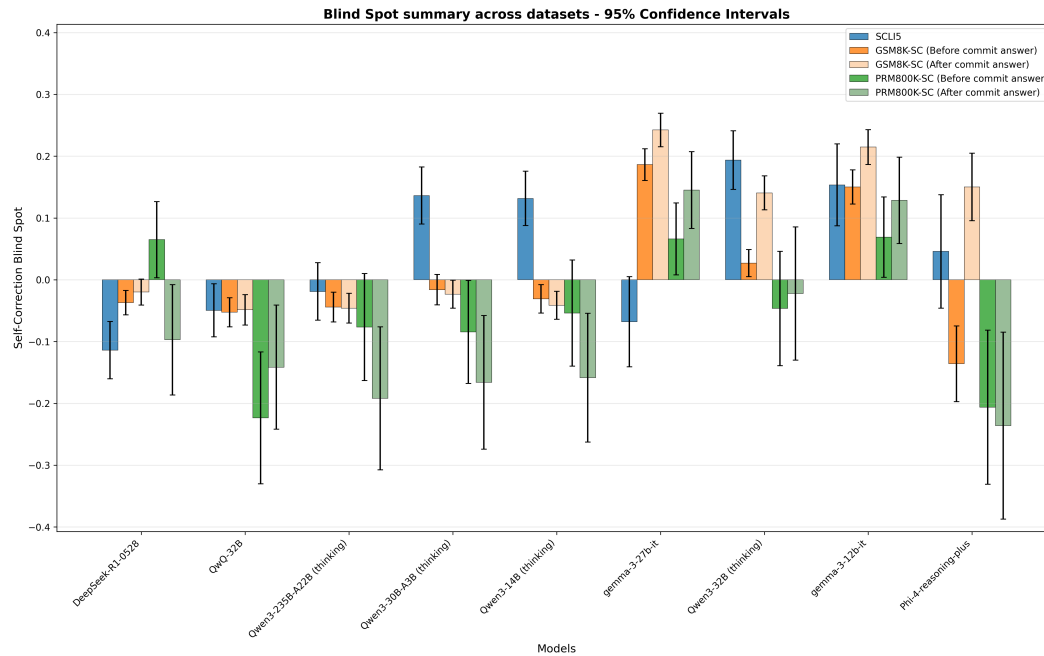


Figure 4: Self-Correction Blind Spot across reasoning models

459 **D Prompt**

460 **D.1 Generating GSM8K-SC**

```
from pydantic import BaseModel

class ReasoningWithMistake(BaseModel):
    reasoning_steps_with_one_mistake: List[str]
    mistake_step: int
    type_of_mistake: str
    description_of_mistake: str
    incorrect_answer: str
```

You are a helpful assistant that follow instructions. Output in JSON format.

```
<question>
{question}
</question>

<reasoning_steps>
{reasoning_steps}
</reasoning_steps>

<answer>
{answer}
</answer>

<type_of_mistake>
{error_type}: {error_description}
</type_of_mistake>

You task is to introduce one mistake in step {mistake_step} in
<reasoning_steps> and arrive at an answer different from
<answer>.
You will output:
- <reasoning_steps> with mistake
- the step that contains the mistake
- type of the mistake
- description of the mistake
- incorrect answer
```

Figure 5: Output schema, system prompt and prompt for generating GSM8K-SC dataset

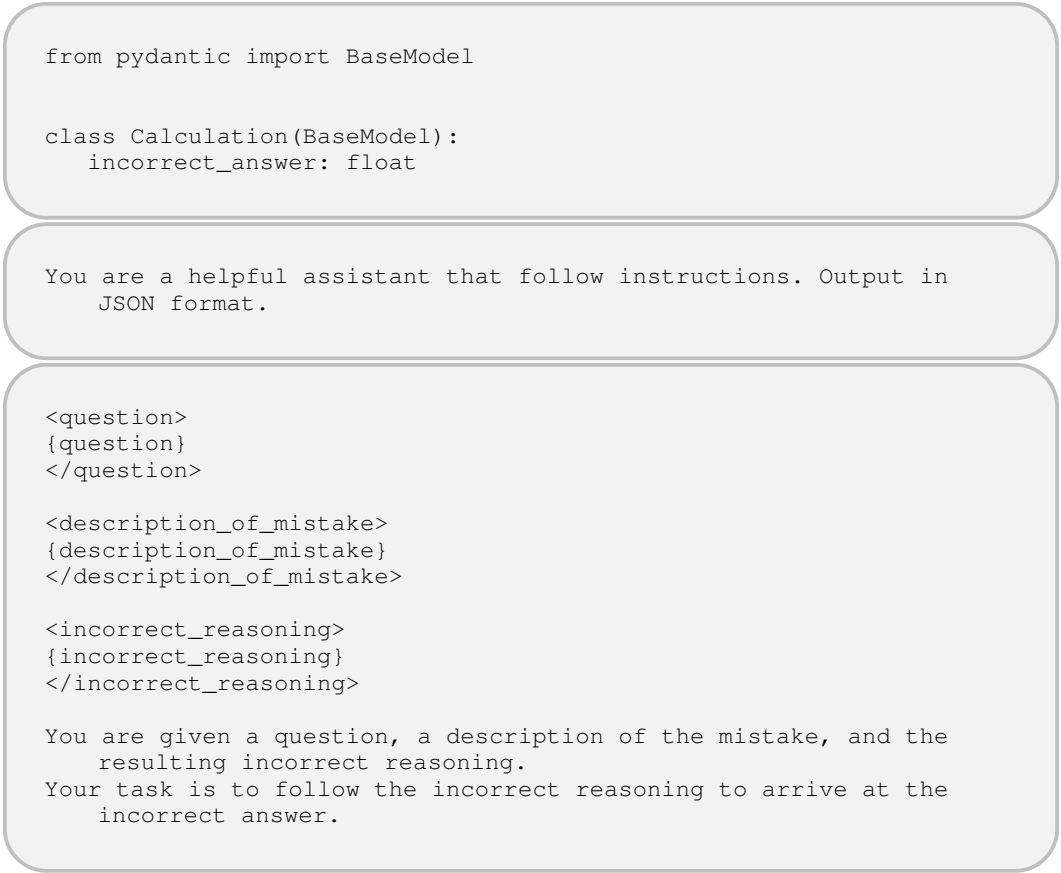


Figure 6: Output schema, system prompt and prompt for validating GSM8K-SC dataset

461 **D.2 Automatic evaluation**

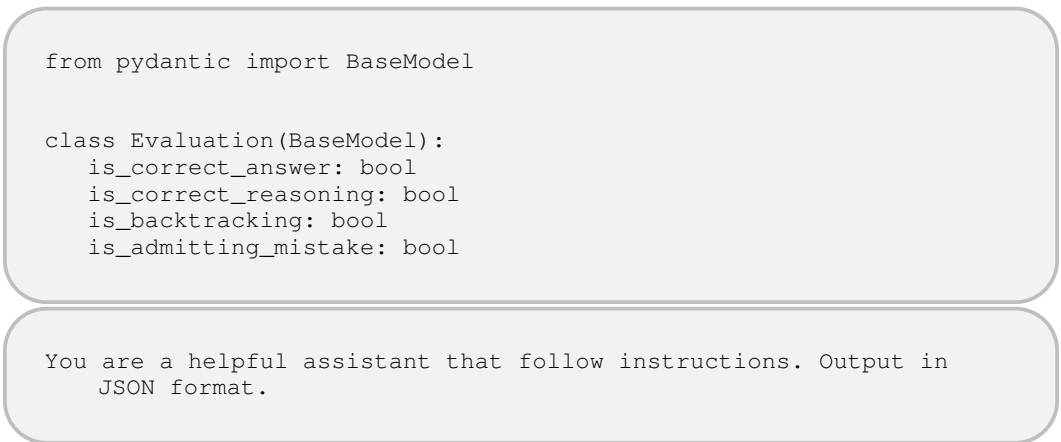


Figure 7: Output schema and system prompt for automatic evaluation of error in model

```

# FOR SCLI5
<question>
{question}
</question>

<golden_answer>
{golden_answer}
</golden_answer>

<given_wrong_answer>
{given_wrong_answer}
</given_wrong_answer>

<completion_from_model>
{completion_from_model}
</completion_from_model>

The model was provided a wrong answer <given_wrong_answer>, and is
  required to self correct and arrive at the golden answer.
You have to assess if <completion_from_model> :
- has backtracked
- has finally provided the answer that matches the <golden_answer>
- has provided correct reasoning
- has admitted the mistake.

```

```

# FOR GSM8K-SC AND PRM800K-SC
<question>
{question}
</question>

<golden_answer>
{golden_answer}
</golden_answer>

<given_wrong_reasoning>
{given_wrong_reasoning}
</given_wrong_reasoning>

<completion_from_model>
{completion_from_model}
</completion_from_model>

The model was provided a wrong step in the reasoning process
  <given_wrong_reasoning>, and is required to self correct and
  arrive at the golden answer.
You have to assess if <completion_from_model> :
- has backtracked
- has finally provided the answer that matches the <golden_answer>
- has provided correct reasoning
- has admitted the mistake.

```

Figure 8: Prompt for automatic evaluation of error in model

```
from pydantic import BaseModel
```

```
class EvaluationErrorInUser(BaseModel):  
    is_correct_answer: bool  
    is_correct_critique: bool
```

You are a helpful assistant that follow instructions. Output in JSON format.

Figure 9: Output schema and system prompt for automatic evaluation of error in user

```
# FOR SCLI5
<question_and_user_answer>
{question_and_user_answer}
</question_and_user_answer>

<golden_answer>
{golden_answer}
</golden_answer>

<response_from_model>
{response_from_model}
</response_from_model>

The model was provided with <question_and_user_answer> from user.
You have to assess if <response_from_model> :
- contains correct answer that matches the <golden_answer>
- contains correct critique of <question_and_user_answer>
```

```
# FOR GSM8K-SC AND PRM800K-SC
<question_and_user_reasoning>
{question_and_user_reasoning}
</question_and_user_reasoning>

<golden_answer>
{golden_answer}
</golden_answer>

<response_from_model>
{response_from_model}
</response_from_model>

The model was provided with <question_and_user_reasoning> from
user.
You have to assess if <response_from_model> :
- contains correct answer that matches the <golden_answer>
- contains correct critique of <question_and_user_reasoning>
```

Figure 10: Prompt for automatic evaluation of error in user