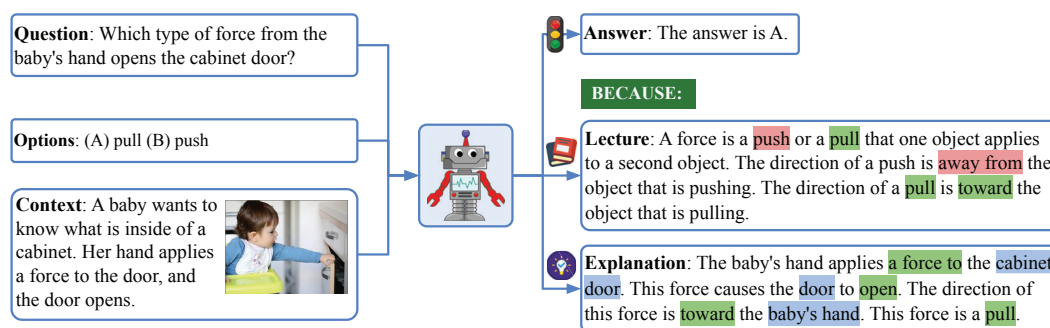

Learn to Explain: Multimodal Reasoning via Thought Chains for Science Question Answering

Pan Lu^{1,3}, Swaroop Mishra^{2,3}, Tony Xia¹, Liang Qiu¹, Kai-Wei Chang¹,
Song-Chun Zhu¹, Oyvind Tafjord³, Peter Clark³, Ashwin Kalyan³



(b) We show that CoT benefits large language models in both few-shot and fine-tuning learning by improving model performance and reliability via generating explanations. (c) We further explore the upper bound of GPT-3 and show that CoT helps language models learn from fewer data.

2 Related Work

Visual question answering. Since the task of visual question answering (VQA) was first proposed in [2], there have been plenty of VQA datasets [56, 58, 23, 11

Statistic	Number
Total questions	21,208
Questions with text context	10,220 (48.2%)
Questions with image context	10,332 (48.7%)
* Image of natural format	≈2,960 (14.0%)
* Image of diagram format	≈7,372 (34.8%)
Questions with both contexts	6,532 (

Biology Genes to traits Classification Adaptations Traits and heredity Ecosystems Classification Scientific names Heredity Ecological interactions Cells Plants Animals Plant reproduction	Physics Materials Magnets Velocity and forces Force and motion Particle motion and energy Heat and thermal energy States of matter Kinetic and potential energy Mixture	Geography State capitals Geography Maps Oceania
--	---	--

4 Baselines and Chain-of-Thought Models

In this section, we establish baselines and develop two chain-of-thought models on SCIENCEQA.

4.1 Baselines

Heuristic baselines. The first heuristic baseline is *random chance*: we randomly select one from the multiple options. Each trial is completed on the whole test set, and we take three different trials

Question: question : I_i^{ques}
 Options: (A) option : I_{i1}^{opt} (B) option : I_{i2}^{opt} (C) option : I_{i3}^{opt}
 Context: context : I_i^{cont}
 Answer: The answer is answer : I_i^a . BECAUSE: lecture : I_i^{lect} explanation : I_i^{exp}

Question: question : I_t^{ques}
 Options

Model	Learning	Format	NAT	SOC	LAN	TXT	IMG	NO	G1-6	G7-12	Avg
Random chance	-	M→A	40.28	46.13	29.25	47.45	40.08	33.66	39.35	40.67	39.83

Similarity [49] to evaluate the generated lectures and explanations, as shown in Table 4. The Similarity metric computes the cosine-similarity of semantic embeddings between two sentences based on the Sentence-BERT network [49]. The results show that UnifiedQA_{BASE} (CoT) generates the most similar explanations to the given ones. However, it's commonly agreed that automatic evaluation of generated texts only provides a partial view and has to be complemented by a human study. By asking annotators to rate the relevance, correctness, and completeness of generated explanations, we find that the explanations generated by GPT-3 (CoT) conform best to human judgment.

Prompt type	Sampling	Acc. (%)
QCML*→A	Random	73.59
QCML*→AE	Random	74.32
QCME*→A	Random	94.03 _{18.86} ↑
QCMLE*→A	Random	94.13 _{18.96} ↑

- [3] Jonathan Bragg, Arman Cohan, Kyle Lo, and Iz Beltagy. Flex: Unifying evaluation for few-shot nlp. *Advances in Neural Information Processing Systems (NeurIPS)*, 34, 2021.
- [4] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. Language models are few-shot learners. *Advances in neural information processing systems (NeurIPS)*, 33:1877–1901, 2020.
- [

- [20] Tushar Khot, Peter Clark, Michal Guerquin, Peter Alexander Jansen, and Ashish Sabharwal. Qasc: A dataset for question answering via sentence composition. *ArXiv*, abs/1910.11473, 2020.
- [21] Jin-Hwa Kim, Jaehyun Jun, and Byoung-Tak Zhang. Bilinear attention networks. In *Advances in Neural Information Processing Systems (NeurIPS)*, pages 1571–1581, 2018.
- [22] Wonjae Kim, Bokyoung Son, and Ildoo

- [38] Swaroop Mishra, Matthew Finlayson, Pan Lu, Leonard Tang, Sean Welleck, Chitta Baral, Tanmay Rajpurohit, Oyvind Tafjord, Ashish Sabharwal, Peter Clark, and Ashwin Kalyan. Lila: A unified benchmark for mathematical reasoning. In *The 2022 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2022.
- [39] Swaroop Mishra, Daniel Khashabi, Chitta Baral, Yejin Choi, and Hannaneh Hajishirzi. Reframing instructional prompts to g

- [56] Peng Zhang, Yash Goyal, Douglas Summers-Stay, Dhruv Batra, and Devi Parikh. Yin and Yang: Balancing and answering binary visual questions. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016.
- [57] Zihao Zhao, Eric Wallace, Shi Feng, Dan Klein, and Sameer Singh. Calibrate before use: Improving few-shot performance of language models. In *International Conference on Machine Learning (ICML)*, pages 12697–12706. PMLR, 2021.
- [58] Yuke Zhu

- (b) Did you describe any potential participant risks, with links to Institutional Review Board (IRB) approvals, if applicable? [N/A]
- (c) Did you include the estimated hourly wage paid to participants and the total amount spent on participant compensation? [Yes] We included the monetary compensation details in Appendix B.2 and B.3.