

Status Empty

Assignee Empty

한 번에 요구 사항을 프롬프트에 넣어야 하는 까닭

20만 건의 대화가 증명한 것: AI 는 대화를 못 한다

LLM을 "대화형 도구"로 사용하지만, 정작 LLM은 대화에 최적화 되어 있지 않습니다.

Microsoft Research와 Salesforce의 공동 연구는 "LLMs Get Lost In Multi-Turn Conversation" 에서 LLM의 치명적 약점을 정량적으로 증명했습니다.

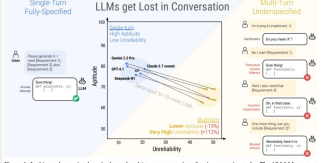


Figure 1: In this work, we simulate single- and multi-turn conversations for six generation tasks. The 15 LLMs we test perform much worse in multi-turn settings (~35%) explained by some loss in aptitude, and large losses in reliability. Aptitude is defined as performance in best-case conversation simulation, and unreliability as the gap between best- and worst-case performance. In short, we find that LLMs get lost in multi-turn, underspecified conversation.

Microsoft Research와 Salesforce의 공동 연구 "LLMs Get Lost In Multi-Turn Conversation"은 이 사실을 정량적으로 증명했습니다. 15개 주요 LLM, 20만 건 이상의 대화 시뮬레이션 결과, 멀티턴 대화에서 LLM 성능은 평균 39% 하락했습니다. GPT-4.1, Gemini 2.5 Pro, Claude 3.7 Sonnet 등 최신 모델도 예외가 없었고, o3-DeepSeek-R1 같은 추론 특화 모델조차 동일하게 무너졌습니다.

Lost in Conversation Experiment

Model	FULL						CONCAT						SHARDED						Overall	
	Code	Database	Actions	Data-to-text	Math	Summary	Code	Database	Actions	Data-to-text	Math	Summary	Code	Database	Actions	Data-to-text	Math	Summary	Full/Concat	Full/Sharded
GPT-4.1	27.4	64.1	82.9	11.7	61.9	7.8	21.2	40.3	83.0	11.7	62.6	6.5	21.7	25.8	45.3	13.3	37.4	3.4	91.6	62.5
Gemini 2.5 Pro	18.8	54.8	96.1	17.2	80.0	-	16.3	48.3	89.8	14.3	80.1	-	14.4	22.4	13.8	8.0	46.3	-	86.3	30.3
Claude 3.7 Sonnet	41.8	83.0	83.3	29.8	73.9	11.6	36.3	76.3	88.2	30.1	76.1	9.2	31.5	31.8	55.9	18.6	47.1	1.6	91.6	32.6
o3-mini	75.9	89.3	94.1	35.9	88.1	14.9	66.7	90.7	92.1	31.2	88.0	11.5	69.3	40.2	52.4	18.8	58.7	7.2	95.8	36.2
o3-1.5-200B	72.0	91.1	93.0	34.1	91.7	15.8	52.3	87.8	91.0	32.0	91.8	14.7	51.6	35.4	71.8	22.4	61.3	30.3	93.2	64.2
o3-1.5-70B	53.2	87.6	82.7	23.9	89.2	-	48.4	79.6	76.0	26.6	90.4	-	39.1	31.1	34.1	23.3	52.5	-	99.8	61.7
o3-1.5-30B	72.0	91.9	98.3	27.7	94.3	24.3	61.6	88.1	88.4	31.2	91.9	21.3	68.8	38.8	72.0	27.9	68.0	4.9	97.3	60.4
o3-1.5-7B	73.9	92.7	98.0	35.2	96.3	13.7	60.3	81.3	88.3	28.1	92.9	11.7	65.4	27.1	68.9	26.1	67.0	12.3	91.8	66.1
o3-1.5-1.7B	86.4	92.0	89.8	40.2	81.6	30.7	67.2	83.3	91.3	29.4	80.0	30.4	55.8	35.4	68.2	20.7	63.1	26.5	98.1	64.1
o3-1.5-0.7B	78.0	93.9	93.4	45.6	85.8	29.3	76.2	81.3	86.0	31.3	87.2	28.9	65.6	36.8	33.3	35.1	70.0	21.6	100.4	63.9
o3-1.5-0.3B	98.4	92.1	97.0	27.0	95.5	26.1	90.1	88.9	97.0	30.7	92.9	34.6	79.9	31.5	47.5	28.0	67.3	17.2	105.6	60.8
o3-1.5-0.1B	88.4	93.6	96.1	42.1	93.8	23.8	82.8	91.7	91.1	32.1	91.9	23.9	61.3	42.3	48.8	28.8	67.9	20.6	96.3	57.9
o3-1.5-0.05B	97.0	96.3	98.4	31.2	90.6	29.1	92.5	95.3	89.2	31.9	88.4	29.4	68.3	51.3	42.6	31.0	66.1	26.1	99.3	63.8
o3-1.5-0.01B	98.6	93.0	94.7	34.6	91.7	26.5	88.7	88.3	98.3	34.4	89.7	26.8	72.6	46.8	42.9	28.6	70.7	33.3	97.8	61.8
o3-1.5-0.001B	97.4	97.3	97.8	34.8	90.2	31.2	95.7	94.9	98.1	26.9	89.3	31.8	68.2	40.8	36.3	46.2	64.3	24.9	100.1	64.5

Table 1: Averaged Performance ($\bar{P}$) of LLMs on six tasks (Code, Database, Actions, Data-to-text, Math, and Summary). For each task, conversations are simulated in three settings: FULL, CONCAT, and SHARDED. Models are sorted in ascending order of average FULL scores across tasks. Background color indicates the level of degradation from the FULL setting. The last two columns average the performance drops from the CONCAT and SHARDED compared to the FULL in percentages across the six tasks.

멀티턴 기반의 에이전트나 플랫폼을 사용하다 보면 대화가 길어질 수록 답변 품질이 떨어지는 것을 체감하게 되는데, 이유를 데이터 기반으로 설명할 수 있습니다.

멀티턴 대화에서 품질이 떨어지는 4가지 원인

1. 정보가 부족한 초반에 성급하게 답변 시도
2. 이전의 틀린 답변에 과도하게 의존
3. 중간 턴의 정보를 잊어 버림
4. 지나치게 장황한 응답으로 잘못된 가정 삽입

Lost in Conversation 현상

대화 초반, 정보가 부족한 상태에서 성급하게 답변을 시도하기 때문 입니다. 초기에 세운 가정이 틀리면 이후 턴에서 그 틀린 답변에 오 히려 더 의존하는 악순환이 생깁니다. Chain of Thought 같은 단일 추론 기법은 물론, Tree of Thought-Graph of Thought 같은 복합 추론 기법에서도 마찬가지입니다. 한 번 잘못 추론하면 이후 답변도 연쇄적으로 틀려버립니다.

고품질 답변을 위한 프롬프트 엔지니어링 방법

1 한 턴에 모든 요구사항을 전달하기

여러 턴에 걸쳐 조건을 추가할수록 답변 품질은 급격히 저하됩니다. 한 번에 많은 맥락을 프롬프트에 담아내려면 구조화는 필수입니다.

2 대화가 꼬이면 이어가지 말고 새로 시작하기

"지금까지 말한 내용을 정리해줘"라고 요청한 뒤, 그 결과를 복사해 새 대화에 붙여넣기만 해도 품질이 눈에 띄게 회복됩니다.

3 긴 대화 시 매 턴마다 요약 쌓기

매 턴마다 이전 정보를 반복 포함시키는 Snowball 방식은 답변 고 임을 15~20% 회복시키는 효과를 보였습니다.

4 Temperature 를 낮춰도 근본적인 해결책은 아님

Temperature를 0으로 설정하면 일관성과 신뢰성은 높아지지만, 멀티턴 환경에서는 0으로 설정해도 여전히 높은 불안정성이 관찰됩니다.

LLM의 벤치마크의 높은 점수와 실제 사용 경험 사이의 간극, 이 논문이 그 간극을 숫자로 보여줍니다. 프롬프트 엔지니어링은 모델의 한계를 이해는 데서 시작합니다.

#LLM #AI #conversation#promptengineering

#AIResearch #GenerativeAI #prompting

LLMs Get Lost In Multi-Turn Conversation

Large Language Models (LLMs) are conversational interfaces. As such, LLMs have the potential to assist their...

arxiv.org



Subscribe to 'sujin-prompt-engineer'

안녕하세요,
슬래시페이지 구독을 하시면, 이따금씩 발행하는 프롬프트와 프
롬프트 엔지니어링에 관한 글을 이메일로 받아보실 수 있어요.

구독하시겠어요? 😊