

Logical forms complement probability in understanding language model (and human) performance

Anonymous ACL submission

Abstract

With the increasing interest in using large language models (LLMs) for planning in natural language, understanding their behaviors becomes an important research question. This work conducts a systematic investigation of LLMs’ ability to perform logical reasoning in *natural language*. We introduce a controlled dataset of hypothetical and disjunctive syllogisms in propositional and modal logic and use it as the testbed for understanding LLM performance. Our results lead to novel insights in predicting LLM behaviors: in addition to the probability of input (Gonen et al., 2023; McCoy et al., 2024), logical forms should be considered as important factors. In addition, we show similarities and discrepancies between the logical reasoning performances of humans and LLMs by collecting and comparing behavioral data from both.

1 Introduction

Logical reasoning is a fundamental aspect of building AI systems for reliable decision-making (Kautz et al., 1992, *inter alia*)—given a set of premises, an AI system should be able to deduce valid conclusions. With the advent of large language models (LLMs; Touvron et al., 2023; Jiang et al., 2023; AI@Meta, 2024, *inter alia*), there has been a surge of interest in using these models to assist planning and decision-making (Huang et al., 2022, *inter alia*); therefore, understanding the logical reasoning capabilities becomes crucial in understanding the reliability and potential of LLMs in planning. While recent work has shown that LLMs exhibit decent performance on logical reasoning problems (Liu et al., 2020; Ontanon et al., 2022; Wan et al., 2024, *inter alia*), there is still a lack of fine-grained understanding of the logical forms—among many argument forms presented in natural language (Shieber, 1993), do LLMs perform equally well, or do they exhibit preferences for

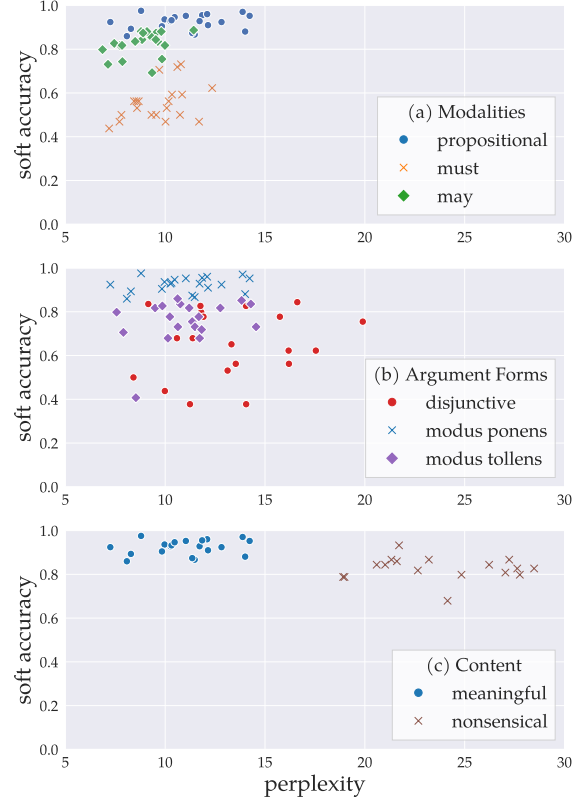


Figure 1: Illustration of the fact that perplexity does not serve as a reliable indicator of logical reasoning performance; and therefore, neither does probability. The distributions of the probabilities assigned to the ground-truth answer (i.e., soft accuracy; Y-axis) by Llama-3-70B are plotted against the perplexity of the corresponding example question (X-axis) and grouped by (a) modality, (b) argument forms, and (c) logic interpretation content. Each group consists of 20 randomly selected examples with other factors controlled.

certain argument forms? Do more complex components of logical forms, such as modalities, matter for LLM performance?

In this work, we investigate the logical reasoning capabilities of LLMs by assessing their performance on different logical forms. We curate a dataset of natural language statements and questions based on several logical forms in both proposi-

tional and modal logic, which is designed to mirror reasoning in daily communication. An example is shown in §3.3. We then conduct a series of controlled experiments to analyze the performance of a set of LLMs on the dataset. Although our findings generally align with those by Gonen et al. (2023) and McCoy et al. (2024), who suggest that LLMs excel on examples with high probability, our results indicate that logical form, including but not limited to modalities and argument forms, is a crucial complementary factor in predicting the performance of LLMs (Figure 1). Additionally, with meaningful real-world interpretations, we find that:

1. LLMs are still far from being perfect in atomic-level propositional and modal logic reasoning.
2. LLMs prefer an affirmative answer under the modality of possibility, whereas they prefer a negative answer under the modality of necessity.
3. In line with the recent results on categorical syllogisms (Eisape et al., 2024), we verify on hypothetical and disjunctive syllogisms that LLMs achieve better performance on certain logical forms that humans perform well. However, some logical forms receive favor from LLMs, while the phenomena lack support from human intuition or human behavioral data.

This paper is structured as follows. After reviewing related work (§2), we describe the dataset synthesis process (§3). We report the LLM reasoning results on our data (§4), compare them with human performance (§5). We conclude by discussing the implications of our results and the limitations (§6).

2 Related Work

Logical reasoning benchmarks. Prior LLM logical reasoning benchmarks (Liu et al., 2020; Han et al., 2022, *inter alia*) focus on complex, multi-hop reasoning problems with manually annotated problems, making cross-problem comparisons challenging. Recent work has introduced benchmarks with synthesized natural-language questions using predefined logical formulas and substitution rules (Saparov and He, 2022; Saparov et al., 2023; Parmar et al., 2024; Wan et al., 2024, *inter alia*). Compared to them, our work uniquely incorporates modal logic, which has been largely unexplored in existing benchmarks—while Holliday and Mandelkern (2024) present a case study, our approach offers two key advances: controlled knowledge bias in logic interpretations (§3.3) and a more rigorous statistical evaluation framework (§4.1).

Propositional and modal logic reasoning in language models. Recent work has explored training and finetuning language models specifically for logical reasoning (Clark et al., 2021; Hahn et al., 2021; Tafjord et al., 2022). Our work differs in two key aspects: (1) we evaluate general-purpose language models through prompting, a cost-efficient setup that has been widely adopted in recent years, and we focus on propositional and alethic modal logic rather than temporal (Hahn et al., 2021) or epistemic (Sileo and Lernould, 2023) logic;¹ (2) unlike studies comparing LLM and human performance on categorical syllogisms (Eisape et al., 2024, *inter alia*),² we focus on hypothetical and disjunctive syllogisms with considerations of modality.

Human logic reasoning. Work on human reasoning capabilities has informed studies of LLM logical reasoning: Eisape et al. (2024) compared LLM syllogistic reasoning with human behavior results (Ragni et al., 2019) under the framework of the Mental Models Theory (Johnson-Laird, 1983); Lampinen et al. (2024) found similar content effects in human and LLM reasoning, supporting the need to control for common-sense knowledge in benchmarks (§3.2); Belem et al. (2024) studied human and LLM perception of uncertainty at a lexical level. Compared to them, we focus on the propositional and modal logic reasoning process and contribute new behavioral data.

3 Dataset

We curate a dataset of natural-language multiple-choice questions to measure the logical inference performance of LLMs. Starting from propositional and modal logical forms as templates (§3.1), we assign meanings (e.g., real-world interpretations) to each variable and translate templates into natural-language Yes/No questions (§3.2). A subsidiary visualization of the process is shown in Figure A1.

3.1 Background: Propositional and Modal Logic

Propositional logic studies the relation between propositions. In this framework, each proposition is typically represented by a variable, and multi-

¹Technically, any logic that involves non-truth-functional operators, including first-order logic, temporal logic, and epistemic logic, can be viewed as a modal logic; however, we adopt the most restrictive sense of *modal logic* (Ballarín, 2023) and use it interchangeably with *alethic modal logic*.

²We refer readers to Zong and Lin (2024) for a more comprehensive review of categorical syllogisms.

ple propositions combine with logical connectives (e.g., $\vee$ and $\rightarrow$) to form compound propositions.

In propositional logic, a proposition can be evaluated as either true or false; however, this system can be overly simplistic when dealing with the complexity of real-world events. Consider the statement *Alice is not eating*, while it is true in a world where Alice is not eating, it may become false in a hypothetical *possible world* where Alice is indeed eating. This idea, known as possible world semantics (Kripke, 1959), provides a framework for more nuanced statements about event possibilities, such as *Alice may be eating* and *Alice must be eating*. The former statement can be understood as there *exists* a possible world where Alice is eating, and the latter can be understood as in *all* possible worlds, Alice is eating.³ Normal modal logic (Kripke, 1963) formalizes this idea and extends propositional logic to reason about event necessity and possibility. In the Backus–Naur form, a normal modal logic system $\mathcal{L}$ can be written as

$$\mathcal{L} : \varphi := p \mid \neg\varphi \mid \Box\varphi \mid \Diamond\varphi \mid \varphi \vee \psi \mid \varphi \wedge \psi \mid \varphi \rightarrow \psi, \quad (1)$$

where p is a propositional variable that serves as an atom in $\mathcal{L}$, $\neg$ is the negation operator, $\Box$ is the necessity operator (*must*), $\Diamond$ is the possibility operator (*may*), $\vee$ is logical disjunction (*or*), $\wedge$ is logical conjunction (*and*), and $\rightarrow$ is the logical implication operator (*if...then*). φ denotes the syntactic category of a formula in $\mathcal{L}$. The right-hand side of Eq. (1) describes all possible logical formulas under the system $\mathcal{L}$: for example, if $\varphi \in \mathcal{L}$, the rules imply that $\neg\varphi \in \mathcal{L}$, $\Box\varphi \in \mathcal{L}$, and so on. Following the convention in logic, the operator precedence is $\{\neg, \Box, \Diamond\} \succ \{\vee, \wedge\} \succ \{\rightarrow\}$.

Indeed, the operators $(\neg, \Box, \rightarrow)$ forms a functional complete set of operators under $\mathcal{L}$. Suppose φ and ψ are variables that represent logical formulas. The logical or ($\vee$) and logical and ($\wedge$) operators can be rewritten with logical not ($\neg$) and logical implication ($\rightarrow$), as follows:

$$\begin{aligned} \varphi \vee \psi &\Leftrightarrow \neg\varphi \rightarrow \psi, \\ \varphi \wedge \psi &\Leftrightarrow \neg(\varphi \rightarrow \neg\psi). \end{aligned} \quad (2)$$

Possibility operator $\Diamond$ can also be derived from the necessity operator.

$$\Diamond\varphi \Leftrightarrow \neg\Box\neg\varphi \quad (3)$$

³The possible world semantics, therefore, connects the notion of *necessity* and *possibility* to the universal and existential quantification ($\forall, \exists$) under first-order logic.

Deduction and sequent. Given a formula set Γ as premises, if a deduction to a conclusion φ exists using axiom schemata and inference rules under the normal modal logic, we say the premises *infer* the conclusion, and the deduction can be represented as a logic *sequent* $\Gamma \vdash \varphi$. If a formula set Γ do not infer the conclusion, we denote it as $\Gamma \not\vdash \varphi$ and call it a *non-entailment*.

3.2 Translating Logic to Natural Language

An *interpretation* maps propositional variables to concrete meanings. For example, under the interpretation that p is “*Jane is eating apples*” and q is “*John is eating oranges*”, the logical formula $p \vee q$ becomes “*Jane is eating apples or John is eating oranges*.”

Choices of interpretation, i.e., the concrete content of the sentence, should not affect the underlying logical reasoning process. However, in natural-language utterances, reasoning can be influenced by various confounding factors. Knowledge bias is a common pitfall. For example, given the logical form $\{\Box p \rightarrow \Box\neg q, \Box p\} \vdash \Box\neg q$, regardless of p ’s interpretation, if we interpret $\neg q :=$ “*Cats are not animals*” then the conclusion will be “*It is certain that cats are not animals*.” But common-sense knowledge suggests that “*it is certain that cats are animals*” ($\Box q$), which logically contradicts the existing premise set.⁴ Such bias will complicate logical reasoning (Lampinen et al., 2024) and should be avoided in data curation. Besides, each variable should have independent interpretation, as detailed in Appendix B.2.

After being assigned interpretations, each logical form is further articulated as a yes-no question on whether the conclusion can be inferred from the premises. To mitigate the ambiguity in natural language, we design heuristic rules to translate logic forms into less ambiguous English, which are detailed in Appendix B.2. For the exact wordings we used, see Table A1 in Appendices. If a valid deduction exists ($\vdash$) for the logical form, the ground truth answer is Yes, otherwise No. The answer is solely determined by the logical form and is independent of the interpretation.

3.3 Involved Logical Forms

Translated logical forms can have varying degrees of naturalness. For example, the *necessitation rule* $\{\varphi\} \vdash \Box\varphi$, which translates to “ φ is true; therefore,

⁴This confounding factor affects the examples in Table 18 of Han et al. (2022).

it is certain that φ is true,” appears to be unnatural due to redundancy.⁵ Based on the relationship between $\vee$ and $\rightarrow$ in Eq. (2), we use hypothetical and disjunctive syllogisms with four basic variants:

$$\begin{aligned} \{\varphi \vee \psi, \neg\varphi\} &\vdash \psi, & (\vee^L) \\ \{\neg\varphi \rightarrow \psi, \neg\varphi\} &\vdash \psi, & (\rightarrow^L; \text{modus ponens}) \\ \{\varphi \vee \psi, \neg\psi\} &\vdash \varphi, & (\vee^R) \\ \{\neg\varphi \rightarrow \psi, \neg\psi\} &\vdash \varphi. & (\rightarrow^R; \text{modus tollens}) \end{aligned}$$

Despite the semantic similarity, these logical forms translate to different natural-language questions. For example, taking the interpretations of $\varphi := \text{Jane is watching a show}$ and $\psi := \text{John is reading a book}$, $\vee^L$ translates to

Consider the following statements:

Jane is watching a show or John is reading a book.

Jane isn’t watching a show.

Question: Based on these statements, can we infer that *John is reading a book*?

With the same interpretation, $\rightarrow^L$ ’s translation of the first statement is *If Jane isn’t watching a show, then John is reading a book.*

According to the commutativity of disjunction operator, we group $\vee^L$ and $\vee^R$ together as *disjunctive syllogism*, alongside two hypothetical syllogism groups, modus ponens ($\rightarrow^L$) and modus tollens ($\rightarrow^R$). All the logical forms shown above are valid sequents with ground-truth answer Yes. To balance the dataset, we introduce some logic fallacies that generate questions with ground-truth label No. By flipping the second premises and the conclusions, we obtain the following fallacies:

$$\begin{aligned} \{\varphi \vee \psi, \psi\} &\not\vdash \neg\varphi, & (\vee_{\not\vdash}^L) \\ \{\neg\varphi \rightarrow \psi, \psi\} &\not\vdash \neg\varphi, & (\rightarrow_{\not\vdash}^L) \\ \{\varphi \vee \psi, \varphi\} &\not\vdash \neg\psi, & (\vee_{\not\vdash}^R) \\ \{\neg\varphi \rightarrow \psi, \varphi\} &\not\vdash \neg\psi, & (\rightarrow_{\not\vdash}^R) \end{aligned}$$

where $\vee_{\not\vdash}^L$ and $\vee_{\not\vdash}^R$ are grouped as *affirming the disjunction*, $\rightarrow_{\not\vdash}^L$ and $\rightarrow_{\not\vdash}^R$ corresponds to *affirming the consequent* and *denying the antecedent*, respectively. In our dataset, we require the formulas φ and ψ to the form of $\mathcal{M}p$ and $\mathcal{M}q$, where p and q are propositional variables, each assigned with an interpretation. Both variables are constrained under the same modality $\mathcal{M}$, which can be necessity ($\square$),

⁵Nevertheless, we report the experiment results on necessity rule in Appendix C.1.

possibility ($\Diamond$) or no modality ($\emptyset$). Pairing with four rules and theorem–fallacy variations, we have a total of $3 \times 4 \times 2 = 24$ forms.

3.4 Involved Logic Interpretations

For logic interpretations, we generate a set of verb phrases by prompting the CodeLlama 2 model (Rozière et al., 2024), and select 204 of them manually. and combine them with top-200 popular baby names in the US into subject-verb-object pairs,⁶ such as (*Ray, make, a pizza*). We randomly generate 1000 interpretations with two pairs each. The same set of interpretations is applied to variables p, q in each logic sequent’s natural language template. In total, there are $24 \times 1000 = 24000$ question, with samples shown in Table A1.

4 Experiment

4.1 Metrics and Investigated Models

Hu and Levy (2023) have suggested that the standard approach of greedily decoding yes-no strings (Dentella et al., 2023) may underestimate the competence of a language model; therefore, we adopt a probability-based metric to evaluate the model performance. In our evaluation protocol, the predicted likelihood of the tokens Yes and No, conditioned on the prompt s —denoted as $p(\text{Yes} \mid s)$ and $p(\text{No} \mid s)$, respectively—serve as the soft labels for yes-no answers. The soft accuracy $\hat{p}$ on the single example with ground-truth answer $y \in \{\text{Yes}, \text{No}\}$ is defined as the relative probability of y :

$$\hat{p} = \frac{p(\text{No} \mid s) \mathbb{1}[y = \text{No}] + p(\text{Yes} \mid s) \mathbb{1}[y = \text{Yes}]}{p(\text{No} \mid s) + p(\text{Yes} \mid s)},$$

where $\mathbb{1}[\cdot]$ is the indicator function that returns 1 if the condition is true and 0 otherwise. This relative probability can also be viewed as the confidence score of the model on the ground-truth answer. The soft accuracy Acc_{soft} of a model on the entire dataset $\mathcal{D}$ is defined as the average soft accuracy over all examples,

$$Acc_{\text{soft}} = \frac{1}{|\mathcal{D}|} \sum_{i=1}^{|\mathcal{D}|} \hat{p}_i.$$

We use a zero-shot setting to investigate the general performance of the models’ logical inference capabilities—while adding detailed instructions or few-shot demonstrations may increase the absolute performance, they are at the cost of introducing

⁶<https://www.ssa.gov/oact/babynames/names.zip>

Model	Overall		Leaderboard		Modality			Argument Form					
	(Rank)		(Rank)		$\emptyset$	$\square$	$\diamond$	$\vee_{\vdash}^{L,R}$	$\rightarrow_{\vdash}^L$	$\rightarrow_{\vdash}^R$	$\vee_{\nvdash}^{L,R}$	$\rightarrow_{\nvdash}^L$	$\rightarrow_{\nvdash}^R$
mistral-7b	0.645	(4)	0.145	(7)	0.464	0.496	0.974	0.877	0.663	0.280	0.434	0.653	0.939
mistral-8x7b	0.724	(1)	0.193	(5)	0.698	0.601	0.874	0.963	0.873	0.023	0.757	0.648	0.813
llama-2-7b	0.335	(10)	0.094	(10)	0.262	0.207	0.538	0.444	0.147	0.315	0.208	0.451	0.468
llama-2-13b	0.513	(9)	0.110	(9)	0.488	0.362	0.688	0.418	0.581	0.393	0.631	0.436	0.591
llama-2-70b	0.611	(5)	0.127	(8)	0.616	0.471	0.746	0.446	0.845	0.518	0.775	0.389	0.694
llama-3-8b	0.565	(6)	0.239	(3)	0.598	0.460	0.639	0.526	0.470	0.332	0.664	0.625	0.716
llama-3-70b	0.714	(2)	0.362	(1)	0.745	0.554	0.843	0.606	0.773	0.515	0.882	0.661	0.788
yi-34b	0.518	(8)	0.226	(4)	0.457	0.413	0.683	0.346	0.498	0.205	0.685	0.638	0.737
phi-2	0.532	(7)	0.155	(6)	0.469	0.456	0.673	0.670	0.757	0.522	0.365	0.402	0.510
phi-3-mini	0.690	(3)	0.272	(2)	0.657	0.536	0.877	0.839	0.974	0.475	0.664	0.462	0.604
OpenAI-o1	0.926	N/A	N/A	N/A	1.000	0.773	1.000	0.895	1.000	0.775	0.919	1.000	1.000
Gemini-1.5-Pro	0.859	N/A	N/A	N/A	0.831	0.748	0.997	1.000	1.000	0.919	0.661	0.991	0.638
human	0.595	N/A	N/A	N/A	0.589	0.566	0.640	0.691	0.901	0.628	0.594	0.225	0.411

Table 1: Overall and break-down accuracies of different models, as well as their HuggingFace OpenLLM Leaderboard performance and relative ranking (Fourrier et al., 2024). Each argument form category denotes the union of the fine-grained categories specified in the superscripts and subscripts—for example, $\vee_{\vdash}^{L,R}$ denotes the entire disjunctive syllogism group. **Boldfaced** values indicate the row-wise maximum for each factor. Note that due to technical limitations of commercial LLMs, results from OpenAI-o1 (OpenAI, 2024) and Gemini-1.5-pro (Team et al., 2024) are greedy-decoding based evaluation on 2,000 random samples that serve as references, and are therefore not directly comparable to other probability-based evaluations. Human results are detailed in §5.

possibly undesired confounding factors or behaviors, such as simply copy-pasting the answers in the examples.

We evaluate on the following models with open-sourced weights: mistral-7b-v0.2 and -8x7b (Jiang et al., 2023, 2024); llama-2-7b, -13b and -70b (Touvron et al., 2023); 3.1 version of llama-3-8b and -70b (AI@Meta, 2024); yi-34b (01.AI, 2024); phi-2 and phi-3-mini (Microsoft, 2023, 2024).⁷

4.2 Results: Performance w.r.t. Logical Forms

We evaluate the aforementioned models with the probability-based protocol (Table 1). Generally, models that rank higher in the leaderboard also achieve higher soft accuracy on our dataset. The break-down accuracies on modalities and argument forms reveal that:

- (Modality) All models consistently perform better on the possibility ($\diamond$) than necessity ($\square$) or plain propositional logic.
- (Argument Forms) The pattern is more diverse, yet most of the models struggle the most on modus tollens ($\rightarrow_{\vdash}^R$) within logic sequents (i.e., questions with ground-truth answers Yes), and affirming the consequent ($\rightarrow_{\nvdash}^L$) within fallacies.

⁷Our evaluation protocol technically requires the conditional probabilities of specified answers given a prompt, which are not supported by most commercial models; however, we report the greedy-decoding accuracy of these models on a sample subset for reference.

4.2.1 Analysis on Logic Sequents

To systematically analyze the effect on model performance of each factor of interest, as well as cross-validating the observations above, we fit a linear mixed-effects model (Raudenbush, 2002) to the soft accuracy data on valid logic sequents (i.e. with ground truth of Yes) across different LLMs and logical forms,

$$Acc_{soft} \sim Modality + ArgForm + Perplexity + (1 + Perplexity | LLM), \quad (4)$$

with the linear fixed effects of (i.) modality, (ii.) argument form, and (iii.) input perplexity. Individual probability, coupled with a constant term, is modeled as a random effect to account for potential model-specific biases. Here, *Perplexity* denotes the perplexity of the input text ($x_1 x_2 \dots x_N$), which is defined as the exponential of the token-wise average negative log-likelihood of the text given a specific language model:

$$Perplexity = \exp \left(-\frac{1}{N} \sum_{i=1}^N \log p(x_i | x_{<i}) \right)$$

The mixed-effects model yields a marginal R^2 of 0.342 and a conditional R^2 of 0.543, suggesting a reasonable predictive power. The likelihood ratio test on the full regression model vs. the null regression model without each of the fixed effects

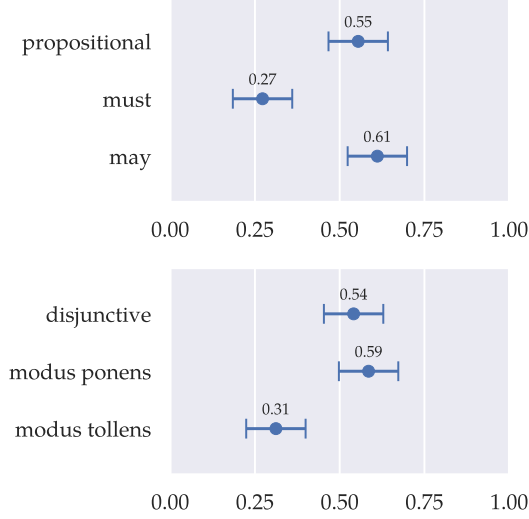


Figure 2: Estimated marginal means of logical form factors in the mixed-effects model of Eq. (4), along with their 95% confidence intervals.

Hypothesis	<i>p</i> -value
propositional < may	< 0.001
must < propositional	< 0.001
must < may	< 0.001
disjunctive < modus ponens	< 0.001
modus tollens < modus ponens	< 0.001
modus tollens < disjunctive	< 0.001

Table 2: Hypothesis testing results on the effect of logical form factors on soft accuracy (Figure 2).

yields a significant result ($p < 0.001$), suggesting the importance of all these factors in determining the model performance.

Fixed effects. In line with Gonen et al. (2023) and McCoy et al. (2024), we find a negative correlation between perplexity and soft accuracy ($p < 0.001$); however, the correlation is weak ($\rho = -0.09$), which suggests the necessity of the complementary factors below in predicting LLM performance.

For different modalities and argument forms, we estimate their marginal means on soft accuracy (Figure 2), and perform pairwise hypothesis testing on the estimated coefficients (Table 2). The results generally align with the general observations on the full dataset. The only exception is that modus ponens ($\rightarrow^L$), instead of disjunctive syllogism ($\vee$), appears to be the easiest argument form (i.e., the one with the highest soft accuracy) among all.

Random effects. We analyze the per-LLM random effects on the soft accuracy (Figure 3). All the model-specific mixed effects of perplexity are negative, suggesting the negative correlation be-

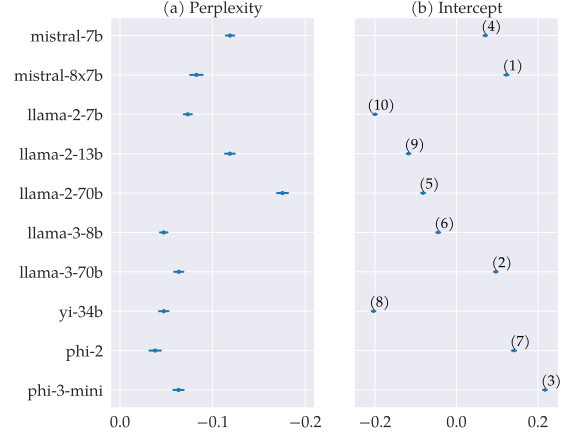


Figure 3: Illustration of per-model random effects on soft accuracy in the mixed-effects model of Eq. (4) with 99.9% confidence intervals. (a) Mixed effects (i.e., the sum of fixed and random effects) of perplexity. (b) Intercept random effects (i.e., constant term per model on soft accuracy), with the model performance rank (Table 1) annotated in parentheses.

tween perplexity and soft accuracy is consistent across models (Figure 3a). While the intercept random effects are not perfectly aligned with the model performance—since the perplexity random effects may introduce confounding factors—higher-ranked models generally tend to have higher intercept random effects (Figure 3b), which cross-validates the general performance ranking.

4.2.2 Extended Analysis on the Negative Perplexity–Performance Correlation

We further investigate the negative correlation between perplexity and model performance through a controlled experiment: we create a mirror dataset of the same size, keeping all the logical formulas while interpreting them with nonsensical words. For example, the formula $\Diamond(\varphi \vee \psi)$ may be interpreted as *it’s possible that Neva is balaring a monterey or Lucille is sweeling prandates*, where the underlined words and phrases are nonsensical. Intuitively, the perplexity of the problems in this mirror dataset should be much higher than that of the primary dataset problems (§3) under any reasonably trained language model.

We analyze the correlation between perplexity and model performance (Figure 4). As desired, the perplexity of problems with nonsensical words are indeed much higher than that of the primary dataset (≈ 30 vs. ≈ 10). The significant portion of horizontal and inclined lines in the figures again suggests that perplexity is not a reliable predictor of model performance. Meanwhile, the overall

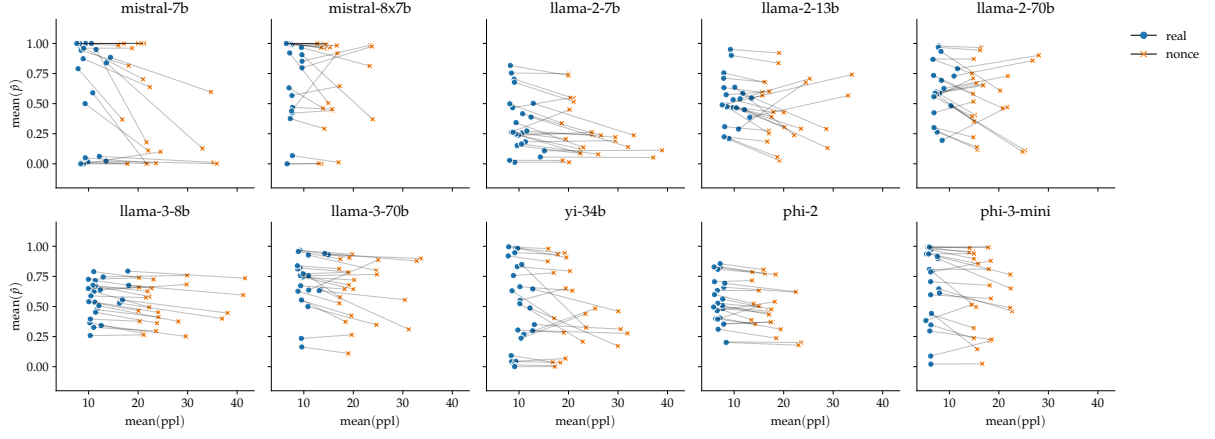


Figure 4: Correlation between mean perplexity and mean confidence score on each logic sequent. Each point represents an average over a group of 1000 prompts that share the same underlying logic sequent. Two connected dots share the same logic formula.

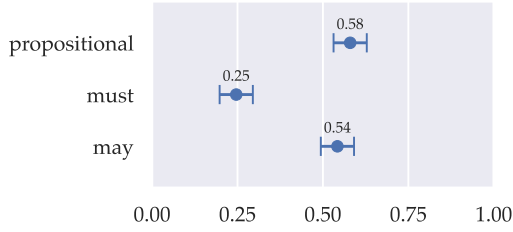


Figure 5: Estimated marginal means of the factors in the mixed-effects model of Eq. (5) with 95% confidence intervals. Higher coefficients indicate a higher tendency to affirm the claim.

parallelism of the lines echos our results that logical forms are important factors for such prediction.

4.2.3 The Affirmation Bias over Modalities

One key argument of Dentella et al. (2023) is that large language models exhibit a bias towards affirming the claim, i.e., answering Yes more frequently than No. We investigate this phenomenon by fitting a mixed-effects model

$$\frac{P(\text{Yes} | s)}{P(\text{Yes} | s) + P(\text{No} | s)} \sim \text{Modality} + \text{ArgForm} + \text{Perplexity} + (1 + \text{Perplexity} | \text{LLM}), \quad (5)$$

which has the same structure as Eq. (4), except the dependent variable being the relative probability of answering Yes conditioned on input text s .

We present the estimated marginal means of the factors in the mixed-effects model (Figure 5). While our results confirm the affirmation bias on propositional logic, such bias is slightly less pronounced on the possibility modality ($\diamond$, around 0.03), and the models even show a bias towards rejecting claims under the necessity modality ($\square$).

5 Human Experiments

LLMs are trained on text produced by humans and are able to generate plausible text; therefore, there have been interests in using LLMs as human models (Eisape et al., 2024; Misra and Kim, 2024, *inter alia*). Following this line of work, we conduct a human behavioral experiment to ground the LLM reasoning behavior. Using samples from our primary dataset, we collected 710 responses from adults fluent in English through Prolific.⁸ More experiment details can be found in Appendix A.2.

The average human accuracy on each group is shown in the last row of Table 1.⁹ Aligned with our LLM results (§4), on modalities the overall human results also show an accuracy order of ($\diamond > \emptyset > \square$), and on argument forms, modus ponens ($\rightarrow^L$) is the most accurately answered pattern.

To further investigate the interactions of logic factors, we fit a generalized linear mixed-effects model (Bates et al., 2015) to verify the effect of modality and argument forms on human logic reasoning accuracy (Eq. (6) and Figure 6).

$$\text{logit}(\text{Acc}) \sim \text{Modality} + \text{ArgForm} + \text{Rt} + (1 + \text{Rt} | \text{ParticipantID}), \quad (6)$$

where Acc is the binary accuracy of human responses, and Rt is the response time. The generalized mixed-effects model yields a marginal R^2 of 0.121 yet a 0.419 conditional R^2 , indicating a diverse response pattern across participants. The likelihood ratio test on the full model against the null model shows that only the effect of argument form

⁸<https://prolific.com>

⁹Human responses are binary classes, so correct and incorrect responses are coded as 1 and 0, respectively.

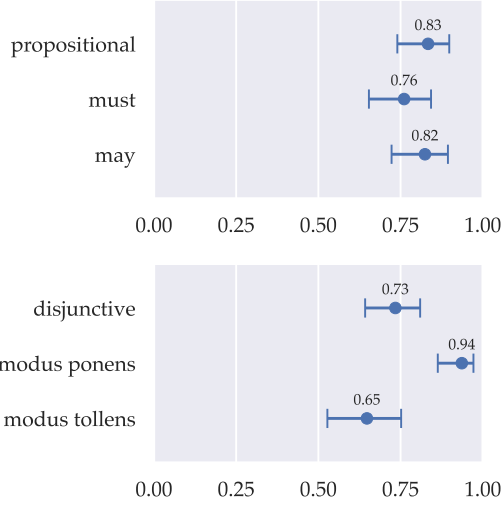


Figure 6: Estimated marginal means of logical form factors in the generalized mixed-effects model of Eq. (6), along with their 95% confidence intervals.

is significant ($\chi^2(2) = 25.6, p < 0.001$). However, in accordance with the overall performance, we find modus ponens ($\rightarrow^L$) has a significantly higher effect than other two valid argument forms. This confirms that logical forms can also have a significant impact on human reasoning accuracy, which is consistent with the LLM results, although the effect sizes are not the same.

6 Conclusion and Discussion

We present an analysis of hypothetical and disjunctive syllogisms on propositional and modal logic and systematically analyze the LLM performance on the dataset. Our analysis provides novel insights on explaining and predicting LLM performance: in addition to the perplexity or probability of the input text, the underlying logic forms play an important role in determining the performance of LLMs. In addition, we compare the behaviors of LLMs and humans using the same data through human behavioral experiments. We discuss the implications of our results as follows.

Probability in language models. Probability and perplexity are often used as intrinsic evaluation metrics for language models. While Gonen et al. (2023) and McCoy et al. (2024) show that probability and perplexity correlate well with LLM performance, literature in program synthesis with LLMs shows little correlation between probability and execution-based evaluation results (Li et al., 2022; Shi et al., 2022). This work does not necessarily contradict either line but rather provides complementary factors for analyzing LLM performance.

We argue that probability may have become an overloaded term in analyzing LLMs. Low probability may be due to one or more of the following non-exhaustive reasons: (1) out-of-context content, (2) ungrammatical language, or (3) grammatical but semantically awkward content (cf. the mirror dataset in §4.2.2), (4) reasonable but rare content. We hypothesize that the probability of language models may not be essentially able to capture all these nuanced differences, and call for encoding and decoding algorithms—such as Meister et al. (2023)—that can better decompose the probability into finer-grained and explainable components.

Comparing humans and LLMs. What is our goal for building LLMs? To achieve better performance on practical tasks or to build a more human-like model? Our results, together with Eisape et al. (2024), suggest that these two goals may not be perfectly aligned by revealing a mixture of similarity and discrepancy between LLMs and humans—for example, while LLMs exhibit higher benchmark performance than humans on our dataset and show the same argument form preferences with humans (Figures 2 and 6), they also show systematic biases that we do not find significant in human reasoning (e.g., disfavoring the necessity modality, §4.2.3). While there has been positive evidence of using LLMs as human models in psycholinguistic studies (Misra and Kim, 2024, *inter alia*), our results suggest executing such approaches cautiously.

On the relation between modality and performance. Our results show that there is a significant difference in performance between necessity and possibility modalities, with the former much lower than the latter (Table 1). Part of the reason for this is that LLMs have a significant tendency to say “No” to necessity modality (Figure 5).

On the one hand, our results extend the conclusion of Dentella et al. (2023) that LLMs generally respond positively—LLM behaviors may be significantly affected by finer-grained factors, including but not necessarily limited to the modality involved in the input. On the other hand, while LLMs systematically tend to answer “No” to questions in necessity modality, we do not find related evidence in human experiments, which leads us to hypothesize that such rejection bias comes from either the model architecture or the training strategies, such as the reinforcement learning with human feedback (RLHF; Ouyang et al., 2022) protocol. We leave this as an open question for future research.

Limitations

This work comes with two major limitations:

1. While we have verified that our data has a low perplexity (9.82 ± 2.47 under mistral-7b; much lower than that of the data by Wan et al. (2024), 25.44), and, therefore, are similar enough to natural language utterances, the synthetic language cannot fully substitute natural language in daily life. Our dataset and analysis are not comprehensive enough to cover many nuanced examples that may appear in real communication, especially when context-dependent understanding is crucial to conveying communication goals.
2. Despite more than 7,000 languages worldwide, as a first step, our material only covers English. This narrow focus is due to the languages the authors are proficient in and the coverage of the language models. We acknowledge the importance of extending the scope of this work to a more comprehensive set of languages and leave the extension as an immediate follow-up step.

In addition, the sample size of human experiments is somewhat limited. We leave more comprehensive human behavioral data collection and analysis to future work.

Ethics Statement

While this work involves human logical reasoning experiments, we have ensured that (1) the data are generated procedurally following templates listed in the paper and (2) there is no harmful content in the atomic logical interpretations, reviewed by all the authors. In addition, we have ensured that all participants are paid a fair wage through the Prolific platform. Instructions and consent forms delivered to the participants can be found in the Appendix A.2. The institutional ethics review board has approved the data collection process.

This work contributes to the understanding of LLMs. We do not foresee risk beyond the minimal risk posed by LLM evaluation work. We acknowledge that using LLMs in real-world scenarios could significantly impact human behaviors, raising the need for model transparency, safety, security, and interpretability. We will open-source the synthetic logical reasoning dataset upon publication.

References

01.AI. 2024. Yi: Open Foundation Models by 01.AI.

AI@Meta. 2024. The Llama 3 Herd of Models.

Roberta Ballarín. 2023. Modern origins of modal logic. In Edward N. Zalta and Uri Nodelman, editors, *The Stanford Encyclopedia of Philosophy*, Fall 2023 edition. Metaphysics Research Lab, Stanford University.

Douglas Bates, Martin Mächler, Ben Bolker, and Steve Walker. 2015. *Fitting Linear Mixed-Effects Models Using lme4*. *Journal of Statistical Software*, 67(1).

Catarina G. Belem, Markelle Kelly, Mark Steyvers, Sameer Singh, and Padhraic Smyth. 2024. *Perceptions of Linguistic Uncertainty by Language Models and Humans*.

Peter Clark, Øyvind Taffjord, and Kyle Richardson. 2021. Transformers as soft reasoners over language. In *IJCAI*.

Vittoria Dentella, Fritz Günther, and Evelina Leivada. 2023. Systematic testing of three language models reveals low language accuracy, absence of response stability, and a yes-response bias. *Proceedings of the National Academy of Sciences*, 120(51):e2309583120.

Tiwalayo Eisape, Michael Tessler, Ishita Dasgupta, Fei Sha, Sjoerd Steenkiste, and Tal Linzen. 2024. A systematic comparison of syllogistic reasoning in humans and language models. In *NAACL*.

Clémentine Fourrier, Nathan Habib, Alina Lozovskaya, Konrad Szafer, and Thomas Wolf. 2024. Open llm leaderboard v2. https://huggingface.co/spaces/open-llm-leaderboard/open_llm_leaderboard.

Hila Gonen, Srini Iyer, Terra Blevins, Noah Smith, and Luke Zettlemoyer. 2023. Demystifying prompts in language models via perplexity estimation. In *Findings of ACL: EMNLP*.

Christopher Hahn, Frederik Schmitt, Jens U Kreber, Markus Norman Rabe, and Bernd Finkbeiner. 2021. Teaching temporal logics to neural networks. In *ICLR*.

Simeng Han, Hailey Schoelkopf, Yilun Zhao, Zhenting Qi, Martin Riddell, Luke Benson, Lucy Sun, Ekaterina Zubova, Yujie Qiao, Matthew Burtell, David Peng, Jonathan Fan, Yixin Liu, Brian Wong, Malcolm Sailor, Ansong Ni, Linyong Nan, Jungo Kasai, Tao Yu, Rui Zhang, Shafiq Joty, Alexander R. Fabbri, Wojciech Kryscinski, Xi Victoria Lin, Caiming Xiong, and Dragomir Radev. 2022. *FOLIO: Natural Language Reasoning with First-Order Logic*.

Wesley H. Holliday and Matthew Mandelkern. 2024. *Conditional and Modal Reasoning in Large Language Models*. ArXiv:2401.17169.

Jennifer Hu and Roger Levy. 2023. *Prompting is not a substitute for probability measurements in large language models*. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language*

642	<i>Processing</i> , pages 5040–5060, Singapore. Association	R. Thomas McCoy, Shunyu Yao, Dan Friedman,	697
643	for Computational Linguistics.	Mathew D. Hardy, and Thomas L. Griffiths. 2024.	698
644	Wenlong Huang, Pieter Abbeel, Deepak Pathak, and	Embers of autoregression show how large language	699
645	Igor Mordatch. 2022. Language models as zero-shot	models are shaped by the problem they are trained	700
646	planners: Extracting actionable knowledge for em-	to solve . <i>Proceedings of the National Academy of</i>	701
647	odied agents. In <i>ICML</i> , pages 9118–9147. PMLR.	<i>Sciences</i> , 121(41):e2322420121.	702
648	Albert Q. Jiang, Alexandre Sablayrolles, Arthur Men-	Clara Meister, Tiago Pimentel, Gian Wiher, and Ryan	703
649	sch, Chris Bamford, Devendra Singh Chaplot, Diego	Cotterell. 2023. Locally typical sampling. <i>Transac-</i>	704
650	de las Casas, Florian Bressand, Gianna Lengyel, Guil-	<i>tions of the Association for Computational Linguis-</i>	705
651	laume Lample, Lucile Saulnier, L��lio Renard Lavaud,	<i>tics</i> , 11:102–121.	706
652	Marie-Anne Lachaux, Pierre Stock, Teven Le Scao,	Microsoft. 2023. Phi-2: The surprising power of small	707
653	Thibaut Lavril, Thomas Wang, Timoth��e Lacroix,	language models . <i>Microsoft Research Blog</i> .	708
654	and William El Sayed. 2023. Mistral 7B .	Microsoft. 2024. Phi-3 Technical Report: A Highly	709
655	Albert Q. Jiang, Alexandre Sablayrolles, Antoine	Capable Language Model Locally on Your Phone .	710
656	Roux, Arthur Mensch, Blanche Savary, Chris	Kanishka Misra and Najoung Kim. 2024. Generat-	711
657	Bamford, Devendra Singh Chaplot, Diego de las	ing novel experimental hypotheses from language	712
658	Casas, Emma Bou Hanna, Florian Bressand, Gi-	models: A case study on cross-dative generalization.	713
659	anna Lengyel, Guillaume Bour, Guillaume Lam-	<i>arXiv preprint arXiv:2408.05086</i> .	714
660	ple, L��lio Renard Lavaud, Lucile Saulnier, Marie-	Santiago Ontanon, Joshua Ainslie, Vaclav Cvicek, and	715
661	Anne Lachaux, Pierre Stock, Sandeep Subramanian,	Zachary Fisher. 2022. LogicInference: A new	716
662	Sophia Yang, Szymon Antoniak, Teven Le Scao,	Dataset for Teaching Logical Inference to seq2seq	717
663	Th��ophile Gervet, Thibaut Lavril, Thomas Wang,	Models. In <i>ICLR2022 Workshop on the Elements of</i>	718
664	Timoth��e Lacroix, and William El Sayed. 2024. Mix-	<i>Reasoning: Objects, Structure and Causality</i> .	719
665	tral of Experts .	OpenAI. 2024. Learning to reason with	720
666	Philip Nicholas Johnson-Laird. 1983. <i>Mental models:</i>	llms. https://openai.com/index/	721
667	<i>Towards a cognitive science of language, inference,</i>	learning-to-reason-with-llms/ .	722
668	<i>and consciousness</i> . 6. Harvard University Press.	Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida,	723
669	Henry A Kautz, Bart Selman, et al. 1992. Planning as	Carroll Wainwright, Pamela Mishkin, Chong Zhang,	724
670	satisfiability. In <i>ECAI</i> , volume 92, pages 359–363.	Sandhini Agarwal, Katarina Slama, Alex Ray, et al.	725
671	Citeseer.	2022. Training language models to follow instruc-	726
672	Saul A. Kripke. 1959. A Completeness Theorem	tions with human feedback. <i>Advances in neural in-</i>	727
673	in Modal Logic . <i>The Journal of Symbolic Logic</i> ,	<i>formation processing systems</i> , 35:27730–27744.	728
674	24(1):1–14.	Mihir Parmar, Nisarg Patel, Neeraj Varshney, Mutsumi	729
675	Saul A. Kripke. 1963. Semantical Analysis of Modal	Nakamura, Man Luo, Santosh Mashetty, Arindam	730
676	Logic I Normal Modal Propositional Calculi . <i>Mathe-</i>	Mitra, and Chitta Baral. 2024. LogicBench: Towards	731
677	<i>matical Logic Quarterly</i> , 9(5-6):67–96.	systematic evaluation of logical reasoning ability of	732
678	Andrew K Lampinen, Ishita Dasgupta, Stephanie C Y	large language models . In <i>Proceedings of the 62nd</i>	733
679	Chan, Hannah R Sheahan, Antonia Creswell, Dhar-	<i>Annual Meeting of the Association for Computational</i>	734
680	shan Kumaran, James L McClelland, and Felix	<i>Linguistics (Volume 1: Long Papers)</i> .	735
681	Hill. 2024. Language models, like humans, show	Marco Ragni, Hannah Dames, Daniel Brand, and Nico-	736
682	content effects on reasoning tasks . <i>PNAS Nexus</i> ,	las Riesterer. 2019. When Does a Reasoner Respond:	737
683	3(7):pgae233.	Nothing Follows?: 41st Annual Meeting of the Cog-	738
684	Yujia Li, David Choi, Junyoung Chung, Nate Kushman,	nitive Science Society. <i>Proceedings of the 41st An-</i>	739
685	Julian Schrittwieser, R��mi Leblond, Tom Eccles,	<i>nual Conference of the Cognitive Science Society</i> ,	740
686	James Keeling, Felix Gimeno, Agustin Dal Lago,	pages 2640–2645.	741
687	et al. 2022. Competition-level code generation with	Stephen W Raudenbush. 2002. Hierarchical linear mod-	742
688	alphacode. <i>Science</i> , 378(6624):1092–1097.	els: Applications and data analysis methods. <i>Ad-</i>	743
689	Jian Liu, Leyang Cui, Hanmeng Liu, Dandan Huang,	<i>vanced Quantitative Techniques in the Social Sci-</i>	744
690	Yile Wang, and Yue Zhang. 2020. LogiQA: A Chal-	<i>ences Series/SAGE</i> .	745
691	lenge Dataset for Machine Reading Comprehension	Baptiste Rozi��re, Jonas Gehring, Fabian Gloeckle, Sten	746
692	with Logical Reasoning . In <i>Proceedings of the</i>	Sootla, Itai Gat, Xiaoqing Ellen Tan, Yossi Adi,	747
693	<i>Twenty-Ninth International Joint Conference on Ar-</i>	Jingyu Liu, Romain Sauvestre, Tal Remez, J��r��my	748
694	<i>tificial Intelligence</i> , pages 3622–3628, Yokohama,	Rapin, Artyom Kozhevnikov, Ivan Evtimov, Joanna	749
695	Japan. International Joint Conferences on Artificial	Bitton, Manish Bhatt, Cristian Canton Ferrer, Aaron	750
696	Intelligence Organization.	Grattafiori, Wenhan Xiong, Alexandre D��fossez,	751

752	Jade Copet, Faisal Azhar, Hugo Touvron, Louis Mar-
753	tin, Nicolas Usunier, Thomas Scialom, and Gabriel
754	Synnaeve. 2024. Code Llama: Open Foundation
755	Models for Code .
756	Abulhair Saparov and He He. 2022. Language Mod-
757	els Are Greedy Reasoners: A Systematic Formal
758	Analysis of Chain-of-Thought. In <i>The Eleventh Inter-</i>
759	<i>national Conference on Learning Representations</i> .
760	Abulhair Saparov, Richard Yuanzhe Pang, Vishakh Pad-
761	makumar, Nitish Joshi, Mehran Kazemi, Najoung
762	Kim, and He He. 2023. Testing the General Deduc-
763	tive Reasoning Capacity of Large Language Models
764	Using OOD Examples. <i>Advances in Neural Informa-</i>
765	<i>tion Processing Systems</i> , 36:3083–3105.
766	Freda Shi, Daniel Fried, Marjan Ghazvininejad, Luke
767	Zettlemoyer, and Sida I Wang. 2022. Natural lan-
768	guage to code translation with execution. In <i>Proce-</i>
769	<i>edings of the 2022 Conference on Empirical Methods</i>
770	<i>in Natural Language Processing</i> .
771	Stuart M. Shieber. 1993. The problem of logical form
772	equivalence . <i>Computational Linguistics</i> , 19(1):179–
773	190.
774	Damien Sileo and Antoine Lerneuld. 2023.
775	MindGames: Targeting theory of mind in large lan-
776	guage models with dynamic epistemic modal logic .
777	In <i>Findings of the Association for Computational</i>
778	<i>Linguistics: EMNLP 2023</i> , pages 4570–4577.
779	Oyvind Tafjord, Bhavana Dalvi, and Peter Clark. 2022.
780	Entailer: Answering questions with faithful and truth-
781	ful chains of reasoning. In <i>EMNLP</i> .
782	Gemini Team, Petko Georgiev, Ving Ian Lei, Ryan
783	Burnell, Libin Bai, Anmol Gulati, Garrett Tanzer,
784	Damien Vincent, Zhufeng Pan, Shibo Wang, et al.
785	2024. Gemini 1.5: Unlocking multimodal under-
786	standing across millions of tokens of context. <i>arXiv</i>
787	<i>preprint arXiv:2403.05530</i> .
788	Hugo Touvron, Louis Martin, and Kevin Stone. 2023.
789	Llama 2: Open Foundation and Fine-Tuned Chat
790	Models.
791	Yuxuan Wan, Wenxuan Wang, Yiliu Yang, Youliang
792	Yuan, Jen-tse Huang, Pinjia He, Wenxiang Jiao, and
793	Michael R. Lyu. 2024. A & B == B & A: Trigger-
794	ing Logical Reasoning Failures in Large Language
795	Models .
796	Yimei Xiang. 2019. Two types of higher-order readings
797	of wh-questions. <i>Proceedings of the 22nd Amsterdam</i>
798	<i>Colloquium</i> .
799	Shi Zong and Jimmy Lin. 2024. Categorical syllogisms
800	revisited: A review of the logical reasoning abilities
801	of llms for analyzing categorical syllogism. <i>arXiv</i>
802	<i>preprint arXiv:2406.18762</i> .

A	Additional Experiment Details	803
A.1	LLM Experiment Details	804
	All LLMs used are obtained from Hugging Face	805
	checkpoints. Time and compute power require-	806
	ments vary, the largest llama-3-70b model takes	807
	around 2 hours on NVIDIA A6000 GPU to obtain	808
	all results in §4.	809
A.2	Human Experiment Details	810
	Participant consent. We use the following lan-	811
	guage to obtain consent from participants, where	812
	our institution name is replaced with <i>the Anony-</i>	813
	<i>mous Institution</i> to protect the anonymity of sub-	814
	mission.	815
	<i>This study is part of a scientific research project</i>	816
	<i>at the Anonymous Institution. Your decision to com-</i>	817
	<i>plete this study is voluntary. There is no way for us</i>	818
	<i>to identify you. The only information we will have,</i>	819
	<i>in addition to your responses, is the demographic</i>	820
	<i>information you provided to Prolific and the time at</i>	821
	<i>which you completed the survey. The results of the</i>	822
	<i>research may be presented at scientific meetings or</i>	823
	<i>published in scientific journals. Clicking on the but-</i>	824
	<i>ton below indicates that you are at least 18 years</i>	825
	<i>of age and agree to complete this study voluntarily.</i>	826
	<i>Press the button below to start the experiment.</i>	827
	Participant instructions. We use keys <i>F</i> and <i>J</i> ,	828
	which are roughly symmetric on a standard English	829
	keyboard, to collect participant responses. Half of	830
	the participants see the following instruction:	831
	<i>In this study, you will be presented with two</i>	832
	<i>statements followed by a question. Your task is to</i>	833
	<i>answer either Yes or No to the question, based on</i>	834
	<i>the information provided in the statements. Please</i>	835
	<i>respond quickly and accurately by pressing "F" for</i>	836
	<i>Yes, and "J" for No.</i>	837
	To mitigate the possible bias introduced by the	838
	dominant hand, we have the other half of the par-	839
	ticipants see instruction with reversed keys: <i>In this</i>	840
	<i>study, you will be presented with two statements fol-</i>	841
	<i>lowed by a question. Your task is to answer either</i>	842
	<i>Yes or No to the question, based on the information</i>	843
	<i>provided in the statements. Please respond quickly</i>	844
	<i>and accurately by pressing "F" for No, and "J" for</i>	845
	<i>Yes.</i>	846
	Participant wage. We offer participants an	847
	hourly wage of 1.5 times Prolific’s minimum wage.	848
	The duration is determined by the median comple-	849
	tion time among all participants.	850

B Extra Details of the Dataset

B.1 Data Synthesis Pipeline

To better explain the data synthesis process, we provide a detailed visualization of our pipeline in Figure A1.

B.2 Considerations in Translating Logical Form to Natural Language

During the interpretation process, another key point is to assign independent interpretations to variables. Deciding the dependency also involves common sense knowledge. For example, consider the premises $\neg p \rightarrow q$ and q . If we interpret $p :=$ “Jane is inside the house” and $q :=$ “Jane is out” to proposition variables p and q , the two variables are possibly not independent. According to common sense, “Jane is not inside the house” ($\neg p$) correlates with or is even equivalent to “Jane is out” (q). Logically, $\{\neg p \rightarrow q, q\} \not\models \neg p$; however, with the extra premise $\neg p \leftrightarrow q$ given by common sense, people may conclude that $\neg p$.¹⁰

Besides, natural language is ambiguous—one sentence in natural language can come from multiple logical forms under the same interpretation. We use present tense and progressive aspect to encourage a reading of imaginary ongoing events, corresponding to the alethic modality. Such events are less likely to induce LLM’s or human’s individual bias, as they are unrelated to factual knowledge or moral judgements. Also, we always use two full verb phrases, ruling out sentences like “Jane is eating apples or oranges,” so the two events are less likely to be mutually exclusive. In this way, we can reduce the ambiguity of the questions in our dataset.

B.3 Data Samples

All logic forms and corresponding natural language sentences can be found in Table A1.

The exact prompt format is as follows:

```
Consider the following statements:\n
Jane is watching a show or John is reading a book.\n
Jane isn't watching a show.\n
Question: Based on these statements, can we infer that John\n
is reading a book?\n
Answer:<eof>
```

¹⁰This confounding factor affects the examples in Figure 10 of Holliday and Mandelkern (2024).

C Additional Experiments

C.1 Extra Experiment: Introduction Rule of Modality

We report the results on the necessitation rule and its variants here, as these rules are obscure and verbose to be articulated in natural language:

$$\begin{aligned} \{\varphi\} &\vdash \Box\varphi, & (\text{necessitation rule}) \\ \{\varphi\} &\vdash \Diamond\varphi, \\ \{\varphi\} &\vdash \varphi. \end{aligned}$$

Its natural language form is as follows:

- Jane is watching a show.
- ($\Box$) Can we infer that it’s certain that Jane is watching a show?
- ($\Diamond$) Can we infer that it’s possible that Jane is watching a show?
- ($\emptyset$) Can we infer that Jane is watching a show?

All three variants are paired with 1000 logic interpretations. As they are all rules of inference, the ground truth answer is always Yes. Overall accuracy is shown in Table A2, where across all LLMs, the necessitation rule has the lowest accuracy. This echoes the necessity modality’s tendency to be rejected discussed in §4.2.3.

We further fit a linear mixed-effects model similar to Eq. (4), except that the argument form effect is now constant across all data points. The mixed-effects model yields a marginal R^2 of 0.391 and a conditional R^2 of 0.745. Estimated marginal means shows that the accuracy on $\emptyset$ is 0.171 less than $\Diamond$, but 0.371 higher than $\Box$, with both differences significant at $p < 0.0001$. This further suggests that modality serves as an important factor on logic reasoning performance.

C.2 Extra Experiment: Distribution of Modalities

Besides the necessitation rule, *distribution axiom* is the other fundamental axiom in normal modal logic. It can be transformed into the rule shown in Eq. (A1), and plugging in the definition of $\vee$ in Eq. (2) gives the rule shown in Eq. (A2). Notice that Eq. (A2) closely resembles rule $\vee^L$ ’s variant with necessity, as shown in Eq. (A3), except the different scope of the necessity operator and the position of the negation operator. Moving the negation operator out of the necessity operator will result in

Validity	Modality	Argument Form	Logical Form	Natural Language
$\vdash$	$\emptyset$	$\vee^L$	$\{p \vee q, \neg p\} \vdash q$	Jane is <u>watching a show</u> or John is <u>reading a book</u> . Jane isn't <u>watching a show</u> . Can we infer that John is <u>reading a book</u> ?
	$\emptyset$	$\vee^R$	$\{p \vee q, \neg q\} \vdash p$	Jane is <u>watching a show</u> or John is <u>reading a book</u> . John isn't <u>reading a book</u> . Can we infer that Jane is <u>watching a show</u> ?
	$\emptyset$	$\rightarrow^L$	$\{\neg p \rightarrow q, \neg p\} \vdash q$	If Jane isn't <u>watching a show</u> , then John is <u>reading a book</u> . Jane isn't <u>watching a show</u> . Can we infer that John is <u>reading a book</u> ?
	$\emptyset$	$\rightarrow^R$	$\{\neg p \rightarrow q, \neg q\} \vdash p$	If Jane isn't <u>watching a show</u> , then John is <u>reading a book</u> . John isn't <u>reading a book</u> . Can we infer that Jane is <u>watching a show</u> ?
	$\square$	$\vee^L$	$\{\square p \vee \square q, \neg \square p\} \vdash \square q$	It's certain that Jane is <u>watching a show</u> or it's certain that John is <u>reading a book</u> . It's uncertain whether Jane is <u>watching a show</u> . Can we infer that it's certain that John is <u>reading a book</u> ?
	$\square$	$\vee^R$	$\{\square p \vee \square q, \neg \square q\} \vdash \square p$	It's certain that Jane is <u>watching a show</u> or it's certain that John is <u>reading a book</u> . It's uncertain whether John is <u>reading a book</u> . Can we infer that it's certain that Jane is <u>watching a show</u> ?
	$\square$	$\rightarrow^L$	$\{\neg \square p \rightarrow \square q, \neg \square p\} \vdash \square q$	If it's uncertain whether Jane is <u>watching a show</u> , then it's certain that John is <u>reading a book</u> . It's uncertain whether Jane is <u>watching a show</u> . Can we infer that it's certain that John is <u>reading a book</u> ?
	$\square$	$\rightarrow^R$	$\{\neg \square p \rightarrow \square q, \neg \square q\} \vdash \square p$	If it's uncertain whether Jane is <u>watching a show</u> , then it's certain that John is <u>reading a book</u> . It's uncertain whether John is <u>reading a book</u> . Can we infer that it's certain that Jane is <u>watching a show</u> ?
	$\diamond$	$\vee^L$	$\{\diamond p \vee \diamond q, \neg \diamond p\} \vdash \diamond q$	It's possible that Jane is <u>watching a show</u> or it's possible that John is <u>reading a book</u> . It's impossible that Jane is <u>watching a show</u> . Can we infer that it's possible that John is <u>reading a book</u> ?
	$\diamond$	$\vee^R$	$\{\diamond p \vee \diamond q, \neg \diamond q\} \vdash \diamond p$	It's possible that Jane is <u>watching a show</u> or it's possible that John is <u>reading a book</u> . It's impossible that John is <u>reading a book</u> . Can we infer that it's possible that Jane is <u>watching a show</u> ?
	$\diamond$	$\rightarrow^L$	$\{\neg \diamond p \rightarrow \diamond q, \neg \diamond p\} \vdash \diamond q$	If it's impossible that Jane is <u>watching a show</u> , then it's possible that John is <u>reading a book</u> . It's impossible that Jane is <u>watching a show</u> . Can we infer that it's possible that John is <u>reading a book</u> ?
	$\diamond$	$\rightarrow^R$	$\{\neg \diamond p \rightarrow \diamond q, \neg \diamond q\} \vdash \diamond p$	If it's impossible that Jane is <u>watching a show</u> , then it's possible that John is <u>reading a book</u> . It's impossible that John is <u>reading a book</u> . Can we infer that it's possible that Jane is <u>watching a show</u> ?
$\nVdash$	$\emptyset$	$\vee^L$	$\{p \vee q, q\} \nVdash \neg p$	Jane is <u>watching a show</u> or John is <u>reading a book</u> . John is <u>reading a book</u> . Can we infer that Jane isn't <u>watching a show</u> ?
	$\emptyset$	$\vee^R$	$\{p \vee q, p\} \nVdash \neg q$	Jane is <u>watching a show</u> or John is <u>reading a book</u> . Jane is <u>watching a show</u> . Can we infer that John isn't <u>reading a book</u> ?
	$\emptyset$	$\rightarrow^L$	$\{\neg p \rightarrow q, q\} \nVdash \neg p$	If Jane isn't <u>watching a show</u> , then John is <u>reading a book</u> . John is <u>reading a book</u> . Can we infer that Jane isn't <u>watching a show</u> ?
	$\emptyset$	$\rightarrow^R$	$\{\neg p \rightarrow q, p\} \nVdash \neg q$	If Jane isn't <u>watching a show</u> , then John is <u>reading a book</u> . Jane is <u>watching a show</u> . Can we infer that John isn't <u>reading a book</u> ?
	$\square$	$\vee^L$	$\{\square p \vee \square q, \square q\} \nVdash \neg \square p$	It's certain that Jane is <u>watching a show</u> or it's certain that John is <u>reading a book</u> . It's certain that John is <u>reading a book</u> . Can we infer that it's uncertain whether Jane is <u>watching a show</u> ?
	$\square$	$\vee^R$	$\{\square p \vee \square q, \square p\} \nVdash \neg \square q$	It's certain that Jane is <u>watching a show</u> or it's certain that John is <u>reading a book</u> . It's certain that Jane is <u>watching a show</u> . Can we infer that it's uncertain whether John is <u>reading a book</u> ?
	$\square$	$\rightarrow^L$	$\{\neg \square p \rightarrow \square q, \square q\} \nVdash \neg \square p$	If it's uncertain whether Jane is <u>watching a show</u> , then it's certain that John is <u>reading a book</u> . It's certain that John is <u>reading a book</u> . Can we infer that it's uncertain whether Jane is <u>watching a show</u> ?
	$\square$	$\rightarrow^R$	$\{\neg \square p \rightarrow \square q, \square p\} \nVdash \neg \square q$	If it's uncertain whether Jane is <u>watching a show</u> , then it's certain that John is <u>reading a book</u> . It's certain that Jane is <u>watching a show</u> . Can we infer that it's uncertain whether John is <u>reading a book</u> ?
	$\diamond$	$\vee^L$	$\{\diamond p \vee \diamond q, \diamond q\} \nVdash \neg \diamond p$	It's possible that Jane is <u>watching a show</u> or it's possible that John is <u>reading a book</u> . It's possible that John is <u>reading a book</u> . Can we infer that it's impossible that Jane is <u>watching a show</u> ?
	$\diamond$	$\vee^R$	$\{\diamond p \vee \diamond q, \diamond p\} \nVdash \neg \diamond q$	It's possible that Jane is <u>watching a show</u> or it's possible that John is <u>reading a book</u> . It's possible that Jane is <u>watching a show</u> . Can we infer that it's impossible that John is <u>reading a book</u> ?
	$\diamond$	$\rightarrow^L$	$\{\neg \diamond p \rightarrow \diamond q, \diamond q\} \nVdash \neg \diamond p$	If it's impossible that Jane is <u>watching a show</u> , then it's possible that John is <u>reading a book</u> . It's possible that John is <u>reading a book</u> . Can we infer that it's impossible that Jane is <u>watching a show</u> ?
	$\diamond$	$\rightarrow^R$	$\{\neg \diamond p \rightarrow \diamond q, \diamond p\} \nVdash \neg \diamond q$	If it's impossible that Jane is <u>watching a show</u> , then it's possible that John is <u>reading a book</u> . It's possible that Jane is <u>watching a show</u> . Can we infer that it's impossible that John is <u>reading a book</u> ?

Table A1: Samples of all logical forms and corresponding natural language sentences.

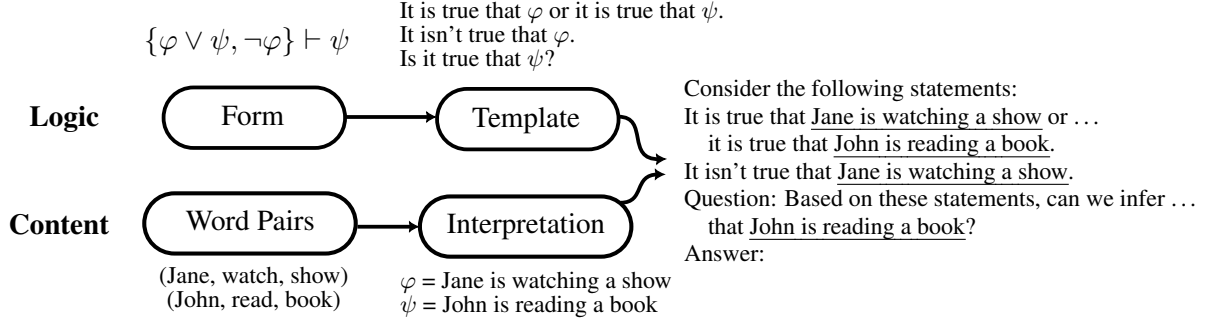


Figure A1: The data synthesis pipeline: for each variable in logic forms (§3.1) we assign meanings to them to obtain the natural language question-answering pairs (§3.2).

	$\emptyset$	$\square$	$\diamond$
mistral-7b	0.998	0.885	0.999
mistral-8x7b	0.957	0.540	0.987
llama-2-7b	0.768	0.013	0.920
llama-2-13b	0.368	0.004	0.829
llama-2-70b	0.511	0.051	0.834
llama-3-8b	0.398	0.225	0.783
llama-3-70b	0.674	0.384	0.794
yi-34b	0.960	0.382	0.999
phi-2	0.814	0.226	0.892
phi-3-mini	0.992	0.925	0.994

Table A2: Overall accuracy of the necessitation rule and its modality variants on each model.

Modality	Argument Form	Logical Form
$\emptyset$	base	$\varphi \vee \psi, \neg\varphi \vdash \psi$
$\square$	base	$\square\varphi \vee \square\psi, \neg\square\varphi \vdash \square\psi$
$\square$	theorem	$\square(\varphi \vee \psi), \square\neg\varphi \vdash \square\psi$
$\square$	<u>spurious</u>	$\square(\varphi \vee \psi), \neg\square\varphi \not\vdash \square\psi$
$\diamond$	base	$\diamond\varphi \vee \diamond\psi, \neg\diamond\varphi \vdash \diamond\psi$
$\diamond$	theorem	$\diamond(\varphi \vee \psi), \diamond\neg\varphi \vdash \diamond\psi$
$\diamond$	spurious	$\diamond(\varphi \vee \psi), \neg\diamond\varphi \vdash \diamond\psi$

Table A3: Logical forms and their ground truth to study the distribution of modalities. Only the spurious form of the necessity modality (marked by underline) has a ground truth of false.

a fallacy (Eq. A4).

$$\{\square(\varphi \rightarrow \psi), \square\varphi\} \vdash \square\psi, \quad (\text{A1})$$

$$\{\square(\varphi \vee \psi), \square\neg\varphi\} \vdash \square\psi, \quad (\text{A2})$$

$$\{\square\varphi \vee \square\psi, \neg\square\varphi\} \vdash \square\psi, \quad (\text{A3})$$

$$\{\square(\varphi \vee \psi), \neg\square\varphi\} \not\vdash \square\psi. \quad (\text{A4})$$

We say (A2) to (A4) are of argument form theorem, base and spurious, respectively. See Table A3 for the logical forms and their ground truth we used to study the distribution of modalities. The natural language form is as follows:

(theorem) It's certain that if Freddy is not going shopping, then Coy is making dinner.
 (spurious) It's certain that Freddy is not going shopping.
 It's uncertain whether Freddy is going shopping.
 Can we infer that it's certain that Coy is making dinner?

This group of rules and fallacies comes from the fact that the necessity modality $\square$ is not distributive to disjunction, i.e. $\square(\varphi \vee \psi) \not\vdash \square\varphi \vee \square\psi$ (Xiang, 2019, Ex. 5). In contrast, the possibility

modality $\diamond$ is distributive to disjunction. This particular case could have served as a material to test the LLM's knowledge of the asymmetry between the two modalities, yet in §4.2.3 we showed that there is a bias towards rejection on the necessity modality. As the false case of the disjunction is on the necessity modality, this bias confounds the experiment.

We fit a linear mixed-effects model similar to Eq. (4) to the data,

$$Acc_{soft} \sim \text{Modality} \times \text{ArgForm} + \text{Perplexity} + (1 + \text{Perplexity} \mid \text{LLM}),$$

with an interaction term between the modality and argument form. On the theorem form compared to the base form, the necessity modality $\square$ has a 0.173 higher estimated marginal means with $p < 0.0001$ significance, yet the possibility modality $\diamond$ has a 0.071 lower estimated marginal means. On the spurious form compared to the base form, the $\square$ has a 0.312 higher means, and the $\diamond$ has no significant difference. On both forms, $\diamond \succ \square$ in terms of accuracy still holds at a slight margin of 0.110 and 0.047 respectively.

To verify whether on $\square$ the performance increase on spurious form is due to the rejection bias, we fit a linear mixed-effects model with the relative probability of answering Yes as dependent variable. Results show that on spurious form compared to the base form, the effect of $\square$'s tendency to answer Yes is only 0.060 lower, indicating the rejection bias of the base form is still present. Therefore, we hypothesize that the LLM's performance on recognizing the fallacy of necessity distribution over disjunction is hindered by the rejection bias on the necessity modality.