

A Survey on Enhancing Large Language Models with Symbolic Reasoning

Anonymous ACL submission

Abstract

Reasoning is one of the fundamental human abilities, central to problem-solving, decision-making, and planning. With the development of large language models (LLMs), significant attention has been paid to enhancing and understanding their reasoning capabilities. Most existing works attempt to directly use LLMs for natural-language-based reasoning. However, due to the inherent semantic ambiguity and complexity of natural language, LLMs often struggle with complex problems, leading to challenges such as hallucinations and inconsistent reasoning. Therefore, techniques for constructing formal language representations, most of which are symbolic languages, have emerged. In this work, we focus on symbolic reasoning in LLMs and provide a comprehensive review of the related research. This includes the types of symbolic languages used, different symbolic reasoning tasks and related benchmarks, and typical techniques for enhancing LLMs' symbolic reasoning abilities. Our goal is to offer a thorough review of symbolic reasoning in LLMs, highlighting key findings and challenges while providing a reference for future research in this area.

1 Introduction

Reasoning involves deriving conclusions or solutions from limited information by applying logical analysis, pattern recognition, and knowledge integration. Researchers have found that once large language models (LLMs) reach a certain threshold in terms of parameters and training data, they can exhibit reasoning capabilities (Wei et al., 2022), leading researchers to explore a variety of techniques to enhance the reasoning capabilities of LLMs (Zhang et al., 2022; Yao et al., 2024; Ning et al., 2023). However, LLMs face significant challenges in their reasoning processes, with hallucination being one of the most notable. This issue becomes especially prominent in complex reasoning tasks, where LLMs may fail to account for all

relevant factors, omit key information, or fall into logical traps.

To address these challenges and further enhance the reasoning capabilities of LLMs, researchers have begun to explore an innovative framework that enables LLMs to focus on question comprehension and symbolic representation generation, while delegating the execution of reasoning steps to external solvers (Olausson et al., 2023; Pan et al., 2023; Gao et al., 2023). This framework stems from classic neuro-symbolic approaches (Andreas et al., 2016; Liang et al., 2017; Ebrahimi et al., 2021) and effectively alleviates the limitations of LLMs in reasoning tasks by integrating the strengths of symbolic representation and specialized solvers.

As long as LLMs generate accurate symbolic representations, external solvers can take over and efficiently solve complex reasoning tasks based on these representations. This technique not only reduces the burden on the language model but also fully leverages the professional advantages of the external solver in dealing with specific problems, achieving complementary advantages and collaborative work between the LLMs and the external solver. Symbolic languages, with their precision, interpretability, and logicity, bring significant advantages to the reasoning process (Olausson et al., 2023; Lyu et al., 2023a; Chen et al., 2022).

Research on LLMs reasoning has progressed significantly, and numerous review articles have discussed it from different perspectives. Some papers provide an overall review of reasoning tasks (Plaat et al., 2024; Xu et al., 2025; Huang and Chang, 2022). Others focus on specific reasoning tasks and deeply analyze the technological advancements within particular fields (Lu et al., 2022). Yu et al. conduct a comprehensive review of natural language reasoning (Yu et al., 2024). However, despite its potential to greatly enhance LLMs reasoning, symbolic reasoning remains underexplored in terms of systematic review and analysis, leaving a gap in understanding its full impact and utility in complex reasoning tasks.

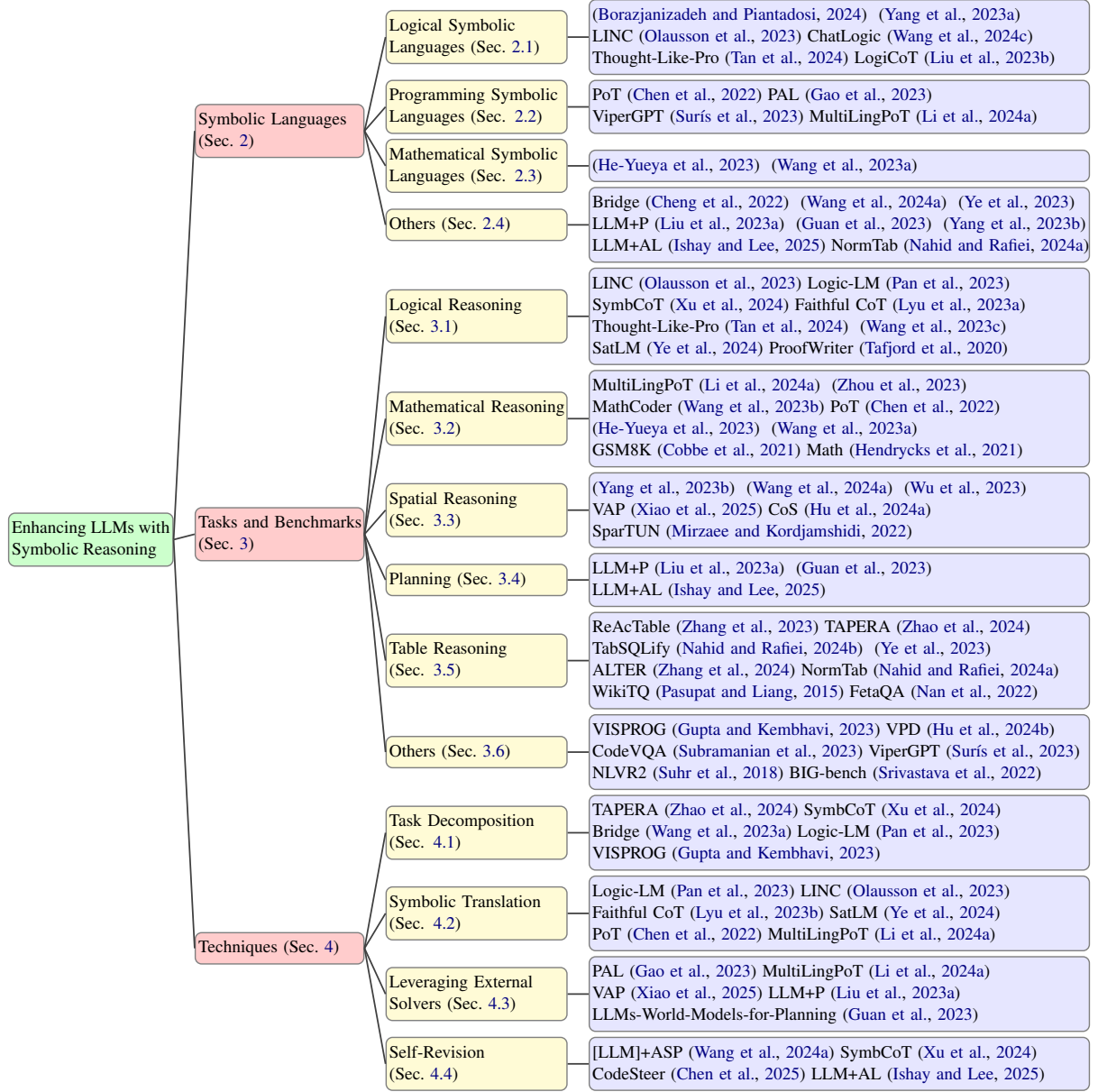


Figure 1: The structure of this survey.

To address this research gap, we present this comprehensive survey that systematically examines how to use symbolic language to enhance the reasoning capabilities of LLMs. As illustrated in Figure 1, we systematically organize these studies along three dimensions: symbolic languages, tasks and benchmarks, and techniques. Through an extensive synthesis and forward-looking analysis of existing research, this paper aims to establish a clear conceptual framework for future researchers, thereby facilitating further advancements and breakthroughs in the application of symbolic languages for LLMs reasoning.

2 Symbolic Languages

Symbolic languages serve as vital tools within the realms of artificial intelligence and logic. When faced with a variety of reasoning tasks, selecting the appropriate symbolic language becomes particularly crucial. In the following text, we will introduce various symbolic languages.

2.1 Logical Symbolic Languages

In Principia Mathematica (Newton, 1934), the language of logical symbols is rigorously defined, encompassing both the symbols and the rules. Common logical symbols include conjunction ($\wedge$), disjunction ($\vee$), negation ($\neg$), and quantifiers ($\forall$, $\exists$). These symbols form the foundation of logical expressions and, through precise rules, ensure the

rigor and consistency of logical reasoning. Compared to natural language, the language of logical symbols adheres to strict rules of reasoning. As long as the premises are correct, it ensures that the derivation of conclusions is error-free.

It is pointed out in LINC (Olausson et al., 2023) that utilizing LLMs to transform natural language into first-order logic and then processing them with Prover9 can help improve the accuracy of logical reasoning tasks. The programming language Prolog, based on first-order logic and known for its powerful logical reasoning capabilities, is also widely applied in reasoning tasks (Borazjanizadeh and Piantadosi, 2024; Yang et al., 2023a; Tan et al., 2024; Wang et al., 2024c).

2.2 Programming Symbolic Languages

A programming language is a formal language specifically designed for writing computer programs, aimed at enabling computers to execute tasks with precision and efficiency. Nowadays, there is a wide variety of programming languages, yet they share some common characteristics. First, programming languages establish rigorous syntactic rules. Moreover, they commonly incorporate fundamental logical constructs, including selection and loop structures. The advantage of programming languages lies in their capability to handle various complex reasoning tasks with the assistance of data structures and algorithms.

Different programming languages possess unique advantages. For example, Matlab is better at matrix operations than Java, while Python provides a rich library for number theory (Li et al., 2024a). These advantages make Matlab and Python excel in mathematical reasoning tasks. The experimental results conducted by Luo et al. (Luo et al., 2024) indicate that Java performs the best in table reasoning tasks, and C++ demonstrates higher efficiency in mathematical applications and spatial reasoning tasks.

2.3 Mathematical Symbolic Language

Mathematical Symbolic Language is the cornerstone of mathematical expression and communication. It utilizes a standardized set of symbols (such as $+$, $-$, $\times$, $\div$) and rules to precisely describe mathematical concepts, relationships, and operations (Newton, 1934). Mathematical Symbolic Language is characterized by its exceptional conciseness and precision, making it highly effective in enhancing the accuracy of LLMs when tackling reasoning tasks, particularly those involving mathematical reasoning.

When faced with mathematical problems, He et al. (He-Yueya et al., 2023) and Wang et al. (Wang et al., 2023a) advocate LLMs should transform these problems into mathematical equations and subsequently solve them using external solvers.

2.4 Others

There are other symbolic languages specifically tailored for specialized domains. Answer Set Programming (ASP) is a declarative programming paradigm. In [LLM]+ASP (Yang et al., 2023b), prompts are employed to steer LLMs in transforming natural language to ASP code and then invoking a solver to obtain faithful results. Wang et al. (Wang et al., 2024a) introduces DSPy (Declarative Self-improving Language Programs in Python) to conduct self-refinement on the prompts to better obtain the ASP code.

Planning Domain Definition Language (PDDL) is a formal language for describing planning problems. The approach proposed by LLM+P (Liu et al., 2023a) uses PDDL to formally describe planning tasks and relies on prompts to generate the problem code. LLMs-World-Models-for-Planning (Guan et al., 2023) takes the LLMs as an intermediate layer. It also brings in human experts to modify the generated PDDL code. LLM+AL (Ishay and Lee, 2025) expounds on the limitations of the PDDL and proposes Action Language (AL) as an alternative. This language is designed to be more flexible than PDDL, capable of expressing complex causal relationships, temporal constraints, and uncertain events, while maintaining the rigor of formalization.

SQL is a declarative language specifically designed for managing and manipulating relational databases, widely used in table reasoning (Nahid and Rafiei, 2024a; Ye et al., 2023; Cheng et al., 2022).

3 Tasks and Benchmarks

This section focuses on the symbolic reasoning tasks of LLMs and their associated datasets and benchmarks. Each task faces unique challenges and has specific requirements. Detailed information about these datasets is presented in Table 1.

3.1 Logical Reasoning

Logical reasoning tasks (Pan et al., 2023; Xu et al., 2024; Lyu et al., 2023a) require LLMs to infer the truth value of a conclusion from a set of rules and conditions. This critically depends on the LLMs' capability to understand logical rules and conditions while consistently upholding rigor throughout

Domains	Benchmarks	Size	Representative Works
Logical Reasoning	ProofWriter (Tafjord et al., 2020) PrOntoQA (Saparov and He, 2022) FOLIO (Han et al., 2022) AR-LSAT (Zhong et al., 2022) LogiQA (Liu et al., 2020) CLUTRR (Sinha et al., 2019)	500K 10K 1435 2046 8678 10K	(Yang et al., 2023a; Lee and Hwang, 2024; Pan et al., 2023; Xu et al., 2024) (Pan et al., 2023; Xu et al., 2024; Tan et al., 2024) (Li et al., 2024b; Kalyanpur et al., 2024; Liu et al., 2025) (Pan et al., 2023; Xu et al., 2024; Wang et al., 2024b) (Liu et al., 2025; Li et al., 2024b; Bao et al., 2024) (Ye et al., 2024; Yang et al., 2023b)
Mathematical Reasoning	GSM8K (Cobbe et al., 2021) Math (Hendrycks et al., 2021) AQuA (Ling et al., 2017) SVAMP (Patel et al., 2021) ASDiv (Miao et al., 2020) MAWPS (Koncel-Kedziorski et al., 2016) ALGEBRA (He-Yueya et al., 2023)	8.5K 12.5K 100K 1K 2305 3320 222	(Borazjanizadeh and Piantadosi, 2024; Lyu et al., 2023a; Chen et al., 2022; Gao et al., 2023) (Zhou et al., 2023; Wang et al., 2023b; Li et al., 2024a; Tan et al., 2024) (Lyu et al., 2023a; Chen et al., 2022; Leang et al., 2024) (Lyu et al., 2023a; Chen et al., 2022; Gao et al., 2023; Leang et al., 2024) (Lyu et al., 2023a; Gao et al., 2023) (Gao et al., 2023) (He-Yueya et al., 2023; Wang et al., 2023a)
Spatial Reasoning	StepGame (Shi et al., 2022) SparTUN (Mirzaee and Kordjamshidi, 2022)	6.1K 50K	(Wang et al., 2024a; Yang et al., 2023b) (Wang et al., 2024a)
Planning	International Planning Competition (IPC) domains (Seipp et al., 2022)	-	(Liu et al., 2023a; Guan et al., 2023)
Table Reasoning	WikiTQ (Pasupat and Liang, 2015) FetaQA (Nan et al., 2022) TabFact (Chen et al., 2020) WikiSQL (Zhong et al., 2017)	22033 10K 117854 80654	(Zhang et al., 2023; Nahid and Rafiei, 2024b; Mouravieff et al., 2024; Cheng et al., 2022; Zhang et al., 2024) (Ye et al., 2023; Zhang et al., 2023; Nahid and Rafiei, 2024b) (Ye et al., 2023; Nahid and Rafiei, 2024b; Wang et al., 2024d; Cheng et al., 2022; Zhang et al., 2024) (Nahid and Rafiei, 2024b)
Others	BIG-bench (Srivastava et al., 2022) Fruit Shop (Hu et al., 2023) VQAv2 (Goyal et al., 2017)	- 70 1105904	(Borazjanizadeh and Piantadosi, 2024; Gao et al., 2023; Li et al., 2023) (Hu et al., 2023) (Hu et al., 2024b)

Table 1: Common datasets and benchmarks used to evaluate the symbolic reasoning capabilities of LLMs. We classify these datasets into different domains, mark their sizes, and note some representative works that used them for evaluation.

the reasoning process. However, due to the ambiguity and complexity of natural language, LLMs are prone to generating hallucinations and errors. Symbolic reasoning effectively mitigates this ambiguity by employing precise symbolic representations and formal rules to express concepts and relationships (Olausson et al., 2023; Tan et al., 2024; Ye et al., 2024).

The content of the datasets used for logical reasoning tasks mainly includes rules, conditions, conclusions, and conclusion validity labels. ProofWriter (Tafjord et al., 2020) is a synthetic dataset dedicated to multi-step logical reasoning. PrOntoQA (Saparov and He, 2022) is a logical reasoning dataset constructed based on formal ontology structures. FOLIO (Han et al., 2022) is a natural language inference benchmark constructed based on first-order logic. AR-LSAT (Zhong et al., 2022) serves as a benchmark formulated around the analytical reasoning questions of the Law School Admission Test. LogiQA (Liu et al., 2020) is a logical reasoning question-answering dataset constructed based on legal and daily scenarios.

3.2 Mathematical Reasoning

Mathematical tasks (Zhou et al., 2023; He-Yueya et al., 2023; Wang et al., 2023a) encompass a wide range of problems, including algebra, geometry, and calculus. During the reasoning process, precise analysis of mathematical conditions is essential, often accompanied by extensive computations. This poses significant challenges to LLMs, demanding both strong analytical understanding and accurate computational capabilities (Chen et al., 2022). Therefore, by guiding LLMs to focus on generating symbolic representations (such as mathemati-

cal formulas or Python code) and leveraging these representations to delegate specific computational tasks to external solvers, not only is the burden on LLMs reduced, but the accuracy of mathematical reasoning is also significantly enhanced (Wang et al., 2023b; Chen et al., 2022).

The content of the datasets used for mathematical reasoning tasks consists of problem descriptions, rationales, and answers. The ASDiv (Miao et al., 2020), SVAMP (Patel et al., 2021), and GSM8K (Cobbe et al., 2021) are primarily composed of primary-school-level math word problems. Math (Hendrycks et al., 2021) mainly consists of challenging math competition problems. AQuA (Ling et al., 2017) contains multiple-choice math reasoning questions and detailed explanations. MAWPS (Koncel-Kedziorski et al., 2016) is a benchmark that contains various math word problems and their solutions. ALGEBRA (He-Yueya et al., 2023) mainly focuses on algebraic problems.

3.3 Spatial Reasoning

Spatial reasoning (Hu et al., 2024a; Xiao et al., 2025) is a fundamental aspect of human cognition, enabling human to interact effectively with the environment. It plays a crucial role in tasks that involve understanding and reasoning about the spatial relationships between objects and their movements. However, spatial reasoning tasks pose a significant challenge for LLMs. It is usually hard for LLMs to understand the complex relationship descriptions formed by natural language in spatial reasoning. Utilizing specialized symbolic expressions to model spatial relationships helps to achieve reliable results (Yang et al., 2023b; Wang et al., 2024a).

The data format of the spatial reasoning datasets are composed of scene descriptions, questions, and answers. StepGame (Shi et al., 2022) is designed to test the ability to multi-hop spatial reasoning, SparTUN (Mirzaee and Kordjamshidi, 2022) is built upon the NLVR (Natural Language for Visual Reasoning) images.

3.4 Planning

At the heart of planning tasks (Liu et al., 2023a; Guan et al., 2023; Ishay and Lee, 2025) is crafting a viable path from the initial state to the goal. This requires guiding the system or environment towards the desired outcome based on the initial state, through a series of carefully selected actions. Symbolic languages enable the flexible construction of planning algorithms tailored to specific needs, allowing for precise definition and optimization of state transitions. This significantly enhances the executability and reliability of plans.

Planning datasets contain a scenario in which some decisions need to be made by robots. Some works conduct experiments in the domains (e.g., GRIPPERS, TYREWORLD) of the International Planning Competition (IPC) (Seipp et al., 2022).

3.5 Table Reasoning

The task of table reasoning (Zhang et al., 2023; Zhao et al., 2024; Nahid and Rafiei, 2024b; Ye et al., 2023; Zhang et al., 2024) aims to enable models to generate corresponding results as answers based on task requirements when receiving one or more tables as input (Wang et al., 2024d).

The complexity and huge volume of information in table may obscure key details, potentially weakening the decision-making ability of LLMs (Liu et al., 2023c). Compared to relying on natural language processing, employing domain-specific symbolic languages designed for structured data, such as SQL, can effectively reduce the burden on LLMs.

The content of the datasets used for table reasoning tasks consists of tables, questions, and answers. WikiTQ (Pasupat and Liang, 2015) is constructed from Wikipedia tables. FetaQA (Nan et al., 2022) is built based on multi-source factual tables. TabFact (Chen et al., 2020) is a fact-checking dataset constructed from Wikipedia tables. WikisQL (Zhong et al., 2017) is a large-scale text-to-SQL dataset built upon Wikipedia tables.

3.6 Others

Given that multi-modal reasoning tasks (Surís et al., 2023; Hu et al., 2024b) typically span multiple cate-

gories discussed previously, they cannot be strictly classified into a single specific category. The main challenge in multi-modal reasoning lies in how to efficiently integrate information from different modalities. Leveraging symbolic languages enable seamless interaction and collaborative cooperation among different modalities, thereby enhancing the overall reasoning capabilities (Gupta and Kembhavi, 2023).

The content of the datasets used for multi-modal reasoning tasks consists of visual information and textual information. The NLVR2 dataset (Suhr et al., 2018) consists of images and their corresponding descriptive sentences. The CoVR dataset (Ventura et al., 2024) is a dataset for composite video retrieval tasks.

Some datasets are challenging to precisely match with specific tasks. We provide a brief introduction here. BIG-bench (Srivastava et al., 2022) encompasses over 200 tasks designed to test various reasoning capabilities of LLMs. Some benchmarks are designed for testing the performance of the framework they proposed to solve some real-world problems (Hu et al., 2023). There are corresponding datasets available for testing some works that apply VLMs (Goyal et al., 2017).

4 Techniques

In general, four typical techniques are used to enhance the reasoning capabilities of LLMs, including task decomposition, symbolic translation, leveraging external solvers, and self-revision. The typical processes of enhancing LLMs with symbolic reasoning are shown in Figure 2. Many works adopt task decomposition as a fundamental technique. Symbolic translation and leveraging external solvers are usually used together. The mapping relationship between techniques and tasks is shown in Table 2.

4.1 Task Decomposition

By decomposing problems into smaller and manageable sub-problems, LLMs can handle these problems effectively, which makes the problem-solving process clear and trackable. Task decomposition can serve as a fundamental technique, followed by the deployment of symbolic languages. Bridge (Wang et al., 2023a) improves the accuracy of generating equations by decomposing complex problems into independent sub-problems. VIS-PROG (Gupta and Kembhavi, 2023) addresses complex tasks by decomposing them into multiple modules, each processing using visual models, Python, and other tools.

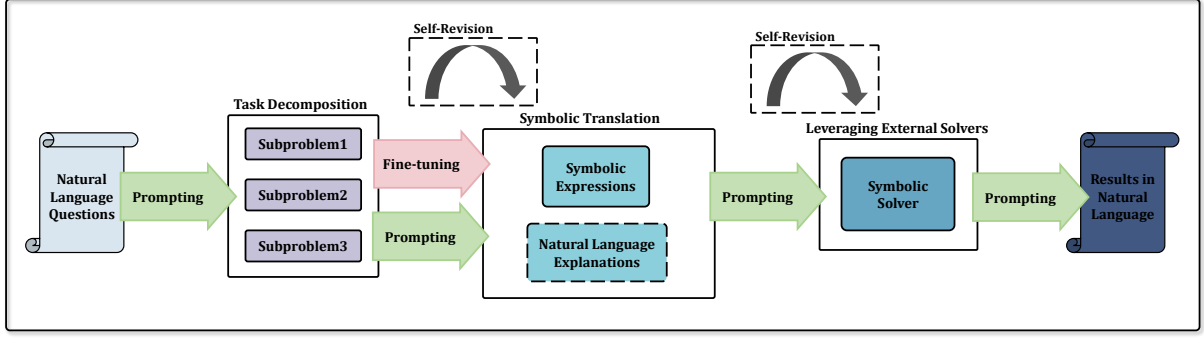


Figure 2: Typical processes of enhancing LLMs with symbolic reasoning.

Task decomposition is widely used in table reasoning. The implementation forms of task decomposition in table reasoning include problem decomposition and table decomposition. Chain-of-Table (Wang et al., 2024d) framework decomposes tables by dynamically generating a chain of table operations, and presents intermediate reasoning results in the form of structured tables. Aiming at long-form table question answering, TAPERA (Zhao et al., 2024) decomposes complex questions into sub-problems through three modules: QA content planner, executable table reasoner, and answer generator. In TabSQLify (Nahid and Rafiei, 2024b) and ReAcTable (Zhang et al., 2023), symbolic languages have been used to decompose the large tables into small sub-tables containing only the essential information required to answer questions. DATER (Ye et al., 2023) and ALTER (Zhang et al., 2024) adopt a dual decomposition mechanism, extracting sub-tables relevant to the question from large tables and breaking down complex problems into logical sub-problems. The issue of table structure has been addressed by NormTab (Nahid and Rafiei, 2024a), which optimizes tables through value normalization and structural normalization, and decomposes tables into sub-tables.

4.2 Symbolic Translation

In most of the works, symbolic translation appears together with leveraging external solvers. The process of these two techniques is first to use LLMs to translate the natural language into symbolic expressions, then select an appropriate solver to address the problem. There are some differences in the implementation forms of symbolic translation. Some works use prompts to guide LLMs to conduct symbolic translation, while some works adopt fine-tuning to enhance the translation capabilities of LLMs.

Translating natural language into logical expressions is a prevalent practice in logical reasoning.

Logic-LM (Pan et al., 2023) and LINC (Olausson et al., 2023) generate a symbolic representation for the input problem with LLMs via in-context learning. Faithful CoT (Lyu et al., 2023b) prompts LLMs to translate the problems into a reasoning chain, which interleaves natural language comments and symbolic language programs. SatLM (Ye et al., 2024) uses LLMs to parse natural language problems into declarative task specifications (e.g., logical formulas), which have relative solvers.

Programming symbolic languages and mathematical symbolic languages can represent quantitative relationships and are suitable for solving mathematical problems. PoT (Chen et al., 2022) prompts LLMs to generate Python code. Considering the differences between programming symbolic languages, MultiLingPoT (Li et al., 2024a) fine-tunes LLMs to generate code in four distinct programming symbolic languages. Declarative-math-word-problem (He-Yueya et al., 2023) translates natural language problems into systems of equations according to set principles. Bridge (Wang et al., 2023a) erases information irrelevant to equation generation from the sub-problems and transforms them into equations.

Spatial reasoning and planning have domain-specific symbolic languages. CoS has been proposed by (Hu et al., 2024a), leveraging specific symbols to simplify the spatial-relationship information expressed in natural language in CoT. [LLM]+ASP and DSPy-based LLM+ASP (Yang et al., 2023b; Wang et al., 2024a) introduces ASP (Answer Set Programming), translates the natural language into ASP code. VAP (Xiao et al., 2025) guides the LLMs to call upon Multi-Modal Large Language Model (MLLM) to understand image information and assist in plan generation.

Compared to other tasks, planning is more dependent on symbolic languages because it is necessary to ensure that the plans are executable and reliable.

Tasks	Papers	Techniques
Logical Reasoning	Logic-LM (Pan et al., 2023) SymbCoT (Xu et al., 2024) LINC (Olausson et al., 2023) AMR-LDA (Bao et al., 2024) Faithful CoT (Lyu et al., 2023b) SatLM (Ye et al., 2024) LoT (Liu et al., 2025) AMR-LDA (Bao et al., 2024) Thought-Like-Pro (Tan et al., 2024)	Symbolic Translation, Leveraging External Solvers Self-Revision Symbolic Translation, Leveraging External Solvers Symbolic Translation Symbolic Translation, Leveraging External Solvers Symbolic Translation, Leveraging External Solvers Symbolic Translation Symbolic Translation Symbolic Translation
Mathematical Reasoning	PoT (Chen et al., 2022) CSV (Zhou et al., 2023) MultiLingPoT (Li et al., 2024a) Declarative-math-word-problem (He-Yueya et al., 2023) Bridge (Wang et al., 2023a)	Symbolic Translation, Leveraging External Solvers Leveraging External Solvers, Self-Revision Symbolic Translation, Leveraging External Solvers Symbolic Translation Task Decomposition, Symbolic Translation
Spatial Reasoning	CoS (Hu et al., 2024a) [LLM]+ASP (Yang et al., 2023b) DSPy-based LLM+ASP (Wang et al., 2024a) VAP (Xiao et al., 2025)	Symbolic Translation Symbolic Translation, Leveraging External Solvers Symbolic Translation, Leveraging External Solvers, Self-Revision Symbolic Translation, Leveraging External Solvers
Planning	LLM+P (Liu et al., 2023a) LLMs-World-Models-for-Planning (Guan et al., 2023) LLM+AL (Ishay and Lee, 2025)	Symbolic Translation, Leveraging External Solvers Symbolic Translation, Leveraging External Solvers, Self-Revision Symbolic Translation, Leveraging External Solvers, Self-Revision
Table Reasoning	Chain-of-Table (Wang et al., 2024d) TAPERA (Zhao et al., 2024) TabSQLify (Nahid and Rafiei, 2024b) ReAcTable (Zhang et al., 2023) DATER (Ye et al., 2023) ALTER (Zhang et al., 2024) NormTab (Nahid and Rafiei, 2024a)	Task Decomposition Task Decomposition, Leveraging External Solvers Task Decomposition Task Decomposition Task Decomposition Task Decomposition Task Decomposition
Others	VISPROG (Gupta and Kembhavi, 2023) CodeVQA (Subramanian et al., 2023) CodeSteer (Chen et al., 2025) ViperGPT (Surís et al., 2023) PAL (Gao et al., 2023)	Task Decomposition Leveraging External Solvers Self-Revision Leveraging External Solvers Leveraging External Solvers

Table 2: Techniques used in different reasoning tasks.

Symbolic languages enable flexible construction of planning algorithms, allowing for the precise definition and optimization of state transitions. The approach proposed by LLM+P (Liu et al., 2023a) converts natural language into PDDL using LLMs. LLMs-World-Models-for-Planning (Guan et al., 2023) incorporates corrections for errors in the content translated by LLMs. LLM+AL (Ishay and Lee, 2025) introduces an AL called BC+ (Babb and Lee, 2020) as an alternative to PDDL and prompts LLMs to translate natural language into AL.

Logical expansion based on symbolic representation provides a more complete semantic expression, reducing the risk of semantic conversion errors in LLMs (Liu et al., 2025). LoT (Liu et al., 2025) utilizes LLMs to extract sentences containing conditional reasoning relationships from the input and implement logical expansion using Python. According to AMR-LDA (Bao et al., 2024), the Abstract Meaning Representation (AMR) graph is used to carry out logical expansion by applying the principle of logical equivalence. Imitation learning has been leveraged by Thought-Like-Pro (Tan et al., 2024) to enable LLMs to mimic the reasoning trajectories generated by the Prolog logic engine.

4.3 Leveraging External Solvers

Leveraging external solvers is a technique that uses solvers to address formal symbolic expressions. This technique has effectively addressed the limitations of LLMs in terms of precise computation. By guiding the LLMs to generate content in the input format required by these solvers, reliable results can be obtained.

For logical reasoning problems, after obtaining the symbolic expressions, using a solver to derive the answers is prevalent. LINC (Olausson et al., 2023) focuses on first-order logic and Logic-LM (Pan et al., 2023) adopts various solvers to address different kinds of problems. Faithful CoT (Lyu et al., 2023b) incorporates external solvers and ensures that the answer is the deterministic result of executing the reasoning chain. SatLM (Ye et al., 2024) adopts relative solvers to derive answers based on logical formulas.

Programming symbolic languages are suitable for solving mathematical problems. PoT (Chen et al., 2022) executes the generated Python code to solve mathematical problems. PAL (Gao et al., 2023) and CSV (Zhou et al., 2023) use generated Python code as intermediate reasoning steps. Mul-

tiLingPoT (Li et al., 2024a) considers multiple programming symbolic languages, selects an appropriate one from four different options.

Spatial reasoning and planning require considering complex spatial relationships, and have exclusive symbolic languages and solvers. [LLM]+ASP (Yang et al., 2023b) and DSPy-based LLM+ASP (Wang et al., 2024a) obtain results using an ASP solver. Visual Question Answering (VQA) is transformed into a modular code generation task, where CodeVQA (Subramanian et al., 2023) and ViperGPT (Surís et al., 2023) prompt the LLMs to generate Python code that invokes the APIs of visual language models (VLMs) to process images. LLM+P (Liu et al., 2023a) and LLMs-World-Models-for-Planning (Guan et al., 2023) employ a planner to generate a plan after the symbolic translation. LLM+AL (Ishay and Lee, 2025) invokes a BC+ solver to obtain plans.

4.4 Self-Revision

For some complex problems, due to issues such as hallucinations, using LLMs to conduct symbolic translation or leveraging external solvers may result in errors. Self-Revision can be used multiple times within the framework to correct these errors. Logic-LM (Pan et al., 2023) designs a module to modify inaccurate logical formulations using the error messages from the symbolic solver as feedback. SymbCoT (Xu et al., 2024) designs a Verifier in the framework. For the found invalid logic, the Verifier refines the reasoning steps. CSV (Zhou et al., 2023) proposes "code-based explicit self-verification", which introduces an additional verification stage. Iteratively refining the generated ASP programs and employing the DSPy framework, DSPy-based LLM+ASP (Wang et al., 2024a) manages complex workflows and optimizes prompts effectively. By fine-tuning a small model CodeSteerLLM as an assistant, combined with a symbolic checker and a self-answer verifier, CodeSteer (Chen et al., 2025) guides LLMs to switch between text reasoning and code generation and continuously refines the results.

For the planning task, it is prone to result in factual errors and syntax errors when generating plans. Introducing a self-revision module is a prevalent practice in planning. LLMs-World-Models-for-Planning (Guan et al., 2023) combines the PDDL verification tool with the feedback from human domain experts to correct model errors. LLM+AL (Ishay and Lee, 2025) introduces multiple verification processes to ensure the executability and correctness of the plan.

5 Future Directions

Customized Symbolic Languages and Solvers.

Recent research relies mainly on existing formal symbolic languages to assist in reasoning, without innovating the grammar for specific tasks (Olausson et al., 2023; Yang et al., 2023b; Liu et al., 2023a). However, these existing formal symbolic languages cannot fully cover all application scenarios in the real world (Ishay and Lee, 2025) and are unable to entirely meet the requirements. Therefore, customizing language parsers and related symbolic grammar according to application scenarios may be the focus of future work.

Multi-Symbolic Language Integration. Reasoning tasks are inherently complex and typically cannot be effectively addressed using a single language alone. Multi-Symbolic language integration can fully leverage the strengths of each symbolic language, but current methods are superficial in integration and lack flexibility (Li et al., 2024a; Pan et al., 2023; Zhang et al., 2023). We encourage researchers to explore better ways of multilingual integration to resolve the grammatical and semantic conflicts between symbolic languages.

Inherent Structured Languages for LLMs.

When leveraging symbolic languages for reasoning in LLMs, many errors stem from how these models generate task-specific symbolic expressions. The currently proposed methods of generating code in multiple steps (Zhou et al., 2023) and human-like debugging (Zhong et al., 2024) cannot fully resolve the problem. We encourage researchers to explore alternatives that either replace symbolic languages altogether or imbue LLMs with inherent structured and rigorous reasoning capabilities. A possible direction is to design a neuro-symbolic framework that enables LLMs to implicitly align free-form reasoning processes with symbolic representations.

6 Conclusion

In this paper, we conduct a systematic review of the symbolic reasoning in LLMs. We summarize the types of symbolic languages used in reasoning, illustrate different reasoning tasks and relative benchmarks, and discuss the typical techniques for enhancing the symbolic reasoning capabilities of LLMs. We also discuss the promising future directions of symbolic reasoning. We hope this paper can offer a comprehensive and valuable overview of the present status of the field and facilitate further advancements in the application of symbolic language for LLM reasoning.

7 Limitations

While this survey aims to provide a comprehensive overview of the integration of symbolic reasoning with large language models (LLMs), it has its limitations, which we acknowledge to provide a balanced perspective on our work.

First, the field of enhancing LLMs with symbolic reasoning is evolving at an unprecedented pace. New methodologies, frameworks, and applications are being published frequently, making it challenging to capture the most recent advancements. Despite our rigorous efforts to include the latest research up to the submission deadline, some cutting-edge developments may have emerged during the final stages of this survey’s preparation.

Second, in an effort to provide a broad overview of the field, this survey categorizes the research from three perspectives: symbolic languages, tasks and benchmarks, and techniques. While this approach offers a structured framework for understanding the landscape, the breadth of coverage inevitably comes at the expense of depth in certain areas. Some technical nuances, domain-specific challenges, and emerging sub-fields may have been underexplored or oversimplified.

Finally, the majority of reasoning benchmarks are collected and categorized from the experimental sections of mainstream industry works, potentially leading to insufficient coverage of niche or domain-specific reasoning tasks.

Despite these limitations, we believe this survey provides a valuable foundation for understanding the current state of research and identifying future directions in symbolic reasoning with LLMs. We encourage researchers to build upon this work and address the gaps identified here to further advance the field.

References

Jacob Andreas, Marcus Rohrbach, Trevor Darrell, and Dan Klein. 2016. Neural module networks. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 39–48.

Joseph Babb and Joohyung Lee. 2020. Action language +. *Journal of Logic and Computation*, 30(4):899–922.

Qiming Bao, Alex Yuxuan Peng, Zhenyun Deng, Wanjun Zhong, Gaël Gendron, Timothy Pistotti, Neşet Tan, Nathan Young, Yang Chen, Yonghua Zhu, Paul Denny, Michael Witbrock, and Jiamou Liu. 2024. [Abstract Meaning Representation-based logic-driven data augmentation for logical reasoning](#). In *Findings of the Association for Computational Linguistics:*

ACL 2024, pages 5914–5934, Bangkok, Thailand. Association for Computational Linguistics.

Nasim Borazjanizadeh and Steven T Piantadosi. 2024. Reliable reasoning beyond natural language. *arXiv preprint arXiv:2407.11373*.

Wenhu Chen, Xueguang Ma, Xinyi Wang, and William W Cohen. 2022. Program of thoughts prompting: Disentangling computation from reasoning for numerical reasoning tasks. *arXiv preprint arXiv:2211.12588*.

Wenhu Chen, Hongmin Wang, Jianshu Chen, Yunkai Zhang, Hong Wang, Shiyang Li, Xiyu Zhou, and William Yang Wang. 2020. [TabFact: A large-scale dataset for table-based fact verification](#). *Preprint*, arXiv:1909.02164.

Yongchao Chen, Yilun Hao, Yueying Liu, Yang Zhang, and Chuchu Fan. 2025. CodeSteer: Symbolic-augmented language models via code/text guidance. *arXiv preprint arXiv:2502.04350*.

Zhoujun Cheng, Tianbao Xie, Peng Shi, Chengzu Li, Rahul Nadkarni, Yushi Hu, Caiming Xiong, Dragomir Radev, Mari Ostendorf, Luke Zettlemoyer, et al. 2022. Binding language models in symbolic languages. *arXiv preprint arXiv:2210.02875*.

Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, et al. 2021. Training verifiers to solve math word problems. *arXiv preprint arXiv:2110.14168*.

Monireh Ebrahimi, Aaron Eberhart, Federico Bianchi, and Pascal Hitzler. 2021. Towards bridging the neuro-symbolic gap: Deep deductive reasoners. *Applied Intelligence*, 51:6326–6348.

Luyu Gao, Aman Madaan, Shuyan Zhou, Uri Alon, Pengfei Liu, Yiming Yang, Jamie Callan, and Graham Neubig. 2023. Pal: Program-aided language models. In *International Conference on Machine Learning*, pages 10764–10799. PMLR.

Yash Goyal, Tejas Khot, Douglas Summers-Stay, Dhruv Batra, and Devi Parikh. 2017. Making the V in VQA Matter: Elevating the role of image understanding in visual question answering. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 6904–6913.

Lin Guan, Karthik Valmeekam, Sarath Sreedharan, and Subbarao Kambhampati. 2023. Leveraging Pre-trained Large Language Models to Construct and Utilize World Models for Model-based Task Planning. *Advances in Neural Information Processing Systems*, 36:79081–79094.

Tanmay Gupta and Aniruddha Kembhavi. 2023. Visual Programming: Compositional Visual Reasoning without Training. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 14953–14962.

733	Simeng Han, Hailey Schoelkopf, Yilun Zhao, Zhen-	Chengshu Li, Jacky Liang, Andy Zeng, Xinyun Chen,	787
734	ing Qi, Martin Riddell, Wenfei Zhou, James Coady,	Karol Hausman, Dorsa Sadigh, Sergey Levine, Li Fei-	788
735	David Peng, Yujie Qiao, Luke Benson, et al. 2022.	Fei, Fei Xia, and Brian Ichter. 2023. Chain of Code:	789
736	FOLIO: Natural Language Reasoning with First-	Reasoning with a Language Model-Augmented Code	790
737	Order Logic. <i>arXiv preprint arXiv:2209.00840</i> .	Emulator. <i>arXiv preprint arXiv:2312.04474</i> .	791
738	Joy He-Yueya, Gabriel Poesia, Rose E Wang, and	Nianqi Li, Zujie Liang, Siyu Yuan, Jiaqing Liang,	792
739	Noah D Goodman. 2023. Solving math word prob-	Feng Wei, and Yanghua Xiao. 2024a. MultiL-	793
740	lems by combining language models with symbolic	ingPoT: Enhancing Mathematical Reasoning with	794
741	solvers. <i>arXiv preprint arXiv:2304.09102</i> .	Multilingual Program Fine-tuning. <i>arXiv preprint</i>	795
742	Dan Hendrycks, Collin Burns, Saurav Kadavath, Akul	<i>arXiv:2412.12609</i> .	796
743	Arora, Steven Basart, Eric Tang, Dawn Song, and Ja-	Qingchuan Li, Jiatong Li, Tongxuan Liu, Yuting Zeng,	797
744	cob Steinhardt. 2021. Measuring Mathematical Prob-	Mingyue Cheng, Weizhe Huang, and Qi Liu. 2024b.	798
745	lem Solving With the Math Dataset. <i>arXiv preprint</i>	Leveraging LLMs for Hypothetical Deduction in Log-	799
746	<i>arXiv:2103.03874</i> .	ical Inference: A Neuro-Symbolic Approach. <i>arXiv</i>	800
747	Chenxu Hu, Jie Fu, Chenzhuang Du, Simian Luo, Junbo	<i>preprint arXiv:2410.21779</i> .	801
748	Zhao, and Hang Zhao. 2023. ChatDB: Augmenting	Chen Liang, Jonathan Berant, Quoc Le, Kenneth For-	802
749	LLMs with Databases as Their Symbolic Memory.	bus, and Ni Lao. 2017. Neural symbolic machines:	803
750	<i>arXiv preprint arXiv:2306.03901</i> .	Learning semantic parsers on freebase with weak	804
751	Hanxu Hu, Hongyuan Lu, Huajian Zhang, Yun-Ze Song,	supervision. In <i>Proceedings of the 55th Annual Meet-</i>	805
752	Wai Lam, and Yue Zhang. 2024a. Chain-of-Symbol	<i>ing of the Association for Computational Linguistics</i>	806
753	Prompting For Spatial Reasoning in Large Language	(<i>Volume 1: Long Papers</i>), pages 23–33.	807
754	Models. In <i>First Conference on Language Modeling</i> .	Wang Ling, Dani Yogatama, Chris Dyer, and Phil Blun-	808
755	Yushi Hu, Otilia Stretcu, Chun-Ta Lu, Krishnamurthy	som. 2017. Program Induction by Rationale Genera-	809
756	Viswanathan, Kenji Hata, Enming Luo, Ranjay Kr-	tion: Learning to Solve and Explain Algebraic Word	810
757	ishna, and Ariel Fuxman. 2024b. Visual Program	Problems . In <i>Proceedings of the 55th Annual Meet-</i>	811
758	Distillation: Distilling Tools and Programmatic Rea-	<i>ing of the Association for Computational Linguistics</i>	812
759	soning into Vision-Language Models. In <i>Proceed-</i>	(<i>Volume 1: Long Papers</i>), pages 158–167, Vancouver,	813
760	<i>ings of the IEEE/CVF Conference on Computer Vi-</i>	Canada. Association for Computational Linguistics.	814
761	<i>sion and Pattern Recognition</i> , pages 9590–9601.	Bo Liu, Yuqian Jiang, Xiaohan Zhang, Qiang Liu,	815
762	Jie Huang and Kevin Chen-Chuan Chang. 2022. To-	Shiqi Zhang, Joydeep Biswas, and Peter Stone.	816
763	wards Reasoning in Large Language Models: A sur-	2023a. LLM+P: Empowering Large Language	817
764	vey. <i>arXiv preprint arXiv:2212.10403</i> .	Models with Optimal Planning Proficiency. <i>arXiv</i>	818
765	Adam Ishay and Joohyung Lee. 2025. LLM+AL:	<i>preprint arXiv:2304.11477</i> .	819
766	Bridging Large Language Models and Action Lan-	Hanmeng Liu, Zhiyang Teng, Leyang Cui, Chaoli	820
767	guages for Complex Reasoning about Actions. <i>arXiv</i>	Zhang, Qiji Zhou, and Yue Zhang. 2023b. Logi-	821
768	<i>preprint arXiv:2501.00830</i> .	CoT: Logical Chain-of-Thought Instruction Tuning.	822
769	Aditya Kalyanpur, Kailash Karthik Saravanakumar, Vic-	In <i>Proc. of EMNLP Findings</i> , pages 2908–2921.	823
770	tor Barres, Jennifer Chu-Carroll, David Melville, and	Jian Liu, Leyang Cui, Hanmeng Liu, Dandan Huang,	824
771	David Ferrucci. 2024. LLM-ARC: Enhancing LLMs	Yile Wang, and Yue Zhang. 2020. LogiQA: A	825
772	with an Automated Reasoning Critic. <i>arXiv preprint</i>	Challenge Dataset for Machine Reading Compre-	826
773	<i>arXiv:2406.17663</i> .	hension with Logical Reasoning. <i>arXiv preprint</i>	827
774	Rik Koncel-Kedziorski, Subhro Roy, Aida Amini, Nate	<i>arXiv:2007.08124</i> .	828
775	Kushman, and Hannaneh Hajishirzi. 2016. MAWPS:	Tianyang Liu, Fei Wang, and Muhao Chen. 2023c. Re-	829
776	A Math Word Problem Repository. In <i>Proceedings of</i>	thinking Tabular Data Understanding with Large Lan-	830
777	<i>the 2016 conference of the north american chapter of</i>	guage Models. <i>arXiv preprint arXiv:2312.16702</i> .	831
778	<i>the association for computational linguistics: human</i>	Tongxuan Liu, Wenjiang Xu, Weizhe Huang, Yuting	832
779	<i>language technologies</i> , pages 1152–1157.	Zeng, Jiaying Wang, Xingyu Wang, Hailong Yang,	833
780	Joshua Ong Jun Leang, Aryo Pradipta Gema, and	and Jing Li. 2025. Logic-of-thought: Injecting logic	834
781	Shay B Cohen. 2024. CoMAT: Chain of Mathemat-	into contexts for full reasoning in large language	835
782	ically Annotated Thought Improves Mathematical	models . In <i>Proceedings of the 2025 Conference of the</i>	836
783	Reasoning. <i>arXiv preprint arXiv:2410.10336</i> .	<i>Nations of the Americas Chapter of the Association</i>	837
784	Jinu Lee and Wonseok Hwang. 2024. SymBa: Sym-	<i>for Computational Linguistics: Human Language</i>	838
785	bolic Backward Chaining for Multi-step Natural Lan-	<i>Technologies (Volume 1: Long Papers)</i> , pages 10168–	839
786	guage Reasoning. <i>arXiv preprint arXiv:2402.12806</i> .	10185, Albuquerque, New Mexico. Association for	840
		Computational Linguistics.	841

954	capabilities of language models. <i>arXiv preprint arXiv:2206.04615</i> .		
955			
956	Sanjay Subramanian, Medhini Narasimhan, Kushal		
957	Khangaonkar, Kevin Yang, Arsha Nagrani, Cordelia		
958	Schmid, Andy Zeng, Trevor Darrell, and Dan Klein.		
959	2023. Modular visual question answering via code		
960	generation. <i>arXiv preprint arXiv:2306.05392</i> .		
961	Alane Suhr, Stephanie Zhou, Ally Zhang, Iris Zhang,		
962	Huajun Bai, and Yoav Artzi. 2018. A corpus for		
963	reasoning about natural language grounded in pho-		
964	tographs. <i>arXiv preprint arXiv:1811.00491</i> .		
965	Dídac Surís, Sachit Menon, and Carl Vondrick. 2023.		
966	ViperGPT: Visual inference via python execution		
967	for reasoning. In <i>Proceedings of the IEEE/CVF In-</i>		
968	<i>ternational Conference on Computer Vision</i> , pages		
969	11888–11898.		
970	Oyvind Taffjord, Bhavana Dalvi Mishra, and Peter		
971	Clark. 2020. ProofWriter: Generating implications,		
972	proofs, and abductive statements over natural lan-		
973	guage. <i>arXiv preprint arXiv:2012.13048</i> .		
974	Xiaoyu Tan, Yongxin Deng, Xihe Qiu, Weidi Xu,		
975	Chao Qu, Wei Chu, Yinghui Xu, and Yuan Qi.		
976	2024. THOUGHT-LIKE-PRO: Enhancing Reason-		
977	ing of Large Language Models through Self-Driven		
978	Prolog-based Chain-of-Thought. <i>arXiv preprint</i>		
979	<i>arXiv:2407.14562</i> .		
980	Lucas Ventura, Antoine Yang, Cordelia Schmid, and		
981	Gül Varol. 2024. Covr: Learning composed video		
982	retrieval from web video captions. In <i>Proceedings</i>		
983	<i>of the AAAI Conference on Artificial Intelligence</i> ,		
984	volume 38, pages 5270–5279.		
985	Dingzirui Wang, Longxu Dou, Wenbin Zhang, Junyu		
986	Zeng, and Wanxiang Che. 2023a. Exploring equa-		
987	tion as a better intermediate meaning represen-		
988	tation for numerical reasoning. <i>arXiv preprint</i>		
989	<i>arXiv:2308.10585</i> .		
990	Ke Wang, Houxing Ren, Aojun Zhou, Zimu Lu, Sichun		
991	Luo, Weikang Shi, Renrui Zhang, Linqi Song,		
992	Mingjie Zhan, and Hongsheng Li. 2023b. Math-		
993	Coder: Seamless code integration in llms for en-		
994	hanced mathematical reasoning. <i>arXiv preprint</i>		
995	<i>arXiv:2310.03731</i> .		
996	Rong Wang, Kun Sun, and Jonas Kuhn. 2024a. Dspy-		
997	based neural-symbolic pipeline to enhance spatial		
998	reasoning in llms. <i>arXiv preprint arXiv:2411.18564</i> .		
999	Ruocheng Wang, Eric Zelikman, Gabriel Poesia, Yewen		
1000	Pu, Nick Haber, and Noah D Goodman. 2023c. Hy-		
1001	pothesis Search: Inductive reasoning with language		
1002	models. <i>arXiv preprint arXiv:2309.05660</i> .		
1003	Siyuan Wang, Zhongyu Wei, Yejin Choi, and Xiang		
1004	Ren. 2024b. Symbolic working memory enhances		
1005	language models for complex rule application. <i>arXiv</i>		
1006	<i>preprint arXiv:2408.13654</i> .		
	Zhongsheng Wang, Jiamou Liu, Qiming Bao, Hongfei	1007	
	Rong, and Jingfeng Zhang. 2024c. ChatLogic: Inte-	1008	
	grating logic programming with large language mod-	1009	
	els for multi-step reasoning. In <i>2024 International</i>	1010	
	<i>Joint Conference on Neural Networks (IJCNN)</i> , pages	1011	
	1–8. IEEE.	1012	
	Zilong Wang, Hao Zhang, Chun-Liang Li, Julian Mar-	1013	
	tin Eisenschlos, Vincent Perot, Zifeng Wang, Lesly	1014	
	Miculicich, Yasuhisa Fujii, Jingbo Shang, Chen-Yu	1015	
	Lee, et al. 2024d. Chain-of-Table: Evolving tables	1016	
	in the reasoning chain for table understanding. <i>arXiv</i>	1017	
	<i>preprint arXiv:2401.04398</i> .	1018	
	Jason Wei, Yi Tay, Rishi Bommasani, Colin Raffel,	1019	
	Barret Zoph, Sebastian Borgeaud, Dani Yogatama,	1020	
	Maarten Bosma, Denny Zhou, Donald Metzler, et al.	1021	
	2022. Emergent abilities of large language models.	1022	
	<i>arXiv preprint arXiv:2206.07682</i> .	1023	
	Xiaoqian Wu, Yong-Lu Li, Jianhua Sun, and Cewu Lu.	1024	
	2023. Symbol-LLM: Leverage language models for	1025	
	symbolic system in visual human activity reason-	1026	
	ing . In <i>Advances in Neural Information Processing</i>	1027	
	<i>Systems</i> , volume 36, pages 29680–29691. Curran As-	1028	
	sociates, Inc.	1029	
	Ziyang Xiao, Dongxiang Zhang, Xiongwei Han, Xi-	1030	
	aojin Fu, Wing Yin Yu, Tao Zhong, Sai Wu, Yuan	1031	
	Wang, Jianwei Yin, and Gang Chen. 2025. Enhanc-	1032	
	ing llm reasoning via vision-augmented prompting.	1033	
	<i>Advances in Neural Information Processing Systems</i> ,	1034	
	37:28772–28797.	1035	
	Fengli Xu, Qian Yue Hao, Zefang Zong, Jingwei Wang,	1036	
	Yunke Zhang, Jingyi Wang, Xiaochong Lan, Jiahui	1037	
	Gong, Tianjian Ouyang, Fanjin Meng, et al. 2025.	1038	
	Towards large reasoning models: A survey of rein-	1039	
	forced reasoning with large language models. <i>arXiv</i>	1040	
	<i>preprint arXiv:2501.09686</i> .	1041	
	Jundong Xu, Hao Fei, Liangming Pan, Qian Liu, Mong-	1042	
	Li Lee, and Wynne Hsu. 2024. Faithful logical rea-	1043	
	soning via symbolic chain-of-thought. <i>arXiv preprint</i>	1044	
	<i>arXiv:2405.18357</i> .	1045	
	Sen Yang, Xin Li, Leyang Cui, Lidong Bing, and Wai	1046	
	Lam. 2023a. Neuro-symbolic integration brings	1047	
	causal and reliable reasoning proofs. <i>arXiv preprint</i>	1048	
	<i>arXiv:2311.09802</i> .	1049	
	Zhun Yang, Adam Ishay, and Joohyung Lee. 2023b.	1050	
	Coupling large language models with logic program-	1051	
	ming for robust and general reasoning from text.	1052	
	<i>arXiv preprint arXiv:2307.07696</i> .	1053	
	Shunyu Yao, Dian Yu, Jeffrey Zhao, Izhak Shafran,	1054	
	Tom Griffiths, Yuan Cao, and Karthik Narasimhan.	1055	
	2024. Tree of Thoughts: Deliberate problem solving	1056	
	with large language models. <i>Advances in Neural</i>	1057	
	<i>Information Processing Systems</i> , 36.	1058	
	Xi Ye, Qiaochu Chen, Isil Dillig, and Greg Durrett.	1059	
	2024. SatLM: Satisfiability-aided language models	1060	
	using declarative prompting. <i>Advances in Neural</i>	1061	
	<i>Information Processing Systems</i> , 36.	1062	

- Yunhu Ye, Binyuan Hui, Min Yang, Binhua Li, Fei Huang, and Yongbin Li. 2023. Large Language Models are Versatile Decomposers: Decomposing evidence and questions for table-based reasoning. In *Proceedings of the 46th International ACM SIGIR Conference on Research and Development in Information Retrieval*, pages 174–184.
- Fei Yu, Hongbo Zhang, Prayag Tiwari, and Benyou Wang. 2024. Natural language reasoning, a survey. *ACM Computing Surveys*, 56(12):1–39.
- Han Zhang, Yuheng Ma, and Hanfang Yang. 2024. Alter: Augmentation for large-table-based reasoning. *arXiv preprint arXiv:2407.03061*.
- Yunjia Zhang, Jordan Henkel, Avriella Floratou, Joyce Cahoon, Shaleen Deep, and Jignesh M Patel. 2023. ReAcTable: Enhancing react for table question answering. *arXiv preprint arXiv:2310.00815*.
- Zhuosheng Zhang, Aston Zhang, Mu Li, and Alex Smola. 2022. Automatic chain of thought prompting in large language models. *arXiv preprint arXiv:2210.03493*.
- Yilun Zhao, Lyuhao Chen, Arman Cohan, and Chen Zhao. 2024. TaPERA: enhancing faithfulness and interpretability in long-form table qa by content planning and execution-based reasoning. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 12824–12840.
- Li Zhong, Zilong Wang, and Jingbo Shang. 2024. Ldb: A large language model debugger via verifying runtime execution step-by-step. *arXiv preprint arXiv:2402.16906*.
- Victor Zhong, Caiming Xiong, and Richard Socher. 2017. Seq2SQL: Generating structured queries from natural language using reinforcement learning. *arXiv preprint arXiv:1709.00103*.
- Wanjun Zhong, Siyuan Wang, Duyu Tang, Zenan Xu, Daya Guo, Yining Chen, Jiahai Wang, Jian Yin, Ming Zhou, and Nan Duan. 2022. [Analytical reasoning of text](#). In *Findings of the Association for Computational Linguistics: NAACL 2022*, pages 2306–2319, Seattle, United States. Association for Computational Linguistics.
- Aojun Zhou, Ke Wang, Zimu Lu, Weikang Shi, Sichun Luo, Zipeng Qin, Shaoqing Lu, Anya Jia, Linqi Song, Mingjie Zhan, et al. 2023. Solving challenging math word problems using gpt-4 code interpreter with code-based self-verification. *arXiv preprint arXiv:2308.07921*.