

Leveraging Large Language Models to Generate Answer Set Programs

Adam Ishay¹, Zhun Yang¹, Joohyung Lee^{1,2}

¹Arizona State University

²Samsung Research

{aishay, zyang90, joolee}@asu.edu

Abstract

Large language models (LLMs), such as GPT-3 and GPT-4, have demonstrated exceptional performance in various natural language processing tasks and have shown the ability to solve certain reasoning problems. However, their reasoning capabilities are limited and relatively shallow, despite the application of various prompting techniques. In contrast, formal logic is adept at handling complex reasoning, but translating natural language descriptions into formal logic is a challenging task that non-experts struggle with. This paper proposes a neuro-symbolic method that combines the strengths of large language models and answer set programming. Specifically, we employ an LLM to transform natural language descriptions of logic puzzles into answer set programs. We carefully design prompts for an LLM to convert natural language descriptions into answer set programs in a step by step manner. Surprisingly, with just a few in-context learning examples, LLMs can generate reasonably complex answer set programs. The majority of errors made are relatively simple and can be easily corrected by humans, thus enabling LLMs to effectively assist in the creation of answer set programs.

1 Introduction

Transformer-based large language models (LLMs) have recently shown remarkable success in many downstream tasks, demonstrating their general reasoning capability across diverse problems. However, while LLMs excel in generating System 1 thinking, they struggle with System 2 thinking, resulting in output that is often inconsistent and incoherent (Nye et al. 2021). This is because LLMs are basically trained to predict subsequent words in a sequence and do not appear to have a deep understanding of concepts such as cause and effect, logic, and probability, which are essential for reasoning.

To address the issue, Nye et al. (2021) propose a dual-system model that combines the strengths of LLMs and symbolic logic to achieve improved performance on reasoning tasks. They leverage an LLM to generate a System 1 proposal and employ symbolic computation to filter these proposals for consistency and soundness.

We are interested in situations where problems are described in natural language and solving them requires deep reasoning. A system needs to take into account linguistic variability and be able to perform symbolic reasoning. We

take logic puzzles as the testbed as they are well-suited for this purpose.

We first note that GPT-3 (Brown et al. 2020) and GPT-4¹ by themselves struggle with solving logic puzzles, despite various prompts we tried. On the other hand, we find that they can convert the natural language descriptions of the puzzles into declarative answer set programming languages (Lifschitz 2008; Brewka, Niemelä, and Truszczyński 2011) surprisingly well. Even the errors these LLMs make are mostly simple for humans to correct. We hope that our finding will ease the efforts of writing answer set programs and expand the application of answer set programming to a broader audience.

The remainder of this paper is organized as follows. Section 2 offers a brief overview of related work on automated solving of logic puzzles. Sections 3 and 4 delve into the proposed approach in detail. Section 5 presents experimental results and performance evaluations of the approach. Section 6 shows more examples demonstrating the generalizability of our method.

The code is available at <https://github.com/azreasoners/gpt-asg-rules>.

2 Preliminaries

2.1 Large Language Models (LLMs)

LLMs have significantly improved natural language processing, achieving strong performance on a variety of tasks using few-shot learning (Brown et al. 2020). However, LLMs remain weak at tasks that involve complex reasoning (Creswell, Shanahan, and Higgins 2022; Valmeekam et al. 2022), and scaling model size alone is not enough to achieve good performance (Rae et al. 2021). It has been shown that various prompting methods improve accuracy on reasoning tasks (Wei et al. 2022; Zhou et al. 2022; Creswell, Shanahan, and Higgins 2022). Nye et al. (2021) present a dual-system model which uses an LLM as a semantic parser and couples it with a custom symbolic module to achieve performance gains on reasoning tasks. This framework combines the strengths of LLMs for parsing complex natural language and symbolic logic for handling

¹Throughout the paper, we use GPT-3 to refer to the “text-davinci-003” model and GPT-4 to refer to the “gpt-4-0314” (released March, 2023) model in the OpenAI API.

complex reasoning. However, the authors had to use hand-engineered set of constraints for the latter part. To our knowledge, our work is the first to use LLMs to generate logic rules to solve complex reasoning tasks.

2.2 Automated Logic Puzzle Solving

Works focused on solving logic puzzles typically involve a mapping from natural language to logic formalism. This process often includes problem simplification techniques, such as tailoring the puzzle to a specific domain, restricting natural language input to a certain form, or assuming additional inputs like enumerated types. Lev et al. (2004) employ a specialized automated multi-stage parsing process to convert natural language text into an intermediate form called Semantic Logic, which is then converted into First Order Logic to finally evaluate on law school admissions tests (LSAT) and the Graduate Records Examination (GRE). Shapiro (2011) manually encodes the “Jobs Puzzle” in a few different logical formalisms and compare them. Puzzler (Milicevic, Near, and Singh 2012) uses a general link parser to translate puzzles into the Alloy language for solving, primarily through an automated process, albeit with assumed types. LogicSolver (Nordstrom 2017) follows a similar approach to Puzzler but replaces Alloy with a custom solver and conducts a more comprehensive evaluation.

Several works utilize translations into the language of answer set programming (ASP) (Lifschitz 2008; Brewka, Niemelä, and Truszczyński 2011). Schwitter (2013) addresses the “Jobs Puzzle” by representing the problem using controlled natural language (Schwitter 2010), which can be further turned into ASP. Baral and Dzifcak (2012) employ a λ -calculus-based approach and trains a model that converts a manually simplified version of natural language clues into ASP rules for solving Zebra puzzle-type logic puzzles. Mitra and Baral (2015) train a maximum entropy-based model to extract relations for each clue, which are then converted into a common ASP rule format, where a stable model corresponds to the puzzle solution. LGPSolver (Jabrayilzade and Tekir 2020) uses DistilBERT, a transformer-based model, as a classifier that can distinguish between representative rule types. With the clue classification, the authors use a hand-crafted clue to Prolog translation (as opposed to ASP) and compute the solution. The works mentioned involve some combination of manual processing and/or brittle problem-specific translations. Our work distinguishes itself by being both fully automated and featuring a general pipeline, leveraging the extensive translation capacity available from LLMs.

2.3 Generate-Define-Test with ASP

ASP programs are typically written following the Generate-Define-Test structure, which generates potential solutions (*Generate*) and eliminates invalid ones based on certain constraints (*Test*). The *Generate* portion usually includes choice rules, while the *Test* portion consists of a set of constraints that prune out invalid solutions. An additional part of the program, the *Define* portion, includes necessary auxiliary predicates that are used in the *Test* portion.

3 Method

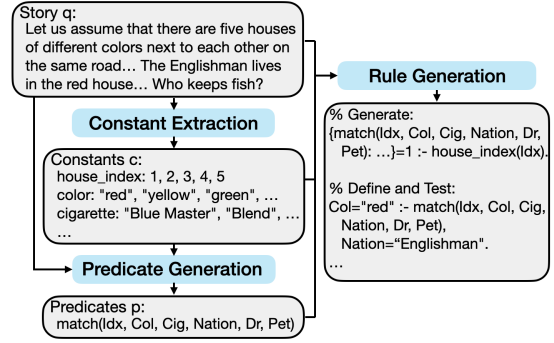


Figure 1: Flow of Generating Answer Set Programs from Logic Puzzle in English

In order to find a solution to a logic puzzle, we utilize GPT-3 to convert the puzzle into an answer set program so that the stable model (a.k.a answer set) encodes the solution.² Although GPT-3 exhibits strong capabilities, we discovered that it cannot generate a correct answer set program without being guided by carefully engineered prompts. These prompts instruct GPT-3 to reliably extract constants and generate accurate predicates and rules. In this paper, we detail our prompt engineering efforts.

Figure 1 illustrates the structure of our pipeline, which utilizes GPT-3 step by step to generate an ASP program. Similar to how a human would approach the task, our pipeline first extracts the relevant objects and their categories. Then, it generates a predicate that describes the relations among the objects from different categories. Using the generated information, the pipeline further constructs an ASP program in the style of Generate-Define-Test.

Let $\mathcal{F}_c$ and $\mathcal{F}_p$ denote the *Constant Extraction* and *Predicate Generation* steps in Figure 1. Let $\mathcal{F}_{r1}$ and $\mathcal{F}_{r2}$ represent the two parts of the *Rule Generation* step, i.e., the *Generate* part and the *Define&Test* part, respectively. Our pipeline can be modeled by the following equations that map a puzzle story q to an ASP program $\Pi = \Pi_{generate} \cup \Pi_{define_test}$.

$$\begin{aligned}
 c &= \mathcal{F}_c(q) & p &= \mathcal{F}_p(q, c) \\
 \Pi_{generate} &= \mathcal{F}_{r1}(c, p) & \Pi_{define_test} &= \mathcal{F}_{r2}(q, c, p).
 \end{aligned}$$

Here, c and p denote extracted objects and generated predicates. Each step $\mathcal{F}_*$ is realized by GPT-3 with 2-shot prompting, i.e., only 2 examples in each prompt.

3.1 Constant Extraction

The first step in the pipeline is to extract constants or entities from the given story along with their corresponding categories. To accomplish this, we invoke GPT-3 using **Prompt C**, which consists of three parts: instruction, examples, and a query.

Prompt C:

²Though this section mostly mentions GPT-3, GPT-4 can be used instead.

```

1  Given a problem, extract all different
   constants and their categories in the form "
   category: constant_1; constant_2; ...;
   constant_n". Here, the format of each constant
   is turned into either an integer or a string
   surrounded by double quotes, e.g., "some name".
2
3  Problem 1:
4  Consider N-Queens Puzzle on a chessboard of
   size 8x8. The goal is to assign 8 queens on
   the chessboard so that no two queens can share
   the same row, column, or diagonal.
5
6  Constants:
7  index_of_row: 1; 2; 3; 4; 5; 6; 7; 8.
8  index_of_column: 1; 2; 3; 4; 5; 6; 7; 8.
9
10 Problem 2:
11 "Against the Grain" offers hand-made wooden
   furniture at reasonable prices. Each item is
   made by an in-house employee. Using only the
   clues that follow, match each item to the
   employee who crafted it, and determine its
   price and the type of wood used to make it.
   Remember, as with all grid-based logic puzzles
   , no option in any category will ever be used
   more than once.
12 1. Bonita's piece costs $325.
13 2. The item made of poplar costs more than
   Yvette's piece.
14 3. Tabitha's item costs 50 dollars less than
   the piece made of sandalwood.
15 4. The $275 item is either the piece made of
   ash or Yvette's item.
16
17 Constants:
18 employee: "Bonita"; "Yvette"; "Tabitha".
19 price: 225; 275; 325.
20 wood_type: "ash"; "poplar"; "sandalwood".
21
22 Problem 3:
23 <story>
24
25 Constants:

```

Line 1 provides a general instruction for the task of extracting objects and directing GPT-3 to generate them in the form of “category: constant₁; ...; constant_n.” Then, two examples follow: Lines 6-8 for Problem 1 specified in Lines 3-4, and Lines 17-20 for Problem 2 specified in Lines 10-15. By replacing Line 23 (<story>) with a new example story and invoking GPT-3 with the above prompt, a new list of categories and constants for that story is generated, as with the previous two examples.

The above two examples are chosen to cover two cases of object extraction. For the N-Queens problem, the constants 1, ..., 8 are not described in the Problem 1 statement (Line 4) but can be inferred. For the second puzzle, however, all constants in Lines 18-20 are mentioned in the example story provided in Lines 11-15.

The second puzzle is also intentionally selected to give an example for GPT-3 so that certain constants (e.g., \$225) can be turned into valid integers (e.g., 225) so that arithmetic

can be applied correctly later when generating rules later on, while others should be surrounded by double quotes. We experimented with various prompts to instruct GPT-3 to generate all non-numeric constants in lowercase and replace special characters with underscores. However, GPT-3 was unable to strictly adhere to these instructions and consequently made more errors.

3.2 Predicate Generation

The next step in the pipeline is to generate predicates p that describe the relations among the extracted constants. We use GPT-3 on the **Prompt P** below.

Prompt P:

```

1  Given a problem and some categorized constants
   of the form "category: constant_1; constant_2;
   ...; constant_n", generate the minimum number
   of predicates to define the relations among
   the categories of constants. Each generated
   predicate is of the form "predicate(X1, X2,
   ..., Xn)" where X1, X2, ..., Xn are different
   variables and each variable X belongs to one
   of the categories. For each category, there
   must exist at least one variable of some
   predicate that belongs to this category.
2
3  Problem 1:
4  (Lines 4-8 from Prompt C: Omitted)
5
6  Predicates:
7  % The categories in Constants include
   index_of_row and index_of_column. We use
   different variables Ir and Ic to represent
   index_of_row and index_of_column.
8  % We assign a queen at row Ir and column Ic,
   where Ir belongs to index_of_row and Ic
   belongs to index_of_column.
9  assign(Ir, Ic)
10
11 Problem 2:
12 (Lines 11-20 from Prompt C: Omitted)
13
14 Predicates:
15 % The categories in Constants include employee,
   price, and wood_type. We use different
   variables E, P, and W to represent employee,
   price, and wood_type.
16 % We match an employee E with price P and wood
   type W, where E belongs to employee, P belongs
   to price, and W belongs to wood_type.
17 match(E, P, W)
18
19 Problem 3:
20 <story>
21
22 Constants:
23 <constants>
24
25 Predicates:

```

Line 1 is a general instruction describing the task of predicate generation, and that the generated predicates should follow the form of “predicate(X₁, ..., X_n)” where each X_i is a distinct variable that represents a category of constants.

Again, the two examples follow. Lines 3–4 are a copy of the first example in Lines 3–8 of **Prompt C** (where we omit Lines 4–8 from **Prompt C** to reduce the space). Lines 6–9 continue the first example, where it now generates the predicates with variables as arguments following the instruction. It also contains two comments (starting with symbol %). The first comment in Line 7 recalls the categories of constants and assigns a different variable to each category. The second comment in Line 8 gives the English reading of the predicate and variables, and emphasizes the link between each variable and a category of constants. Similarly, Lines 11–17 present the second example.

Next, the story and constants are given for the third problem and GPT-3 is prompted to generate the predicate for that example, given the general instruction and the preceding two examples.

Given the extracted constants c and generated predicates p , the next step in the pipeline is to generate ASP rules Π , consisting of the *Generate* part and the *Define&Test* part.

3.3 Rule Generation: Generate

The *Generate* part of an ASP program defines all possible mappings of constants from different categories. This is done by choice rules. In this step, an ASP program $\Pi_{generate}$ is obtained by calling GPT-3 with **Prompt R1**.

Prompt R1:

```

1  Given some categorized constants in the form "
   category: constant_1; constant_2; ...;
   constant_n" and some predicates about the
   relation among different categories of
   constants, write ASP (Answer Set Programming)
   rules to generate the search space of possible
   relations.
2
3  Constants:
4  employee: "Bonita"; "Yvette"; "Tabitha".
5  price: 225; 275; 325.
6  wood_type: "ash"; "poplar"; "sandalwood".
7
8  Predicates:
9  % The categories include employee, price, and
   wood_type. We use different variables E, P,
   and W to represent employee, price, and
   wood_type.
10 % We match an employee E with price P and wood
   type W, where E belongs to employee, P belongs
   to price, and W belongs to wood_type.
11 match(E, P, W)
12
13 ASP Rules:
14 % Define the constants in each category.
15 employee("Bonita"; "Yvette"; "Tabitha").
16 price(225; 275; 325).
17 wood_type("ash"; "poplar"; "sandalwood").
18
19 % For each employee E, it matches with exactly
   1 price P and 1 wood type W.
20 {match(E, P, W) : price(P), wood_type(W)}=1 :-
   employee(E).
21
22 Constants:
```

```

23 <constants>
24
25 Predicates:
26 <predicates>
27
28 ASP rules:
```

In the above prompt, `<constants>` and `<predicates>` are to be replaced for a new example. GPT-3 generates facts and choice rules following the last line of the prompt.

The task in this step is to write facts and choice rules based on the generated constants and predicates. Since this step doesn't require the details of the story, we omit the story from the prompt to avoid unnecessary noisy information being included in the prompt. Each example only consists of constants, predicates, and ASP rules to be generated, i.e., facts and choice rules.

Similar to the previous prompts, Line 1 is a general instruction, Lines 3–20 provide an example, and Lines 22–28 are for the queried example. The example ASP rules in Lines 14–20 contain comments (Lines 14 and 19), which will also be generated for the queried example and help to gather semantic information before generating a rule.

3.4 Rule Generation: Define and Test

The *Define&Test* part of an ASP program contains constraints that “weed out” the stable models that do not correspond to valid answers. This step takes as input the puzzle story q , constants c , and predicates p : semantically, the ASP rules represent the content in story q while, syntactically, the ASP rules must be formed by the extracted constants c and generated predicates p . The ASP program $\Pi_{define-test}$ is obtained by calling GPT-3 with **Prompt R2**.

Prompt R2:

```

1  Consider the constraint in the following form
2  <C1>; <C2>; ...; <Cm> :- <L1>, <L2>, ..., <Ln>.
3  which says that if the conjunction "<L1> and <
   L2> and ... and <Ln>" is true, then the
   disjunction of comparisons "<C1> or <C2> or
   ... or <Cm>" must be true.
4
5  One can also add a restriction that "exactly k
   of <C1>, <C2>, ..., <Cm> is true" by using the
   following form
6  {<C1>; <C2>; ...; <Cm>}=k :- <L1>, <L2>, ..., <
   Ln>.
7
8  Given a problem, extract all constraints from
   the clues in the problem using only the
   provided constants and predicates.
9
10 Problem 1:
11 "Against the Grain" offers hand-made wooden
   furniture at reasonable prices. Each item is
   made by an in-house employee. Using only the
   clues that follow, match each item to the
   employee who crafted it, and determine its
   price and the type of wood used to make it.
   Remember, as with all grid-based logic puzzles
   , no option in any category will ever be used
   more than once.
```

```

12 1. Bonita's piece costs $325.
13 2. The item made of poplar costs more than
   Yvette's piece.
14 3. Tabitha's item costs 50 dollars less than
   the piece made of sandalwood.
15 4. The $275 item is either the piece made of
   ash or Yvette's item.
16
17 Constants:
18 employee: "Bonita"; "Yvette"; "Tabitha".
19 price: 225; 275; 325.
20 wood_type: "ash"; "poplar"; "sandalwood".
21
22 Predicates:
23 % The categories include employee, price, and
   wood_type. We use different variables E, P,
   and W to represent employee, price, and
   wood_type.
24 % We match an employee E with price P and wood
   type W, where E belongs to employee, P belongs
   to price, and W belongs to wood_type.
25 match(E, P, W)
26
27 Constraints:
28 % No option in any category will ever be used
   more than once.
29 {E1=E2; P1=P2; W1=W2}=0 :- match(E1,P1,W1),
   match(E2,P2,W2), (E1,P1,W1) != (E2,P2,W2) .
30
31 % 1. Bonita's piece costs $325.
32 P=325 :- match(E,P,W), E="Bonita".
33
34 % 2. The item made of poplar costs more than
   Yvette's piece.
35 P1>P2 :- match(E1,P1,W1), match(E2,P2,W2), W1="
   poplar", E2="Yvette".
36
37 % 3. Tabitha's item costs 50 dollars less than
   the piece made of sandalwood.
38 P1=P2-50 :- match(E1,P1,W1), match(E2,P2,W2),
   E1="Tabitha", W2="sandalwood".
39
40 % 4. The $275 item is either the piece made of
   ash or Yvette's item.
41 {W="ash"; E="Yvette"}=1 :- match(E,P,W), P=275.
42
43 (Problem 2 omitted)
44
45 Problem 3:
46 <story>
47
48 Constants:
49 <constants>
50
51 Predicates:
52 <predicates>
53
54 Constraints:

```

In the above prompt, `<story>` is a new puzzle, and `<constants>`, `<predicates>` are generated by GPT-3 for that story using **Prompt C** and **Prompt P** in Section 3.1 and 3.2.

Lines 1–8 are a general instruction describing the task of $\Pi_{define.test}$ generation and provides two rule forms for the

target ASP rules. The first rule form

$$C_1; C_2; \dots; C_m \leftarrow L_1, L_2, \dots, L_n$$

says that “ C_1 or ... or C_m is true if L_1 and ... and L_n are true.” Here, each L_i is a literal and each C_i is a comparison in the input language of CLINGO, e.g., $A > B$, $A = B + 3$, etc. The second rule form

$$\{C_1; C_2; \dots; C_m\} = k \leftarrow L_1, L_2, \dots, L_n$$

additionally restricts that “exactly k of $\{C_1, \dots, C_m\}$ must be true.” In principle, the first rule form is enough to represent various constraints. However, since the second rule form is syntactically closer to certain complex sentences related to cardinality, e.g., “either ... or ...”, “neither ... nor ...”, or “no ... is ...”, etc, we found that GPT-3 works much better when we also include the second rule form.

4 Optional Enhancements to the Pipeline

Section 3 presented a general pipeline that automatically writes an ASP program for a puzzle in natural language using LLM. This section explains two optional enhancements that strengthen its robustness.

4.1 Constant Formatting

In the Constant Extraction step (Section 3.1), GPT-3 may extract the names of the objects as they appear in the puzzle story, such as \$225, Sue Simpson, and 8:30 AM, which do not conform to the syntax of the input language of answer set solver CLINGO. Also, GPT-3 applies arithmetic computations (e.g., $L1=L2+3$) to constants surrounded by double quotes (e.g., $L2$ is constant “9 inches”) instead of constants that are integers (e.g., $L2$ is constant 9).

A rule-based post-processing could be applied to turn them into the right syntax, but alternatively, we employ GPT-3 to generate syntactically correct forms. We found that this method requires significantly less efforts and is more general because GPT-3 applies the constant formatting correctly even for unforeseen formats using some “common sense,” which is lacking in the rule-based approach. We use the following prompt for this.

The *Constant Formatting* step is done by calling GPT-3 with the following prompt, where `<constants>` at the end of the prompt is replaced by the original (extracted) constants c obtained by the Constant Extraction step (Section 3.1). The GPT-3 response in this step is the updated constants c , serving as an input to other steps in the pipeline.

```

1 Given categorized constants of the form "
   category: constant_1; constant_2; ...;
   constant_n", format the category and constants
   such that:
2 each category consists of only lowercase
   letters and underscores, and
3 each constant is either an integer or a string
   surrounded by double quotes, e.g., "United
   States".
4
5 There are two ways below to format constants
   and we must use the same way for all constants
   of the same category.

```



```

6 1. Turn all constants of the same category into
   integers with no space or special character.
7 2. Add double quotes around all constants of
   the same category.
8 Note that the 1st way has a higher priority,
   meaning that we must turn all constants of the
   same category into integers whenever possible
   . For example, twice or second can be turned
   into 2, September can be turned into 9,
   September 5th can be turned into 5 if all
   dates are in September, but 9:30am can only be
   turned into "9:30am" since no integer can
   represent 9:30am.
9
10 Original constants:
11 Employees: Bonita; Yvette; Tabitha.
12 Prices: $225; $275; $325.
13 Wood types: ash; poplar; sandalwood.
14
15 Formatted constants:
16 employee: "Bonita"; "Yvette"; "Tabitha".
17 price: 225; 275; 325.
18 wood_type: "ash"; "poplar"; "sandalwood".
19
20 Original constants:
21 months: January; April; October; December.
22 times: 8:30AM; 10:30AM; 2:30PM; 3:30PM.
23 durations: 1 day; 3 days; 11 days; 12 days.
24
25 Formatted constants:
26 month: 1; 4; 10; 12.
27 time: "8:30AM"; "10:30PM"; "2:30PM"; "3:30PM".
28 duration: 1; 3; 11; 12.
29
30 Original constants:
31 <constants>
32
33 Formatted constants:

```

4.2 Sentence Paraphrasing

Sometimes sentences may need to be paraphrased before an LLM can correctly generate rules from them. The *Sentence Paraphrasing* step provides the opportunity to not only simplify or formalize the sentences from the original question but also add the hidden information assumed to underlie the question. For example, the following sentence

```

1 Of the person who won the prize in
   bioengineering and Sue Simpson, one won in
   1976 and the other won in 1968.

```

is one clue in the example question in Section 3. The correct translation requires an LLM to turn the above sentence into at least 3 ASP rules, which would be hard for the current LLMs (e.g., GPT-3). Instead, we can ask GPT-3 to first paraphrase such kind of sentence into simpler ones below.

```

1 The person who won the prize in bioengineering
   and Sue Simpson are different.
2 The person who won the prize in bioengineering
   won in 1976 or won in 1968.
3 Sue Simpson won in 1976 or won in 1968.

```

The Sentence Paraphrasing step is done by calling GPT-3 with the following prompt, where `<sentences>` at the end of the prompt is replaced by the numbered sentences in the queried puzzle story q , and the GPT-3 response in text is used to replace the original sentences in q . This prompt is dedicated to the logic puzzles from Puzzle Baron and only paraphrases one kind of sentence in the form “of A and B, one is C and the other is D.”

```

1 Copy a sequence of numbered sentences.
2
3 If a sentence is of the form "N. Of A and B,
   one is C and the other is D", replace it with
   3 sentences below.
4 N.1 A and B are different.
5 N.2 A is C or D.
6 N.3 B is C or D.
7
8 For every sentence, if it is not of the form "N
   . Of ... and ...", simply copy it without
   replacement. An easy way to determine if a
   sentence is not of the above form is to check
   if its first word is not of.
9
10 In the following example, one sentence is of
   the above form.
11 Given:
12 1. The squad from Grenada ended with 2 silver
   medals.
13 2. Of the team from Oman and the team that won
   10 silver medals, one finished with 2 gold
   medals and the other finished with 1 gold
   medal.
14 Copy:
15 1. The squad from Grenada ended with 2 silver
   medals.
16 2.1 The team from Oman and the team that won 10
   silver medals are different.
17 2.2 The team from Oman finished with 2 gold
   medals or finished with 1 gold medal.
18 2.3 The team that won 10 silver medals finished
   with 2 gold medals or finished with 1 gold
   medal.
19
20 In the following example, no sentence is of the
   above form.
21 Given:
22 1. Tabitha's item costs 50 dollars less than
   the piece made of sandalwood.
23 2. The $275 item is either the piece made of
   ash or Yvette's item.
24 Copy:
25 1. Tabitha's item costs 50 dollars less than
   the piece made of sandalwood.
26 2. The $275 item is either the piece made of
   ash or Yvette's item.
27
28 Given:
29 <sentences>
30 Copy:

```

5 Experiments

We tested the above pipeline on the logic puzzles dataset from (Mitra and Baral 2015). Since the constants are pro-

Method	train set	test set
(Mitra and Baral 2015)	–	71%
Zero-shot GPT-3	0%	2%
Few-shot GPT-3	4%	3%
Zero-shot GPT-4	12%	21%
Few-shot GPT-4	6%	7%
GPT-3 Generated ASP Rules	86%	81%
GPT-4 Generated ASP Rules	92%	92%

Table 1: Accuracy on 50 train and 100 test puzzles. GPT-3 refers to the model named “text-davinci-003” in the OpenAI API, while GPT-4 is the model named “gpt-4.”

Step	Count	
	GPT-3	GPT-4
constant formatting	3	1
paraphrasing	2	4
constraint generation (syntax)	3	0
constraint generation (semantics)	13	3

Table 2: Mistakes on 100 test puzzles at different pipeline steps.

vided in the dataset as necessary information to solve each puzzle, we apply Constant Formatting directly on the given constants to generate constants c .

The dataset consists of 50 training examples and 100 testing examples. When designing our prompts, we only consult the training examples and not the testing examples. Table 1 shows the performance of our approach to zero-shot GPT-3/GPT-4, few-shot GPT-3/GPT-4, and a fully-supervised learning system LOGICIA (Mitra and Baral 2015).³ In the few-shot setting, we use the first two examples in the training set as the few-shot examples. GPT-3 with zero-shot and few-shot settings didn’t perform well, while zero-shot GPT-4 could solve 21% of the test puzzles correctly, which is significantly better than GPT-3’s performance. However, this is much lower than our method’s 81%. Interestingly, while the few-shot setting slightly improves over the zero-shot for GPT-3, this is quite different with GPT-4. This is likely because GPT-4 with the zero-shot setting was instructed to solve the puzzles in a step by step manner. However, for the few-shot setting, the examples only include the problem and solution, which may have discouraged GPT-4 from working through the puzzles in steps.

Besides the fact that the direct execution of the LLMs results in low performance, it is hard to understand why they fail to solve puzzles; in other words, the results are hard to interpret.

On the other hand, in our method, although the LLMs can still produce incorrect ASP programs, the errors are much less and it is easy to identify the failures by examining the generated ASP programs. Table 2 summarizes the number of mistakes we encountered on 100 test data at each pipeline step. Below we provide a representative example for each

³For GPT-3/GPT-4, to avoid randomness, we use a temperature of 0 (deterministic) and a top P value of 1 (default setting).

type of errors from Table 2 and highlight the mistakes in red.⁴ The examples are taken from GPT-3, although each type also occurs with GPT-4, except for the syntax error.

- **Failure to convert certain constants into integers during constant formatting.** The only mistake happens when, given constants “points: 181 points; 184 points; 187 points; 190 points”, GPT-3 generated

```
1 points: "181 points"; "184 points";
2 "187 points"; "190 points".
```

while it should generate constants in integers below to be able to apply arithmetic operations.

```
1 points: 181; 184; 187; 190.
```

- **Addition of wrong clues during paraphrasing.** One such error happens when paraphrasing the sentence “2. The book published in 2011 is either Orlando’s project or Jan’s project”. GPT-3 generated the following three sentences.

```
1 % 2.1 The book published in 2011 is either
   Orlando's project or Jan's project.
2 % 2.2 Orlando's project is the book
3   published in 2011 or the 200 page
4   project.
5 % 2.3 Jan's project is the book
6   published in 2011 or the 200 page
7   project.
```

The first sentence is simply a copy of the original sentence while the last two should not be there since they cannot be inferred from the original sentence.

- **Syntax error during constraint generation.** One typical syntax error is an incorrect variable name in a generated rule. For example, given the sentence “The study on the Liberty starts 2 months after Eunice’s project”, GPT-3 generates the following rule

```
1 M=M1+2 :- match(S,M,Sh), match(S1,M1,Sh1),
   S="Eunice", Sh="Liberty".
```

while the variable S in the second line should have been $S1$.

- **Semantic error during constraint generation.** One typical semantic error is caused by a wrong equation. For example, given the sentence “the \$35,000 structure is 15 sq ft smaller than the \$29,000 home”, GPT-3 generated

```
1 S1=S2+15 :- match(C1,S1,P1), match(C2,S2,P2)
   , P1=35000, P2=29000.
```

while the equation in the head should have been $S1=S2-15$.

While our pipeline doesn’t achieve 100% accuracy on generated ASP programs, most failed puzzles only have one

⁴We have provided a breakdown of the errors we encountered on Github at https://github.com/azreasoners/gpt-asp-rules/tree/main/error_analysis.

mistake and such a mistake is easy to correct. This means that our pipeline could serve as a good suggestion tool to prepare draft ASP programs for users. For example, compared to designing all the ASP programs for 50+100 puzzles manually, it would save a significant amount of time to only check the correctness of the automatically generated rules for the programs that don't have a single stable model.

6 More Examples

Previous approaches that automate logic puzzle solving either only predict constants and relations (Mitra and Baral 2015) or treat rule generation as a classification problem on a small set of rule templates (Jabrayilzade and Tekir 2020). In comparison, our method is generative, where rules are generated in an open-ended manner under the guidance of a few examples.

While it's hard to apply the previous methods to other domains without substantial changes, applying our pipeline to new domains requires only minor adjustments on the prompts. To apply our pipeline to other domains, we make a slight adjustment by turning the last sentence in Line 11 of **Prompt R2** into a numbered clue "0. No option in any category will ever be used more than once.", since it was specific to grid logic puzzles.

In the following part of this section, we show how our pipeline can be further applied to generate ASP programs for Sudoku and the Jobs Puzzle.

6.1 Sudoku

If we describe Sudoku problem with the following story q

- 1 The Sudoku problem is a classic computer science problem that involves filling a 9x9 grid with digits from 1 to 9.
- 2 1. The same number cannot appear twice in the same row.
- 3 2. The same number cannot appear twice in the same column.
- 4 3. The same number cannot appear twice in the 3 x3 box. (You can use division / to check if 2 cells are in the same box.)

our pipeline generates the following ASP program Π .

```

1 % Define the constants in each category.
2 index_of_row(1; 2; 3; 4; 5; 6; 7; 8; 9).
3 index_of_column(1; 2; 3; 4; 5; 6; 7; 8; 9).
4 number(1; 2; 3; 4; 5; 6; 7; 8; 9).
5 % For each row Ir and column Ic, it assigns
  exactly 1 number N.
6 {assign(Ir, Ic, N): number(N)}=1 :-
  index_of_row(Ir), index_of_column(Ic).
7
8 % 1. The same number cannot appear twice in the
  same row.
9 {N1=N2}=0 :- assign(Ir,Ic1,N1), assign(Ir,Ic2,
  N2), (Ic1,N1) != (Ic2,N2).
10
11 % 2. The same number cannot appear twice in the
  same column.
12 {N1=N2}=0 :- assign(Ir1,Ic,N1), assign(Ir2,Ic,
  N2), (Ir1,N1) != (Ir2,N2).
```

```

13
14 % 3. The same number cannot appear twice in the
  3x3 box.
15 {N1=N2}=0 :- assign(Ir1,Ic1,N1), assign(Ir2,Ic2,
  N2), (Ir1/3,Ic1/3) != (Ir2/3,Ic2/3), (Ir1,
  Ic1,N1) != (Ir2,Ic2,N2).
```

This ASP program Π is almost correct except that the red part in Line 16 of Π should be

```

1 ((Ir1-1)/3, (Ic1-1)/3) = ((Ir2-1)/3, (Ic2-1)/3)
```

since the row and column indices start from 1. This formula seems too difficult for GPT-3 to notice and generate unless some examples are provided. On the other hand, if we slightly adjust Lines 7–8 of **Prompt C** (Section 3.1) to make the indices start from 0, then the generated ASP program Π becomes correct as Lines 2–3 of Π are changed to the following facts.

```

1 index_of_row(0; 1; 2; 3; 4; 5; 6; 7; 8).
2 index_of_column(0; 1; 2; 3; 4; 5; 6; 7; 8).
```

GPT-4 also fails to generate the last rule correctly, although it makes a different mistake.

6.2 Jobs Puzzle

The Jobs Puzzle studied in (Schwitter 2013) asks one to assign 8 different jobs to 4 people while satisfying the given constraints. The full puzzle q is shown below.

- 1 1. There are four people: Roberta, Thelma, Steve, and Pete.
- 2 2. Among them, they hold eight different jobs.
- 3 3. Each holds exactly two jobs.
- 4 4. The jobs are: chef, guard, nurse, telephone operator, police officer (gender not implied), teacher, actor, and boxer.
- 5 5. The job of nurse is held by a male.
- 6 6. The husband of the chef is the telephone operator.
- 7 7. Roberta is not a boxer.
- 8 8. Pete has no education past the ninth grade.
- 9 9. Roberta, the chef, and the police officer went golfing together.
- 10 Question: Who holds which jobs?

This puzzle was considered a challenge for logical expressibility and automated reasoning (Shapiro 2011).

To apply our method to the Jobs Puzzle, some paraphrasing was needed before the *Define&Test* part of rule generation. We manually paraphrased the above puzzle to the following

- 1 There are four people: Roberta, Thelma, Steve, and Pete. Among them, they hold eight different jobs. Each holds exactly two jobs. The jobs are: chef, guard, nurse, telephone operator, police officer (gender not implied), teacher, actor, and boxer.
- 2 5. The job of nurse is held by a male.
- 3 6. The husband of the chef is the telephone operator, which means the chef is a female and the telephone operator is a male.
- 4 7. Roberta is not a boxer.


```

5  8. Pete has no education past the ninth grade.
   (Three jobs teacher, nurse, and police officer
   require higher education.)
6  9. Roberta, the chef, and the police officer
   went golfing together, meaning they are
   different people.
7  10.1 The same job cannot be assigned to
   different people.
8  10.2 Roberta is a female.
9  10.3 Thelma is a female.
10 10.4 Steve is a male.
11 10.5 Pete is a male.
12 10.6 The actor is a male.
13 Question: Who holds which jobs?

```

by turning clues 1–4 as background story, clarifying clues 6, 8, and 9, and adding a few hidden clues numbered 10.X at the end.

As for the prompts, we only need to update Line 1 of **Prompt R1** to the following to allow for $\{ \dots \} = k$ in a rule.

```

1  Given a problem as the background information,
   turn the provided categorized constants and
   predicates into rules.
2
3  Tasks:
4  1. Write a rule of the form "category(
   constant_1; constant_2; ...; constant_n)" to
   represent each row of categorized constants in
   the form "category: constant_1; constant_2;
   ...; constant_n".
5  2. Write a rule of the form "{predicate(X_1,
   X_2, ..., X_n): category_2(X_2), ...,
   category_n(X_n)}=k :- category_1(X_1)" for
   each predicate of the form "predicate(X_1, X_2,
   ..., X_n)".

```

Finally, GPT-3 generates the following ASP program:

```

1  % Task 1. Define the constants in each category
   .
2  person("Roberta"; "Thelma"; "Steve"; "Pete").
3  job("chef"; "guard"; "nurse"; "telephone
   operator"; "police officer"; "teacher"; "actor
   "; "boxer").
4  gender("male"; "female").
5  % Task 2. For each person P, there is exactly 2
   jobs J and gender G.
6  {assign(P, J, G): job(J), gender(G)}=2 :-
   person(P).
7  % Done tasks 1 and 2.
8  % 5. The job of nurse is held by a male.
9  G="male" :- assign(P, J, G), J="nurse".
10 % 6. The husband of the chef is the telephone
   operator, which means the chef is a female and
   the telephone operator is a male.
11 G1="female" :- assign(P1, J1, G1), assign(P2, J2,
   G2), J1="chef", J2="telephone operator".
12 G2="male" :- assign(P1, J1, G1), assign(P2, J2, G2)
   , J1="chef", J2="telephone operator".
13 % 7. Roberta is not a boxer.
14 J!="boxer" :- assign(P, J, G), P="Roberta".
15 % 8. Pete has no education past the ninth grade
   . (Three jobs teacher, nurse, and police
   officer require higher education.)

```

```

16 {J="teacher"; J="nurse"; J="police officer"}=0
   :- assign(P, J, G), P="Pete".
17 % 9. Roberta, the chef, and the police officer
   went golfing together, meaning they are
   different people.
18 P1!=P2 :- assign(P1, J1, G1), assign(P2, J2, G2),
   assign(P3, J3, G3), P1="Roberta", J2="chef", J3
   ="police officer".
19 P2!=P3 :- assign(P1, J1, G1), assign(P2, J2, G2),
   assign(P3, J3, G3), P1="Roberta", J2="chef", J3
   ="police officer".
20 P1!=P3 :- assign(P1, J1, G1), assign(P2, J2, G2),
   assign(P3, J3, G3), P1="Roberta", J2="chef", J3
   ="police officer".
21 % 10.1 The same job cannot be assigned to
   different people.
22 {P1=P2}=0 :- assign(P1, J, G1),
23      assign(P2, J, G2), (P1, G1) != (P2, G2) .
24 % 10.2 Roberta is a female.
25 G="female" :- assign(P, J, G), P="Roberta".
26 % 10.3 Thelma is a female.
27 G="female" :- assign(P, J, G), P="Thelma".
28 % 10.4 Steve is a male.
29 G="male" :- assign(P, J, G), P="Steve".
30 % 10.5 Pete is a male.
31 G="male" :- assign(P, J, G), P="Pete".
32 % 10.6 The actor is a male.
33 G="male" :- assign(P, J, G), J="actor".

```

which is almost correct with a single mistake in translating clue 10.1. If we just replace this constraint (in red) with

```

1  P1=P2 :- assign(P1, J, G1), assign(P2, J, G2) .

```

the corrected ASP program has exactly one stable model, which is the correct solution to the Jobs Puzzle.

Similarly, GPT-4 also failed to generate a completely correct ASP program. It also couldn't generate a correct rule for constraint 10.1, and furthermore failed to produce the gender category in constant extraction step **Prompt C**, missing "gender: "male"; "female"."

7 Conclusion

LLMs are a relatively recent technology that have shown to be disruptive. Despite their wide range of applications, their responses are not always reliable and cannot be trusted.

Automatic rule generation is a difficult problem. However, by using LLMs as a front-end to answer set programming, we can utilize their linguistic abilities to translate natural language descriptions into the declarative language of answer set programs. Unlike previous methods that use algorithmic or machine learning techniques, we find that a pre-trained large language model with a good prompt can generate reasonably accurate answer set programs. We present a pipeline with general steps that systematically build an ASP program in a natural way. This method not only leads to higher accuracy but also makes the results interpretable.

We expect this type of work to expand the application of KR methods that may appear unfamiliar to non-experts. We also anticipate that this pipeline will serve as a suggestion tool to help users prepare valid constants, useful predicates, or draft ASP programs.

Acknowledgements

We are grateful to the anonymous referees for their useful comments. This work was partially supported by the National Science Foundation under Grant IIS-2006747.

References

- Baral, C., and Dzifcak, J. 2012. Solving puzzles described in english by automated translation to answer set programming and learning how to do that translation. In *Proceedings of the Thirteenth International Conference on Principles of Knowledge Representation and Reasoning*, 573–577.
- Brewka, G.; Niemelä, I.; and Truszczyński, M. 2011. Answer set programming at a glance. *Communications of the ACM* 54(12):92–103.
- Brown, T.; Mann, B.; Ryder, N.; Subbiah, M.; Kaplan, J. D.; Dhariwal, P.; Neelakantan, A.; Shyam, P.; Sastry, G.; Askell, A.; et al. 2020. Language models are few-shot learners. *Advances in neural information processing systems* 33:1877–1901.
- Creswell, A.; Shanahan, M.; and Higgins, I. 2022. Selection-inference: Exploiting large language models for interpretable logical reasoning. *arXiv preprint arXiv:2205.09712*.
- Jabrayilzade, E., and Tekir, S. 2020. LGPSolver - solving logic grid puzzles automatically. In *Findings of the Association for Computational Linguistics: EMNLP 2020*, 1118–1123.
- Lev, I.; MacCartney, B.; Manning, C. D.; and Levy, R. 2004. Solving logic puzzles: From robust processing to precise semantics. In *Proceedings of the 2nd Workshop on Text Meaning and Interpretation*, 9–16.
- Lifschitz, V. 2008. What is answer set programming? In *Proceedings of the AAAI Conference on Artificial Intelligence*, 1594–1597. MIT Press.
- Milicevic, A.; Near, J. P.; and Singh, R. 2012. Puzzler: An automated logic puzzle solver. Technical report, Massachusetts Institute of Technology (MIT).
- Mitra, A., and Baral, C. 2015. Learning to automatically solve logic grid puzzles. In *Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing*, 1023–1033.
- Nordstrom, R. 2017. LogicSolver - Solving logic grid puzzles with part-of-speech tagging and first-order logic. Technical report, University of Colorado, Colorado Springs.
- Nye, M.; Tessler, M.; Tenenbaum, J.; and Lake, B. M. 2021. Improving coherence and consistency in neural sequence models with dual-system, neuro-symbolic reasoning. *Advances in Neural Information Processing Systems* 34:25192–25204.
- Rae, J. W.; Borgeaud, S.; Cai, T.; Millican, K.; Hoffmann, J.; Song, F.; Aslanides, J.; Henderson, S.; Ring, R.; Young, S.; et al. 2021. Scaling language models: Methods, analysis & insights from training gopher. *arXiv preprint arXiv:2112.11446*.
- Schwitter, R. 2010. Controlled natural languages for knowledge representation. In *Coling 2010: Posters*, 1113–1121.
- Schwitter, R. 2013. The jobs puzzle: Taking on the challenge via controlled natural language processing. *Theory and Practice of Logic Programming* 13(4-5):487–501.
- Shapiro, S. C. 2011. The jobs puzzle: A challenge for logical expressibility and automated reasoning. In *AAAI spring symposium: logical formalizations of commonsense reasoning*.
- Valmeekam, K.; Olmo, A.; Sreedharan, S.; and Kambhampati, S. 2022. Large language models still can’t plan (a benchmark for LLMs on planning and reasoning about change). In *NeurIPS 2022 Foundation Models for Decision Making Workshop*.
- Wei, J.; Wang, X.; Schuurmans, D.; Bosma, M.; brian ichter; Xia, F.; Chi, E. H.; Le, Q. V.; and Zhou, D. 2022. Chain of thought prompting elicits reasoning in large language models. In Oh, A. H.; Agarwal, A.; Belgrave, D.; and Cho, K., eds., *Advances in Neural Information Processing Systems*.
- Zhou, D.; Schärli, N.; Hou, L.; Wei, J.; Scales, N.; Wang, X.; Schuurmans, D.; Bousquet, O.; Le, Q.; and Chi, E. 2022. Least-to-most prompting enables complex reasoning in large language models. *arXiv preprint arXiv:2205.10625*.

A Prompts in the Pipeline

In this section, we list all prompts used in our pipeline.

Prompt C:

```
1  Given a problem, extract all different
   constants and their categories in the form "
   category: constant_1; constant_2; ...;
   constant_n". Here, the format of each constant
   is turned into either an integer or a string
   surrounded by double quotes, e.g., "some name
   ".
2
3  Problem 1:
4  Consider N-Queens Puzzle on a chessboard of
   size 8x8. The goal is to assign 8 queens on
   the chessboard so that no two queens can share
   the same row, column, or diagonal.
5
6  Constants:
7  index_of_row: 1; 2; 3; 4; 5; 6; 7; 8.
8  index_of_column: 1; 2; 3; 4; 5; 6; 7; 8.
9
10 Problem 2:
11 "Against the Grain" offers hand-made wooden
   furniture at reasonable prices. Each item is
   made by an in-house employee. Using only the
   clues that follow, match each item to the
   employee who crafted it, and determine its
   price and the type of wood used to make it.
   Remember, as with all grid-based logic puzzles
   , no option in any category will ever be used
   more than once.
12 1. Bonita's piece costs $325.
13 2. The item made of poplar costs more than
   Yvette's piece.
14 3. Tabitha's item costs 50 dollars less than
   the piece made of sandalwood.
15 4. The $275 item is either the piece made of
   ash or Yvette's item.
16
17 Constants:
18 employee: "Bonita"; "Yvette"; "Tabitha".
19 price: 225; 275; 325.
20 wood_type: "ash"; "poplar"; "sandalwood".
21
22 Problem 3:
23 <question: query>
24
25 Constants:
```

Prompt P:

```
1  Given a problem and some categorized constants
   of the form "category: constant_1; constant_2;
   ...; constant_n", generate the minimum number
   of predicates to define the relations among
   the categories of constants. Each generated
   predicate is of the form "predicate(X1, X2,
   ..., Xn)" where X1, X2, ..., Xn are different
   variables and each variable X belongs to one
   of the categories. For each category, there
   must exist at least one variable of some
   predicate that belongs to this category.
2
```

```
3  Problem 1:
4  Consider N-Queens Puzzle on a chessboard of
   size 8x8. The goal is to assign 8 queens on
   the chessboard so that no two queens can share
   the same row, column, or diagonal.
5
6  Constants:
7  index_of_row: 1; 2; 3; 4; 5; 6; 7; 8.
8  index_of_column: 1; 2; 3; 4; 5; 6; 7; 8.
9
10 Predicates:
11 % The categories in Constants include
   index_of_row and index_of_column. We use
   different variables Ir and Ic to represent
   index_of_row and index_of_column.
12 % We assign a queen at row Ir and column Ic,
   where Ir belongs to index_of_row and Ic
   belongs to index_of_column.
13 assign(Ir, Ic)
14
15 Problem 2:
16 "Against the Grain" offers hand-made wooden
   furniture at reasonable prices. Each item is
   made by an in-house employee. Using only the
   clues that follow, match each item to the
   employee who crafted it, and determine its
   price and the type of wood used to make it.
   Remember, as with all grid-based logic puzzles
   , no option in any category will ever be used
   more than once.
17 1. Bonita's piece costs $325.
18 2. The item made of poplar costs more than
   Yvette's piece.
19 3. Tabitha's item costs 50 dollars less than
   the piece made of sandalwood.
20 4. The $275 item is either the piece made of
   ash or Yvette's item.
21
22 Constants:
23 employee: "Bonita"; "Yvette"; "Tabitha".
24 price: 225; 275; 325.
25 wood_type: "ash"; "poplar"; "sandalwood".
26
27 Predicates:
28 % The categories in Constants include employee,
   price, and wood_type. We use different
   variables E, P, and W to represent employee,
   price, and wood_type.
29 % We match an employee E with price P and wood
   type W, where E belongs to employee, P belongs
   to price, and W belongs to wood_type.
30 match(E, P, W)
31
32 Problem 3:
33 <question: query>
34
35 <constants: query>
36
37 Predicates:
```

Prompt R1:

```
1  Given some categorized constants in the form "
   category: constant_1; constant_2; ...;
   constant_n" and some predicates about the
```

```

1 relation among different categories of
2 constants, write ASP (Answer Set Programming)
3 rules to generate the search space of possible
4 relations.
5
6 Constants:
7 number: 1; 2; 3; 4; 5; 6; 7; 8.
8
9 Predicates:
10 % The categories include number. Note that we
11 must use different variables in each predicate
12 .
13 % We assign a queen at row N1 and column N2,
14 where N1 belongs to number and N2 belongs to
15 number.
16 assign(N1, N2)
17
18 ASP Rules:
19 % Define the constants in each category.
20 number(1; 2; 3; 4; 5; 6; 7; 8).
21 % For each row N1, there is exactly 1 queen
22 assigned at some column N2.
23 {assign(N1, N2) : number(N2)}=1 :- number(N1).
24
25 Constants:
26 employee: "Bonita"; "Yvette"; "Tabitha".
27 price: 225; 275; 325.
28 wood_type: "ash"; "poplar"; "sandalwood".
29
30 Predicates:
31 % The categories include employee, price, and
32 wood_type. We use different variables E, P,
33 and W to represent employee, price, and
34 wood_type.
35 % We match an employee E with price P and wood
36 type W, where E belongs to employee, P belongs
37 to price, and W belongs to wood_type.
38 match(E, P, W)
39
40 ASP Rules:
41 % Define the constants in each category.
42 employee("Bonita"; "Yvette"; "Tabitha").
43 price(225; 275; 325).
44 wood_type("ash"; "poplar"; "sandalwood").
45 % For each employee E, it matches with exactly
46 1 price P and 1 wood type W.
47 {match(E, P, W) : price(P), wood_type(W)}=1 :-
48 employee(E).
49
50 <constants: query>
51
52 <predicates: query>
53
54 ASP rules:

```

Prompt R2:

```

1 Consider the constraint in the following form
2 <C1>; <C2>; ...; <Cm> :- <L1>, <L2>, ..., <Ln>.
3 which says that if the conjunction "<L1> and <
4 L2> and ... and <Ln>" is true, then the
5 disjunction of comparisons "<C1> or <C2> or
6 ... or <Cm>" must be true.
7
8 One can also add a restriction that "exactly k

```

```

of <C1>, <C2>, ..., <Cm> is true" by using the
following form
6 {<C1>; <C2>; ...; <Cm>}=k :- <L1>, <L2>, ..., <
7 Ln>.
8
9 Given a problem, extract all constraints from
10 the clues in the problem using only the
11 provided constants and predicates.
12
13 Problem 1:
14 "Against the Grain" offers hand-made wooden
15 furniture at reasonable prices. Each item is
16 made by an in-house employee. Using only the
17 clues that follow, match each item to the
18 employee who crafted it, and determine its
19 price and the type of wood used to make it.
20 Remember, as with all grid-based logic puzzles
21 , no option in any category will ever be used
22 more than once.
23
24 1. Bonita's piece costs $325.
25 2. The item made of poplar costs more than
26 Yvette's piece.
27 3. Tabitha's item costs 50 dollars less than
28 the piece made of sandalwood.
29 4. The $275 item is either the piece made of
30 ash or Yvette's item.
31
32 Constants:
33 employee: "Bonita"; "Yvette"; "Tabitha".
34 price: 225; 275; 325.
35 wood_type: "ash"; "poplar"; "sandalwood".
36
37 Predicates:
38 % The categories include employee, price, and
39 wood_type. We use different variables E, P,
40 and W to represent employee, price, and
41 wood_type.
42 % We match an employee E with price P and wood
43 type W, where E belongs to employee, P belongs
44 to price, and W belongs to wood_type.
45 match(E, P, W)
46
47 Constraints:
48 % No option in any category will ever be used
49 more than once.
50 {E1=E2; P1=P2; W1=W2}=0 :- match(E1,P1,W1),
51 match(E2,P2,W2), (E1,P1,W1) != (E2,P2,W2).
52
53 % 1. Bonita's piece costs $325.
54 P=325 :- match(E,P,W), E="Bonita".
55
56 % 2. The item made of poplar costs more than
57 Yvette's piece.
58 P1>P2 :- match(E1,P1,W1), match(E2,P2,W2), W1="
59 poplar", E2="Yvette".
60
61 % 3. Tabitha's item costs 50 dollars less than
62 the piece made of sandalwood.
63 P1=P2-50 :- match(E1,P1,W1), match(E2,P2,W2),
64 E1="Tabitha", W2="sandalwood".
65
66 % 4. The $275 item is either the piece made of
67 ash or Yvette's item.
68 {W="ash"; E="Yvette"}=1 :- match(E,P,W), P=275.
69
70

```

```

43 Problem 2:
44 The Winter Olympics have just wrapped up in
   Norway. Using only the clues that follow,
   determine the number of gold, silver and
   bronze medals won by each country. Remember,
   as with all grid-based logic puzzles, no
   option in any category will ever be used more
   than once.
45 1. The four teams are the team from Bolivia,
   the team that won 3 gold medals, the team that
   won 6 silver medals, and the team from
   Argentina.
46 2. The team from Oman and the team that won 10
   silver medals are different.
47 3. The team from Oman finished with 2 gold
   medals or finished with 1 gold medal.
48 5. The squad that won 3 gold medals, the squad
   that won 6 silver medals and the squad from
   Bolivia were all different teams.
49 6. Neither the team from Argentina nor the
   squad that won 2 silver medals is the squad
   that won 4 gold medals.
50 8. The squad that won 2 gold medals is either
   the squad that won 6 silver medals or the team
   from Grenada.
51
52 Constants:
53 country: "Argentina"; "Bolivia"; "Grenada"; "
   Oman".
54 silver_medals: 2; 6; 10; 11.
55 gold_medals: 1; 2; 3; 4.
56
57 Predicates:
58 % The categories include country, silver_medals
   , and gold_medals. We use different variables
   C, S, and G to represent country,
   silver_medals, and gold_medals.
59 % We assign a country C with silver medals S
   and gold medals G, where C belongs to country,
   S belongs to silver_medals, and G belongs to
   gold_medals.
60 assign(C, S, G)
61
62 Constraints:
63 % No option in any category will ever be used
   more than once.
64 {C1=C2; S1=S2; G1=G2}=0 :- assign(C1,S1,G1),
   assign(C2,S2,G2), (C1,S1,G1) != (C2,S2,G2).
65
66 % 1. The four teams are the team from Bolivia,
   the team that won 3 gold medals, the team that
   won 6 silver medals, and the team from
   Argentina.
67 {C="Bolivia"; G=3; S=6; C="Argentina"}=1 :-
   assign(C,S,G).
68
69 % 2. The team from Oman and the team that won
   10 silver medals are different.
70 C1!=C2 :- assign(C1,S1,G1), assign(C2,S2,G2),
   C1="Oman", S2=10.
71
72 % 3. The team from Oman finished with 2 gold
   medals or finished with 1 gold medal.
73 {G=2; G=1}=1 :- assign(C,S,G), C="Oman".
74

```

```

75 % 5. The squad that won 3 gold medals, the
   squad that won 6 silver medals and the squad
   from Bolivia were all different teams.
76 C1!=C2 :- assign(C1,S1,G1), assign(C2,S2,G2),
   assign(C3,S3,G3), G1=3, S2=6, C3="Bolivia".
77 C2!=C3 :- assign(C1,S1,G1), assign(C2,S2,G2),
   assign(C3,S3,G3), G1=3, S2=6, C3="Bolivia".
78 C1!=C3 :- assign(C1,S1,G1), assign(C2,S2,G2),
   assign(C3,S3,G3), G1=3, S2=6, C3="Bolivia".
79
80 % 6. Neither the team from Argentina nor the
   squad that won 2 silver medals is the squad
   that won 4 gold medals.
81 {C="Argentina"; S=2}=0 :- assign(C,S,G), G=4.
82 C1!=C2 :- assign(C1,S1,G1), assign(C2,S2,G2),
   C1="Argentina", S2=2.
83
84 % 8. The squad that won 2 gold medals is either
   the squad that won 6 silver medals or the
   team from Grenada.
85 {S=6; C="Grenada"}=1 :- assign(C,S,G), G=2.
86
87 Problem 3:
88 <question: query>
89
90 <constants: query>
91
92 <predicates: query>
93
94 Constraints:

```

In Prompt R2, since the maximum number of tokens (including the tokens in prompt and response) allowed by GPT-3 is 4096, we deleted 2 clues (numbered 4 and 7) in the question description of the second logic puzzle (lines 44-50) and their corresponding ASP rules (lines 63-85).

B A Few Example Answer Set Programs Generated by GPT-3

B.1 Puzzle 1 in test set

Given the question q below,

- 1 The local foodie club met at Chez Martin last night for their monthly meal. Match each person to their choice of wine and entree, and determine how much each owed at the end of the night. Remember, as with all grid-based logic puzzles, no option in any category will ever be used more than once.
- 2 1. The person who had the port paid 1 dollar more than Kurt.
- 3 2. Of the person who paid \$25 and the person who paid \$24, one was Priscilla and the other had the shiraz.
- 4 3. Of the person who paid \$27 and Priscilla, one had the chianti and the other had the port.
- 5 4. Isabel paid \$25.

our pipeline generates the following ASP program

```

1 % Define the constants in each category.
2 wine("chianti"; "port"; "riesling"; "shiraz").
3 price(24; 25; 26; 27).

```



```

4 name("Isabel"; "Kurt"; "Priscilla"; "Robin").
5 % For each person N, it matches with exactly 1
  wine W and 1 price P.
6 {match(W, P, N): wine(W), price(P)}=1 :- name(N
  ).
7
8 % No option in any category will ever be used
  more than once.
9 {W1=W2; P1=P2; N1=N2}=0 :- match(W1,P1,N1),
  match(W2,P2,N2), (W1,P1,N1) != (W2,P2,N2) .
10
11 % 1. The person who had the port paid 1 dollar
  more than Kurt.
12 P1=P2+1 :- match(W1,P1,N1), match(W2,P2,N2), W1
  ="port", N2="Kurt".
13
14 % 2.1 The person who paid $25 and the person
  who paid $24 are different.
15 N1!=N2 :- match(W1,P1,N1), match(W2,P2,N2), P1
  =25, P2=24.
16
17 % 2.2 The person who paid $25 was Priscilla or
  had the shiraz.
18 {N="Priscilla"; W="shiraz"}=1 :- match(W,P,N),
  P=25.
19
20 % 2.3 The person who paid $24 was Priscilla or
  had the shiraz.
21 {N="Priscilla"; W="shiraz"}=1 :- match(W,P,N),
  P=24.
22
23 % 3.1 The person who paid $27 and Priscilla are
  different.
24 N1!=N2 :- match(W1,P1,N1), match(W2,P2,N2), P1
  =27, N2="Priscilla".
25
26 % 3.2 The person who paid $27 had the chianti
  or had the port.
27 {W="chianti"; W="port"}=1 :- match(W,P,N), P
  =27.
28
29 % 3.3 Priscilla had the chianti or had the port
  .
30 {W="chianti"; W="port"}=1 :- match(W,P,N), N="
  Priscilla".
31
32 % 4. Isabel paid $25.
33 P=25 :- match(W,P,N), N="Isabel".

```

which has exactly 1 stable model and this stable model matches with the solution.

```

1 match("chianti",24,"Priscilla")
2 match("shiraz",25,"Isabel")
3 match("riesling",26,"Kurt")
4 match("port",27,"Robin")

```

B.2 Puzzle 2 in test set

Given the question q below,

```

1 A group of friends has decided to try several
  different weight-loss diets and exercises to
  see who amongst them can lose the most weight
  in 3 months. Using only the clues below, match
  each person to their total number of pounds

```

```

lost, and determine the diet type they used
and their preferred type of exercise. Remember
, as with all grid-based logic puzzles, no
option in any category will ever be used more
than once.
2 1. Celia used the gluten-free diet.
3 2. The friend who lost 3 pounds used the low-
  fat diet.
4 3. Mandy lost 2 fewer pounds than Raymond.
5 4. Mandy lost 4 more pounds than Tom.
6 5. The dieter who used the vegan diet is either
  Mandy or the friend who lost 3 pounds.

```

our pipeline generates the following ASP program

```

1 % Define the constants in each category.
2 pounds_lost(3; 5; 7; 9).
3 diet("dairy-free"; "gluten-free"; "low-fat"; "
  vegan").
4 name("Celia"; "Mandy"; "Raymond"; "Tom").
5 % For each name N, it matches with exactly 1
  pounds lost P1 and 1 diet D.
6 {match(N, P1, D): pounds_lost(P1), diet(D)}=1
  :- name(N) .
7
8 % No option in any category will ever be used
  more than once.
9 {N1=N2; P11=P12; D1=D2}=0 :- match(N1,P11,D1),
  match(N2,P12,D2), (N1,P11,D1) != (N2,P12,D2) .
10
11 % 1. Celia used the gluten-free diet.
12 D="gluten-free" :- match(N,P1,D), N="Celia".
13
14 % 2. The friend who lost 3 pounds used the low-
  fat diet.
15 P1=3 :- match(N,P1,D), D="low-fat".
16
17 % 3. Mandy lost 2 fewer pounds than Raymond.
18 P11=P12-2 :- match(N1,P11,D1), match(N2,P12,D2)
  , N1="Mandy", N2="Raymond".
19
20 % 4. Mandy lost 4 more pounds than Tom.
21 P11=P12+4 :- match(N1,P11,D1), match(N2,P12,D2)
  , N1="Mandy", N2="Tom".
22
23 % 5. The dieter who used the vegan diet is
  either Mandy or the friend who lost 3 pounds.
24 {N="Mandy"; P1=3}=1 :- match(N,P1,D), D="vegan
  ".

```

which has exactly 1 stable model and this stable model matches with the solution.

```

1 match("Tom",3,"low-fat")
2 match("Celia",5,"gluten-free")
3 match("Mandy",7,"vegan")
4 match("Raymond",9,"dairy-free")

```

B.3 Variants of Sudoku

Continuing the process in Section 6, we can generate the ASP program for variants of Sudoku by adding 1 or 2 more clues to the puzzle description q . Below are the newly generated constraints for the added clues in each variant.

Anti-Knight Sudoku

```

1 % 4. The same number cannot appear twice in a
  knight move.
2 {N1=N2}=0 :- assign(Ir1,Ic1,N1), assign(Ir2,Ic2
  ,N2), |Ir1-Ir2|+|Ic1-Ic2|=3, (Ir1,Ic1,N1)!=(
  Ir2,Ic2,N2) .

```

Sudoku-X

```

1 % 4. The same number cannot appear twice among
  the cells whose row index is equal to column
  index.
2 {N1=N2}=0 :- assign(Ir1,Ic1,N1), assign(Ir2,Ic2
  ,N2), Ir1=Ic1, Ir2=Ic2, (Ir1,Ic1,N1)!=(Ir2,Ic2
  ,N2) .
3
4 % 5. The same number cannot appear twice among
  the cells whose row and column indices sum up
  to 8.
5 {N1=N2}=0 :- assign(Ir1,Ic1,N1), assign(Ir2,Ic2
  ,N2), Ir1+Ic1=8, Ir2+Ic2=8, (Ir1,Ic1,N1)!=(Ir2
  ,Ic2,N2) .

```

Offset Sudoku

```

1 % 4. The same number cannot appear twice among
  the cells with the same relative position in
  3*3 boxes.
2 {N1=N2}=0 :- assign(Ir1,Ic1,N1), assign(Ir2,Ic2
  ,N2), Ir1\3=Ir2\3, Ic1\3=Ic2\3, (Ir1,Ic1,N1)
  !=(Ir2,Ic2,N2) .

```

C Additional GPT-4 Analysis

C.1 Representative Example Errors (for GPT-4)

- **Failure to convert certain constants into integers during constant formatting.** Given the constants: “times: 8:00am; 9:00am; 10:00am; 11:00am”, GPT-4 generated

```

1 times("8:00am"; "9:00am"; "10:00am";
2 "11:00am") .

```

instead of the correct generation

```

1 times(8;9;10;11) .

```

- **Addition of wrong clues during paraphrasing.** Given the clue “2. The conductor working on June 12 is either the conductor departing from Buttonwillow or Greg.”, GPT-4 generates the three sentences

```

1 2.1 The conductor working on June 12
  and
2 Greg are different.
3 2.2 The conductor working on June 12 is
  either the conductor departing from
  Buttonwillow or Greg.
4 Greg is either the conductor departing
5 from Buttonwillow or the conductor
6 working on June 12

```

The second sentence is a copy of the original, while 2.1 and 2.3 cannot be inferred and are therefore wrong.

- **Semantic error during constraint generation.** The sentence “Vasquez will leave sometime after Macdonald.” is parsed by GPT-4 into

```

1 M1<M2 :- schedule(D1,M1,Du1), schedule(D2,M2
  ,Du2), D1="Vasquez", D2="Macdonald" .

```

which is incorrect, the less than sign should be changed to greater than:

```

1 M1>M2 :- schedule(D1,M1,Du1), schedule(D2,M2,
  Du2), D1="Vasquez", D2="Macdonald" .

```

There are no syntax errors encountered with GPT-4.

C.2 Error Subtypes

We further break down the paraphrasing error into two types, (p1) a sentence representing an exclusive disjunction is incorrectly translated into additional sentences. For example, “3. The card with an APR of 11% is either the card with the \$4,000 credit limit or the one with the \$20,000 credit limit.”, is incorrectly translated into

```

1 3.1 The card with an APR of 11% and the card
  with the $4,000 credit limit are different.
2 3.2 The card with an APR of 11% is either the
  card with the $4,000 credit limit or the one
  with the $20,000 credit limit.
3 3.3 The card with the $4,000 credit limit is
  either the card with an APR of 11% or the one
  with the $20,000 credit limit.

```

and (p2) a sentence representing that four things are different is incorrectly translated into two incorrect sentences. For example, the statement “5. The four people are Deep Shadow, the superhero who started in 2007, the hero who started in 2009 and Matt Minkle.”, is incorrectly translated into

```

1 5.1 Deep Shadow, the superhero who started in
  2007, the hero who started in 2009, and Matt
  Minkle are different.
2 5.2 Deep Shadow is either the superhero who
  started in 2007 or the hero who started in
  2009 or Matt Minkle.
3 5.3 The superhero who started in 2007 is either
  Deep Shadow or the hero who started in 2009
  or Matt Minkle.
4 5.4 The hero who started in 2009 is either Deep
  Shadow or the superhero who started in 2007
  or Matt Minkle.

```

Constraint Generation (semantic errors) is further broken into four subtypes. The first (c1) has to do with an incorrect comparison between times. For example, the statement “Tricia came in a half-hour after Ora.” is incorrectly translated into

```

1 % 2. Tricia came in a half-hour after Ora.
2 T1="7:00am", T2="7:30am" :- order(C1,B1,T1),
  order(C2,B2,T2), C1="Ora", C2="Tricia" .
3 T1="7:30am", T2="8:00am" :- order(C1,B1,T1),
  order(C2,B2,T2), C1="Ora", C2="Tricia" .
4 T1="8:00am", T2="8:30am" :- order(C1,B1,T1),
  order(C2,B2,T2), C1="Ora", C2="Tricia" .

```

The second (c2) is when an incorrect operator is used (e.g. “+” in place of “-”). For example, the statement “% 1. Vasquez will leave sometime after Macdonald.” is incorrectly translated into

```
1 M1<M2 :- schedule(D1,M1,Du1), schedule(D2,M2,
    Du2), D1="Vasquez", D2="Macdonald".
```

the third (c3) is a disjunction in the head of a rule which should not be there. For example, the statement “% 3. The 11-year-old bird has a wingspan 8 inches shorter than Charlie.” is incorrectly translated into

```
1 A=11, W=W1-8 :- assign(N,W,A), assign(N1,W1,A1)
    , N="Charlie".
```

and last (c4) belongs to semantic errors which do not fit into any of the previous types and only occur once.

Error	Subtype	Count	
		GPT-3	GPT-4
Constant Formatting		3	1
Paraphrasing		2	3
	p1	1	3
	p2	1	1
Cons. Gen. (syntax)		3	0
Cons. Gen. (semantics)		13	4
	c1	3	2
	c2	4	1
	c3	3	0
	c4	3	0

Table 3: Errors on 100 test puzzles on GPT-3 and GPT-4, broken down by subtype.

Table 3 shows the count for the errors and error subtypes encountered for GPT-3 and GPT-4. We find that unlike GPT-3, GPT-4 does not make any syntax errors, however it increases its rate of paraphrasing errors.

C.3 Sudoku and Jobs Puzzle

Sudoku GPT-4 correctly generates all rules except for the the constraints here:

```
1 % 1. The same number cannot appear twice in the
    same row.
2 {Irl=Ir2; Icl=Ic2; N1=N2}=0 :- assign(Irl,
    Icl,N1), assign(Ir2,Ic2,N2), (Irl,Icl,N1)!=(
    Ir2,Ic2,N2), Irl=Ir2.
3 % 2. The same number cannot appear twice in the
    same column.
4 {Irl=Ir2; Icl=Ic2; N1=N2}=0 :- assign(Irl,
    Icl,N1), assign(Ir2,Ic2,N2), (Irl,Icl,N1)!=(
    Ir2,Ic2,N2), Icl=Ic2.
5 % 3. The same number cannot appear twice in the
    3x3 box.
6 {Irl=Ir2;
    Icl=Ic2; N1=N2}=0 :- assign(Irl,Icl,N1),
    assign(Ir2,Ic2,N2), (Irl,Icl,N1)!=(Ir2,Ic2,N2)
    , (Irl-1)/3=(Ir2-1)/3, (Icl-1)/3=(Ic2-1)/3.
```

The red portion should not be included, otherwise the program runs correctly. Note that GPT-4 generates the correct (Ir1-1), (Ir2-2), (Icl-1), and (Ic2-1) terms while GPT-3 does not.

Jobs Puzzle In the constant extraction step, GPT-4 fails to generate the gender category. From the problem

```
1 Problem 3:
2 1. There are four people: Roberta, Thelma,
    Steve, and Pete.
3 2. Among them, they hold eight different jobs.
4 3. Each holds exactly two jobs.
5 4. The jobs are: chef, guard, nurse, telephone
    operator, police officer (gender not
6 implied), teacher, actor, and boxer.
7 5. The job of nurse is held by a male.
8 6. The husband of the chef is the telephone
    operator.
9 7. Roberta is not a boxer.
10 8. Pete has no education past the ninth grade.
11 9. Roberta, the chef, and the police officer
    went golfing together.
12 Question: Who holds which jobs?
```

It produces:

```
1 person: "Roberta"; "Thelma"; "Steve"; "Pete".
2 job: "chef"; "guard"; "nurse"; "telephone
    operator"; "police officer"; "teacher"; "actor
    "; "boxer"
```

and is missing “gender: "male"; "female".”

Also, on prompt **Prompt R1**, GPT-4 produces the correct output but then continues to produces constraints that it should not:

```
1 % Task 1. Define the constants in each category
    .
2 person("Roberta"; "Thelma"; "Steve"; "Pete").
3 job("chef"; "guard"; "nurse"; "telephone
    operator"; "police officer"; "teacher"; "actor
    "; "boxer").
4 gender("male"; "female").
5 % Task 2. For each person P, they hold exactly
    2 jobs J and have a gender G.
6 {assign(P, J, G): job(J), gender(G)}=2 :-
    person(P).
7 % Additional rules based on the problem
8 constraints:
9 % 5. The job of nurse is held by a
    male.
10 assign(P, "nurse", "male") :-
    person(P).
11 % 6. The husband of the chef is the
    telephone operator.
12 assign(P1, "chef", "female"),
13 assign(P2,
    "telephone operator", "male") :-
14 person(P1), person(P2).
15 % 7. Roberta is not a boxer.
16 :- assign("Roberta", "boxer", _).
17 % 8. Pete has no education past the
    ninth
18 grade.
19 :- assign("Pete", "teacher", _).
20 % 9. Roberta, the chef, and the police
    officer went golfing together.
```

```

23  assign("Roberta", -, -), assign(-,
    "chef", -),
24  assign(-, "police officer", - ) :-
25  person().

```

Since these constraints not supposed to be generated yet, they are not appropriately prompted and hence GPT-4 produces incorrect constraints. Apart from that, GPT-4 produces a similar incorrect rule on rule 10.1 that GPT-3 produces.

```

1  % 10.1 The same job cannot be assigned to
    different people.
2  {P1=P2}=0 :- assign(P1,J1,G1), assign(P2,J2,G2)
    , J1=J2, P1!=P2.

```

which should be

```

1  P1=P2 :- assign(P1,J,G1), assign(P2,J,G2) .

```