

---

# TESTING AND UNDERSTANDING ERRONEOUS PLANNING IN LLM AGENTS THROUGH SYNTHESIZED USER INPUTS

---

Zhenlan Ji, Daoyuan Wu\*, Pingchuan Ma, Zongjie Li, Shuai Wang\*

The Hong Kong University of Science and Technology

Hong Kong SAR, China

{zjiae, daoyuan, pmaab, zligo, shuaiw}@cse.ust.hk

## ABSTRACT

Agents based on large language models (LLMs) have demonstrated effectiveness in solving a wide range of tasks by integrating LLMs with key modules such as planning, memory, and tool usage. Increasingly, customers are adopting LLM agents across a variety of commercial applications critical to reliability, including support for mental well-being, chemical synthesis, and software development. Nevertheless, our observations and daily use of LLM agents indicate that they are prone to making erroneous plans, especially when the tasks are complex and require long-term planning.

In this paper, we propose PDoctor, a novel and automated approach to testing LLM agents and understanding their erroneous planning. As the first work in this direction, we formulate the detection of erroneous planning as a constraint satisfiability problem: *an LLM agent’s plan is considered erroneous if its execution violates the constraints derived from the user inputs*. To this end, PDoctor first defines a domain-specific language (DSL) for user queries and synthesizes varying inputs with the assistance of the Z3 constraint solver. These synthesized inputs are natural language paragraphs that specify the requirements for completing a series of tasks. Then, PDoctor derives constraints from these requirements to form a testing oracle. PDoctor features several design considerations, such as mock tool and input mutation, to enhance testing effectiveness. Its synthesized inputs can also incorporate advanced features like dynamic constraint update to better test the LLM agent’s planning ability. We evaluate PDoctor with three mainstream agent frameworks and two powerful LLMs (GPT-3.5 and GPT-4). The results show that PDoctor can effectively detect diverse errors in agent planning, and provide insights and error characteristics that are valuable to both agent developers (for improving LLM agents) and users (for using contemporary agents). We conclude by discussing potential alternative designs and directions to extend PDoctor.

## 1 Introduction

Large language models (LLMs) have become extremely popular due to their remarkable performance across a wide range of tasks [22, 12, 26, 44, 21, 55, 47], demonstrating an ability to *understand*, *reason*, and *act* in ways akin to human cognition and reasoning. With their extensive parameters and training data, these models have shown proficiency in complex pattern recognition in natural language, leading to advancements in logical reasoning [16, 56], vulnerability detection [47, 46], robot planning [11, 29], and causal inference [31, 24]. This has fueled the development of LLM agents [59, 43, 14, 40], which are designed to interact with the external world and make decisions to achieve complex objectives, leveraging LLMs’ cognitive capabilities for tasks that require planning, memory, and execution of a sequence of actions. These agents, integrating key modules alongside LLMs, excel in processing complex user queries, learning from past interactions, and adapting to new tasks.

To date, the industry has been increasingly commercializing LLM agents in various highly profitable and even mission-critical applications [59, 5, 7, 9, 8, 2, 33], such as healthcare, personal assistance, and financial services. However, despite their great potential and overall enthusiasm, LLM agents often struggle with *erroneous planning*, leading to significant consequences. For example, an LLM agent used to manage a chemical synthesis process may fail to

---

\*Corresponding authors.

produce the desired chemical compound if it makes an erroneous plan, and the involved, possibly expensive, chemical reagents are already consumed during the process. This challenge, highlighted by our observations and experiences, underscores the need for improvements in their design and operation. Errors in planning, particularly in complex, long-term tasks, can result in the misuse of resources or failure to achieve intended outcomes, underscoring the importance of developing more reliable and effective LLM agents.

In this paper, we propose PDoctor, a novel framework for testing and understanding erroneous planning in LLM agents. Our approach is fully automated and can be used to detect erroneous planning in LLM agents using different core LLMs and following different paradigms (see introduction in Sec. 2). As the first work in this direction, PDoctor formulates the occurrence of “erroneous planning” as a constraint satisfiability problem: *an LLM agent’s plan is erroneous if its execution violates the constraints derived from the user inputs*. The constraints can be rigorously checked with moderate cost, thus providing a reliable way to detect erroneous planning.

PDoctor defines a domain specific language (DSL) that captures the semantics of user queries and features a synthesis procedure (with the assistance of Z3) that can generate diverse user requests as the test inputs to the LLM agent. We provide configurable parameters to control the diversity and complexity of the user queries and offer a mutation procedure to further transform each generated user query. Each user query denotes a paragraph that outlines the requirements for conducting a series of tasks. PDoctor accordingly derives constraints from these requirements and checks if the LLM agent’s plan aligns with the satisfiability of these constraints; misalignment flags erroneous planning. We provide a set of design considerations and optimizations (e.g., mock tools employed by the agent) to deliver effective testing, and also augment the synthesized inputs with advanced features like dynamic constraint update to test the LLM agent’s planning capability.

We evaluate PDoctor on three mainstream LLM agent frameworks, ReAct [59], OpenAI Tools (OT) [5], and OpenAI Assistant (OA) [7]. These LLM agents represent different paradigms and are widely used in various applications. We incorporate these frameworks with two widely-used LLM models, GPT-3.5 and GPT-4 [10]. PDoctor can effectively detect thousands of erroneous plans across all these settings. We configure PDoctor to depict the planning capability “upper bound” of different LLM agents and also summarize the detected erroneous plans. These results can provide insights into the common pitfalls of LLM agents and offer guidance for developers and users in daily practice. We also discuss the extension of PDoctor and the potential future work to repair the detected erroneous plans. In summary, the contributions of this paper are as follows:

- We pioneer the effort to test and understand erroneous planning in LLM agents, given their increasing commercialization in reliability-sensitive fields and the potential severe consequences of erroneous planning.
- We formulate the problem of detecting erroneous planning as a constraint satisfiability problem and present a fully automated framework, PDoctor, that employs input synthesis and constraint checking to detect erroneous planning.
- We evaluate PDoctor on mainstream LLM agents using different paradigms and show that our approach can effectively detect thousands of erroneous planning with moderate cost. We also summarize insights from different aspects to benefit developers and users.

**Tool Availability.** We have made PDoctor available at <https://anonymous.4open.science/r/PDoctor-E872> for review purposes. We will continue to maintain the tool and add more documentation to assist users in utilizing it.

## 2 Preliminary

### 2.1 Planning Problem

Planning is a fundamental problem-solving task that involves generating a sequence of actions to achieve a goal [34, 18]. Extensive efforts have been devoted to this area, achieving significant progress in various domains, such as robotics [36, 13, 23] and autonomous vehicles [15, 32]. Formally, the planning problem can be defined as follows:

**Definition 1** (Planning Problem). *Given a set of states  $\mathcal{S}$ , a set of actions  $\mathcal{A}$ , an and a state transition function  $f : \mathcal{S} \times \mathcal{A} \rightarrow \mathcal{S}$ , the planning problem  $\mathcal{P}$  is defined as a tuple  $\langle \mathcal{S}, \mathcal{A}, f, s_0, s_g \rangle$ , where  $s_0 \in \mathcal{S}$  is the initial state and  $s_g \in \mathcal{S}$  is the goal state. The goal of  $\mathcal{P}$  is to find a sequence of actions  $\langle a_1, a_2, \dots, a_n \rangle$  called a plan such that  $f(f(\dots f(s_0, a_1), a_2), \dots, a_n) = s_g$ .*

In the context of LLM agents, the user query can be viewed as a planning problem  $\mathcal{P} = \langle \mathcal{S}, \mathcal{A}, f, s_0, s_g \rangle$ . Likewise, the invocation of a specific tool can be viewed as taking an action  $a$ , where  $a \in \mathcal{A}$ . The goal of LLM agents is to generate a correct answer, i.e., a sequence of actions  $\langle a_1, a_2, \dots, a_n \rangle$ , to satisfy the constraint set  $\mathcal{C}$  implied by the state transition

function  $f$ . The initial state  $s_0$  represents that the agent begins to address the user query, and the goal state  $s_g$  denotes the state in which tasks specified in the user query are completed, i.e., all constraints in  $\mathcal{C}$  are satisfied.

## 2.2 An Example of LLM Agents

| (a) User input                                                                                                                                                    | (c) Agent's action chain                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>User:</b><br>Tell me the current speed of the network. If the speed is zero, help me fix the problem.                                                          | <b>Action:</b><br><span style="color: green;">Invoking: `network_speed_test` with `{}`</span><br><span style="color: blue;">&gt; Speed is 0Mbps.</span><br><span style="color: green;">Invoking: `network_status_check` with `{}`</span><br><span style="color: blue;">&gt; Network is down. Error code: 1.</span><br><span style="color: green;">Invoking: `network_diagnose` with `{'error_code': 1}`</span><br><span style="color: blue;">&gt; DHCP server is down. Need to restart the server.</span><br><span style="color: green;">Invoking: `dhcp_service_restart` with `{}`</span><br><span style="color: blue;">&gt; Done.</span><br><span style="color: green;">Finished.</span> |
| <b>(b) Instruction list</b><br><br><b>Tools:</b><br>network_status_check()<br>network_speed_test()<br>network_diagnose(error_code: int)<br>dhcp_service_restart() |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |

Figure 1: An example illustrating an LLM agent working on a complex user query. Green and blue represent the agent's actions and the execution results of invoked tools, respectively.

Fig. 1 presents a typical example of an LLM agent working on a sample user query (Fig. 1(a)). The agent first *understands* the requirements implied by the user query, then plans (*reasons*) to select the appropriate tools (Fig. 1(b)) to invoke, and finally *acts* to execute each selected tool and collect its response (Fig. 1(c)). In particular, after receiving the output from each invoked tool, the agent repeats the above process until the user query is fully answered. This series of interactions between the agent and the environment, as listed in Fig. 1(c), collectively constitutes a problem-solving procedure and then determines the final response (omitted here) to the user query. LLM agents benefit from powerful LLMs that can behave in a human-like manner to understand, reason, and act simultaneously, and the synergy of these three capabilities.

Nevertheless, given the susceptibility of LLMs that has been exposed in prior research [38, 53, 58, 27], it is not surprising that LLM agents are substantially prone to errors and failures. Intuitively, the agent's action chain can be viewed as an "error amplifier," where a small mistake in the early stage of the action chain could be continuously amplified and propagated in each subsequent step, leading to catastrophic failures in the end. This character further exacerbates the difficulty of providing a correct and stable response to the user query. Moreover, the interaction between the LLM agent and the environment is often complex and dynamic, distinct from the static and isolated settings in which LLMs's abilities have been tested [49, 26, 22, 46].

In general, previous studies that endeavor to evaluate LLMs are often limited to a straightforward linear flow following a "text in and text out" paradigm. In this paradigm, problems are presented in their entirety, and LLMs are expected to generate a textual response based on the given problem. In contrast, LLM agent testing necessitates a focus on the interaction and interplay between the LLM and the environment, where the problems are dynamic and the LLM agent is required to adjust its plan in response to the environment's feedback. Here, the feedback refers to the results of the invoked tools, as shown in the blue text of Fig. 1(c). In line with these dynamic and complex settings, we design PDoctor to dynamically update the constraint sets (see details in Sec. 4.4). This underscores the importance of tool design, in addition to the textual problems, during test case generation. Consequently, it is crucial to develop a systematic approach tailored to the unique characteristics of LLM agents to test and assess their performance in a thorough manner.

## 2.3 Architecture of LLM Agents

With a typical example of LLM agents illustrated in Sec.2.2, we now delve into the architecture of an LLM agent to better understand the mechanisms behind the agent's operation, which is valuable for designing the testing framework.

**Modules.** Fig. 2 presents a typical architecture of an LLM agent. This agent consists of three main components: (1) an understand module, (2) a plan module, and (3) an act (tool usage) module. The understand module is tasked with comprehending the user query and some historical records of previous actions (i.e., in the "memory" component) to extract a concrete task to be accomplished. The plan module is responsible for: (1) decomposing the extracted task into a sequence of sub-tasks that can be accomplished by invoking various tools, (2) managing the state and resources

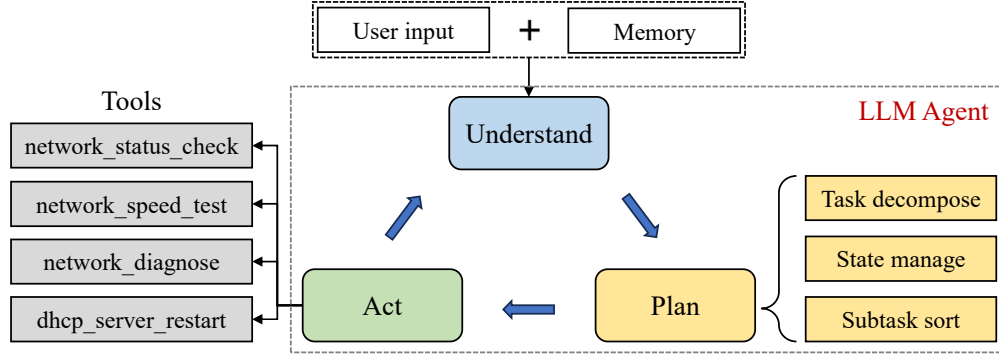


Figure 2: A typical architecture of LLM agents, according to [54].

to support an appropriate plan (e.g., retaining the error code until requested by `network_diagnose`, as shown in Fig. 1), and (3) determining the execution order of the sub-tasks. The act module then invokes the tools according to the generated plan and collects the results of the tool invocations.

**LLM and Prompts.** All the components mentioned above are implemented by an LLM, which serves as the “brain” of the agent. The agent offers prompt templates composing the query context, how the agent should respond, and the accessible tools, whereas the users provide the query content. The formed prompt will be used to guide the LLM to generate the response. For example, ReAct [59] categorizes the core LLM’s response into three types: *observation*, *thought*, and *action*, which correspond to the *understand*, *plan*, and *act* modules in Fig. 2, respectively. Additionally, prompt engineering strategies like role-play are often employed to enhance the agent’s performance under specific scenarios, such as through the `instruction` parameter in OpenAI’s Assistant API [7].

## 2.4 Our Testing Focus

As shown in Fig. 2, different components jointly contribute to the agent’s overall performance. Nevertheless, this paper considers testing the plan module a top priority. This is because extensive efforts [25, 19, 62, 46, 51, 28] have been devoted to evaluating the understand module, which rely on LLMs’ basic natural language understanding capability. Similarly, the act module has also made promising progress, as evidenced by OpenAI’s recent work in incorporating function calling functionality into GPT models [6]. Moreover, Yue et al. [20] have conducted a thorough investigation into LLM agents’ tool usage performance and proposed a benchmark, MetaTool, to evaluate the act module. In contrast, the plan module, which determines how to concretize the agent’s thought into a sequence of actions that can be executed by the agent, has not been systematically studied or tested in existing research. To sum up, we have the following focuses:

- *(Isolating the test of the planning module only)* While maintaining the complexity of the planning problem, we strive to simplify the user queries handled by the understand and act modules (see Fig. 2). By doing this, we can guarantee that any detected agent failure can be solely attributed to an error in the planning module, thereby achieving a simulated “isolation” of the planning module. This approach facilitates more precise testing of LLM agents’ planning ability.
- *(Covering both textual queries and underlying tools)* Unlike previous works [50, 49] that tested the planning ability of LLMs by generating and evaluating textual plans, the planning module in LLM agents is required to decide the order of tool invocation, making the testing of the planning module more complex. We aim to generate specialized test cases that comprise both textual queries and corresponding tools.
- *(Using only valid user queries for testing)* For individual queries, we focus on generating valid user queries to test LLM agents. That said, we are *not* using extreme user queries to stress LLM agents. Based on our observation, extreme queries, such as overly long or complex ones, as well as those with broken or confusing contents, may hinder the LLM agent from generating meaningful outputs. This is not surprising since LLM agents rely on LLMs to generate outputs and are, therefore, sensitive to the quality of the input prompts. Sticking to valid user queries can help us better understand the planning module’s performance in a real-world setting.

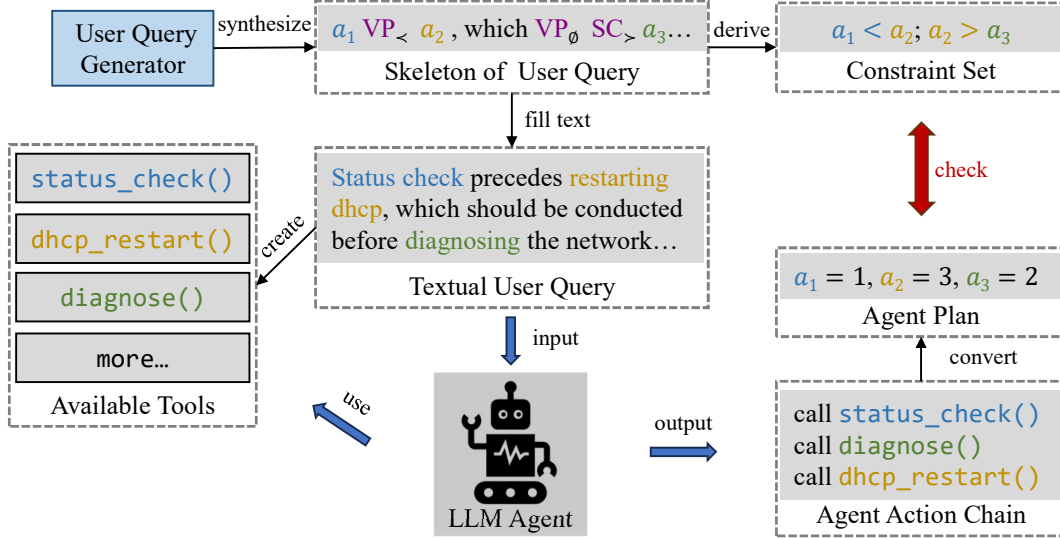


Figure 3: An overall workflow of PDoctor, which synthesizes textual user queries with a list of tools and derives the constraints as the oracle to test the LLM agent.

### 3 Overview

With the testing focus described in Sec. 2.4, we now introduce PDoctor, a novel system for testing and understanding erroneous planning in LLM agents. The basic idea of PDoctor is to synthesize textual user queries, and based on the simultaneously derived formal constraints, we can identify the erroneous planning. Based on this idea, we present the overall design of PDoctor as shown in Fig. 3. At a high level, PDoctor generates test cases comprising textual user queries that represent planning problems, along with a series of tools for invocation by the LLM agent. The testing oracle, automatically derived from the synthesized queries, is used to check the correctness of the LLM agent’s planning. This oracle also facilitates the dissection and characterization of erroneous planning when it occurs. Specifically, PDoctor has the following key components:

① **User Query Generation.** PDoctor performs a natural language synthesis process to generate textual user queries, which are considered as a planning problem  $\mathcal{P}$ , as formulated in Sec. 2.1, for the LLM agent. This process comprises two steps: *skeleton synthesis* and *text filling*. The former step generates a skeleton of the user query, which is a high-level, abstract representation of the planning problem. This enables PDoctor directly to derive the constraints for subsequent testing. The latter step fills in this skeleton with text to form a complete user query that describes the requirements for problem-solving. The design of this synthesis process is motivated by two key considerations: (1) the derived constraints should be equivalent to the semantics of the generated queries, since these constraints are used as the oracle to test the LLM agent; (2) the synthesis process should be flexible and extensible, allowing users to specify the complexity and scenarios of the generated queries. For the first consideration, alternative methods for generating user queries, such as text mutation and LLM-based text generation, are further discussed in Sec. 7.

② **Test Results Check.** After deriving formal constraints from the synthesis process, PDoctor utilizes these constraints to detect erroneous planning of the LLM agent. Specifically, PDoctor maps each action, which corresponds to subtasks specified by the user query in a one-to-one manner, to a specific tool invocation and collects the action chain (i.e., the planning made by the LLM agent) by recording the history of tool invocations. The collected planning is then compared against the constraints to determine whether the planning is correct.

③ **Error Dissection.** In contrast to prior studies that modified inputs based on superficial natural language properties (e.g., word-level replacement) [48, 30], PDoctor synthesizes complete textual user queries from scratch, gaining full control over both the semantics (i.e., the derived constraints) and the structure (i.e., the skeleton) of the user query. This design allows PDoctor to conduct comprehensive mutations and transformations on the generated query. By maintaining the overall semantic meaning, PDoctor is capable of altering words, sub-tasks, the context of the prompt, multiple sentences, and even the entire structure of the prompt. This capability is crucial for dissecting erroneous planning in LLM agents, as it enables PDoctor to pinpoint the root cause of failures by comparing the planning results of the original and mutated prompts.

**Challenges and Key Solutions.** To achieve the design above, however, we need to address some unique challenges:

- C1: Designing a mechanism or specification to support user query synthesis.** Without a mechanism or specification to guide their formulation, the synthesized user queries could become highly varied and uncontrollable. Moreover, to easily derive the corresponding constraints from the synthesized queries, we also need the support of a formal specification. To address this problem, we propose a specialized DSL in Sec. 4.1. With this DSL, PDoctor is able to cover a diverse set of user queries with configurable complexity and diversity. Moreover, DSL-based skeleton synthesis ensures the full control over the structure and semantics of the user queries, enabling PDoctor to exhaustively conduct mutation and transformation on a given prompt, thereby allowing the dissection of erroneous planning in Sec. 4.5.
- C2: Generating syntactically valid and semantically meaningful user queries.** As mentioned in Sec. 2.4, we aim to generate valid user queries to test LLM agents. That said, PDoctor should focus on generating *syntactically valid* and *semantically meaningful* user queries as testing inputs. To address this challenge, the user query skeletons are specified using the aforementioned DSL, which prevents the generation of broken content; see details in Sec. 4.2. Moreover, a constraint solver (Z3 [17] in our case) is used to ensure that the semantics underlying the user queries are meaningful, i.e., the constraints derived from the generated prompt are always satisfiable, as to be illustrated in Sec. 4.3.
- C3: Covering complex planning problems with state management and dynamic problem-solving.** While the original design of PDoctor above already covers planning tests for both textual queries and corresponding tools, it has not yet considered complex planning problems involving state management and dynamic problem-solving. We address this challenge by introducing time and duration constraints, thereby obtaining an extended testing framework for PDoctor, which will be presented in Sec. 4.4.

## 4 Design

### 4.1 Domain-Specific Language

To support user query synthesis, we have designed a specialized DSL tailored for LLM agent planning problems. This DSL is not intended to cover the infinitely vast semantic space of natural language for describing anything conceivable concept. Instead, it narrows its focus exclusively to the semantic space required for the planning problems defined in Sec. 2.1. Therefore, before delving into the specifics of the DSL, we first present how we simplify the semantic space to be explored, which could also avoid ambiguity due to English not being context-free.

**Semantic Simplification.** For two given tasks that cannot be executed in parallel, denoted as task  $A$  and task  $B$ , the planning problem can be boiled down to identify whether  $A < B$  (task  $A$  comes before task  $B$ ), or  $A > B$  (task  $A$  comes after task  $B$ ). Accordingly, constraints can be reduced to a specific description of the execution order of the tasks. For example, “ $A$  should be executed before  $B$ ” is a simplified constraint in this context. Moreover, it is straightforward to generalize this simplification to multiple tasks. Thus, PDoctor simplifies prompt synthesis into generating several sentences, each of which is composed of this kind of simplified constraints.

In addition, each “task” in the user query is simplified to take an action  $a$ . Here  $a \in \mathcal{A}$ , and  $\mathcal{A}$  is the set of actions that the LLM agent can perform (see Sec. 2.1). Take Fig. 3 as an example, conducting network status check is a task and invoking `network_status_check()` is the action. Without necessitating additional decomposition, every task can be completed directly through one single action. This simplification is based on the observation that task decomposition is mainly determined by the context of the task, rather than the performance of the LLM agent itself. For example, LLM agents that are familiar with the network would be able to correctly decompose the task of “fix disconnected network” into a series of sub-tasks like “check network status” and “network diagnosis”, while others that are designed for addressing book management would not. Nonetheless, this task decomposition failure can be easily mitigated by providing the LLM agent with instructive knowledge of network repair. Due to the potential unanticipated impact that task decomposition could have on the performance of the LLM agent, PDoctor avoids generating complex tasks that require decomposition.

**DSL Specifics.** Fig. 4 presents the syntax of this domain-specific language. Overall, each user query is represented as a paragraph  $\mathcal{P}$ . Every task that the LLM agent is expected to perform is detailed in this paragraph; each task requires the LLM agent to take a specific action  $a$ . Paragraphs are composed of one or more sentences  $S$ . And each sentence  $S$  is a sequence of sub-sentences  $s$  separated by a clause conjunction  $j$ . A sub-sentence can be regarded as an independent unit of semantic meaning, containing one or multiple constraints to be obeyed by the LLM agent. In particular, a sub-sentence  $s$ , no matter if it is an independent clause  $c_i$ , or a multi-clause  $m$ , comprises a subject  $s$ , an object  $o$ , and one or more keywords in purple that directly indicate a sequential relationship. These keywords are denoted in a special way, where the capital letter stands for their part-of-speech (POS) tag, and the subscript denotes the sequential relationship. Specifically, **VP** denotes a verb phrase, **P** denotes a preposition, and **SC** denotes a sequential



---

**Syntax**

$$\begin{aligned}
\langle \mathcal{P} \in \text{Planning Problem} \rangle &::= S^+ \\
\langle S \in \text{Sentence} \rangle &::= s \mid S \ j \ s \\
\langle s \in \text{Sub-sentence} \rangle &::= c_i \mid m \\
\langle c_i \in \text{Independent Clause} \rangle &::= s \ (\mathbf{VP}_{<} \mid \mathbf{VP}_{>}) \ o \\
&\quad \mid s \ \mathbf{VP}_{\emptyset} \ (\mathbf{P}_{<} \mid \mathbf{P}_{>}) \ o \\
&\quad \mid (\mathbf{P}_{<} \mid \mathbf{P}_{>}) \ o \ \text{“,”} \ s \ \mathbf{VP}_{\emptyset} \\
\langle m \in \text{Multi-clauses} \rangle &::= c_d \ (\mathbf{SC}_{<} \mid \mathbf{SC}_{>}) \ c'_d \\
&\quad \mid (\mathbf{SC}_{<} \mid \mathbf{SC}_{>}) \ c'_d \ \text{“,”} \ c_d \\
\langle c_d \in \text{Dependent Clause} \rangle &::= s \ \mathbf{VP}_{\emptyset} \\
\langle c_r \in \text{Relative Clause} \rangle &::= \text{“which”} \ (\mathbf{VP}_{<} \mid \mathbf{VP}_{>}) \ o \\
&\quad \mid \text{“which”} \ \mathbf{VP}_{\emptyset} \ (\mathbf{P}_{<} \mid \mathbf{P}_{>}) \ o \\
\langle s \in \text{Subject} \rangle &::= a^+ \\
&\quad \mid a^+ \ \text{“,”} \ c_r \ \text{“,”} \\
\langle o \in \text{Object} \rangle &::= a^+ \\
&\quad \mid a^+ \ \text{“,”} \ c_r \ \text{“,”} \\
\langle j \in \text{Conjunction} \rangle &::= \text{“,”} \\
&\quad \mid \text{“,”} \ (\text{“and”} \mid \text{“but”} \mid \text{“yet”}) \\
&\quad \mid \text{“,”} \ (\text{“while”} \mid \text{“whereas”}) \\
\langle a \in \text{Action} \rangle &::= \text{the set of actions, } \mathcal{A}
\end{aligned}$$

---

Figure 4: The syntax of our DSL, specifically designed for synthesizing user queries and deriving their constraints.

conjunction. Here, “sequential conjunction” is a special type of conjunction that indicates the temporal relationship between two dependent clauses. For the subscript, we use  $<$  to stand for the case where the events in the subject clause happen before the events in the object clause, while  $>$  denotes the opposite.  $\emptyset$  is an exceptional case, as it is merely used to decorate the verb phrase  $\mathbf{VP}$ , indicating that the verb phrase is not associated with any temporal relationship, like “happen” or “occur”.

## 4.2 User Query Synthesis

The DSL designed in Sec. 4.1 moderates natural language complexity while preserving the expressiveness and variety of the user queries it generates. In general, synthesizing a user query is to gradually expand an abstract syntax tree (AST) where each node in the AST is randomly selected from a set of valid nodes according to the grammar of the DSL. That said, the synthesis process is conducted in a top-down manner, with the root node being the paragraph  $P$ , and the leaf nodes being a sequence of non-terminal symbols in the DSL. Without accounting for the expansion of  $a^+$ , the DSL offers a total of 340 possible expanding options, ensuring the diversity of the synthesized user queries.

Based on this carefully designed DSL, the synthesis of user queries mainly comprises three steps: synthesizing the skeleton, translating the skeleton into NL (natural language) prompts, and simultaneously deriving the constraints from the synthesized skeleton, as illustrated in Fig. 3.

**Skeleton Synthesis.** Alg. 1 presents the algorithmic procedure for prompt skeleton synthesis. The algorithm takes the action set  $\mathcal{A}$ , the maximum number of sentences  $N$ , and the maximum iteration  $K$  as inputs, and outputs the generated paragraph  $P$  and its corresponding constraint set  $\mathcal{C}$ . In general, for each sentence generation, the algorithm iteratively generates a sentence by randomly expanding the sentence  $S$  according to the DSL and then substituting non-terminals with natural language words (line 6). For action symbols in the skeleton, they are substituted with actions from  $\mathcal{A}$ . Constraints of this sentence are then derived and used to extend the constraint set  $\mathcal{C}$  (line 7). If the extended constraint set  $\mathcal{C}'$  is satisfiable, the sentence will be added to paragraph  $P$ , and the constraint set  $\mathcal{C}$  will get updated (lines 8–11). Otherwise, the algorithm will try to generate another sentence by iterating the process until the maximum iteration  $K$  is reached, at which point the generation process will be terminated since it indicates that finding a satisfiable sentence is infeasible (lines 12–17). It is worth noting that this algorithm ensures the generated user query must be satisfiable,

**Algorithm 1:** Skeleton Synthesis.**Input:** Action Set  $A = \{a_1, a_2, \dots, a_n\}$ , Max Sentence Number  $N$ , Max Iteration  $K$ **Output:** Generated Paragraph  $P$ , Constraint Set  $C$ 


---

```

1  $P \leftarrow \emptyset$ 
2  $C \leftarrow \text{initConstraints}(A)$  // Initialize constraints.
3 for  $n \in \{1, \dots, N\}$  do
    /* Iteratively generate sentences to form a paragraph. */
4      $\text{iter} \leftarrow 0$ 
5     while  $\text{iter} < K$  do
6          $s \leftarrow \text{genSentence}(A, C)$ 
7          $C' \leftarrow C \cup \text{getConstraints}(s)$ 
8         if  $\text{SATSolver}(C') == \text{SAT}$  then
            /* Verify the satisfiability of the constraints after adding the new sentence. */
9              $P \leftarrow P \cup s$ 
10             $C \leftarrow C'$ 
11            break
12        end
13         $\text{iter} \leftarrow \text{iter} + 1$ 
14    end
15    if  $\text{iter} \geq K$  then
16        break
17    end
18 end
19 return  $P, C$ 

```

---

as the constraints are checked after each sentence generation. This way, the synthesized user query can effectively be used for testing the LLM agent’s planning ability.

**Text Filling.** To “translate” the skeleton into natural language user queries, PDoctor fills text into the slots marked by the terminal symbols (e.g., various **keywords** and actions) in the skeleton. Except for action symbols, there are seven alternative natural language words on average for each terminal symbol. For example, “happen”, “occur”, “be executed”, etc., for the verb phrase  $VP_\emptyset$ , and “after”, “behind”, “later than”, etc., for the preposition  $P_\succ$ . For the action symbols, we substitute them with the daily activities of various jobs. More specifically, we first instruct the LLM to provide a list of common jobs (e.g., “teacher”, “software developer”, etc.) by querying the LLM with the prompt “Please provide a list of 50 typical jobs. Ensure that these 50 positions span a variety of industries.”. For each provided job, we further query the LLM with the prompt “Please list 20 activities in noun phrase format that a [role] may need to do in a day.”, where “[role]” is replaced with the job name. Action symbols in the same paragraph are substituted with daily activities from the same job, and the job becomes the *Topic* of the user query. By changing the topic, PDoctor is capable of smoothly altering the context of the user queries, which is beneficial for the error dissection in Sec. 4.5. The design of text filling facilitates the diversity and realism of the synthesized user queries, in contrast to the previous works that are limited to a few fixed, artificial scenarios or templates [49, 42].

**Deriving Constraints.** Simultaneously, PDoctor derives the constraints from each synthesized user query. Specifically, each action  $a$  is represented as an integer variable  $\mathbf{a}$ , denoting the execution order of the action. Then, the subsequent relationship between two actions, denoted as  $a_1$  and  $a_2$ , can be expressed in a comparative form, like  $\mathbf{a}_1 < \mathbf{a}_2$  or  $\mathbf{a}_1 > \mathbf{a}_2$ . The former indicates that action  $a_1$  should be executed before action  $a_2$ , while the latter denotes the opposite. Therefore, considering a sub-sentence  $s$ , “network diagnosis comes after network status check” with the skeleton  $a_1 VP_\succ a_2$ , the constraint derived from this sub-sentence is  $\mathbf{a}_1 > \mathbf{a}_2$ .

### 4.3 Test Results Check

Once we feed the synthesized user query to the LLM agent, it will generate a textual response that describes the planning of the tasks. We can check the correctness of the planning by comparing the action sequence extracted from the response to the constraints derived from the user query. It is worth noting that some may argue for employing metrics like the BLEU score [37] and BertScore [64] to assess the similarity between the generated response and the ground truth. However, this is less feasible in our context because (i) these metrics may be biased or expensive to compute, and (ii) the effectiveness of planning testing should be assessed based on whether the action sequence satisfies the constraints, rather than on the similarity between the generated response and the ground truth. Our approach of



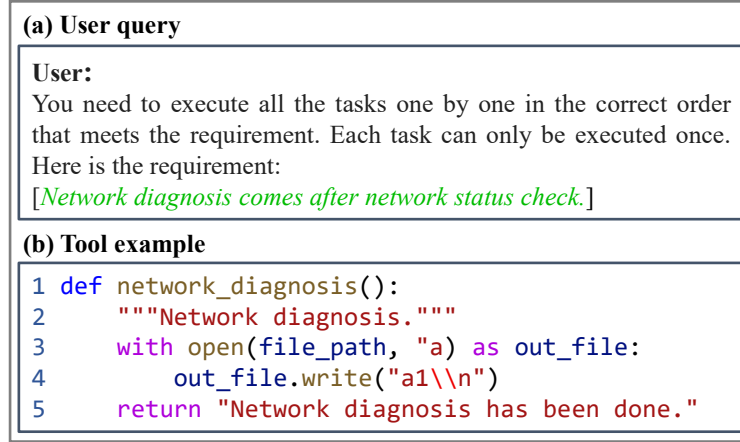


Figure 5: An example illustrating our “mock test” approach to extract the action sequence of LLM agents. *Green italic* denotes the synthesized part, which describes the constraints for conducting a series of tasks.

comparing the action sequence to the constraints, instead, is a more direct and effective way to validate the LLM agent’s planning.

A new challenge arises, however, as extracting the action sequence from the LLM agent’s textual response is non-trivial, requiring a precise understanding of the LLM response meaning. Existing work typically instructs the LLM to respond in a structural format through the adoption of few-shot prompts [29, 49]. We argue, nevertheless, that this practice is still not ideal, as LLMs may fail to comply with the format requirements. For example, the LLM may generate a response that is not structured as expected, or the response may contain irrelevant information that complicates the extraction process.

**Mock Tool Design.** PDoctor conducts a “mock test” approach to extracting the action sequence of LLM agents. Specifically, for each action  $a$  in the synthesized user query, PDoctor creates a “simulated” tool, with the tool’s name and description designed to be consistent with the textual content of the action (e.g., “network diagnosis” in Fig. 5). The internal logic of the tool is designated as a simple file-write module that directly writes the ID of the corresponding action  $a$  into a log file. This design is based on the fact that LLM agents treat tools in a black-box manner: they select tools based on the tool’s name and description, rather than the internal implementation. Therefore, the literal task described by  $a$  is not conducted in the environment, and the tool merely returns a string that deceives the LLM agent into believing that the action has been completed. Fig. 5 illustrates an example of this design. For the synthesized user query presented in Fig. 5(a), PDoctor creates a set of tools corresponding to the actions in the user query, where Fig. 5(b) depicts one of the tools. The `network_diagnosis` does not perform an actual network diagnosis in the environment. Instead, it merely writes the tool’s ID, “a1”, into the log file and deceives the LLM agent into believing that the network has been diagnosed.

**Result Checking.** After the agent finishes processing the synthesized user query, the log file is read to collect the action sequence. Take the case in Fig. 3 as an example, the action sequence is “[ $a_1, a_2, a_3$ ]”. Then, for each action  $a$  in the action sequence, the corresponding variable  $\mathbf{a}$  is assigned with the execution order of the action, i.e.,  $\mathbf{a}_1 = 1$ ,  $\mathbf{a}_2 = 3$ , and  $\mathbf{a}_3 = 2$ . Afterward, a plan is identified as erroneous if there exists any constraint unsatisfied by the assigned values of the variables. Still referring to the example in Fig. 3, the assigned values of variables  $\mathbf{a}_1$ ,  $\mathbf{a}_2$ , and  $\mathbf{a}_3$  are checked against the logic (and  $\wedge$ ) of the constraint set  $\{\mathbf{a}_1 < \mathbf{a}_2, \mathbf{a}_2 > \mathbf{a}_3\}$ . Since none of the constraints (and their  $\wedge$ ) is unsatisfied, this planning is identified as correct. Overall, we view the integration of this check design with the mock tools as a novel approach to rigorously examine the planning of LLM agents while avoiding the complexity of understanding LLMs’ textual response.

#### 4.4 Extended Testing Framework

Although the aforementioned test framework is effective in evaluating the planning ability of LLM agents, real-world scenarios would be more complex and challenging. Referring to the example in Fig. 1, the agent is expected to dynamically adjust its planning based on the execution results of the tools. Moreover, some tools may require several specific inputs, like `error_code` in this example, which requires the LLM agent to retain the state information across different tasks. Motivated by these observations, we extend the test framework by introducing time and tool duration

## Syntax

$$\langle o \in \text{Object} \rangle ::= t$$

$$\langle t \in \text{Time} \rangle ::= \text{the set of time point, } \mathcal{T}$$

Figure 6: DSL extension for covering time and duration constraints.

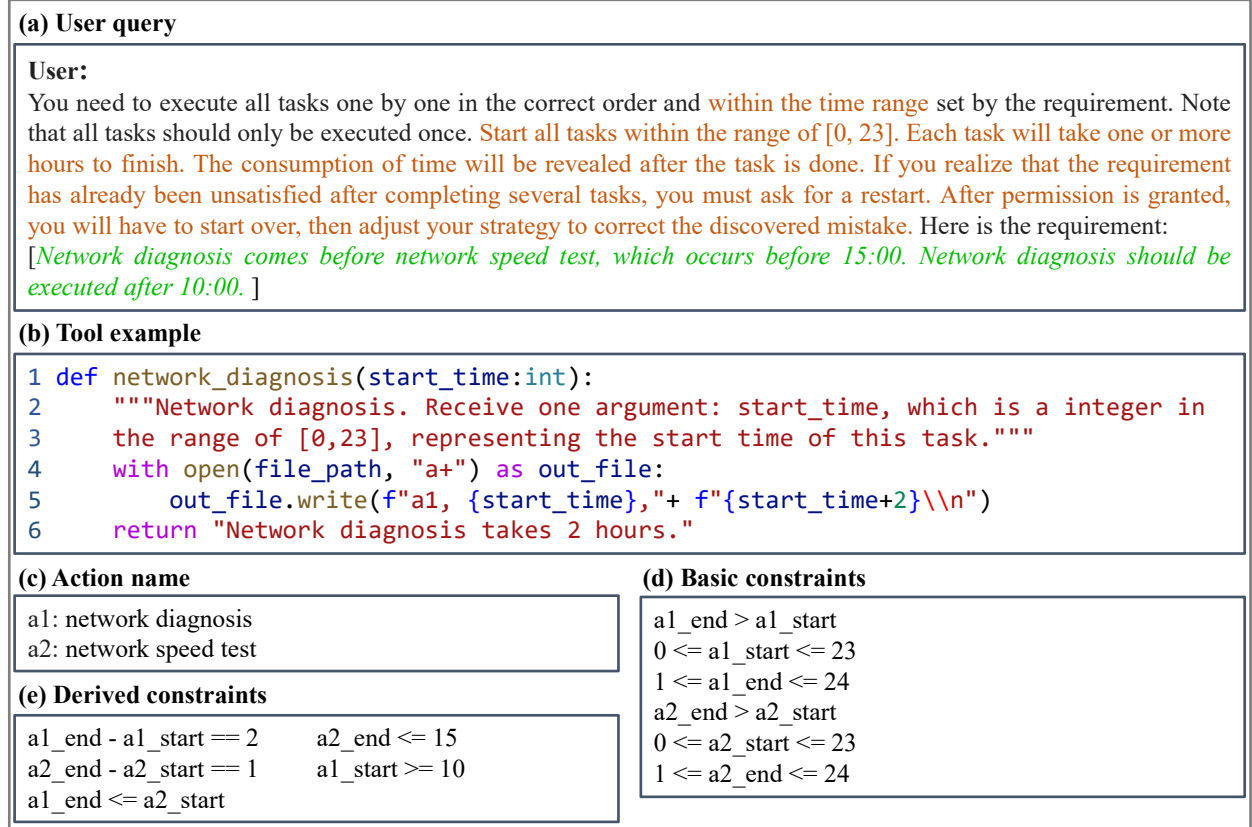


Figure 7: An example illustrating the extended version of the generated test case, including the user query and corresponding tools. Similar to Fig. 5, *green italic* denotes the synthesized part. Newly added instructions are highlighted in **brown**. Sub-figure (c) shows the mapping between the action ID and the action name. (d) and (e) present the basic constraints and the additional constraints derived from the user query, respectively.

concepts into the agent planning. For one thing, time constraints are widely used in real-world planning, such as daily schedules. For another, the introduction of time allows PDoctor to simulate the dynamic planning over global resources, which is a representative complex planning problem in real-world scenarios. In particular, the synthesized user query is extended to include new constraints that regulate the start or end time of given tasks as illustrated in Fig. 6. Accordingly, the mock tools are modified to require the start time point as the input parameter and return the duration of their execution. Fig. 7 illustrates an example of the extended version of the generated test case.

From this figure, an obvious change is the template of the user query. Several new instructions are added, to require the LLM agent to consider the time constraints. Besides, the change in the return value of the mock tool, which is presented in Fig. 7(b), requires the agent to dynamically adjust its planning. Due to the absence of knowledge about the tool execution time, the agent may encounter a situation where the planning is negated by the execution result of tools. Taking the case in Fig. 7 as an example, if the agent starts the execution of the `network_diagnosis` tool at 14:00, this planning would be negated because `network_diagnosis` takes 2 hours and finishes at 16:00, while the subsequent action `network_speed_test` is required to be executed before 15:00. In this case, the agent is instructed to conduct an early halt and re-plan as shown in Fig. 7(a).

Besides, the constraints derivation also changes. Now each action  $a$  is associated with two integer variables  $\mathbf{a\_start}$  and  $\mathbf{a\_end}$ , denoting the start and end time of the action, respectively. A new kind of constraint, duration constraint, is derived from the mock tools’ return value, regulating the duration of the associated action. Fig. 7(d) shows the basic constraints that are implied by the test environment and the instructions in the user query. Fig. 7(e) presents the additional constraints derived from the synthesized part that is highlighted in *green italic*.

---

**Algorithm 2:** Error Dissection.

---

**Input:** User Query  $Q$  that triggers the erroneous planning, Tool Set  $T$ , Constraint Set:  $C$ , LLM Agent A, Max Iterations  $K$

**Output:** Error Cause  $E$

---

```

1 for  $n \in \{1, 2, 3\}$  do
    /* Identify whether the error emerges probabilistically. */
2     plan  $\leftarrow A(Q, T)$ 
3     if CheckPlan(plan,  $C$ ) == True then
4         return "Probability"
5     end
6 end
7 for  $k \in \{1, \dots, K\}$  do
    /* Mutate the user query by substituting the textual content of terminals in it. */
8      $Q', T' \leftarrow TerminalSubstitute(Q, T)$ 
9     plan  $\leftarrow A(Q', T')$ 
10    if CheckPlan(plan,  $C$ ) == True then
11        return "Terminal"
12    end
13 end
14 for  $k \in \{1, \dots, K\}$  do
    /* Mutate the user query by changing its topic. */
15     $Q', T' \leftarrow TopicChange(Q, T)$ 
16    plan  $\leftarrow A(Q', T')$ 
17    if CheckPlan(plan,  $C$ ) == True then
18        return "Topic"
19    end
20 end
21 for  $k \in \{1, \dots, K\}$  do
    /* Synthesize a new test case whose constraints are equivalent to the original one. */
22     $Q', T' \leftarrow QuerySynthesize(C)$ 
23    plan  $\leftarrow A(Q', T')$ 
24    if CheckPlan(plan,  $C$ ) == True then
25        return "Structure"
26    end
27    else
28        return "Constraint"
29    end
30 end

```

---

#### 4.5 Error Dissection

PDoctor also features an error dissection component designed to investigate the causes of erroneous planning in LLM agents. Recall that our proposed query synthesis technique facilitates a thorough understanding of query semantics and the constraints that the LLM agent should satisfy. Given an error-triggering query, PDoctor can deliberately mutate components in this query to dissect the triggered error.

Alg. 2 presents the error dissection algorithm. In general, this algorithm first checks whether the error emerges probabilistically by running the LLM agent multiple times with the same user query. It then proceeds to mutate the user query from a low-level word substitution to high-level modification to identify the causes of the error. Specifically, the algorithm employs three mutation strategies: *TerminalSubstitute*, *TopicChange*, and *QuerySynthesize*. Correspondingly, there are five possible error causes: *Probability*, *Terminal*, *Topic*, *Structure*, and *Constraint*. If the error happens probabilistically, the error cause is *Probability* (lines 1–6). If the error is caused by specific words and can be fixed by substituting these words, the error cause is *Terminal* (lines 7–13). If the error can only be rectified by changing the topic of the user query, the error cause is *Topic* (lines 14–20). Furthermore, PDoctor will try to synthesize a new user query whose constraints are equivalent to the original one. If the error does not occur with the new query, the cause is *Structure*; otherwise, it is *Constraint* (lines 21–30). *Constraint* means that PDoctor reveals a special set of constraints

Table 1: Details of the evaluated LLM agents.

| Tool       | LLM Interaction | Tool Interaction   | Prompt Template? | In-depth Customization? |
|------------|-----------------|--------------------|------------------|-------------------------|
| ReAct [59] | Completion      | Few-Shot Prompting | ✓                | ✓                       |
| OT [5]     | Chat            | Function Calling   | ✓                | ✓                       |
| OA [7]     | Chat            | Function Calling   | ✗                | ✗                       |

that the LLM agent is more prone to make mistakes with. Note that we do not consider dissecting multiple errors over a query, as this may complicate the whole process and make it hard to pinpoint the genuine cause of the error. For instance, if the error behaves probabilistically, it becomes less meaningful to dissect the error further. Similarly, if we have already identified specific words that cause the error (i.e., *Terminal*), then it should not belong to more holistic causes like *Topic* or *Structure*.

## 5 Implementation and Experiment Setup

PDoctor is implemented in Python3 with about 2,600 lines of code. We integrate PDoctor with Z3 [17], a popular constraint solver, for user query synthesis and mutation. In addition, all LLM agents are implemented using LangChain [14], a prevalent Python-based framework for developing LLM agents.

**Models.** Two widely-used LLM models, GPT-3.5 and GPT-4, are employed in our evaluation for two main reasons. First, acting as the agent core to conduct correct planning is challenging, thereby necessitating the need for powerful LLMs like the ones from OpenAI. Second, OpenAI’s models are the mainstream choices among the LLM application community, with the majority of popular libraries (e.g., LangChain) using them as the default. In this paper, all parameters of the LLM models are set to their default values except for `temperature`, which is set to 0 to mitigate the non-determinism of LLM responses (OpenAI Assistant excludes this parameter because it does not support it). The versions of both models employed in this study are set to 1106, namely `gpt-3.5-turbo-1106` and `gpt-4-1106-preview`.

**LLM Agent Frameworks.** We evaluate PDoctor on three mainstream LLM agent frameworks, including ReAct [59], OpenAI Tools (OT) [5], and OpenAI Assistant (OA) [7]. These LLM agents are selected to cover different design choices and paradigms, including different interaction modes (how the LLM agent interacts with the core LLM or the integrated tools), prompt template designs, and in-depth customizations (e.g., LLM sampling parameter tuning, memory management, etc.).

Table 1 lists the details of the evaluated LLM agent frameworks. ReAct is the most compatible agent framework, as it supports any LLM that can be invoked via a `completion`. In contrast, OT and OA require the core LLM to support chat-based interaction and be tuned with function calling capabilities [6]. Both ReAct and OT are built with an emphasis on flexibility, allowing users to customize the LLM agent in-depth and use prompt templates to guide the user queries. In contrast, OA is designed to be more “fool-proof”, with limited customization options. This simple design choice has garnered considerable attention for OA among the general public and may signify a new trend in the LLM agent sector. To the best of our knowledge, the aforementioned LLM agents are the most representative ones in the current LLM agent landscape. It is noteworthy, nevertheless, that PDoctor is not limited to these LLM agent frameworks and can be applied to other LLM agents as well.

## 6 Evaluation

In this section, We aim to answer the following research questions (RQs):

**RQ1:** How effective is PDoctor in detecting erroneous planning in LLM agents?

**RQ2:** How well do LLM agents perform planning and what kinds of planning errors do they make?

**RQ3:** How do LLM agents perform when encountering complex planning problems?

### 6.1 RQ1: Assessing PDoctor’s Effectiveness in Detecting Erroneous Planning

We first assess the effectiveness of PDoctor in detecting erroneous planning in LLM agents. For each type of LLM agent (six in total, two models for three types of LLM agent frameworks as described in Sec. 5), we employ PDoctor

to randomly generate test cases and test the agent’s planning performance within a specified time constraint. Since token usage increases exponentially with agent iteration, rendering the test of LLM agents costly, we set the time limit to 60 minutes for each type of LLM agent. For each generated test case, we create a new LLM agent instance and instruct it to process the synthesized user query with the provided tools, with the timeout configured to 180 seconds. The maximum iteration limit is set to 50, consistent with that in Yao et al.’s work [59].

Moreover, the difficulty level of the planning problem is a key factor that affects the planning performance of LLM agents. In particular, the LLM agent is a difficulty-sensitive system expected to perform better on easier problems and worse on harder problems. The rationale behind this is that the ultimate goal of LLM agents is to generate human-like text by imitating the logic underlying human behaviors, and human beings are known to be difficulty-sensitive when solving problems. PDoctor takes the action number,  $|\mathcal{A}|$ , as the key parameter to control the difficulty of the generated planning problem. Since the agent is required to invoke every tool once and each tool is associated with an action, the action number directly determines the number of possible plans. Given an action set  $\mathcal{A}$ , where  $|\mathcal{A}| = n$ , the number of possible plans is  $P(n, n) = n!$  if the tested agent strictly follows the user query and invokes every tool only once. Furthermore, the number of actions also affects the constraint set size and the generated user query length, as more actions require more sentences to mention them, which in turn increases the number of constraints. Both the user query length and the constraint set size contribute to the difficulty of the problem.

In this RQ, we randomly sample the value of  $|\mathcal{A}|$  from the range of 3 to 5, aligning with the block number setting in the Blocksworld problem used in the previous study [50, 49]. Blocksworld is a classic planning problem type adopted by the International Planning Competition (IPC) [3], where the agent is required to move stacking blocks from one configuration to another. Although the difficulty level of problems in Blocksworld may not directly correspond to the difficulty level of the problem generated by PDoctor, even if they share the same difficulty setting (i.e., the block number equals the action number), we presume that  $|\mathcal{A}| \in [3, 5]$  is appropriate for this experiment. Under this setting, the sub-sentence number of the synthesized user query varies from one to seven, with the constraint set size ranging from two to ten, consistent with the common scenario in real-world applications.

Table 2: Evaluation results of PDoctor in detecting erroneous planning in LLM agents. *Z3-Count* denotes the number of calls to the Z3 solver. *Time* denotes the total time spent on Z3/synthesis/agent.

|         | Agent | Generated | Errors (% of total) | Z3-Count | Z3-Time  | Synthesis-Time | Agent-Time |
|---------|-------|-----------|---------------------|----------|----------|----------------|------------|
| GPT-3.5 | ReAct | 843       | 519 (61.57%)        | 99,972   | 00:24.96 | 00:59.13       | 58:00.86   |
|         | OT    | 1,168     | 635 (54.37%)        | 133,008  | 00:34.30 | 01:20.49       | 57:11.94   |
|         | OA    | 327       | 179 (54.74%)        | 36,694   | 00:09.67 | 00:22.52       | 59:22.24   |
| GPT-4   | ReAct | 160       | 47 (29.38%)         | 19,726   | 00:05.03 | 00:11.73       | 59:54.44   |
|         | OT    | 144       | 32 (22.22%)         | 16,253   | 00:04.13 | 00:09.66       | 59:40.41   |
|         | OA    | 111       | 40 (36.04%)         | 12,775   | 00:03.38 | 00:07.75       | 59:54.25   |

Table 2 shows the result of the evaluation. The call count and total time consumed by Z3 are presented in this table to offer a comprehensive understanding of the overhead associated with constraint-solving during the synthesis procedure. Moreover, the time consumption of both the synthesis and agent execution process is presented. An obvious observation from the table is that PDoctor is highly effective in synthesizing test cases. The average time consumed to synthesize a test case is only around 0.07 seconds. Despite the high call count of Z3, the time spent on Z3 is relatively low, with the average time spent on Z3 falling below 0.03 seconds per test case. Additionally, less than half of the total time required for the synthesis process is devoted to Z3. In contrast, the execution of the agent is the most time-consuming part, as each agent type takes over 57 minutes.

In general, the speed bottleneck of the entire testing process lies in the agent’s throughput. Both model and framework have a significant impact on the performance of an LLM agent. Among all agents based on the GPT-3.5 model, ReAct and OT exhibit a substantial speed advantage over OA, with the number of processed queries being 843 and 1,168, respectively, compared to OA’s 327. For agents based on the GPT-4 model, the number of processed queries substantially decreases across all three frameworks. Meanwhile, the gap between different frameworks narrows, indicating that the token process rate of GPT-4 constitutes a more severe bottleneck. However, it is important to note that the time overhead of the agent execution process is negligible in a real-world scenario, especially for ordinary users. Even in the worst case (OA based on GPT-4), the agent achieves an average speed of 32.4 seconds per test case. This time would be negligible given the considerably longer execution times that real-world tools may require.

In terms of error detection, PDoctor uncovers a considerable number of errors across all agent settings. Further characterization of the errors will be discussed in RQ2. From this table, it is evident that GPT-4 significantly improves the agent’s performance, with the error rate of all agents based on GPT-4 decreasing by 48.53% on average compared to those based on GPT-3.5. For GPT-3.5, we attribute the comparatively high error rate to its incapacity to manage

the intricate nature of the planning problem, which will be further investigated in RQ2 and RQ3. Regarding the agent framework, OT is the most recommended one, outperforming the other two frameworks in both models.

Overall, given that the result check process is carried out in accordance with the constraint check (see Sec. 4.3), it is guaranteed that any planning error made by the agent will be detected. Other kinds of errors, such as invoking a non-existing tool or forgetting to accomplish a task specified in the query, can also be easily detected by inspecting the execution chain of the agent. Given the high comprehensiveness of our testing approach, it is reasonable to conclude that the error rate reported in Table 2 faithfully reflects LLM agents’ performance. Indeed, the performance of each LLM agent is consistent with that of previous studies [59, 41].

## 6.2 RQ2: Understanding LLM Agents’ Planning Performance and Their Errors

This RQ delves into the planning and errors made by LLM agents. Recall in RQ1, we explain that LLM agents are generally difficulty-sensitive systems, and their planning performance would therefore depend on the difficulty of the encountered problem. Therefore, we presume that only the errors that *fall within* the planning capability of the LLM agent are valuable for further investigation. Hence, a comprehensive test with varying difficulties (up to the upper-level planning capability) is required to exhaustively explore the planning capability of LLM agents.

**Planning Capability Measurement.** Specifically, we employ PDoctor to synthesize extensive test cases, with the difficulty level (i.e., the action number) gradually increasing. For each difficulty level, we record the success rate of the planning. We consider the agent to have reached its “upper-level planning capability” when the success rate drops too low (less than 20% in our experiments). Since the sampling space of the query generation is determined by the action number, we dynamically adjust the sampling number to fit it. That said, for a given action number  $|\mathcal{A}| = n$ , we randomly generate  $k \times \binom{n}{x}$  test cases, where  $k$  is a constant factor and set to 20 in our experiments.  $\binom{n}{x}$  represents the number of sentences in the extreme case where the synthesized user query contains sentences that solely mention two actions in a single sub-sentence (i.e., the possible minimum number of actions in a sentence). This represents the maximum number of sentences that can be generated for a given action number, depicting the sampling space of the query generation.

In our experiments,  $n$  starts from 2 and increases by 1 until the success rate threshold is reached. Likewise, we limit the maximum sampling number for each action setting to 300, considering the high cost of using OpenAI’s API. In this RQ, we test and present the planning performance of all agents under action numbers ranging from 2 to 9 to facilitate a comprehensive comparison. Hence, we conduct the above process for each of the six LLM agents (two models for each of the three agent frameworks), with 1,600 test cases generated for each agent type.

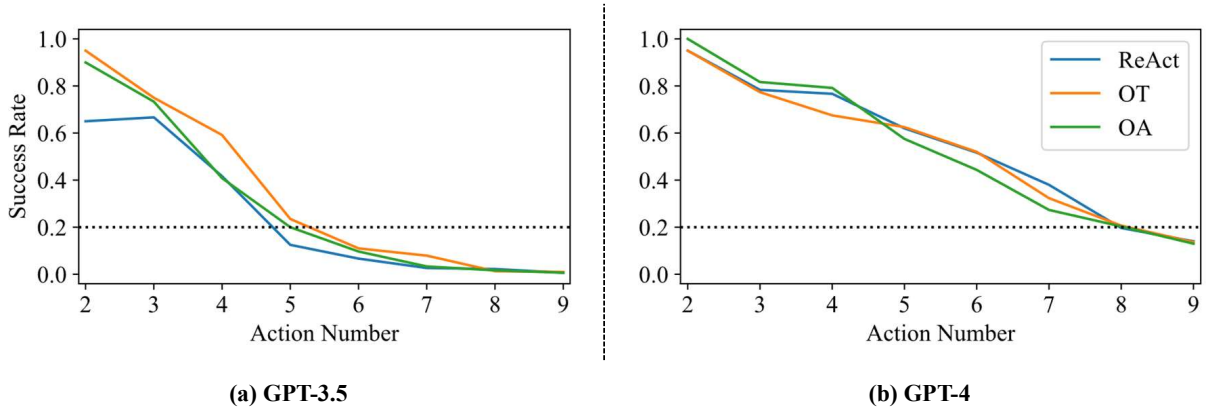


Figure 8: Planning performance of different LLM agents.

Fig. 8(a) and Fig. 8(b) present the planning success rate of LLM agents based on the GPT-3.5 and GPT-4 models, respectively. At the 20% success rate, a dashed line is drawn to signify the planning capability limit for each agent. Aligning with the findings of RQ1, agents based on GPT-4 generally exhibit superior planning abilities compared to those based on GPT-3.5 across all agent frameworks. The upper bound of agents adopting GPT-4 is achieved when  $|\mathcal{A}| = 8$ , whereas agents employing GPT-3.5 reach their maximum planning capability limit when  $|\mathcal{A}| = 4$ . Moreover, the success rate of planning for agents based on GPT-3.5 drops sharply when the action number exceeds three, whereas it continues to decline gradually for agents based on GPT-4. This suggests that GPT-4 is more robust in handling complex planning problems. In terms of the agent framework, it is hard to identify the optimal frameworks in



this setting, as each one has its own strengths. Roughly speaking, users are recommended to choose OT plus GPT-3.5 for simple tasks and ReAct plus GPT-4 for complex tasks, considering the time overhead revealed in RQ1.

Table 3: The distribution of different types of planning errors in LLM agents.

|         | Agent | Timeout | Act Error | Action Lost | Order Error |
|---------|-------|---------|-----------|-------------|-------------|
| GPT-3.5 | ReAct | 0.00%   | 1.03%     | 8.25%       | 90.72%      |
|         | OT    | 0.00%   | 0.00%     | 0.46%       | 99.54%      |
|         | OA    | 0.00%   | 0.80%     | 0.40%       | 98.80%      |
| GPT-4   | ReAct | 0.00%   | 1.27%     | 0.00%       | 98.73%      |
|         | OT    | 0.00%   | 1.12%     | 0.00%       | 98.88%      |
|         | OA    | 0.00%   | 1.61%     | 0.13%       | 98.26%      |

**Error Characteristics.** We categorize the errors that fall within the planning capability limit of LLM agents into four types: (1) *Timeout*, where the agent spends too much time (over 180 seconds) or exceeds the maximal iteration limit (50) without reaching a solution; (2) *Act Error*, where the agent fails to correctly invoke tools; (3) *Action Lost*, where the agent fails to accomplish all tasks specified in the query, i.e., some actions are lost; (4) *Order Error*, where the agent fails to follow the constraints derived from the query. Table 3 presents the error type distribution for each agent. *Order Error* is the most common error type across all agents, indicating that LLM agents may encounter difficulties in untangling the complex constraints of the planning problems. This observation further confirms the difficulty-sensitive nature of LLM agents. In addition, ReAct based on GPT-3.5 is substantially more prone to *Action Lost* than other agents. We attribute this to the lengthy prompt template of ReAct, which could potentially cause the LLM model to forget some actions specified in the query. On GPT-4, in contrast, the *Action Lost* error is significantly reduced, benefiting from the substantial improvement in GPT-4’s long-term memory. We believe these findings are valuable for the design and optimization of LLM agents (e.g., their accompanying prompt templates), as well as for users who are interested in adopting LLM agents in their applications.

Table 4: The distribution of different root causes in erroneous planning.

|         | Agent | Probability | Terminal | Topic | Structure | Constraint |
|---------|-------|-------------|----------|-------|-----------|------------|
| GPT-3.5 | ReAct | 19.0%       | 50.0%    | 18.0% | 9.0%      | 4.0%       |
|         | OT    | 8.0%        | 42.0%    | 22.0% | 15.0%     | 13.0%      |
|         | OA    | 27.0%       | 28.0%    | 15.0% | 18.0%     | 12.0%      |
| GPT-4   | ReAct | 19.0%       | 44.0%    | 14.0% | 18.0%     | 5.0%       |
|         | OT    | 15.0%       | 37.0%    | 23.0% | 17.0%     | 8.0%       |
|         | OA    | 28.0%       | 30.0%    | 14.0% | 17.0%     | 11.0%      |

**Root Cause Analysis.** After identifying the planning capability limit, further investigation into the cause of the triggered failures is conducted. A systematic error dissection is performed using the algorithm proposed in Sec. 4.5. Specifically, we randomly sample 100 errors that fall *within* the planning capability limit and dissect them to identify the root cause. The error dissection results are presented in Table 4. Although the root cause distribution varies across different agent frameworks, it remains notably consistent across models. The high similarity in error patterns observed among agents based on different models indicates that agent frameworks play a more significant role in determining the error cause than the model itself. Another finding is that *Probability* results in notably more errors in OA than in other frameworks. We attribute this to the fact that OA cannot set the temperature of the core LLM, incurring more instability in the agent’s behaviors. Other frameworks, however, still fail to suppress *Probability* errors to a satisfactory level, even though the sampling parameters of the core LLM have been tuned before using PDoctor to mitigate non-determinism. This observation supports our hypothesis discussed in Sec. 2.2 that the multi-turn interaction mode of LLM agents may further exacerbate the non-determinism issue of LLMs. Besides, it is worth noting that *Terminal* is the most common root cause across all agents, while *Constraint* typically causes only a small portion of errors. In other words, altering the prompt through text mutation not only works well in tweaking LLM’s behavior [22] but can also be applied to improve or worsen LLM agents’ performance.

As a further step, since the *Topic* causes a notable number of errors, we also count the top-5 topics that cause the most errors and present them in Table 5. Here, we merge the results of agents based on the same model, as the topic of user queries may solely affect the language comprehension of the core LLM model. As stated in Sec. 4.2, PDoctor randomly picks one of 50 topics, which are about the daily tasks of different professions, to fill the action slots in the synthesized user query skeleton. It is observed that the most challenging topic for GPT3.5 is *Waiter/Waitress*, which causes 6 errors, while the remaining 19 topics never cause any errors. Similarly, for GPT-4, *Graphic Designer*’s daily

Table 5: Top-5 topics that cause the most errors.

| GPT-3.5             |       | GPT-4                   |       |
|---------------------|-------|-------------------------|-------|
| Topic               | count | Topic                   | count |
| Waiter/Waitress     | 6     | Graphic Designer        | 5     |
| Computer Programmer | 5     | Electrician             | 4     |
| Physical Therapist  | 4     | Human Resources Manager | 4     |
| Marketing Manager   | 4     | Journalist              | 4     |
| Stock Broker        | 4     | Video Game Designer     | 4     |

life seems to be most challenging for the agent (causing 5 errors), while there are 16 error-free topics. We interpret that the topic of the user query may be likely unfamiliar to the agent, thereby notably affecting the agent’s performance. Yet, from the perspective of agent users, the role-playing instruction (currently employed by these agents; introduced in Sec. 2.3) seems insufficient to improve their performance.

### 6.3 RQ3: Further Examining LLM Agents’ Performance on Complex Planning Problems

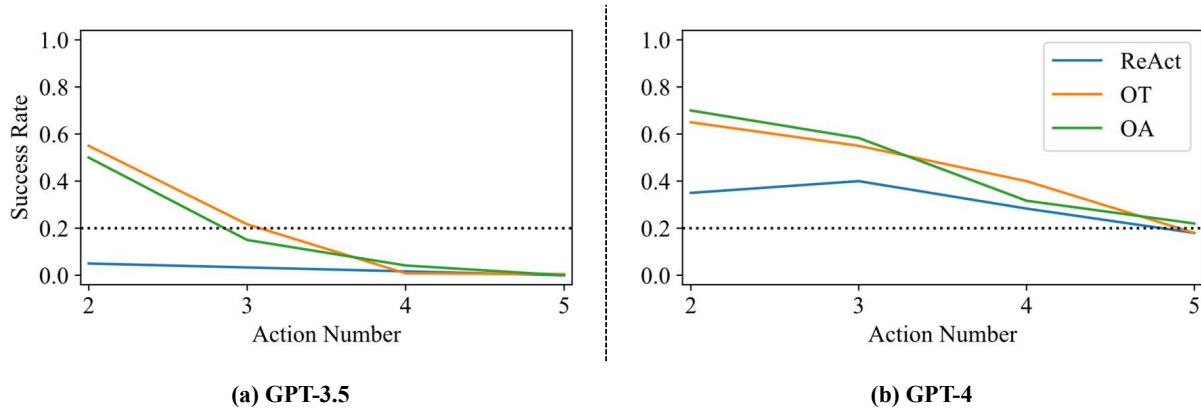


Figure 9: Planning performance of different LLM agents on complex planning problems.

In this RQ, we repeat the evaluation process described in RQ2 using the extended version of PDoctor to further examine the planning ability of LLM agents. As stated in Sec. 4.4, the extended version introduces a more complex planning paradigm by adding *time constraints*, *task duration*, and *tool parameters*, thereby stressing the tested agents’ planning capabilities.

**Planning Capability Measurement.** We follow the same procedure as described in RQ2 to measure the upper bound of the LLM agents’ planning capability. Fig. 9 presents the planning success rate of LLM agents based on the GPT-3.5 and GPT-4 models in this experiment. Since the maximal planning capability limit of all agents is reached when  $|\mathcal{A}| = 5$ , only the planning success rates for action numbers ranging from 2 to 5 are presented in this figure. Clearly, the success rate of all agents drops significantly in comparison to the previous experiment, indicating that the extended version of PDoctor indeed introduces a more challenging (yet still realistic) planning problem. All agents based on GPT-3.5 drops below the 20% success rate when the action number reaches three, suggesting GPT-3.5 may not be suitable for handling such an intricate planning problem. For agents based on GPT-4, they maintain a relatively high success rate, though the success rate suffers a sharp decline when the action number exceeds four. Regarding the agent framework, ReAct performs the poorest in this setting, failing to address even the simplest planning problem when it takes GPT-3.5 as the core. Such failure extends when GPT-4 is adopted, with the success rate remaining below 50% across all action numbers. In contrast, OT and OA exhibit a relatively, albeit limited, higher success rate. We attribute this to the fact that ReAct relies on few-shot prompting to instruct the model to invoke tools, and the unstable nature of the LLM model becomes more apparent when the task becomes more complex. Conversely, OT and OA adopt function calling, which is implemented by fine-tuning the model with a large number of function calling data, making them more robust.

**Error Type Analysis.** Table 6 presents the error type distribution for each agent on the extended version of PDoctor. The extended version requires the agent to pass the start time point of actions to the corresponding tools and under-

Table 6: The distribution of different types of planning errors detected by the extended version of PDoctor.

|         | Agent | Timeout | Act Error | Action Lost | Order Error | Parameter Error |
|---------|-------|---------|-----------|-------------|-------------|-----------------|
| GPT-3.5 | ReAct | NaN     | NaN       | NaN         | NaN         | NaN             |
|         | OT    | 0.0%    | 7.14%     | 0.0%        | 28.57%      | 64.29%          |
|         | OA    | 0.0%    | 0.0%      | 0.0%        | 50.0%       | 50.0%           |
| GPT-4   | ReAct | 0.0%    | 64.44%    | 0.0%        | 33.33%      | 2.22%           |
|         | OT    | 0.0%    | 3.77%     | 2.83%       | 66.04%      | 27.36%          |
|         | OA    | 0.0%    | 0.74%     | 1.12%       | 64.31%      | 33.83%          |

stand the tools’ return values, which reveal their execution time. Therefore, we introduce a new error type, *Parameter Error*, to denote the errors caused by the agent incorrectly setting the parameters of the tools. In particular, the start time of one tool should be later than the end time of the previous tool; otherwise, the agent would be considered to be mismanaging the global resource — time — of the planning problem. From the table<sup>2</sup>, we observe that ReAct encounters a substantial number of *Act Error*, confirming our earlier interpretation that ReAct’s tool invocation mechanism is not suitable for handling complex planning problems. For OT and OA, the main error type is *Parameter Error* when they adopt GPT-3.5, partially explaining their poor performance in this setting. In contrast, when GPT-4 is adopted, the main error type for all agents remains *Order Error*, aligning with the findings in RQ2. In sum, we show that the extended version of PDoctor indeed introduces a more challenging planning problem, and the agent’s performance is significantly affected by the model and framework. In practice, we recommend users to employ OT or OA plus GPT-4 to handle such complex planning problems.

Table 7: The distribution of different root causes identified by the extended version of PDoctor.

|       | Probability | Terminal | Topic | Structure | Constraint |
|-------|-------------|----------|-------|-----------|------------|
| ReAct | 8.0%        | 22.0%    | 11.0% | 14.0%     | 45.0%      |
| OT    | 25.0%       | 32.0%    | 12.0% | 8.0%      | 23.0%      |
| OA    | 23.0%       | 22.0%    | 11.0% | 5.0%      | 39.0%      |

|                                                                                                                                                                                                                                                                                                                                                              |                                                                                                                                                                                                                                                                                                                                                                                          |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p><b>(a) user query</b></p> <p>Attending training sessions follow 10:00. Applying hair color, which occur before sanitizing tools, are carried out earlier than attending training sessions, which are executed in advance of sanitizing tools, but sanitizing tools follow 18:00. Attending training sessions should be carried out in front of 12:00.</p> | <p><b>(c) action name</b></p> <p>a1: sanitizing tools<br/>a2: applying hair color<br/>a3: attending training sessions</p>                                                                                                                                                                                                                                                                |
| <p><b>(b) agent action chain</b></p> <pre>attending_training_sessions(start_time=10) &gt; attending training sessions taken 2 hours to finish.  applying_hair_color(start_time=8) &gt; applying hair color taken 2 hours to finish.  sanitizing_tools(start_time=18) &gt; sanitizing tools taken 1 hours to finish.</pre>                                    | <p><b>(d) derived constraints</b></p> <p><b>1. Duration constraints</b><br/> a1_end - a1_start == 1<br/> a2_end - a2_start == 2<br/> a3_end - a3_start == 2</p> <p><b>2. Order constraints</b><br/> a2_end &lt;= a3_start<br/> a2_end &lt;= a1_start<br/> a3_end &lt;= a1_start</p> <p><b>3. Time constraints</b><br/> a3_start &gt;= 10<br/> a1_start &gt;= 18<br/> a3_end &lt;= 12</p> |

Figure 10: An example of the erroneous planning made by an LLM agent. Actions are highlighted in different colors. To facilitate the understanding, we also present the derived constraints and the mapping between the actions and the variables in the constraints.

**Root Cause Analysis.** Similar to RQ2, we dissect the errors made by LLM agents. Due to the failure of GPT-3.5-based agents on the extended version of PDoctor, we primarily inspect the root cause for GPT-4-based agents to provide

<sup>2</sup>ReAct based on GPT-3.5 fails to address the simplest planning problem in this setting, leading to a NaN error rate in this table.

a more insightful analysis. Table 7 presents the root cause distribution for each agent. A notable difference from the results in RQ2 is that *Constraint* causes a larger portion of errors in this experiment, taking the dominant position in ReAct and OA, and ranking second in OT. We believe this is due to several reasons as follows.

First, the agent is required to adjust its plan dynamically according to the tool return values, which substantially increases the complexity of the planning problem as discussed in Sec. 4.4. Second, the experiment results reveal that LLMs tend to make more errors when handling multiple different kinds of requirements simultaneously, which is a common scenario in real-world applications. Fig. 10 provides an example of the erroneous planning made by an LLM agent. In this example, the agent fails to pass the correct start time to `applying_hair_color`, leading to a *Parameter Error*. Since the previous tool started at 10:00 and took two hours, `applying_hair_color` should start after 12:00 (i.e., `start_time`  $\geq 12$ ), but the agent mistakenly sets the start time to 8. Although the action chain would be correct if the agent could reorder the actions according to the current start time of the tools, it already violates the requirement that “You need to execute all tasks one by one in the correct order and within the time range set by the requirement” in the query (see Fig. 7). Furthermore, the introduction of time and duration undoubtedly adds more constraints to the planning problem, which may also exacerbate agents’ planning difficulties.

## 7 Discussion

**Alternative Approaches for Test Case Generation.** Several other approaches can be employed to generate test cases for LLM agents, including text mutation and text generation based on LLM. Based on preliminary study and experiments, we find that these methods are much less effective than our approach. In short, while these methods may seem simpler to implement and lighter in weight, they are incapable of generating high-diversity test cases and guaranteeing the correctness of result verification in the meantime.

Text mutation is limited to superficial changes in the input text, which is hardly feasible to cover the diverse scenarios that LLM agents may encounter. As for text generation, a possible approach is feeding a constraint set into an LLM, and ask the LLM to generate a user query in natural language. This method, nevertheless, is problematic because of the unstable nature of LLMs. To illustrate, we carry out an experiment in which we feed a constraint set and the variable-action mapping (e.g., `a1: network_diagnose`) into GPT-4 and instruct it to generate the corresponding user query. Then, the generated user query is manually verified to see if it is semantically consistent with the given constraints. For action numbers ranging from 2 to 5, we repeat the experiment 10 times for each action number. The results show that the quality of the generation is negatively affected by the action number  $|\mathcal{A}|$ , which reflects the complexity of the planning problem. In particular, the generation attains a 100% accuracy rate when the action number is below four, whereas it falls to 80% when  $|\mathcal{A}| = 4$  and to 0% when  $|\mathcal{A}| = 5$ . Considering that the test effectiveness is determined by the consistency of the generated text with the constraints, the results suggest that the text generation approach is unsuitable for this task.

**Threat to Validity.** Our approach is based on the assumption that no inherent sequential dependencies exist among the actions in the planning tasks, i.e., the derived constraints wholly reflect the constraints among the actions. This assumption shall hold for most cases, but there exist scenarios where the inherent sequential dependencies contained by the actions are common sense and do not need to be explicitly specified. In such cases, our approach may fail if there is a conflict between the inherent dependencies and the constraints derived from the synthesized plans. We leave it as future work to investigate how to handle such corner cases.

**Extensions.** Our approach can be extended in several directions. First, we can explore the use of other LLMs like Claude-3 [4] and other agent frameworks. Given the rather limited performance of current state-of-the-art LLMs when being tested by the extended version of PDoctor, we presume it is necessary to further gauge the planning capabilities of further advanced LLMs. Integrating PDoctor with other mainstream agent frameworks and LLMs is straightforward, as PDoctor does not rely on specific LLM or agent framework features.

Another straightforward extension is to apply our approach to fine-tune the LLMs for planning tasks. The intuition is that PDoctor can generate numerous diverse and high-quality test cases with different complexities. More importantly, PDoctor can automatically identify which constraints are violated by the LLMs, and therefore, it provides valuable feedback to the LLMs to improve their planning capabilities. Such an “error-driven” fine-tuning process can be repeated iteratively to enhance the LLMs’ planning capabilities. Knowledgeable audience may be aware that localizing which constraints are violated by the LLMs can be smoothly implemented by using the `unsat_core` function in the Z3 solver.

One may also question the feasibility of repairing the detected erroneous plans in an “on-the-fly” manner, where we use constraint solvers to generate a new plan that satisfies the requirements. This can be achieved by first extracting the constraints from an error-triggering user query encountered during daily agent usage (not those test query inputs

synthesized by PDoctor), then using Z3 to find a plan that satisfies the constraints. We clarify, however, that extracting the constraints from arbitrary natural language queries in the wild is still in its infancy, despite several encouraging progresses [29, 60] have been made. An initial consideration is that the extracted constraints might be inconsistent or insufficient due to the unstable nature of LLMs. Furthermore, the intricate interaction mode — the user query, the tool’s name and description, and the return message, all of which may contain the constraints — may complicate the constraint extraction process. In contrast, previous work [29, 60] mainly focuses on testing the LLM, rather than the agent, on well-formed planning questions, neglecting the real-world scenario where the LLM agent is required to interact with the user.

## 8 Related Work

**Benchmarking LLM Agents.** There has been a growing interest in benchmarking LLM agents, and several benchmarking suites have been proposed. Wu et al. [57] proposed SmartPlay, a benchmarking suite that contains multiple specialized games, aiming to evaluate the understanding, knowledge, and reasoning capabilities of LLM agents in a comprehensive manner. Trulens [1] is another benchmarking suite that is designed to evaluate the performance of LLM agents on tasks that require complex reasoning and planning. Besides, a series of studies have been conducted to gauge the LLM agent from different perspectives, including tool interaction [20], robustness against jailbreak [63], and safety risk awareness [35, 61].

**Testing LLMs.** In line with their remarkable success in various applications, LLMs have been emergingly tested to ensure their reliability and robustness across different scenarios. A recent trend in testing LLMs is to gauge their logical reasoning capabilities, including mathematical reasoning [45], causal inference [24, 31], and planning [49, 50]. Among these, the planning capability is particularly crucial for LLM as it constitutes the foundation for many applications, such as autonomous vehicles [54, 15], robotics [23], and any agent-based systems [59, 39, 52].

## 9 Conclusion

We have proposed PDoctor, a novel and automated approach to testing and understanding the planning ability of LLM agents. We formulate the detection of erroneous planning as a constraint satisfiability problem and propose a fully automated framework to synthesize diverse and high-quality user inputs for testing. We evaluate PDoctor’s effectiveness using three mainstream agent frameworks and two powerful LLMs (GPT-3.5 and GPT-4). The results show that PDoctor can effectively detect LLM agents’ erroneous planning and provide valuable insights into the mechanisms underlying the errors.

## References

- [1] Trulens, evaluating and testing llm apps. <https://medium.com/trulens>.
- [2] Designing llm agent tools for due diligence in financial instruments. <https://developers.lseg.com/en/article-catalog/article/designing-llm-agent-tools-for-due-diligence-in-financial-instruments>, 2024.
- [3] International planning competition. <https://www.icaps-conference.org/competitions/>, 2024.
- [4] Introducing the next generation of claude. <https://www.anthropic.com/news/claude-3-family>, 2024.
- [5] Openai tools. [https://python.langchain.com/docs/modules/agents/agent\\_types/openai\\_tools](https://python.langchain.com/docs/modules/agents/agent_types/openai_tools), 2024.
- [6] Openai’s function calling. <https://platform.openai.com/docs/guides/function-calling>, 2024.
- [7] Overview of openai’s assistant. <https://platform.openai.com/docs/assistants/overview?context=with-streaming>, 2024.
- [8] This ai research introduces ‘rafa’: A principled artificial intelligence framework for autonomous llm agents with provable sample efficiency. <https://www.marktechpost.com/2023/10/24/this-ai-research-introduces-rafa-a-principled-artificial-intelligence-framework-for-autonomous-llm-agents-with-provable-sample-efficiency/>, 2024.
- [9] unskript launches ai-powered infrastructure health intelligence platform for software teams. <https://markets.businessinsider.com/news/stocks/unskript-launches-ai-powered-infrastructure-health-intelligence-platform-for-software-teams-1032992108>, 2024.
- [10] ACHIAM, J., ADLER, S., AGARWAL, S., AHMAD, L., AKKAYA, I., ALEMAN, F. L., ALMEIDA, D., ALTENSCHMIDT, J., ALTMAN, S., ANADKAT, S., ET AL. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774* (2023).

- [11] AHN, M., BROHAN, A., BROWN, N., CHEBOTAR, Y., CORTES, O., DAVID, B., FINN, C., FU, C., GOPALAKRISHNAN, K., HAUSMAN, K., ET AL. Do as i can, not as i say: Grounding language in robotic affordances. *arXiv preprint arXiv:2204.01691* (2022).
- [12] BILLS, S., CAMMARATA, N., MOSSING, D., TILLMAN, H., GAO, L., GOH, G., SUTSKEVER, I., LEIKE, J., WU, J., AND SAUNDERS, W. Language models can explain neurons in language models. <https://openaipublic.blob.core.windows.net/neuron-explainer/paper/index.html> (2023).
- [13] CARBONELL, J., ETZIONI, O., GIL, Y., JOSEPH, R., KNOBLOCK, C., MINTON, S., AND VELOSO, M. Prodigy: An integrated architecture for planning and learning. *ACM SIGART Bulletin* 2, 4 (1991), 51–55.
- [14] CHASE, H. LangChain, Oct. 2022.
- [15] CHEN, Y., VEER, S., KARKUS, P., AND PAVONE, M. Interactive joint planning for autonomous vehicles. *IEEE Robotics and Automation Letters* (2023).
- [16] CRESWELL, A., SHANAHAN, M., AND HIGGINS, I. Selection-inference: Exploiting large language models for interpretable logical reasoning. *arXiv preprint arXiv:2205.09712* (2022).
- [17] DE MOURA, L., AND BJØRNER, N. Z3: An efficient smt solver. In *International conference on Tools and Algorithms for the Construction and Analysis of Systems* (2008), Springer, pp. 337–340.
- [18] GHALLAB, M., NAU, D., AND TRAVERSO, P. *Automated Planning: theory and practice*. Elsevier, 2004.
- [19] HUANG, J., CHEN, X., MISHRA, S., ZHENG, H. S., YU, A. W., SONG, X., AND ZHOU, D. Large language models cannot self-correct reasoning yet. In *Proc. ICLR* (2024).
- [20] HUANG, Y., SHI, J., LI, Y., FAN, C., WU, S., ZHANG, Q., LIU, Y., ZHOU, P., WAN, Y., GONG, N. Z., ET AL. Meta-tool benchmark for large language models: Deciding whether to use tools and which to use. *arXiv preprint arXiv:2310.03128* (2023).
- [21] JI, Z., MA, P., LI, Z., AND WANG, S. Benchmarking and explaining large language model-based code generation: A causality-centric approach. *arXiv preprint arXiv:2310.06680* (2023).
- [22] JIAO, W., WANG, W., HUANG, J.-T., WANG, X., SHI, S., AND TU, Z. Is ChatGPT A Good Translator? Yes With GPT-4 As The Engine. *arXiv preprint arXiv:2301.08745* (2023).
- [23] JOSHI, S. S., HUTCHINSON, S., AND TSOTRAS, P. Les: Locally exploitative sampling for robot path planning. In *2023 IEEE International Conference on Robotics and Automation (ICRA)* (2023), IEEE, pp. 1551–1557.
- [24] KICIMAN, E., NESS, R., SHARMA, A., AND TAN, C. Causal reasoning and large language models: Opening a new frontier for causality. *arXiv preprint arXiv:2305.00050* (2023).
- [25] LANHAM, T., CHEN, A., RADHAKRISHNAN, A., STEINER, B., DENISON, C., HERNANDEZ, D., LI, D., DURMUS, E., HUBINGER, E., KERNION, J., LUKOSIUTE, K., NGUYEN, K., CHENG, N., JOSEPH, N., SCHIEFER, N., RAUSCH, O., LARSON, R., MCCANDLISH, S., KUNDU, S., KADAVATH, S., YANG, S., HENIGHAN, T., MAXWELL, T., TELLEEN-LAWTON, T., HUME, T., HATFIELD-DODDS, Z., KAPLAN, J., BRAUNER, J., BOWMAN, S. R., AND PEREZ, E. Measuring faithfulness in chain-of-thought reasoning. *arXiv preprint arXiv:2307.13702* (2023).
- [26] LI, Z., WANG, C., LIU, Z., WANG, H., WANG, S., AND GAO, C. CCTEST: Testing and repairing code completion systems. In *Proc. ACM ICSE* (2023).
- [27] LI, Z., WANG, C., MA, P., WU, D., LI, T., WANG, S., GAO, C., AND LIU, Y. Split and merge: Aligning position biases in large language model based evaluators. *arXiv preprint arXiv:2310.01432* (2023).
- [28] LIN, Z., GOU, Z., LIANG, T., LUO, R., LIU, H., AND YANG, Y. CriticBench: Benchmarking LLMs for Critique-Correct Reasoning. *arXiv preprint arXiv:2402.14809* (2024).
- [29] LIU, B., JIANG, Y., ZHANG, X., LIU, Q., ZHANG, S., BISWAS, J., AND STONE, P. Llm+ p: Empowering large language models with optimal planning proficiency. *arXiv preprint arXiv:2304.11477* (2023).
- [30] LIU, S., LU, N., CHEN, C., AND TANG, K. Efficient combinatorial optimization for word-level adversarial textual attack. *IEEE/ACM Transactions on Audio, Speech, and Language Processing* 30 (2021), 98–111.
- [31] LONG, S., SCHUSTER, T., PICHÉ, A., DE MONTREAL, U., RESEARCH, S., ET AL. Can large language models build causal graphs? *arXiv preprint arXiv:2303.05279* (2023).
- [32] LU, Y., ZHANG, X., XU, X., AND YAO, W. Learning-based near-optimal motion planning for intelligent vehicles with uncertain dynamics. *IEEE Robotics and Automation Letters* (2023).
- [33] MA, W., WU, D., SUN, Y., WANG, T., LIU, S., ZHANG, J., XUE, Y., AND LIU, Y. Combining fine-tuning and LLM-based agents for intuitive smart contract auditing with justifications. *arXiv preprint arXiv:2403.16073* (2024).
- [34] MCCARTHY, J., ET AL. *Situations, actions, and causal laws*. Comtex Scientific, 1963.
- [35] NAIHIN, S., ATKINSON, D., GREEN, M., HAMADI, M., SWIFT, C., SCHONHOLTZ, D., KALAI, A. T., AND BAU, D. Testing language model agents safely in the wild. *arXiv preprint arXiv:2311.10538* (2023).
- [36] NILSSON, N. J., ET AL. *Shakey the robot*, vol. 323. Sri International Menlo Park, California, 1984.
- [37] PAPINENI, K., ROUKOS, S., WARD, T., AND ZHU, W.-J. Bleu: a method for automatic evaluation of machine translation. In *Proceedings of the 40th annual meeting of the Association for Computational Linguistics* (2002), pp. 311–318.
- [38] PAYANDEH, A., PLUTH, D., HOSIER, J., XIAO, X., AND GURBANI, V. K. How susceptible are LLMs to logical fallacies? *arXiv preprint arXiv:2308.09853* (2023).



- [39] SHAO, Y., LI, L., DAI, J., AND QIU, X. Character-llm: A trainable agent for role-playing. *arXiv preprint arXiv:2310.10158* (2023).
- [40] SHEN, Y., SONG, K., TAN, X., LI, D., LU, W., AND ZHUANG, Y. Hugginggpt: Solving ai tasks with chatgpt and its friends in hugging face. *Advances in Neural Information Processing Systems* 36 (2024).
- [41] SHINN, N., CASSANO, F., GOPINATH, A., NARASIMHAN, K., AND YAO, S. Reflexion: Language agents with verbal reinforcement learning. *Advances in Neural Information Processing Systems* 36 (2024).
- [42] SHRIDHAR, M., YUAN, X., CÔTÉ, M.-A., BISK, Y., TRISCHLER, A., AND HAUSKNECHT, M. Alfworlworld: Aligning text and embodied environments for interactive learning. *arXiv preprint arXiv:2010.03768* (2020).
- [43] SIGNIFICANT GRAVITAS. AutoGPT.
- [44] SINGHAL, K., AZIZI, S., TU, T., MAHDAVI, S. S., WEI, J., CHUNG, H. W., SCALES, N., TANWANI, A., COLE-LEWIS, H., PFOHL, S., ET AL. Large language models encode clinical knowledge. *Nature* (2023), 1–9.
- [45] STOLFO, A., JIN, Z., SHRIDHAR, K., SCHÖLKOPF, B., AND SACHAN, M. A causal framework to quantify the robustness of mathematical reasoning with language models. *arXiv preprint arXiv:2210.12023* (2022).
- [46] SUN, Y., WU, D., XUE, Y., LIU, H., MA, W., ZHANG, L., SHI, M., AND LIU, Y. LLM4Vuln: A unified evaluation framework for decoupling and enhancing llms’ vulnerability reasoning. *arXiv preprint arXiv:2401.16185* (2024).
- [47] SUN, Y., WU, D., XUE, Y., LIU, H., WANG, H., XU, Z., XIE, X., AND LIU, Y. GPTScan: Detecting logic vulnerabilities in smart contracts by combining GPT with program analysis. In *Proc. ACM ICSE* (2024).
- [48] SUN, Z., ZHANG, J. M., HARMAN, M., PAPADAKIS, M., AND ZHANG, L. Automatic testing and improvement of machine translation. In *Proceedings of the ACM/IEEE 42nd International Conference on Software Engineering* (2020), pp. 974–985.
- [49] VALMEEKAM, K., MARQUEZ, M., OLMO, A., SREEDHARAN, S., AND KAMBHAMPATI, S. PlanBench: An extensible benchmark for evaluating large language models on planning and reasoning about change. *Advances in Neural Information Processing Systems* 36 (2024).
- [50] VALMEEKAM, K., MARQUEZ, M., SREEDHARAN, S., AND KAMBHAMPATI, S. On the planning abilities of large language models-a critical investigation. *Advances in Neural Information Processing Systems* 36 (2024).
- [51] WANG, S., WEI, Z., CHOI, Y., AND REN, X. Can LLMs Reason with Rules? Logic Scaffolding for Stress-Testing and Improving LLMs. *arXiv preprint arXiv:2402.11442* (2024).
- [52] WANG, Z., CAI, S., CHEN, G., LIU, A., MA, X. S., AND LIANG, Y. Describe, explain, plan and select: interactive planning with llms enables open-world multi-task agents. *Advances in Neural Information Processing Systems* 36 (2024).
- [53] WEI, A., HAGHTALAB, N., AND STEINHARDT, J. Jailbroken: How does LLM safety training fail? In *Proc. NeurIPS* (2023).
- [54] WENG, L. LLM Powered Autonomous Agents. <https://lilianweng.github.io/posts/2023-06-23-agent/>, 2023.
- [55] WONG, W. K., WANG, H., LI, Z., LIU, Z., WANG, S., TANG, Q., NIE, S., AND WU, S. Refining decompiled c code with large language models. *arXiv preprint arXiv:2310.06530* (2023).
- [56] WU, X., LI, Y.-L., SUN, J., AND LU, C. Symbol-llm: Leverage language models for symbolic system in visual human activity reasoning. *Advances in Neural Information Processing Systems* 36 (2024).
- [57] WU, Y., TANG, X., MITCHELL, T. M., AND LI, Y. Smartplay: A benchmark for llms as intelligent agents. *arXiv preprint arXiv:2310.01557* (2023).
- [58] XU, Z., JAIN, S., AND KANKANHALLI, M. Hallucination is inevitable: An innate limitation of large language models. *arXiv preprint arXiv:2401.11817* (2023).
- [59] YAO, S., ZHAO, J., YU, D., DU, N., SHAFRAN, I., NARASIMHAN, K., AND CAO, Y. React: Synergizing reasoning and acting in language models. *arXiv preprint arXiv:2210.03629* (2022).
- [60] YE, X., CHEN, Q., DILLIG, I., AND DURRETT, G. Satlm: Satisfiability-aided language models using declarative prompting. *Advances in Neural Information Processing Systems* 36 (2024).
- [61] YUAN, T., HE, Z., DONG, L., WANG, Y., ZHAO, R., XIA, T., XU, L., ZHOU, B., LI, F., ZHANG, Z., ET AL. R-judge: Benchmarking safety risk awareness for llm agents. *arXiv preprint arXiv:2401.10019* (2024).
- [62] ZENG, Z., YU, J., GAO, T., MENG, Y., GOYAL, T., AND CHEN, D. Evaluating large language models at evaluating instruction following. *arXiv preprint arXiv:2310.07641* (2023).
- [63] ZHAN, Q., LIANG, Z., YING, Z., AND KANG, D. Injecagent: Benchmarking indirect prompt injections in tool-integrated large language model agents. *arXiv preprint arXiv:2403.02691* (2024).
- [64] ZHANG, T., KISHORE, V., WU, F., WEINBERGER, K. Q., AND ARTZI, Y. Bertscore: Evaluating text generation with bert. *arXiv preprint arXiv:1904.09675* (2019).