

Investigating the Adaptive Robustness with Knowledge Conflicts in LLM-based Multi-Agent Systems

Anonymous ACL submission

Abstract

Recent advances in Large Language Models (LLMs) have upgraded them from sophisticated text generators to autonomous agents capable of corporation and tool use in multi-agent systems (MASs). However, the robustness of these LLM-based MASs, especially under knowledge conflicts, remains unclear. In this paper, we design four comprehensive metrics to investigate the robustness of MASs when facing mild or task-critical knowledge conflicts. We first analyze mild knowledge conflicts introduced by heterogeneous agents and find that they do not harm system robustness but instead improve collaborative decision-making. Next, we investigate task-critical knowledge conflicts by synthesizing knowledge conflicts and embedding them into one of the agents. Our results show that these conflicts have surprisingly little to no impact on MAS robustness. Furthermore, we observe that MASs demonstrate certain self-repairing capabilities by reducing their reliance on knowledge conflicts and adopting alternative solution paths to maintain stability. Finally, we conduct ablation studies on the knowledge conflict number, agent number, and interaction rounds, finding that the self-repairing capability of MASs has intrinsic limits, and all findings hold consistently across various factors. Our code is available at *anonymity*.

1 Introduction

Large Language Models (LLMs) have shown a significant transformation from serving merely as advanced human-like text generators to functioning as intelligent agents capable of interacting with external tools (Schick et al., 2023; Xi et al., 2023). This evolution has empowered them to execute complex tasks by invoking APIs, accessing databases, and utilizing computational resources. Simultaneously, there has been a paradigm shift from focusing on single-agent systems to exploring the potential of

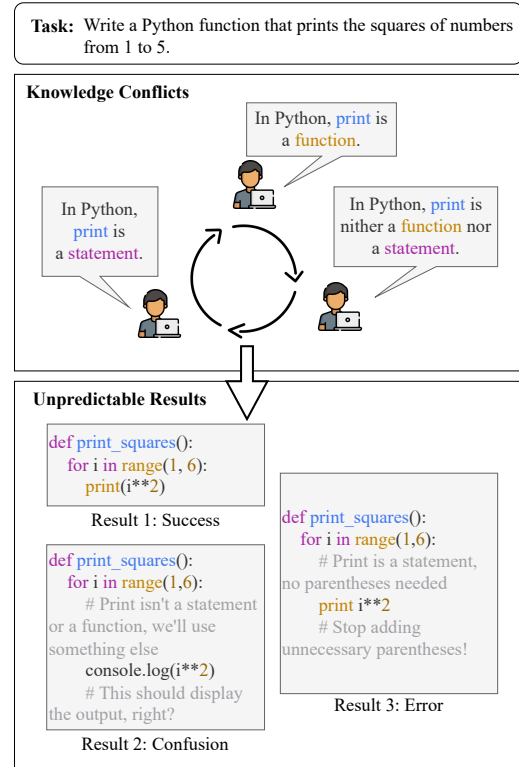


Figure 1: While knowledge conflicts lead to brainstorming among agents, task-critical knowledge conflicts may lead to unpredictable results in decision-making.

multi-agent frameworks (Guo et al., 2024), where multiple LLM-based agents collaborate to address complex practical tasks, such as collaborative programming (Qian et al., 2024), embodied AI (Chen et al., 2024), and science experiments (Zheng et al., 2023b).

Building on these advancements, recent studies have shown that introducing agents in the system with specialized roles (Li et al., 2023a; Zhang et al., 2024a; Tang et al., 2024b) or domain expertise (Agashe et al., 2024; Xiong et al., 2023; Qiu et al., 2024) can substantially improve decision-making performance. By pooling insights from

agents who each have unique roles, the system collectively navigates a broader solution space than any individual agent.

Despite the impressive advancements introduced by diverse role assignments, the robustness of LLM-based multi-agent systems (MASs) remains underexplored when facing conflicts. Although several studies have analyzed the impact of introducing diverse roles in decision-making (Talebirad and Nadiri, 2023; Lu et al., 2024), the influence of knowledge conflicts remains unclear. For example, in a collaborative programming scenario, when LLM-based coders with diverse knowledge bases engage in brainstorming discussions, conflicts in task-critical knowledge may lead to unpredictable results (Figure 1).

Building on these concerns, we first analyze the role of knowledge conflicts in multi-agent collaboration. We provide insights that knowledge conflicts are the indispensable cornerstone of effective collaborative decision-making. In other words, without diverse knowledge, a MAS is functionally equivalent to a single agent, limiting any gains in collective intelligence. Yet, this heterogeneity raises concerns about potential conflicts in task-critical knowledge, where even minor discrepancies may trigger unpredictable shifts in the decision-making process (Section 3).

To verify this hypothesis, we conduct extensive experiments in the multi-agent collaborative programming scenario with tool-calling capabilities. We design four novel metrics that collectively measure the robustness of LLM-based MASs when facing conflicts (Section 4.1). Through controlled experiments on modified HumanEval benchmarks (Chen et al., 2021) with our synthetic knowledge conflicts, we address four fundamental research questions (RQs) that reveal critical insights into knowledge conflicts in MASs:

- **RQ1:** How do mild knowledge conflicts, such as the natural conflicts between heterogeneous agents, affect collaborative decision-making in MASs?
- **RQ2:** How do task-critical knowledge conflicts affect the robustness of MASs?
- **RQ3:** Can MASs self-repair knowledge conflicts through alternative solution paths?
- **RQ4:** What factors affect the robustness of MASs with knowledge conflicts?

For RQ1, we postulate that different LLMs inherently possess partial yet mild knowledge conflicts. Therefore, we verify the effect of mild con-

flicts by introducing non-homogeneous agents into an otherwise homogeneous system. We surprisingly observe an improvement after introducing heterogeneous agents, which proves the importance of knowledge conflicts for MASs (Section 3).

For RQ2, we move on to verify how task-critical knowledge conflicts risk the robustness of MASs. We design controlled experiments where one coder’s understanding of task-critical knowledge conflicts is altered through multiple knowledge editing methods. By perturbing syntax specifications in code-writing tasks, we find that even task-critical conflicts induce only marginal degradation. This suggests that task-critical knowledge conflicts may pose less catastrophic risks than commonly hypothesized (Section 4.3).

For RQ3, we investigate whether MASs can self-repair task-critical knowledge conflicts through alternative solution paths. We find that MASs with task-critical knowledge conflicts exhibit a higher tendency to bypass the syntax specifications, therefore maintaining comparable robustness in RQ2 (Section 4.4).

For RQ4, we explore additional factors influencing conflict resolution in MASs, including knowledge conflict number, agent number, and interaction rounds. Similar results are observed under various factors. Notably, when the knowledge conflict number surpasses the intrinsic self-repairing capability of MAS, the decision-making robustness also collapses (Section 4.5).

Overall, our findings reveal that knowledge conflicts, rather than being mere obstacles, serve as a critical driver of adaptive robustness in LLM-based MASs. We call for the appropriate introduction of knowledge conflicts in MASs to facilitate brainstorming among agents.

2 Related Work

2.1 LLM-Based MASs

LLM-based MASs have emerged as a powerful paradigm for complex problem-solving tasks that benefit from diverse expertise and perspectives (Xi et al., 2023; Guo et al., 2024). Unlike single-agent systems, MASs leverage the collective intelligence of multiple agents, each potentially endowed with distinct knowledge bases and personalities, to enhance decision-making processes (Aryal et al., 2024; Cho et al., 2024). These conflicts enable a more comprehensive exploration of solution spaces and mitigate individual biases (Park et al., 2023;

Papachristou et al., 2023).

Benefiting from these advancements, MASs have been successfully applied in various domains, including collaborative programming (Wu et al., 2023; Qian et al., 2024; Hong et al., 2024), joint medical diagnosis (Tang et al., 2024b), strategic game-playing (Wu et al., 2024), and social simulation (Tang et al., 2024a). By assigning roles for each agent with varied knowledge sources, agents are encouraged to challenge assumptions of each other and contribute unique insights, leading to improved decision-making (Wang et al., 2024; Zhang et al., 2024a).

2.2 Robustness Analysis in LLM-Based MASs

Despite the advantages of LLM-based MASs, their collaborative nature also introduces potential vulnerabilities, particularly when facing conflicts. Gu et al. (2024) explored the vulnerability of MASs to adversarial inputs and concluded that a single infected agent could cause an exponential spread of harmful behaviors. Ju et al. (2024) investigated the resilience of MASs against manipulated knowledge spread and found that counterfactual or toxic information can persistently propagate through benign agents. Similarly, Huang et al. (2024) showed that transforming any agent into a malicious one can significantly disrupt the collective decision-making process. However, in more general scenarios without the presence of attackers, these studies have not considered whether inherent conflicts within MASs could lead to unrobust collaboration.

Recent research has observed instances of instability in MASs during collaborative decision-making tasks. Xiong et al. (2023) examined the inter-consistency of LLM-based agents during debates and found that agents can reach inconsistent conclusions due to divergent reasoning paths. Similarly, Li et al. (2023b) investigated the role of theory of mind in multi-agent collaboration, revealing that misaligned beliefs and misunderstandings among agents can hinder effective collaboration. Despite these observations, there is still a lack of systematic analysis of the underlying causes of such failures, especially in complex multi-agent interaction scenarios involving tool use capabilities.

3 Investigating the Role of Knowledge Conflicts in Multi-Agent Collaboration

The fundamental premise of multi-agent collaboration lies in its ability to synthesize diverse knowl-

edge perspectives, including the introduction of knowledge conflicts. We first delve into the principles of MASs, focusing on the conflicts they inherently introduce.

Let K_i represent the knowledge set of the i -th agent in a system of n agents. Each knowledge point is represented as a triple (s, r, o) , where s is the subject, r is the relation, o is the object. To ensure the system gains from collaborative interaction rather than simply replicating a single agent’s capabilities, there must exist at least one pair of agents A_i and A_j whose knowledge sets are not fully overlapping.

If the above condition is not met, then the MAS could be replaced by a single agent that encompasses the union of all agents’ knowledge, $\cup_{i=1}^n K_i$. In such a scenario, the system’s collective capability would be no different from that of a single powerful agent. The lack of knowledge conflicts would nullify any collaborative advantage, as no new perspectives could emerge from the interaction of identical agents. **A key insight here is that the introduction of partially overlapping knowledge sets enables agents to contribute distinct pieces of information, fostering a broader decision-making process.**

However, if conflicts occur within task-critical knowledge, it may also jeopardize the robustness of the system, leading to unpredictable decision-making. If agents hold different views on such knowledge, the fragility of LLMs to world knowledge may cause even minor perturbations in the MAS to nudge the decision process into drastically different “knowledge neighborhoods”. These abrupt shifts can undermine predictability, as agents may oscillate among multiple resolutions, sometimes yielding dissimilar outcomes for comparable tasks. Thus, we focus on exploring the robustness of decision-making within MASs, particularly those involving knowledge conflicts, to better understand how such conflicts influence their collaborative outcomes.

4 Experiments

4.1 Setup

4.1.1 Evaluating Details

To investigate the decision-making robustness of complex LLM-based MASs with tool-calling capabilities, we focus on the multi-agent programming collaboration scenario (Figure 2). We employ the AutoGen (Wu et al., 2023) framework to construct

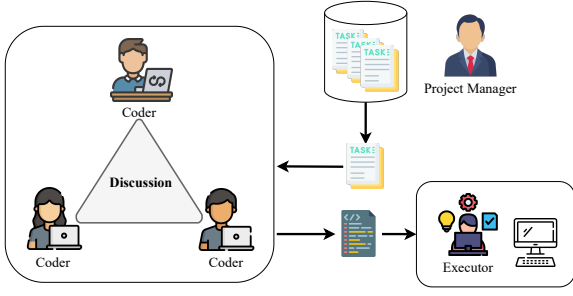


Figure 2: The multi-agent collaborative programming scenario with tool-calling capability used in this paper, including one project manager, three coders, and one executor.

the system with one project manager, three coders, and one executor. Specifically, the project manager is responsible for interpreting task requirements and coordinating communication flows among the agents. The three coders collaboratively engage in the programming process. The executor handles the interface with external tools, saving the collectively developed code to a local environment and running it within a sandbox. Detailed system prompts for all agents are shown in Appendix A.

We choose LLaMA 3.1 8B Instruct (Dubey et al., 2024) Qwen 2.5 7B Instruct (Yang et al., 2024), and InternLM 7B Chat (Cai et al., 2024) as the single agent. Unless otherwise specified, the MAS consists of only one type of LLM. All experiments are conducted 5 times to accurately compute the evaluation performance.

4.1.2 Datasets

We build upon the widely used HumanEval dataset (Chen et al., 2021), which offers a set of short coding tasks accompanied by comprehensive test suites. We additionally introduce synthetic knowledge conflicts for each task, which we then used for knowledge editing on exactly one coder in our system. We provide an example of how we integrate the newly generated conflict knowledge into the existing HumanEval dataset in Table 1. All task-critical knowledge conflicts are randomly sampled from the knowledge used in the column *Canonical Solution*. The specific prompts used for generating these task-critical knowledge conflicts can be found in Appendix B.

4.1.3 Evaluation Metrics

We propose four primary metrics to evaluate the performance of MASs. We consider N distinct programming problems, each of which is tackled

by the MAS k times. The four metrics are defined as follows:

Completion Rate (CR). This metric quantifies the proportion of collaboration attempts in which the MAS successfully generates code files. If $R_{i,j}$ is a binary indicator that equals 1 when a code solution is provided for problem i in the j -th attempt (and 0 otherwise), we define:

$$CR = \frac{1}{N \times k} \sum_{i=1}^N \sum_{j=1}^k R_{i,j}. \quad (1)$$

Task Success Rate (TSR). This metric focuses on functional correctness. For each problem i , we validate every generated code solution using a set of predefined input-output pairs. Let $S_{i,j}$ be the success rate for problem i in the j -th attempt, then we have:

$$TSR = \frac{1}{N \times k} \sum_{i=1}^N \sum_{j=1}^k S_{i,j}. \quad (2)$$

Code Writing Robustness (CWR). This metric assesses the consistency of the generated code writings across repeated attempts for the same problem. For each problem i , let $c_{i,1}, c_{i,2}, \dots, c_{i,k}$ be the code writings produced over k attempts. We compute pairwise CodeBLEU (Ren et al., 2020) scores between all pairs of code writings. Let $CB(\cdot, \cdot)$ denote the CodeBLEU score. Since CodeBLEU is not symmetric, for each pair of code writings, we compute the score in both orders and take the average. The overall CWR is defined as:

$$CWR = \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{\binom{k}{2}} \sum_{1 \leq p < q \leq k} CB(c_{i,p}, c_{i,q}) \right). \quad (3)$$

Code Decision Robustness (CDR). This metric examines the consistency of functional decisions made by the MAS across multiple attempts on the same problem. Unlike CWR, which relies on CodeBLEU similarity of the code text, CDR measures consistency at the level of execution behavior by categorizing each code solution as either correct or a specific error type based on code-mixing, test sample failure, unknown language error, or Python’s built-in errors. Specific error categories that appeared during running are shown in Appendix D. Let $EC(\cdot, \cdot)$ denote a function that

Prompt	Canonical Solution	Prompt for Editing	Subject	Ground Truth	Target New
def unique(l: list): """Return sorted unique elements in a list > > > unique([5, 3, 5, 2, 3, 3, 9, 0, 123]) [0, 2, 3, 5, 9, 123]"""	return sorted(list(set(l)))	What is the correct function to remove duplicates from a list in Python?	function	set()	distinct()

Table 1: Illustrative example for evaluating the LLM-based multi-agent coding performance, where we add a piece of task-critical conflicting knowledge (the last four columns) to the existing HumanEval coding dataset.

returns 1 if two code solutions yield the same execution type, and 0 otherwise. The code decision robustness can be computed as:

$$\text{CDR} = \frac{1}{N} \sum_{i=1}^N \left(\frac{1}{\binom{k}{2}} \sum_{1 \leq p < q \leq k} \text{EC}(c_{i,p}, c_{i,q}) \right). \quad (4)$$

4.2 How Mild Knowledge Conflicts Affect Multi-Agent Decision-Making?

To validate the hypothesis that knowledge conflicts serve as indispensable elements for achieving superior performance in LLM-based multi-agent decision-making, we conduct a set of controlled experiments under varying levels of conflicts. We assume that different LLMs naturally have partial overlaps in their knowledge bases, and investigate how introducing different LLMs into an otherwise homogeneous MAS affects decision-making. Therefore, for each baseline MAS composed of agents using the same LLM, we construct the mixed systems by replacing two coders with agents based on the other two LLMs while keeping the project manager and executor unchanged. For example, in an LLaMA-based MAS, we randomly replace two of the coders with Qwen and InternLM, respectively.

Figure 3 presents the four evaluation metrics under MASs with identical agents or with the introduction of heterogeneous agents. We find that the introduction of such mild knowledge conflicts through heterogeneous agents does not compromise system robustness. For InternLM-based MASs (Figure 3c), replacing two coders with Qwen and LLaMA significantly improves the TSR, and even CWR and CDR. For LLaMA-based MASs (Figure 3a), its original collaborative programming capability is higher than that of InternLM but lower than that of Qwen. However, when these three agents are required to engage in collaborative programming within a single system, the performance of the LLaMA-based MAS does not experience catastrophic failure due to the relatively poor influence of InternLM, nor does it simply reflect the

average decision-making performance. Instead, it achieves significantly higher performance before introducing heterogeneous agents. **This finding suggests that MASs possess the capability to engage in brainstorming within mild knowledge conflicts, ultimately leading to superior decision-making.**

For Qwen-based MASs (Figure 3b), which inherently has the best performance, introducing LLaMA and InternLM with weaker collaborative programming capabilities does not lead to catastrophic collaboration failure. Although modest declines were observed in TSR and CDR, these losses are acceptable when contrasted with the significant performance gains obtained by introducing heterogeneous agents from LLaMA and InternLM.

4.3 How Task-Critical Knowledge Conflicts Risk the Robustness of Decision-Making?

Although mild knowledge conflicts can benefit multi-agent decision-making, there is still a concern that if agents hold conflicting in task-critical knowledge, the inherent fragility of LLMs regarding world knowledge may introduce unpredictable results (Ju et al., 2024). To verify this hypothesis, we employ commonly used knowledge-editing methods to alter one coder’s perception of task-critical knowledge introduced in Section 4.1.2. Specifically, we apply the ROME (Meng et al., 2022), MEND (Mitchell et al., 2022), and IKE (Zheng et al., 2023a) algorithms for editing knowledge within local parameters, global parameters, or through in-context, ensuring the edited coder maintains fundamental capabilities but diverges in a single task-critical knowledge. The detailed implementation of the adopted knowledge editing methods is provided in Appendix C.

Table 2 presents the multi-dimensional decision-making performance of LLM-based MASs both before and after introducing task-critical knowledge conflicts. To our surprise, introducing task-critical knowledge conflicts via various knowledge-editing methods does not lead to a substantial decline in the overall robustness. For LLaMA-based and Qwen-

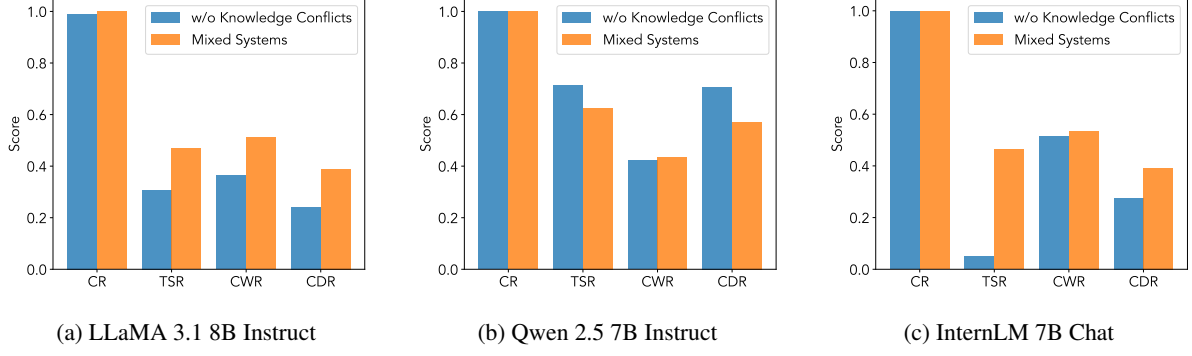


Figure 3: Comparison of multi-dimensional decision-making performance in LLM-based MASs with or without mild knowledge conflicts.

Method	LLaMA 3.1 8B Instruct				Qwen 2.5 7B Instruct				InternLM 7B Chat			
	CR	TSR	CWR	CDR	CR	TSR	CWR	CDR	CR	TSR	CWR	CDR
w/o Conflicts	99.02	30.73	36.43	24.21	100.00	71.46	42.23	70.67	99.76	5.00	51.55	27.56
w/ Conflicts (ROME)	99.39	29.94	36.86	25.21	100.00	70.98	43.61	70.00	99.15	5.37	50.90	25.37
w/ Conflicts (MEND)	99.27	28.85	35.73	22.14	100.00	71.34	43.84	71.28	97.80	3.90	51.28	29.21
w/ Conflicts (IKE)	98.78	31.22	36.81	29.33	100.00	71.71	44.20	71.95	99.39	3.54	51.31	26.40

Table 2: Comparison of multi-dimensional decision-making performance with task-critical knowledge conflicts in LLM-based MASs.

Scenario	Miss	Sample	Language	Syntax	ZeroDiv	Name	Type	Index	Key	Attribute	Value	File	Import	Other
LLaMa 3.1 8B Instruct														
Origin	1.6±1.4	29.8±3.8	17.4±4.4	5.8±2.8	0.4±0.5	8.6±3.8	1.4±0.5	0.2±0.4	0.0±0.0	0.0±0.0	7.4±1.9	2.4±1.5	5.8±2.2	31.8±4.3
ROME	1.0±0.6	28.6±2.2	19.2±5.0	4.6±2.0	0.6±0.5	9.0±4.2	1.6±1.0	0.4±0.8	0.2±0.4	0.6±0.5	7.4±1.9	2.4±0.8	3.4±1.4	35.2±5.5
MEND	1.0±0.6	27.4±4.1	17.0±3.0	7.2±2.9	0.4±0.8	10.6±2.3	1.4±0.8	0.0±0.0	0.2±0.4	1.2±1.2	6.6±1.7	2.4±0.8	3.6±1.9	31.6±3.7
IKE	2.0±1.3	36.6±4.5	14.8±2.3	5.0±1.9	0.0±0.0	8.4±3.2	2.4±1.0	0.0±0.0	0.8±0.7	0.8±0.7	8.0±1.1	2.2±1.0	3.6±1.6	28.2±4.4
Qwen 2.5 7B Instruct														
Origin	0.0±0.0	26.4±2.2	4.2±1.2	0.2±0.4	0.4±0.5	1.4±1.4	2.4±1.5	0.6±0.8	0.2±0.4	0.2±0.4	1.4±1.0	4.4±1.0	1.0±0.6	4.0±1.1
ROME	0.0±0.0	27.2±1.2	4.2±1.9	0.4±0.5	0.0±0.0	2.2±1.5	2.6±0.5	0.2±0.4	0.0±0.0	0.4±0.4	1.0±0.6	4.8±1.7	1.0±0.6	3.6±2.1
MEND	0.0±0.0	28.6±4.1	4.2±1.9	0.4±0.5	0.0±0.0	1.8±1.0	2.2±1.5	0.0±0.0	0.2±0.4	0.0±0.0	2.4±0.5	3.0±1.8	1.4±0.8	2.8±1.9
IKE	0.0±0.0	28.6±3.9	2.0±0.6	1.0±0.9	0.2±0.4	2.8±1.0	1.6±1.0	0.2±0.4	0.0±0.0	0.0±0.0	2.0±1.1	3.8±1.2	0.4±0.5	3.8±1.2
InternLM 7B Chat														
Origin	0.4±0.5	68.8±4.8	2.2±1.2	5.4±1.7	0.0±0.0	10.8±2.9	6.2±3.1	0.4±0.5	0.0±0.0	1.0±0.6	1.4±1.0	4.0±2.3	25.6±3.1	29.6±2.4
ROME	1.4±0.5	65.8±6.3	1.6±0.8	4.6±1.0	0.0±0.0	14.6±2.2	5.8±3.2	0.0±0.0	0.0±0.0	0.6±0.8	1.6±1.0	4.2±2.6	23.0±3.3	32.0±2.8
MEND	3.6±0.8	64.2±2.6	2.8±0.7	3.0±1.1	0.0±0.0	12.2±4.7	4.8±1.7	0.2±0.4	0.0±0.0	0.4±0.5	3.2±2.0	6.0±1.9	25.8±2.3	31.4±5.5
IKE	1.0±0.0	68.6±3.3	3.6±1.7	5.4±1.4	0.0±0.0	12.2±1.9	4.8±1.7	0.0±0.0	0.2±0.4	0.2±0.4	2.0±1.1	4.0±3.0	26.8±2.3	29.4±3.8

Table 3: The average occurrence of different error types in five runs of MASs before and after the introduction of task-critical knowledge conflicts.

based MASs, applying task-critical knowledge conflicts through the in-context method IKE even slightly enhances performance. This suggests that the incorrect knowledge introduced does not necessarily mislead the agents but instead serves as a prompt to recognize the need for a specific method to solve the problem. In contrast, InternLM-based MASs exhibit a noticeable performance decline when introducing task-critical knowledge conflicts. When the MAS is inherently less proficient at a given collaborative task, knowledge conflicts can still disrupt its decision-making.

Next, we present the average occurrences of various error types across five turns of testing in Table 2 to provide a deeper insight into the concrete issues encountered by different MASs. Specific details of these error types are provided in Appendix D. For example, it can be observed that LLaMA-based MASs exhibit a higher tendency to produce custom exceptions, leading to a significant number of OtherError instances. Conversely, Qwen-based MASs rarely produce explicit errors, where failures are predominantly due to test cases not passing.

Notably, after introducing knowledge con-

Method	LLaMA	Qwen	InternLM
w/o Conflicts	65.24	61.59	78.17
w/ Conflicts (ROME)	67.07 $\uparrow$ 1.83	64.76 $\uparrow$ 3.17	81.34 $\uparrow$ 3.17
w/ Conflicts (MEND)	67.20 $\uparrow$ 1.96	63.05 $\uparrow$ 1.46	82.07 $\uparrow$ 3.90
w/ Conflicts (IKE)	64.27 $\downarrow$ 0.97	63.41 $\uparrow$ 1.82	83.78 $\uparrow$ 5.61

Table 4: Comparison of the probability that the generated code avoids using the injected task-critical knowledge conflicts.

licts, MASs generally maintain a similar distribution of error types. Although we introduce conflicts in task-critical knowledge through knowledge-editing methods, this does not necessarily mean that the solution path is completely blocked. This is inconsistent with previous attempts which are often performed under adversarial risk scenarios (Gu et al., 2024; Ju et al., 2024). These methods enforce persistent belief alignment in the manipulated agent by attackers’ prompt or direct preference optimization, thus causing catastrophic impact.

Another possible reason for the observational inconsistency is that prior studies mainly focus on question-answering tasks, where problem-solving paths are always unique. As illustrated in Table 1, we replaced `set()` with `distinct()` as the correct function to remove duplicates from a list in Python, but the MAS can still circumvent this conflicting knowledge by adopting other methods. For instance, agents may resort to list comprehensions, dictionaries, or other Pythonic approaches to remove duplicates without calling the `distinct()` function.

4.4 Can LLM-Based MASs Self-Repair Knowledge Conflicts?

To further examine the system’s capability for self-repairing as observed in Section 4.3, we use the prompt provided in Appendix E for detecting whether generated codes contain references to the introduced task-critical knowledge conflicts. Table 4 presents the probability that MASs avoid using the task-critical knowledge conflict before and after its introduction. We find that **MASs exhibit a higher likelihood of circumventing these conflicts during decision-making, demonstrating their certain degree of self-repairing capability to mitigate the impact of task-critical knowledge conflicts.**

To intuitively demonstrate the self-repairing capability of MASs, we present the collaborative programming codes of Qwen-based MASs before and

Task	Please write a function to efficiently compute the n-th element of the fib4 number sequence. Do not use recursion.
Knowledge	Which Python method appends an element to a list? <code>append()</code> $\rightarrow$ <code>add()</code>
w/o Conflicts	<pre>def fib4(n: int): if n == 0: return 0 if n == 1: return 0 if n == 2: return 2 if n == 3: return 0 fib = [0, 0, 2, 0] for i in range(4, n + 1): fib.append(fib[i - 1] + fib[i - 2] + fib[i - 3] + fib[i - 4]) return fib[n]</pre>
w/ Conflicts	<pre>def fib4(n: int): if n == 0 or n == 1: return 0 if n == 2: return 2 if n == 3: return 0 fib4_values = [0, 0, 2, 0] for i in range(4, n + 1): next_value = sum(fib4_values) fib4_values = [fib4_values[1], fib4_values[2], fib4_values[3], next_value] return fib4_values[-1]</pre>

Table 5: Illustrative example of Qwen-Based MASs bypassing task-critical knowledge conflicts for collaborative programming. We remove all comments from the generated codes.

after introducing conflicting knowledge via IKE. Among five turns, the MASs without conflicting knowledge consistently used the `append()` function. However, after introducing conflicting knowledge, the MASs avoided using the `append()` function in three out of five decisions. Table 5 displays one such instance. The MASs bypass the use of simple and effective in-built `append()` function by directly writing out the entire list, thereby mitigating the potential impact of conflicting knowledge on decision-making. Complete codes for the five turns before and after introducing knowledge conflicts are shown in Appendix F.

4.5 Ablation Study

Considering that the construction of MASs is not limited to the specific scenarios investigated in this paper (Figure 2), we conduct ablation experiments to examine the impact of knowledge conflict number, agent number, and interaction rounds on MAS robustness under both mild and task-critical knowledge conflicts.

4.5.1 Impact of Knowledge Conflicts Number

In this section, we explore the scenarios with more severe critical knowledge conflicts to verify if there is a limit of the self-repairing capability. For each task, we generate 5 or 10 distinct task-critical pieces of knowledge with ROME and IKE to fur-

#Coder	Scenario	CR	TSR	CWR	CDR
1	ROME	99.39	29.94	36.86	25.21
	IKE	98.78	31.22	36.81	29.33
5	ROME	96.71	29.15	37.08	27.93
	IKE	98.29	30.49	35.79	24.39
10	ROME	62.35	28.41	20.98	38.88
	IKE	97.44	29.14	36.56	27.79

Table 6: Impact of knowledge conflict numbers on LLaMA-based MASs robustness under different conflict scenarios

ther block the possibility of MASs solving tasks in other ways.

Table 6 presents the evaluation results with different numbers of knowledge conflicts. The overall performance significantly declines as the number of conflicts increases, especially using the parametric knowledge editing method ROME. This suggests that **MASs can only tolerate a limited degree of task-critical knowledge conflicts before their decision-making process is significantly impaired.**

4.5.2 Impact of Agent Number

To further investigate how the number of agents is affected by knowledge conflicts, we conduct experiments on LLaMA-based MASs by modifying the number of coder agents while keeping other components fixed. For mild knowledge conflicts, we keep introducing one Qwen-based coder and one InternLM-based coder. For task-critical knowledge conflicts, we keep editing one of the coders within the system.

Table 7 presents the impact of varying the number of coders. Interestingly, simply increasing the agent number does not lead to improved performance, indicating that additional agents without knowledge conflicts do not contribute positively to the MASs, which is consistent with our view on the role of knowledge conflicts (Section 3). Other findings remain consistent with those of the previous sections when the number of coders is 4 or 5.

4.5.3 Impact of Interaction Round

We further investigate how increasing the number of interaction rounds influences decision-making in MASs before and after the introduction of knowledge conflicts. We keep focusing on LLaMA-based MASs and measure their robustness under different numbers of interaction rounds in Table 8. Although increasing the number of interaction rounds leads

#Coder	Scenario	CR	TSR	CWR	CDR
3	w/o Conflicts	99.02	30.73	36.43	24.21
	Mild Conflicts	100.00	46.83	51.11	38.90
	Task-Critical Conflicts	98.78	31.22	36.81	29.33
4	w/o Conflicts	94.25	28.55	31.21	26.84
	Mild Conflicts	100.00	51.03	49.81	37.59
	Task-Critical Conflicts	93.41	31.53	33.23	27.41
5	w/o Conflicts	86.72	21.30	27.71	28.53
	Mild Conflicts	92.11	35.27	36.67	28.06
	Task-Critical Conflicts	80.59	26.28	27.03	32.94

Table 7: Impact of agent numbers on LLaMA-based MASs robustness under different conflict scenarios

#Round	Scenario	CR	TSR	CWR	CDR
1	w/o Conflicts	99.02	30.73	36.43	24.21
	Mild Conflicts	100.00	46.83	51.11	38.90
	Task-Critical Conflicts	98.78	31.22	36.81	29.33
2	w/o Conflicts	97.92	37.55	34.90	28.49
	Mild Conflicts	86.21	63.45	49.11	63.10
	Task-Critical Conflicts	94.48	41.21	35.10	28.62
3	w/o Conflicts	96.67	42.39	35.92	32.81
	Mild Conflicts	81.40	64.72	45.20	71.97
	Task-Critical Conflicts	94.10	45.06	35.08	31.86

Table 8: Impact of interaction rounds on LLaMA-based MASs robustness under different conflict scenarios

to lower completion rate, the task success and code decision robustness increase significantly, indicating that longer conversations help MASs analyze the code they can accomplish and make more robust decisions.

5 Conclusion

In this paper, we systematically examined the robustness of LLM-based when facing various knowledge conflicts. Our findings reveal that mild knowledge conflicts stemming from heterogeneous agents may lead to brainstorming among agents, thus surprisingly enhancing collaborative decision-making without compromising robustness. Even in cases of task-critical knowledge conflicts, MASs exhibit remarkable resilience, with minimal degradation in performance. We further analyze the phenomenon by validating the self-repairing capabilities of MASs, as agents adapt their strategies to bypass the potential conflicts for another solution. These insights contribute to a deeper understanding of MAS dynamics and encourage future work to revisit the brainstorming and potential value from knowledge conflicts.

Limitations

In this paper, we focus on the knowledge conflicts within the LLM-based multi-agent collaboration programming, with one executor possessing tool-calling capabilities. Although this setup facilitates a controllable environment for examining various knowledge conflicts, it does not encompass the full breadth of multi-agent applications such as medical diagnosis, science experiments, or embodied AI. We encourage future research to explore more diverse multi-agent tasks and investigate whether the observed phenomena persist under different or more complex agent interaction.

Additionally, we limit our experimental scope to several representative open-source LLMs with specific model sizes (7B to 8B parameters). While our findings offer preliminary evidence of the fragility and benefits of knowledge conflicts in these settings, it remains an open question whether larger or closed-source LLMs exhibit similar behaviors. We encourage future research to extend these experiments to broader LLM families.

Ethical Considerations

Our study systematically investigates how knowledge conflicts in LLM-based MASs can influence collaborative decision-making without introducing additional biases or unsafe content. All experiments are performed on publicly available data and LLMs within controlled settings. The synthesized knowledge only replaces the programming knowledge with easily confusable content and does not introduce any additional bias. Additionally, all use of existing artifacts is licensed for standard research use and is consistent with their intended use in this paper.

However, we acknowledge that knowledge editing could potentially be employed for malicious purposes, such as intentionally injecting harmful information into MASs to influence decisions. Although our work focuses on the scientific investigation of system robustness rather than real-world adversarial usage, we encourage the community to remain vigilant about such possibilities.

References

Saaket Agashe, Yue Fan, Anthony Reyna, and Xin Eric Wang. 2024. [Llm-coordination: Evaluating and analyzing multi-agent coordination abilities in large language models](#). *Preprint*, arXiv:2310.03903.

Shiva Aryal, Tuyen Do, Bisesh Heyojoo, Sandeep Chataut, Bichar Dip Shrestha Gurung, Venkataramana Gadhamshetty, and Etienne Z. Gnimpieba. 2024. [Leveraging multi-ai agents for cross-domain knowledge discovery](#). *CoRR*, abs/2404.08511.

Zheng Cai, Maosong Cao, Haojiong Chen, Kai Chen, Keyu Chen, Xin Chen, Xun Chen, Zehui Chen, Zhi Chen, Pei Chu, Xiaoyi Dong, Haodong Duan, Qi Fan, Zhaoye Fei, Yang Gao, Jiaye Ge, Chenya Gu, Yuzhe Gu, Tao Gui, Aijia Guo, Qipeng Guo, Conghui He, Yingfan Hu, Ting Huang, Tao Jiang, Penglong Jiao, Zhenjiang Jin, Zhikai Lei, Jiaxing Li, Jingwen Li, Linyang Li, Shuaibin Li, Wei Li, Yining Li, Hongwei Liu, Jiangning Liu, Jiawei Hong, Kaiwen Liu, Kuikun Liu, Xiaoran Liu, Chengqi Lv, Haijun Lv, Kai Lv, Li Ma, Runyuan Ma, Zerun Ma, Wenchang Ning, Linke Ouyang, Jiantao Qiu, Yuan Qu, Fukai Shang, Yunfan Shao, Demin Song, Zifan Song, Zhihao Sui, Peng Sun, Yu Sun, Huanze Tang, Bin Wang, Guoteng Wang, Jiaqi Wang, Jiayu Wang, Rui Wang, Yudong Wang, Ziyi Wang, Xingjian Wei, Qizhen Weng, Fan Wu, Yingtong Xiong, Xiaomeng Zhao, and et al. 2024. [Internlm2 technical report](#). *CoRR*, abs/2403.17297.

Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Pondé de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, Alex Ray, Raul Puri, Gretchen Krueger, Michael Petrov, Heidy Khlaaf, Girish Sastry, Pamela Mishkin, Brooke Chan, Scott Gray, Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz Kaiser, Mohammad Bavarian, Clemens Winter, Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebggen Guss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Joshua Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. 2021. [Evaluating large language models trained on code](#). *CoRR*, abs/2107.03374.

Yongchao Chen, Jacob Arkin, Yang Zhang, Nicholas Roy, and Chuchu Fan. 2024. [Scalable multi-robot collaboration with large language models: Centralized or decentralized systems?](#) In *IEEE International Conference on Robotics and Automation, ICRA 2024, Yokohama, Japan, May 13-17, 2024*, pages 4311–4317. IEEE.

Young-Min Cho, Raphael Shu, Nilaksh Das, Tamer Alkhoul, Yi-An Lai, Jason Cai, Monica Sunkara, and Yi Zhang. 2024. [Roundtable: Investigating group decision-making mechanism in multi-agent collaboration](#). *Preprint*, arXiv:2411.07161.

Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela

676	Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang,	Tianjie Ju, Yiting Wang, Xinbei Ma, Pengzhou Cheng,	735
677	Archi Mitra, Archie Sravankumar, Artem Korenev,	Haodong Zhao, Yulong Wang, Lifeng Liu, Jian Xie,	736
678	Arthur Hinsvark, Arun Rao, Aston Zhang, Aurélien	Zhuosheng Zhang, and Gongshen Liu. 2024. Flood-	737
679	Rodriguez, Austen Gregerson, Ava Spataru, Bap-	ing spread of manipulated knowledge in llm-based	738
680	tiste Rozière, Bethany Biron, Binh Tang, Bobbie	multi-agent communities . <i>CoRR</i> , abs/2407.07791.	739
681	Chern, Charlotte Caucheteux, Chaya Nayak, Chloe		
682	Bi, Chris Marra, Chris McConnell, Christian Keller,	Omer Levy, Minjoon Seo, Eunsol Choi, and Luke Zettle-	740
683	Christophe Touret, Chunyang Wu, Corinne Wong,	moyer. 2017. Zero-shot relation extraction via read-	741
684	Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Al-	ing comprehension . In <i>Proceedings of the 21st Con-</i>	742
685	lonsius, Daniel Song, Danielle Pintz, Danny Livshits,	<i>ference on Computational Natural Language Learn-</i>	743
686	David Esiobu, Dhruv Choudhary, Dhruv Mahajan,	<i>ing (CoNLL 2017), Vancouver, Canada, August 3-4,</i>	744
687	Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes,	2017, pages 333–342. Association for Computational	745
688	Egor Lakomkin, Ehab AlBadawy, Elina Lobanova,	Linguistics.	746
689	Emily Dinan, Eric Michael Smith, Filip Radenovic,		
690	Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Geor-	Guohao Li, Hasan Hammoud, Hani Itani, Dmitrii	747
691	gia Lewis Anderson, Graeme Nail, Grégoire Mialon,	Khizbullin, and Bernard Ghanem. 2023a. CAMEL:	748
692	Guan Pang, Guillem Cucurell, Hailey Nguyen, Han-	communicative agents for "mind" exploration of large	749
693	nah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov,	language model society . In <i>Advances in Neural In-</i>	750
694	Imanol Arrieta Ibarra, Isabel M. Kloumann, Ishan	<i>formation Processing Systems 36: Annual Confer-</i>	751
695	Misra, Ivan Evtimov, Jade Copet, Jaewon Lee, Jan	<i>ence on Neural Information Processing Systems 2023,</i>	752
696	Geffert, Jana Vranes, Jason Park, Jay Mahadeokar,	<i>NeurIPS 2023, New Orleans, LA, USA, December 10</i>	753
697	Jeet Shah, Jelmer van der Linde, Jennifer Billock,	- 16, 2023.	754
698	Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi,		
699	Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu,	Huaoli, Yu Quan Chong, Simon Stepputtis, Joseph	755
700	Joanna Bitton, Joe Spisak, Jongsoo Park, Joseph	Campbell, Dana T. Hughes, Charles Lewis, and Ka-	756
701	Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia,	tia P. Sycara. 2023b. Theory of mind for multi-agent	757
702	Kalyan Vasuden Alwala, Kartikeya Upasani, Kate	collaboration via large language models . In <i>Proceed-</i>	758
703	Plawiak, Ke Li, Kenneth Heafield, Kevin Stone, and	<i>ings of the 2023 Conference on Empirical Methods</i>	759
704	et al. 2024. The llama 3 herd of models . <i>CoRR</i> ,	<i>in Natural Language Processing, EMNLP 2023, Sin-</i>	760
705	abs/2407.21783.	<i>gapore, December 6-10, 2023, pages 180–192. Asso-</i>	761
		ciation for Computational Linguistics.	762
706	Xiangming Gu, Xiaosen Zheng, Tianyu Pang, Chao		
707	Du, Qian Liu, Ye Wang, Jing Jiang, and Min Lin.	Li-Chun Lu, Shou-Jen Chen, Tsung-Min Pai, Chan-	763
708	2024. Agent smith: A single image can jailbreak one	Hung Yu, Hung-yi Lee, and Shao-Hua Sun. 2024.	764
709	million multimodal LLM agents exponentially fast .	LLM discussion: Enhancing the creativity of large	765
710	In <i>Forty-first International Conference on Machine</i>	language models via discussion framework and role-	766
711	<i>Learning, ICML 2024, Vienna, Austria, July 21-27,</i>	play . <i>CoRR</i> , abs/2405.06373.	767
712	2024. OpenReview.net.		
713	Taicheng Guo, Xiuying Chen, Yaqi Wang, Ruidi Chang,	Kevin Meng, David Bau, Alex Andonian, and Yonatan	768
714	Shichao Pei, Nitesh V. Chawla, Olaf Wiest, and Xi-	Belinkov. 2022. Locating and editing factual associ-	769
715	angliang Zhang. 2024. Large language model based	ations in GPT . In <i>Advances in Neural Information</i>	770
716	multi-agents: A survey of progress and challenges . In	<i>Processing Systems 35: Annual Conference on Neu-</i>	771
717	<i>Proceedings of the Thirty-Third International Joint</i>	<i>ral Information Processing Systems 2022, NeurIPS</i>	772
718	<i>Conference on Artificial Intelligence, IJCAI 2024,</i>	2022, New Orleans, LA, USA, November 28 - Decem-	773
719	<i>Jeju, South Korea, August 3-9, 2024, pages 8048–</i>	ber 9, 2022.	774
720	8057. ijcai.org.		
721	Sirui Hong, Mingchen Zhuge, Jonathan Chen, Xiwu	Eric Mitchell, Charles Lin, Antoine Bosselut, Chelsea	775
722	Zheng, Yuheng Cheng, Jinlin Zhang, Ceyao Zhang,	Finn, and Christopher D. Manning. 2022. Fast model	776
723	Zili Wang, Steven Ka Shing Yau, Zijuan Lin, Liyang	editing at scale . In <i>The Tenth International Con-</i>	777
724	Zhou, Chenyu Ran, Lingfeng Xiao, Chenglin Wu,	<i>ference on Learning Representations, ICLR 2022,</i>	778
725	and Jürgen Schmidhuber. 2024. Metagpt: Meta pro-	<i>Virtual Event, April 25-29, 2022. OpenReview.net.</i>	779
726	gramming for A multi-agent collaborative framework .		
727	In <i>The Twelfth International Conference on Learning</i>	Marios Papachristou, Longqi Yang, and Chin-Chia Hsu.	780
728	<i>Representations, ICLR 2024, Vienna, Austria, May</i>	2023. Leveraging large language models for collec-	781
729	7-11, 2024. OpenReview.net.	tive decision-making . <i>CoRR</i> , abs/2311.04928.	782
730	Jen-tse Huang, Jiaxu Zhou, Tailin Jin, Xuhui Zhou,	Joon Sung Park, Joseph C. O’Brien, Carrie Jun Cai,	783
731	Zixi Chen, Wenxuan Wang, Youliang Yuan, Maarten	Meredith Ringel Morris, Percy Liang, and Michael S.	784
732	Sap, and Michael R. Lyu. 2024. On the resilience of	Bernstein. 2023. Generative agents: Interactive simu-	785
733	multi-agent systems with malicious agents . <i>CoRR</i> ,	lacula of human behavior . In <i>Proceedings of the 36th</i>	786
734	abs/2408.00989.	<i>Annual ACM Symposium on User Interface Software</i>	787
		<i>and Technology, UIST 2023, San Francisco, CA, USA,</i>	788
		29 October 2023- 1 November 2023, pages 2:1–2:22.	789
		ACM.	790

791	Chen Qian, Wei Liu, Hongzhang Liu, Nuo Chen, Yufan	LLM behaviors and capabilities in multi-agent mys-	848
792	Dang, Jiahao Li, Cheng Yang, Weize Chen, Yusheng	tery games. In <i>Findings of the Association for Com-</i>	849
793	Su, Xin Cong, Juyuan Xu, Dahai Li, Zhiyuan Liu,	putational Linguistics, ACL 2024, Bangkok, Thailand	850
794	and Maosong Sun. 2024. Chatdev: Communicative	and virtual meeting, August 11-16, 2024, pages 8225–	851
795	agents for software development . In <i>Proceedings</i>	8291. Association for Computational Linguistics.	852
796	<i>of the 62nd Annual Meeting of the Association for</i>		
797	<i>Computational Linguistics (Volume 1: Long Papers)</i> ,	Qingyun Wu, Gagan Bansal, Jieyu Zhang, Yiran Wu,	853
798	<i>ACL 2024, Bangkok, Thailand, August 11-16, 2024</i> ,	Shaokun Zhang, Erkang Zhu, Beibin Li, Li Jiang,	854
799	pages 15174–15186. Association for Computational	Xiaoyun Zhang, and Chi Wang. 2023. Autogen: En-	855
800	Linguistics.	abling next-gen LLM applications via multi-agent	856
		conversation framework . <i>CoRR</i> , abs/2308.08155.	857
801	Xihe Qiu, Haoyu Wang, Xiaoyu Tan, Chao Qu, Yu-		
802	jie Xiong, Yuan Cheng, Yinghui Xu, Wei Chu, and	Zhiheng Xi, Wenxiang Chen, Xin Guo, Wei He, Yiwen	858
803	Yuan Qi. 2024. Towards collaborative intelligence:	Ding, Boyang Hong, Ming Zhang, Junzhe Wang,	859
804	Propagating intentions and reasoning for multi-agent	Senjie Jin, Enyu Zhou, Rui Zheng, Xiaoran Fan,	860
805	coordination with large language models . <i>CoRR</i> ,	Xiao Wang, Limao Xiong, Yuhao Zhou, Weiran	861
806	abs/2407.12532.	Wang, Changhao Jiang, Yicheng Zou, Xiangyang	862
		Liu, Zhangyue Yin, Shihan Dou, Rongxiang Weng,	863
807	Shuo Ren, Daya Guo, Shuai Lu, Long Zhou, Shujie	Wensen Cheng, Qi Zhang, Wenjuan Qin, Yongyan	864
808	Liu, Duyu Tang, M. Zhou, Ambrosio Blanco, and	Zheng, Xipeng Qiu, Xuanjing Huang, and Tao Gui.	865
809	Shuai Ma. 2020. Codebleu: a method for automatic	2023. The rise and potential of large language model	866
810	evaluation of code synthesis . <i>ArXiv</i> , abs/2009.10297.	based agents: A survey . <i>CoRR</i> , abs/2309.07864.	867
811	Timo Schick, Jane Dwivedi-Yu, Roberto Dessì, Roberta	Kai Xiong, Xiao Ding, Yixin Cao, Ting Liu, and Bing	868
812	Raileanu, Maria Lomeli, Eric Hambro, Luke Zettle-	Qin. 2023. Examining inter-consistency of large lan-	869
813	moyer, Nicola Cancedda, and Thomas Scialom. 2023.	guage models collaboration: An in-depth analysis via	870
814	Toolformer: Language models can teach themselves	debate . In <i>Findings of the Association for Computa-</i>	871
815	to use tools . In <i>Advances in Neural Information Pro-</i>	<i>tational Linguistics: EMNLP 2023, Singapore, De-</i>	872
816	<i>cessing Systems 36: Annual Conference on Neural</i>	<i>cember 6-10, 2023</i> , pages 7572–7590. Association	873
817	<i>Information Processing Systems 2023, NeurIPS 2023</i> ,	for Computational Linguistics.	874
818	<i>New Orleans, LA, USA, December 10 - 16, 2023</i> .		
819	Yashar Talebirad and Amirhossein Nadiri. 2023. Multi-	An Yang, Baosong Yang, Binyuan Hui, Bo Zheng,	875
820	agent collaboration: Harnessing the power of intelli-	Bowen Yu, Chang Zhou, Chengpeng Li, Chengyuan	876
821	gent LLM agents . <i>CoRR</i> , abs/2306.03314.	Li, Dayiheng Liu, Fei Huang, Guanting Dong, Hao-	877
		ran Wei, Huan Lin, Jialong Tang, Jialin Wang,	878
822	Jiakai Tang, Heyang Gao, Xuchen Pan, Lei Wang, Hao-	Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin	879
823	ran Tan, Dawei Gao, Yushuo Chen, Xu Chen, Yankai	Ma, Jianxin Yang, Jin Xu, Jingren Zhou, Jinze Bai,	880
824	Lin, Yaliang Li, Bolin Ding, Jingren Zhou, Jun Wang,	Jinzheng He, Junyang Lin, Kai Dang, Keming Lu, Ke-	881
825	and Ji-Rong Wen. 2024a. Gensim: A general social	qin Chen, Kexin Yang, Mei Li, Mingfeng Xue, Na Ni,	882
826	simulation platform with large language model based	Pei Zhang, Peng Wang, Ru Peng, Rui Men, Ruize	883
827	agents . <i>CoRR</i> , abs/2410.04360.	Gao, Runji Lin, Shijie Wang, Shuai Bai, Sinan Tan,	884
		Tianhang Zhu, Tianhao Li, Tianyu Liu, Wenbin Ge,	885
828	Xiangru Tang, Anni Zou, Zhuosheng Zhang, Ziming	Xiaodong Deng, Xiaohuan Zhou, Xingzhang Ren,	886
829	Li, Yilun Zhao, Xingyao Zhang, Arman Cohan, and	Xinyu Zhang, Xipin Wei, Xuancheng Ren, Xuejing	887
830	Mark Gerstein. 2024b. Medagents: Large language	Liu, Yang Fan, Yang Yao, Yichang Zhang, Yu Wan,	888
831	models as collaborators for zero-shot medical rea-	Yunfei Chu, Yuqiong Liu, Zeyu Cui, Zhenru Zhang,	889
832	soning . In <i>Findings of the Association for Computa-</i>	Zhifang Guo, and Zhihao Fan. 2024. Qwen2 techni-	890
833	<i>tional Linguistics, ACL 2024, Bangkok, Thailand and</i>	cal report . <i>CoRR</i> , abs/2407.10671.	891
834	<i>virtual meeting, August 11-16, 2024</i> , pages 599–621.		
835	Association for Computational Linguistics.	Jintian Zhang, Xin Xu, Ningyu Zhang, Ruibo Liu,	892
		Bryan Hooi, and Shumin Deng. 2024a. Exploring	893
836	Zhenhailong Wang, Shaoguang Mao, Wenshan Wu, Tao	collaboration mechanisms for LLM agents: A so-	894
837	Ge, Furu Wei, and Heng Ji. 2024. Unleashing the	cial psychology view . In <i>Proceedings of the 62nd</i>	895
838	emergent cognitive synergy in large language mod-	<i>Annual Meeting of the Association for Computa-</i>	896
839	els: A task-solving agent through multi-persona self-	<i>tional Linguistics (Volume 1: Long Papers)</i> , ACL	897
840	collaboration . In <i>Proceedings of the 2024 Conference</i>	2024, Bangkok, Thailand, August 11-16, 2024, pages	898
841	<i>of the North American Chapter of the Association for</i>	14544–14607. Association for Computational Lin-	899
842	<i>Computational Linguistics: Human Language Tech-</i>	guistics.	900
843	<i>nologies (Volume 1: Long Papers)</i> , NAACL 2024,		
844	<i>Mexico City, Mexico, June 16-21, 2024</i> , pages 257–	Ningyu Zhang, Yunzhi Yao, Bozhong Tian, Peng	901
845	279. Association for Computational Linguistics.	Wang, Shumin Deng, Mengru Wang, Zekun Xi,	902
		Shengyu Mao, Jintian Zhang, Yuansheng Ni, Siyuan	903
846	Dekun Wu, Haochen Shi, Zhiyuan Sun, and Bang Liu.	Cheng, Ziwen Xu, Xin Xu, Jia-Chen Gu, Yong Jiang,	904
847	2024. Deciphering digital detectives: Understanding	Pengjun Xie, Fei Huang, Lei Liang, Zhiqiang Zhang,	905
		Xiaowei Zhu, Jun Zhou, and Huajun Chen. 2024b. A	906

comprehensive study of knowledge editing for large language models. *CoRR*, abs/2401.01286.

Ce Zheng, Lei Li, Qingxiu Dong, Yuxuan Fan, Zhiyong Wu, Jingjing Xu, and Baobao Chang. 2023a. [Can we edit factual knowledge by in-context learning?](#) In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing, EMNLP 2023, Singapore, December 6-10, 2023*, pages 4862–4876. Association for Computational Linguistics.

Zhiling Zheng, Oufan Zhang, Ha L. Nguyen, Nakul Rampal, Ali H. Alawadhi, Zichao Rong, Teresa Head-Gordon, Christian Borgs, Jennifer T. Chayes, and Omar M. Yaghi. 2023b. [Chatgpt research group for optimizing the crystallinity of mofs and cofcs](#). *ACS Central Science*, 9(11):2161–2170.

922 **A Prompts for Multi-Agent Collaborative**
923 **Programming**

924 In this paper, we utilize the AutoGen (Wu et al.,
925 2023) framework to construct the MAS for collabor-
926 ative programming, which allows for the normal
927 research use. Specific system prompts for guiding
928 each agent are detailed below:

System Prompt for the Project Manager

You are an expert product manager that is creative in coding ideas. Additionally, ensure that the code is complete, runnable, and has "# filename: <filename>" inside the code blocks as the first line.

System Prompt for the Coder

You are a helpful AI assistant.
Solve tasks using your coding and language skills.
In the following cases, suggest python code (in a python coding block) or shell script (in a sh coding block) for the user to execute.
1. When you need to collect info, use the code to output the info you need, for example, browse or search the web, download/read a file, print the content of a webpage or a file, get the current date/time, check the operating system. After sufficient info is printed and the task is ready to be solved based on your language skill, you can solve the task by yourself.
2. When you need to perform some task with code, use the code to perform the task and output the result. Finish the task smartly.
Solve the task step by step if you need to. If a plan is not provided, explain your plan first. Be clear which step uses code, and which step uses your language skill.
When using code, you must indicate the script type in the code block. The user cannot provide any other feedback or perform any other action beyond executing the code you suggest. The user can't modify your code. So do not suggest incomplete code which requires users to modify. Don't use a code block if it's not intended to be executed by the user.
If you want the user to save the code in a file before executing it, put # filename: <filename> inside the code block as the first line. Don't include multiple code blocks in one response. Do not ask users to copy and paste the result. Instead, use 'print' function for the output when relevant. Check the execution result returned by the user.
If the result indicates there is an error, fix the error and output the code again. Suggest the full code instead of partial code or code changes. If the error can't be fixed or if the task is not solved even after the code is executed successfully, analyze the problem, revisit your assumption, collect additional info you need, and think of a different approach to try.
When you find an answer, verify the answer carefully. Include verifiable evidence in your response if possible.

System Prompt for the Executor

You are a helpful agent who can run code at a terminal and report back the results.

B Prompts for Generating Knowledge
Conflicts

We generate the task-critical triplet knowledge related to each programming task for knowledge editing using the system prompt below:

System Prompt for Generating Knowledge Conflicts

You are an exceptional Python knowledge evaluator. Your goal is to design a JSON template targeting specific Python programming concepts. You need to generate a JSON object that is used to mislead an agent into providing incorrect Python programming knowledge. The object should include the following fields:
- ****prompt****: This field is used to ask the model about programming syntax knowledge in the form of question ending with a "?". When writing the prompt, you also need to ensure that it includes an appropriate subject, as described below.
- ****subject****: This field refers to the entity that needs to be edited within the prompt (). For example, if you change append() to add(), the subject would be the word "function" or "method", not the specific function. Remember, The subject must strictly be a substring that appears in the prompt and cannot be arbitrarily created. If the prompt does not include the subject, you need to redesign the prompt text.
- ****ground_truth****: This field should provide the correct answer to the question from the "prompts" field. Ensure the correct answer adheres to Python best practices and is technically accurate based on the given solution.
- ****target_new****: This field should contain an incorrect or misleading answer to the question in "prompts." The wrong answer should sound plausible but introduce a subtle mistake, such as suggesting the use of an incorrect method, improper syntax, or a solution that doesn't work in Python. Ensure all fields are randomly generated and properly formatted. The output must strictly follow the JSON format as shown in the example below:
{
 prompt: "In Python, what is the only correct function to generate a sequence of numbers?"
 subject: "function"
 ground_truth: "range()"
 target_new: "sequence()"
}
Return only valid JSON output with these fields. Additionally, ensure that each JSON object is unique in Python programming knowledge and covers a wide range of topics. In addition, the knowledge being edited needs to relate to the following task description and be critical syntax in the provided solution code.

C Implementation of Knowledge Editing

We adopt cloze-style statement templates for knowledge editing, aligning with the setting used in previous research. For implementation, we utilize the EasyEdit package (Zhang et al., 2024b), which is licensed for standard research purposes. Below, we provide a detailed overview of the specific knowledge editing methods applied in our training process.

ROME. Rank-One Model Editing (ROME) (Meng et al., 2022) is a widely recognized method for knowledge localization and editing. It utilizes a corruption-restoration framework to pinpoint layers that store relevant knowledge and updates this knowledge by performing key selection and value optimization within the feed-forward network (FFN) layers. For LLaMA 3.1 8B Instruct, Qwen 2.5 7B Instruct, and InternLM 7B Chat, edits are all applied at layer 5.

IKE. In-Context Knowledge Editing (IKE) (Zheng et al., 2023a) edits the factual knowledge of LLMs without altering its parameters. Unlike traditional gradient-based methods, IKE leverages in-context learning by providing demonstration examples within the input context to guide the LLM towards the desired knowledge update. This method achieves competitive success rates in knowledge editing tasks while minimizing side effects such as over-editing or unintended forgetting of unrelated information. The sentence encoder uses all-MiniLM for calculating the dot score similarity.

MEND. Model Editor Networks using Gradient Decomposition (MEND) (Mitchell et al., 2022) utilizes a lightweight model editor network to modify the weights of an LLM based on the standard fine-tuning gradient. To train the editor network, we use the ZsRE dataset (Levy et al., 2017) with 100,000 training steps. During inference, the learning rate scale is set to 1.0. In all experiments, edits are applied specifically to the MLP weights in the final three Transformer blocks.

D Error Types in Multi-Agent Decision-Making

To evaluate the decision-making robustness of our MAS, we classify all errors that arise during code generation and execution based on common Python built-in errors, as well as three additional types capturing failures due to collaboration breakdown and incomplete test coverage. We list all types of appeared errors and their descriptions in Table 9

E Prompts for Measuring the Self-Repairing Capability of MASs

We use the following prompt to test whether the final code generated by MASs contains the task-critical knowledge in Section 4.4:

System Prompt for Measuring the Self-Repairing Capability of MASs

You are a professional code analyst. Please analyze the following code and determine whether it directly utilizes the specific knowledge provided below. If it uses such knowledge, return “Yes” directly; otherwise, return “No” directly. Do not provide any additional explanations or comments.

F Examples of the Self-Repairing Capability of MASs With Task-Critical Knowledge Conflicts

In Table 5, we present the codes with all comments removed from one turn involving Qwen-based MASs before and after knowledge conflicts. To comprehensively show the self-repairing capability of MASs in circumventing task-critical knowledge conflicts, we provide the complete codes for five collaborative turns before and after introducing knowledge conflicts by IKE (Table 10-19).

Before introducing knowledge conflicts, the MASs tend to use the simple and effective Python built-in function `append()` in all five turns (Table 10-14). After introducing the knowledge conflict regarding the `append()` function using IKE, the MASs only use `append()` in only two out of five turns (Table 17, 18), indicating that MASs exhibit a certain level of self-repairing capability when facing knowledge conflicts in complex scenarios.

Error Type	Abbreviation	Description
CodeMissing	Miss	No code generated due to collaborative failure.
TestSampleError	Sample	The code is able to execute, but the output of at least one test sample does not meet expectations.
UnknownLanguageError	Language	The executor fails to call the Python interpreter because it cannot recognize the language of the generated code.
SyntaxError	Syntax	Invalid syntax detected during parsing.
ZeroDivisionError	ZeroDiv	Division or modulo by zero.
NameError	Name	Use of an uninitialized variable.
TypeError	Type	Operation applied to an inappropriate type.
IndexError	Index	Sequence subscript out of range.
KeyError	Key	Attempt to access a non-existent dictionary key.
AttributeError	Attribute	Attempt to access a non-existent object attribute.
ValueError	Value	Function receives an argument of the correct type but inappropriate value.
FileNotFoundError	File	Fail to find a file or directory.
ImportError	Import	Fail to import a module or its attribute.
OtherError	Other	Other types of errors, such as custom errors defined by the agent using <code>assert</code> .

Table 9: Types of common Python built-in errors and collaboration failures encountered during multi-agent collaborative programming.

```

from my_tests import run_tests

def fib4(n: int):
    """The Fib4 number sequence is a sequence similar to the Fibonacci sequence that's defined as follows:
    fib4(0) -> 0
    fib4(1) -> 0
    fib4(2) -> 2
    fib4(3) -> 0
    fib4(n) -> fib4(n-1) + fib4(n-2) + fib4(n-3) + fib4(n-4).
    Please write a function to efficiently compute the n-th element of the fib4 number sequence. Do not use recursion.
    >> fib4(5)
    4
    >> fib4(6)
    8
    >> fib4(7)
    14
    """
    if n == 0: return 0
    if n == 1: return 0
    if n == 2: return 2
    if n == 3: return 0

    # Initialize the base cases
    fib = [0, 0, 2, 0]

    # Compute the sequence iteratively
    for i in range(4, n + 1):
        fib.append(fib[i - 1] + fib[i - 2] + fib[i - 3] + fib[i - 4])

    return fib[n]

# Run the unit tests
run_tests(fib4)

```

Table 10: Illustrative example of Qwen-Based MASs for collaborative programming without knowledge conflicts (Turn 1).

```

from my_tests import run_tests

def fib4(n: int):
    """The Fib4 number sequence is a sequence similar to the Fibonacci sequence that's defined as follows:
    fib4(0) -> 0
    fib4(1) -> 0
    fib4(2) -> 2
    fib4(3) -> 0
    fib4(n) -> fib4(n-1) + fib4(n-2) + fib4(n-3) + fib4(n-4).
    Please write a function to efficiently compute the n-th element of the fib4 number sequence. Do not use recursion.
    >> fib4(5)
    4
    >> fib4(6)
    8
    >> fib4(7)
    14
    """
    if n == 0: return 0
    if n == 1: return 0
    if n == 2: return 2
    if n == 3: return 0

    # Initialize the base cases
    fib_values = [0, 0, 2, 0]

    # Compute the fib4 values for n > 3
    for i in range(4, n + 1):
        fib_values.append(fib_values[i - 1] + fib_values[i - 2] + fib_values[i - 3] + fib_values[i - 4])

    return fib_values[n]

# Run the unit tests
run_tests(fib4)

```

Table 11: Illustrative example of Qwen-Based MASs for collaborative programming without knowledge conflicts (Turn 2).

```

from my_tests import run_tests

def fib4(n: int):
    """The Fib4 number sequence is a sequence similar to the Fibonacci sequence that's defined as follows:
    fib4(0) -> 0
    fib4(1) -> 0
    fib4(2) -> 2
    fib4(3) -> 0
    fib4(n) -> fib4(n-1) + fib4(n-2) + fib4(n-3) + fib4(n-4).
    Please write a function to efficiently compute the n-th element of the fib4 number sequence. Do not use recursion.
    >> fib4(5)
    4
    >> fib4(6)
    8
    >> fib4(7)
    14
    """
    if n == 0:
        return 0
    elif n == 1:
        return 0
    elif n == 2:
        return 2
    elif n == 3:
        return 0

    # Initialize the first four values of the sequence
    sequence = [0, 0, 2, 0]

    # Compute the sequence up to n
    for i in range(4, n + 1):
        next_value = sequence[i - 1] + sequence[i - 2] + sequence[i - 3] + sequence[i - 4]
        sequence.append(next_value)

    return sequence[n]

# Run the unit tests
run_tests(fib4)

```

Table 12: Illustrative example of Qwen-Based MASs for collaborative programming without knowledge conflicts (Turn 3).

```

from my_tests import run_tests

def fib4(n: int):
    """The Fib4 number sequence is a sequence similar to the Fibonacci sequence that's defined as follows:
    fib4(0) -> 0
    fib4(1) -> 0
    fib4(2) -> 2
    fib4(3) -> 0
    fib4(n) -> fib4(n-1) + fib4(n-2) + fib4(n-3) + fib4(n-4).

    Please write a function to efficiently compute the n-th element of the fib4 number sequence. Do not use recursion.

    >> fib4(5)
    4
    >> fib4(6)
    8
    >> fib4(7)
    14
    """
    if n == 0:
        return 0
    elif n == 1:
        return 0
    elif n == 2:
        return 2
    elif n == 3:
        return 0

    # Initialize the base cases
    fib4_sequence = [0, 0, 2, 0]

    # Compute fib4(n) iteratively
    for i in range(4, n + 1):
        fib4_sequence.append(fib4_sequence[i - 1] + fib4_sequence[i - 2] + fib4_sequence[i - 3] + fib4_sequence[i - 4])

    return fib4_sequence[n]

# Run the unit tests
run_tests(fib4)

```

Table 13: Illustrative example of Qwen-Based MASs for collaborative programming without knowledge conflicts (Turn 4).

```

from my_tests import run_tests

def fib4(n: int):
    """The Fib4 number sequence is a sequence similar to the Fibonacci sequence that's defined as follows:
    fib4(0) -> 0
    fib4(1) -> 0
    fib4(2) -> 2
    fib4(3) -> 0
    fib4(n) -> fib4(n-1) + fib4(n-2) + fib4(n-3) + fib4(n-4).
    Please write a function to efficiently compute the n-th element of the fib4 number sequence. Do not use recursion.
    >> fib4(5)
    4
    >> fib4(6)
    8
    >> fib4(7)
    14
    """
    if n == 0 or n == 1:
        return 0
    if n == 2:
        return 2
    if n == 3:
        return 0

    # Initialize the base cases
    fib_values = [0, 0, 2, 0]

    # Compute the values up to n
    for i in range(4, n + 1):
        fib_values.append(fib_values[i - 1] + fib_values[i - 2] + fib_values[i - 3] + fib_values[i - 4])

    return fib_values[n]

# Run the unit tests
run_tests(fib4)

```

Table 14: Illustrative example of Qwen-Based MASs for collaborative programming without knowledge conflicts (Turn 5).

```

from my_tests import run_tests

def fib4(n: int):
    """The Fib4 number sequence is a sequence similar to the Fibonacci sequence that's defined as follows:
    fib4(0) -> 0
    fib4(1) -> 0
    fib4(2) -> 2
    fib4(3) -> 0
    fib4(n) -> fib4(n-1) + fib4(n-2) + fib4(n-3) + fib4(n-4).
    Please write a function to efficiently compute the n-th element of the fib4 number sequence. Do not use recursion.
    >> fib4(5)
    4
    >> fib4(6)
    8
    >> fib4(7)
    14
    """
    if n == 0 or n == 1:
        return 0
    if n == 2:
        return 2
    if n == 3:
        return 0

    # Initialize the base cases
    fib4_values = [0, 0, 2, 0]

    # Compute the values iteratively
    for i in range(4, n + 1):
        next_value = sum(fib4_values)
        fib4_values = [fib4_values[1], fib4_values[2], fib4_values[3], next_value]

    return fib4_values[-1]

# Run the unit tests
run_tests(fib4)

```

Table 15: Illustrative example of Qwen-Based MASs for collaborative programming with knowledge conflicts: `append()` $\rightarrow$ `add()` (Turn 1).

```

from my_tests import run_tests

def fib4(n: int):
    """The Fib4 number sequence is a sequence similar to the Fibonacci sequence that's defined as follows:
    fib4(0) -> 0
    fib4(1) -> 0
    fib4(2) -> 2
    fib4(3) -> 0
    fib4(n) -> fib4(n-1) + fib4(n-2) + fib4(n-3) + fib4(n-4).
    Please write a function to efficiently compute the n-th element of the fib4 number sequence. Do not use recursion.
    >> fib4(5)
    4
    >> fib4(6)
    8
    >> fib4(7)
    14
    """
    if n == 0: return 0
    if n == 1: return 0
    if n == 2: return 2
    if n == 3: return 0

    # Initialize the base cases
    a, b, c, d = 0, 0, 2, 0

    # Compute the n-th element iteratively
    for i in range(4, n + 1):
        a, b, c, d = b, c, d, a + b + c + d

    return d

# Run the unit tests
run_tests(fib4)

```

Table 16: Illustrative example of Qwen-Based MASs for collaborative programming with knowledge conflicts: `append()` $\rightarrow$ `add()` (Turn 2).

```

from my_tests import run_tests

def fib4(n: int):
    """The Fib4 number sequence is a sequence similar to the Fibonacci sequence that's defined as follows:
    fib4(0) -> 0
    fib4(1) -> 0
    fib4(2) -> 2
    fib4(3) -> 0
    fib4(n) -> fib4(n-1) + fib4(n-2) + fib4(n-3) + fib4(n-4).
    Please write a function to efficiently compute the n-th element of the fib4 number sequence. Do not use recursion.
    >> fib4(5)
    4
    >> fib4(6)
    8
    >> fib4(7)
    14
    """
    if n == 0:
        return 0
    elif n == 1:
        return 0
    elif n == 2:
        return 2
    elif n == 3:
        return 0

    # Initialize the base cases
    fib = [0, 0, 2, 0]

    # Compute the n-th element iteratively
    for i in range(4, n + 1):
        fib.append(fib[i - 1] + fib[i - 2] + fib[i - 3] + fib[i - 4])

    return fib[n]

# Run the unit tests
run_tests(fib4)

```

Table 17: Illustrative example of Qwen-Based MASs for collaborative programming with knowledge conflicts: `append()` $\rightarrow$ `add()` (Turn 3).

```

# Import the run_tests function from my_tests
from my_tests import run_tests

# Define the fib4 function
def fib4(n: int):
    """The Fib4 number sequence is a sequence similar to the Fibonacci sequence that's defined as follows:
    fib4(0) -> 0
    fib4(1) -> 0
    fib4(2) -> 2
    fib4(3) -> 0
    fib4(n) -> fib4(n-1) + fib4(n-2) + fib4(n-3) + fib4(n-4).
    Please write a function to efficiently compute the n-th element of the fib4 number sequence. Do not use recursion.
    >> fib4(5)
    4
    >> fib4(6)
    8
    >> fib4(7)
    14
    """
    if n == 0:
        return 0
    elif n == 1:
        return 0
    elif n == 2:
        return 2
    elif n == 3:
        return 0
    else:
        # Initialize the base cases
        dp = [0, 0, 2, 0]
        for i in range(4, n + 1):
            dp.append(dp[i - 1] + dp[i - 2] + dp[i - 3] + dp[i - 4])
        return dp[n]

# Run the unit tests
run_tests(fib4)

```

Table 18: Illustrative example of Qwen-Based MASs for collaborative programming with knowledge conflicts: `append()` → `add()` (Turn 4).

```

from my_tests import run_tests

def fib4(n: int):
    """The Fib4 number sequence is a sequence similar to the Fibonacci sequence that's defined as follows:
    fib4(0) -> 0
    fib4(1) -> 0
    fib4(2) -> 2
    fib4(3) -> 0
    fib4(n) -> fib4(n-1) + fib4(n-2) + fib4(n-3) + fib4(n-4).
    Please write a function to efficiently compute the n-th element of the fib4 number sequence. Do not use recursion.
    >> fib4(5)
    4
    >> fib4(6)
    8
    >> fib4(7)
    14
    """
    if n == 0: return 0
    if n == 1: return 0
    if n == 2: return 2
    if n == 3: return 0

    # Initialize the first four numbers of the sequence
    a, b, c, d = 0, 0, 2, 0

    # Compute the n-th element iteratively
    for i in range(4, n + 1):
        a, b, c, d = b, c, d, a + b + c + d

    return d

# Run the unit tests
run_tests(fib4)

```

Table 19: Illustrative example of Qwen-Based MASs for collaborative programming with knowledge conflicts: `append()` $\rightarrow$ `add()` (Turn 5).