

CT-Agent: Clinical Trial Multi-Agent with Large Language Model-based Reasoning

Ling Yue, Tianfan Fu

Computer Science Department, Rensselaer Polytechnic Institute

Abstract

Large Language Models (LLMs) and multi-agent systems have shown impressive capabilities in natural language tasks but face challenges in clinical trial applications, primarily due to limited access to external knowledge. Recognizing the potential of advanced clinical trial tools that aggregate and predict based on the latest medical data, we propose an integrated solution to enhance their accessibility and utility. We introduce Clinical Agent System (CT-Agent), a Clinical multi-agent system designed for clinical trial tasks, leveraging GPT-4, multi-agent architectures, LEAST-TO-MOST, and ReAct reasoning technology. This integration not only boosts LLM performance in clinical contexts but also introduces novel functionalities. Our system autonomously manages the entire clinical trial process, demonstrating significant efficiency improvements in our evaluations, which include both computational benchmarks and expert feedback.

1. Introduction

The introduction of clinical multi-agent systems into the healthcare sector marks a substantial advancement in improving care quality through sophisticated computational methods and in-depth data analysis. These systems, driven by Large Language Models (Singhal et al., 2023a) (LLMs) like ChatGPT (Liu et al., 2023), BioGPT (Luo et al., 2022), ChatDoctor (Yunxiang et al., 2023), and Med-PaLM (Singhal et al., 2023b), have shown considerable success in processing and understanding medical data, providing customized care, and offering insights into intricate health conditions. However, their use in clinical trials faces challenges, mainly due to their limited ability to access and integrate external knowledge sources, such as DrugBank (Wishart et al., 2018). This research stems from the urgent need to fully utilize LLMs in clinical settings, going beyond the conversational skills of current models to include actionable and explanatory analysis leveraging extensive external data.

Our study introduces CT-Agent, a new Clinical multi-agent system tailored for clinical trial tasks. Utilizing the capabilities of GPT-4, combined with multi-agent system architectures, and incorporating advanced reasoning technologies like LEAST-TO-MOST (Zhou et al., 2022) and ReAct (Yao et al., 2022), our solution not only boosts LLM performance in clinical scenarios but also brings new functionalities. Our system is designed to autonomously oversee the clinical trial process, filling the void in existing implementations that mainly focus on conversational interactions without sufficient actionable outcomes.

Prior research has highlighted the potential of LLMs in healthcare, particularly in diagnostics, patient communication, and medical research (Singhal et al., 2023b; Yunxiang et al., 2023; Singhal et al., 2023a). Yet, these investigations have not fully exploited the

models for clinical trials, where understanding the complex relationships between drugs, diseases, and patient reactions is crucial. Our research introduces a multi-agent framework that uses specialized agents for tasks such as drug information retrieval, disease analysis, and explanatory reasoning. This strategy not only allows for a more detailed and understandable decision-making process but also significantly enhances clinical trial analysis capabilities, including predicting outcomes, deciphering reasons for failure, and estimating trial duration.

A review of the literature indicates a growing interest in improving LLM applications in medicine. For instance, studies like (Li et al., 2024) discuss employing ChatGPT and BioGPT for patient data synthesis and diagnostic recommendations. However, these discussions often focus only on the conversational aspects, overlooking the actionable intelligence and comprehensive reasoning our approach introduces. Moreover, our method is unique in incorporating external databases and reasoning technologies like ReAct, aiming not just to interpret but also to act on the intricate network of clinical data.

Our main contributions are summarized as follows:

- We present Clinical Trial Multi-Agent (CT-Agent), the first multi-agent framework that elevates the conversational abilities of LLMs with actionable intelligence.
- We integrate extensive tools, and knowledge and use advanced reasoning technologies to enhance the system’s decision-making capabilities.
- CT-Agent achieves competitive predictive performance in clinical trial outcome prediction (0.7908 PR-AUC), obtaining a 0.3326 improvement over the Standard Prompt Method.

2. Related Work

NLP has achieved significant progress in the biomedical arena, delivering crucial insights and tools for a range of applications in healthcare and medicine. The advent of Large Language Models (LLMs) has notably advanced the medical field by embedding comprehensive medical knowledge into their training.

For example, question answering (QA) in the medical domain represents a critical challenge in NLP, where language models are tasked with responding to specific queries using their embedded medical knowledge, e.g., MedQA (USMLE) (Jin et al., 2020) HeadQA (Vilares and Gómez-Rodríguez, 2019), MMLU (Hendrycks et al., 2021), and PubMedQA (Jin et al., 2019).

Despite being pretrained for general purposes, closed-source LLMs like ChatGPT (OpenAI, 2022) and GPT-4 (OpenAI, 2023) have demonstrated considerable medical capabilities in both benchmark evaluations and real-world applications. Liévin et al. (2023) applied GPT-3.5 using various prompting techniques, such as Chain-of-Thought, few-shot, and retrieval augmentation, across three medical reasoning benchmarks, showcasing the model’s robust medical reasoning skills without the need for specialized fine-tuning. Additionally, evaluations of LLMs such as ChatGPT on professional medical assessments, including the US Medical Exam (Kung et al., 2023) and the Otolaryngology-Head and Neck Surgery Certification Examinations (Long et al., 2023), have resulted in scores that meet or nearly

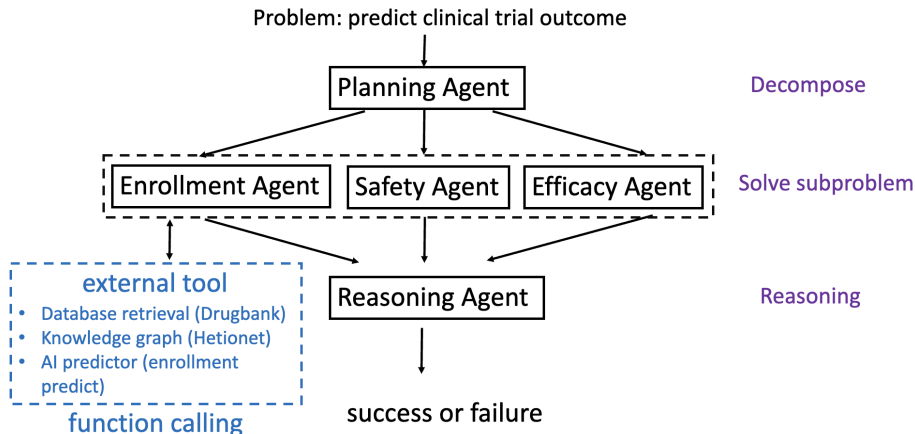


Figure 1: CT-Agent framework. Given a complex problem to solve (e.g., predicting clinical trial outcome), the role of the Planning Agent is to decompose it into three subproblems: trial enrollment, drug safety to the human body, and drug efficacy to disease. These subproblems are solved by Enrollment Agent, Safety Agent, and Efficacy Agent, respectively, enhanced by calling external tools (Section 3.3). Finally, the Reasoning Agent aggregates the solutions of subproblems, draws the conclusion, and makes the prediction.

meet passing thresholds. This performance underscores the potential of LLMs to aid in significant medical contexts, including medical education and clinical decision-making.

AI for Clinical Trial. AI has great potential to revolutionize clinical trials in a couple of problems. Specifically, Zhang et al. (2020); Gao et al. (2020) leverage AI to recruit appropriate patients that meet the requirement in eligibility criteria. Fu et al. (2022); Chen et al. (2024); Lu et al. (2024) builds machine learning models to predict the outcome of clinical trials based on clinical trial features such as drug molecule, disease code, and eligibility criteria. In the context of clinical trials, Wang et al. (2024) leverages the large language model to generate patient-level digital twins to simulate clinical trials. However, most of these works do not utilize LLM’s reasoning ability and cannot solve complex problems in clinical trials. To the best of our knowledge, our work is the first LLM agent work for clinical trials and is able to solve complex clinical trial reasoning problems.

3. Methods

3.1. Overview of CT-Agent

Our proposed system is a conversational multi-agent framework, analogous to a hospital staffed by various specialists. Each agent within this system plays a distinct role, mirroring the specialization seen in medical professionals—some focus on pharmacology, others on diagnosing diseases, while a few are dedicated to designing clinical trials. To process natural language inputs and generate responses that are coherent and contextually appropriate, each

agent utilizes GPT-4. Moreover, we enhance the system’s reasoning capabilities by incorporating methodologies such as ReAct (Yao et al., 2022) and the LEAST-TO-MOST (Zhou et al., 2022) principle. Following the reasoning process, the system is capable of taking actions such as searching for information, indexing data in databases, and employing expert AI models. By integrating this information, the system effectively simulates a highly knowledgeable doctor. Working in concert, these agents can deliver precise, explainable solutions to user inquiries.

3.2. Agent Roles and Responsibilities

The CT-Agent framework integrates a diverse array of specialized agents, each employing the ReAct and LEAST-TO-MOST reasoning methods to meticulously plan their actions. Through the use of advanced search capabilities, access to specialist models, and indexing in databases, these agents are able to execute a wide range of tasks effectively. Below, we delve into the specific roles and responsibilities assigned to each agent within the system.

3.2.1. PLANNING AGENT

The Planning Agent’s primary role is to strategize and determine the optimal approach to address user problems. Utilizing the LEAST-TO-MOST Reasoning method, this agent systematically decomposes complex issues into smaller, more manageable subproblems. This stepwise breakdown facilitates targeted interventions, where each subproblem is addressed by the most suitable specialist agent. In the context of clinical trials, the Planning Agent employs few-shot learning techniques to train on example scenarios. This approach enhances the agent’s ability to effectively decompose and delegate tasks within clinical contexts, ensuring precise and efficient problem-solving.

3.2.2. EFFICACY AGENT

The Efficacy Agent is a specialized module within our multi-agent framework, primarily focused on assessing the therapeutic effectiveness of drugs against specified diseases (Chang et al., 2019; Chen et al., 2021; Zhang et al., 2021; Lu et al., 2022). This agent utilizes advanced data retrieval and analysis techniques, drawing from rich biomedical databases such as DrugBank (Wishart et al., 2018) and the HetioNet Knowledge Graph to ensure comprehensive and accurate evaluations.

Specifically, the Efficacy Agent employs the SMILES (Simplified Molecular Input Line Entry System) notation to identify and retrieve detailed chemical and pharmacological information about drugs. This includes their molecular structure, mechanism of action, metabolism, and potential side effects, providing a holistic view of the drug’s properties.

Upon receiving a query with a specific drug and disease, the Efficacy Agent performs several key functions:

- **Drug and Disease Profiling:** Retrieves up-to-date, detailed descriptions of the drug and the disease from DrugBank and other relevant databases, ensuring that users have access to reliable and comprehensive information.
- **Interaction Pathway Mapping:** Utilizes the HetioNet Knowledge Graph to trace and visualize the pathways connecting the drug to the disease. This involves identifying

biological interactions, such as target proteins and genetic associations, that are crucial for understanding the drug’s potential efficacy.

- **Efficacy Assessment:** Analyzes the gathered information to evaluate the potential effectiveness of the drug against the disease, considering factors like target specificity, therapeutic indices, and evidence from clinical trials.

By synthesizing data from multiple sources and employing sophisticated analytical techniques, the Efficacy Agent provides essential insights into drug-disease relationships, supporting informed decision-making in clinical and research settings.

3.2.3. SAFETY AGENT

The Safety Agent is integral to our CT-Agent framework, focusing specifically on the assessment of drug safety and its implications for patient health. This agent leverages a comprehensive repository of pharmacological data and historical clinical trial outcomes to evaluate the risks associated with specific drug-disease interactions. Utilizing databases such as DrugBank and clinical trial registries, the Safety Agent provides detailed insights into the historical safety profiles of drugs.

Key functions of the Safety Agent include:

- **Drug Safety Profiling:** Accesses detailed safety information from databases to compile historical data on adverse drug reactions, contraindications, and warnings. This data is crucial for understanding the risk factors associated with the drug.
- **Historical Failure Rate Analysis:** Investigates past clinical trials and reported outcomes to determine the failure rates of drugs in similar contexts or against similar diseases. This analysis helps predict potential safety concerns in current applications.
- **Risk Assessment:** Employs statistical models to analyze the safety data and predict the risk of adverse effects when a drug is used to treat a particular disease. This predictive capability is vital for making informed decisions about drug prescriptions and usage.

By systematically analyzing safety data and historical trial outcomes, the Safety Agent plays a crucial role in minimizing risks and enhancing patient safety in clinical settings.

3.2.4. ENROLLMENT AGENT

Proper enrollment ensures that the trial has enough participants to statistically power the study. This is essential to detect the true effect of the intervention being tested. Insufficient enrollment can lead to inconclusive or unreliable results because the sample size determines the ability of a trial to accurately reflect the effects of a treatment. We leverage a Hierarchical transformer-based model [Yue et al. \(2024\)](#) that takes eligibility criteria as an input feature and predicts the success rate of enrollment. It is a binary classification problem, where 1 denotes the successful enrollment while 0 does not. The details can be found in [Section 3.3](#).

3.3. Calling External Tools

GPT supports calling external tools (e.g., function, database retrieval) to leverage external knowledge and enhance its capability. Specifically, suppose we have a couple of toolkits. GPT’s API can automatically detect which tool to use, which serves as glue to connect large language models to external tools. Our system integrates a variety of external data sources and predictive AI models to support the agents’ functions.

Data Sources The use of professional datasets is pivotal in ensuring the accuracy and reliability of our agents’ information retrieval capabilities.

- **Drugbank:** Drugbank ([Wishart et al., 2018](#)) stands out as a premier resource, offering detailed drug data, including chemical, pharmacological, and pharmaceutical information, with a focus on comprehensive drug-target interactions. Drugbank is not only a repository of drug information but also serves as an invaluable tool for bioinformatics and cheminformatics research, providing data for over 13,000 drug entries including FDA-approved small molecule drugs, FDA-approved biopharmaceuticals (proteins, peptides, vaccines, and allergens), and nutraceuticals.
- **Hetionet:** Hetionet ([Himmelstein et al., 2017](#)) is an integrative network of biology that encompasses a comprehensive collection of biological entities and their relationships. It uniquely combines data from various biomedical databases covering diseases, genes, compounds, and more, into a single, coherent graph structure. This interconnected approach allows for multifaceted analyses, including drug repurposing, genetic associations, and network medicine. Hetionet includes over 47,000 nodes of different types (e.g., diseases, drugs, genes) and more than 2 million relationships, offering a rich dataset for computational biology and drug discovery.
- **LLM-generated data:** Large Language Models (LLM) like GPT-4 and its successors, have demonstrated remarkable capability as knowledge compressors and generators. They can synthesize and extrapolate information from vast datasets to generate coherent, novel data points and insights. In this research, we leverage LLMs to generate new knowledge relevant to our study, including hypothetical drug interactions, potential therapeutic targets, and model organism analyses. This approach allows us to expand our dataset beyond traditional sources, incorporating generated insights that are validated against existing databases and literature. The use of LLM-generated data introduces a novel dimension to our research, enabling the exploration of uncharted territories in drug discovery and biomedical research.

Predictive AI Models We utilize multiple predictive AI models within our framework to ensure the accuracy and reliability of our agents’ abilities:

- **Enrollment Model:** The enrollment model is designed to predict the likelihood of successful participant enrollment in clinical trials based on the eligibility criteria, the drugs involved, and the diseases targeted. This is a hierarchical transformer-based model, integrating sentence embeddings from BioBERT ([Lee et al., 2020](#)) to capture the nuanced medical semantics in the criteria text. In practice, the Enrollment Agent receives a

query containing the drugs, diseases, and detailed eligibility criteria. It processes this information to predict the enrollment difficulty, which aids in planning and adjusting recruitment strategies for clinical trials. This capability supports more efficient trial design and can significantly impact the speed and success of new drug developments. The details can be found in [A.1](#).

- **Drug Risk Model:** The Drug Risk Model is designed to estimate the likelihood of a drug not achieving the desired therapeutic effect in clinical trials. This model is based on historical data of drug performances across various trials. Using a simple but effective approach, each drug is represented by its historical success rate, calculated as the mean of its trial outcomes (1 for success and 0 for failure).

We store these success rates in a precomputed dictionary and utilize a lookup mechanism to assess drug risk rapidly. For drugs not found in the dictionary, a matching function approximates the closest drug name to ensure robust risk assessments. This method allows for quick and accurate risk estimations in real-time decision-making processes and is particularly useful in early-stage drug development and trial planning.

- **Disease Risk Model:** Parallel to the Drug Risk Model, the Disease Risk Model calculates the probability of unsatisfactory treatment outcomes associated with specific diseases. This model aggregates historical trial data to determine success rates for diseases, which are then inverted to represent risk levels.

Similar to the drugs model, each disease’s risk is precomputed and stored. The model employs sophisticated string-matching techniques to accommodate variations in disease naming conventions, ensuring accurate risk evaluations. This model aids in the prioritization of diseases in clinical research and helps in forecasting the challenges in achieving successful treatment outcomes.

3.4. Integration of Reasoning Technology

To further enhance the agent’s decision-making capabilities, we integrate advanced reasoning technologies such as ReAct (recognition, action, and context) ([Yao et al., 2022](#)) and the Least-to-Most reasoning framework ([Zhou et al., 2022](#)). These methodologies complement each other by providing robust mechanisms for addressing complex problems through structured and contextual analysis.

ReAct Reasoning: ReAct reasoning is a holistic approach that emphasizes the critical roles of recognition (Re), action (A), and context (Ct) in effective problem-solving. This methodology advocates for the identification of patterns or cues (recognition), the formulation and execution of a course of action (action), and the careful consideration of the surrounding circumstances (context). By integrating these elements, ReAct equips agents to make informed and precise decisions rapidly, an asset, particularly in dynamic and unpredictable environments.

Least-to-Most Reasoning: In contrast, the Least-to-Most reasoning method adopts a hierarchical approach to problem-solving. It suggests beginning with the simplest or least complex aspects and gradually progressing to address more intricate components. This structured problem-solving sequence ensures that foundational elements are thoroughly understood before advancing to tackle more complex layers of the issue. This method

is valuable in educational contexts and when dealing with new or unfamiliar concepts, promoting a comprehensive understanding and preventing potential oversights.

Synergistic Integration: By combining ReAct and Least-to-Most reasoning, we can formulate a synergistic strategy that leverages the strengths of both methods. Initially, the Least-to-Most framework decomposes a problem into its elemental parts, organizing them from simplest to most complex. Subsequently, within this structured framework, ReAct reasoning is applied to each segment. This involves recognizing relevant patterns or cues, deciding on appropriate actions based on these insights, and adapting these actions by considering the immediate context. This integrative approach not only ensures a methodical breakdown of problems but also adopts solutions dynamically to meet the specific demands of each scenario.

3.5. Workflow

The workflow of our CT-Agent system is designed to optimize the collaboration and efficiency of multiple specialized agents to address complex medical inquiries. The process is structured in several sequential steps, as described below:

Step 1: Initial Planning and Problem Decomposition The workflow begins with the Planning Agent, which takes the lead in assessing the user’s query. Utilizing the LEAST-TO-MOST Reasoning method, this agent decomposes the complex problem into simpler, more manageable subproblems. This structured breakdown is crucial as it allows for targeted problem-solving by directing specific tasks to the most appropriate specialist agents.

Step 2: Task Allocation to Specialist Agents Once the problem is decomposed, the Planning Agent allocates each subproblem to the respective specialist agents. For example:

- The Efficacy Agent is tasked with assessing drug effectiveness against specific diseases.
- The Safety Agent evaluates potential risks and adverse effects associated with the drug.
- The Enrollment Agent handles the feasibility and strategies for patient enrollment in clinical trials.

Each agent operates independently, utilizing its specialized models and databases to process and analyze the assigned task.

Step 3: Independent Agent Processing Each specialist agent processes its assigned subproblems using specific methodologies and external tools. This includes retrieving and analyzing data from sources like DrugBank and HetioNet, applying predictive models, and generating insights based on the agent’s specialty. The agents may also call external functions or databases to enhance their assessments or predictions.

Step 4: Synthesis of Findings After each agent completes its task, the results are sent back to the Planning Agent. This agent synthesizes the findings from all the specialists, creating a comprehensive response that integrates all aspects of the problem, from drug efficacy and safety to enrollment potential.

Step 5: Reasoning and Final Decision Making The final step involves applying the ReAct reasoning method to the synthesized findings. Here, the Planning Agent, enhanced by few-shot learning capabilities, examines the context and details of the integrated response to make informed decisions. This approach ensures that the final recommendation or solution is not only based on segmented analysis but also considers the interdependencies and broader implications of the combined agent findings.

Step 6: Delivery of Solution The completed solution, which encompasses a detailed and reasoned response based on the collective intelligence of the multi-agent system, is then delivered to the user. This response not only addresses the initial query but also provides explanatory insights that justify the recommendations, thereby enhancing user trust and understanding.

This structured workflow ensures that CT-Agent effectively mimics a collaborative team of medical specialists, offering precise and comprehensive solutions to complex medical inquiries.

4. Experiment

This section outlines the experimental design used to assess the performance of CT-Agent in the setting of clinical trials. We aim to demonstrate the superior predictive capabilities of our model by comparing it with established baseline methods.

4.1. Baseline Methods

To ensure a comprehensive evaluation, we have selected diverse baseline methods known for their robustness in similar tasks:

1. **Gradient-Boosted Decision Trees (GBDT):** This method integrates embeddings for drugs, diseases, and eligibility criteria derived from BioBERT. The concatenated embeddings are then processed using LightGBM, a popular gradient-boosting framework that is highly efficient and scalable, making it suitable for handling complex datasets typical in clinical trials.
2. **Hierarchical Attention Transformer (HAtten):** Employing BioBERT embeddings for drugs and diseases, this model introduces a hierarchical attention mechanism. It systematically focuses on different granularity levels, from entire paragraphs to specific sentences within the eligibility criteria, enhancing its ability to discern relevant information. The process culminates in a two-layer Multilayer Perceptron (MLP), which aids in refining the decision process.
3. **Standard Prompting:** As a control, this baseline employs large language models (LLMs) GPT-4 (OpenAI, 2022) in their standard configuration. It tests the hypothesis that without tailored adaptations or integrations of external data, the pre-trained knowledge embedded within LLMs can competently perform outcome prediction in clinical trials, albeit potentially less effectively than more specialized approaches.

This comparative analysis will help in highlighting the strengths and potential areas for improvement in CT-Agent, guiding future enhancements in the model’s architecture and its application in clinical trial settings.

4.2. Experimental Setup

Our experimental framework was implemented on a server equipped with an AMD Ryzen 9 3950X CPU, 64GB RAM, and an NVIDIA RTX 3080 Ti GPU. We utilized Python 3.8 for scripting and PyTorch for model implementation and training. For each experiment, we used the same seed to ensure reproducibility.

4.2.1. DATA AND RESOURCES

In our investigation, we employed a variety of datasets and external resources to construct a comprehensive experimental framework. These are detailed as follows:

- **Drug Databases:** We used *DrugBank*, a comprehensive, freely accessible online database containing information on drugs and drug targets. DrugBank is instrumental for acquiring detailed pharmacological information, which supports the pharmacokinetics and molecular mechanisms of drug actions in our models.
- **Knowledge Graphs:** *Hetionet*, an integrative knowledge graph of biomedical information that interlinks biological entities through relationships, was utilized. It provides a structured form of data that helps in understanding complex drug-disease relationships, gene interactions, and more, which are crucial for our predictive analytics.
- **ClinicalTrials.gov:** We extracted data from <https://clinicaltrials.gov/>, which includes information from both completed and ongoing clinical trials. This data is essential for validating our predictive models and for training them to understand clinical outcomes based on past trial data.
- **GPT-Generated Data:** To augment our dataset and test the robustness of our models under varied scenarios, we generated synthetic data. This was done by prompting Generative Pre-trained Transformer (GPT) models with scenarios derived from our existing datasets, thereby creating realistic, hypothetical data scenarios for further testing.

These diverse resources were meticulously integrated into our model, referred to hereafter as CT-Agent. This integration helps in providing a rich and informed context for each predictive task undertaken by our model.

For our experimental validation, we randomly selected 40 training samples from the clinical trial outcome prediction benchmark provided in (Fu et al., 2022). An identical approach was used to select 40 samples from the test set, ensuring that our training and testing datasets were balanced and randomized, thereby providing a robust evaluation framework for our model’s performance.

4.2.2. PROCEDURE

Each agent in CT-Agent was tasked with specific roles, as outlined in the Methods section. The Safety Agent queried drug databases, the Efficacy Agent analyzed disease information, the Enrollment Agent predicted the enrollment difficulty, and the Planning Agent synthesized these findings into actionable insights. We conducted experiments to assess the accuracy of outcome predictions and the precision of failure reason identifications.

Table 1: A real example of CT-Agent on clinical trial outcome prediction.

User Input Problem: I have designed a clinical trial and hope you can help me predict whether this trial can pass. Features contain (1) drug: Aggrenox capsule; (2) disease: cerebrovascular accident; (3) inclusion criteria: ...; (4) exclusion criteria: ...;
Planning (Planning Agent): few-shot learning Thought: Decompose a complex problem into several easier subproblems. Subproblem 1 (Enrollment Agent): Determine the level of enrollment feasibility based on inclusion and exclusion criteria. Subproblem 2 (Safety Agent): Evaluate the safety of the drug “aggrenox capsule” Subproblem 3 (Efficacy Agent): Assess the efficacy of the drug “aggrenox capsule” on the disease “cerebrovascular accident”
Solving Subproblems: Subproblem 1 (Enrollment Agent): Detect the function “enrollment prediction model” should to be called. The predicted enrollment failure rate is 0.3597. Subproblem 2 (Safety Agent): Detect functions “disease historical statistics” and “drug historical statistics” should be called. The historical failure rate of Aggrenox capsules in clinical trials is reported as 1.0. Subproblem 3 (Efficacy Agent): Detect functions “retrieve Hetionet knowledge graph” and “retrieve DrugBank database” should be called.
Reasoning (Reasoning Agent): few-shot learning 1. Enrollment Difficulty (Enrollment Agent): A moderate level of difficulty in enrolling participants. This could pose challenges in recruiting a sufficient number of eligible patients meeting the strict inclusion and exclusion criteria set for the trial, which can impact the overall success rate of the trial. 2. Safety of Aggrenox Capsule (Safety Agent): The historical failure rate of Aggrenox capsules (1.0) indicates a lack of success in previous trials. Evaluating the safety profile of the Aggrenox capsule in the context of cerebral infarction and cerebrovascular accidents is crucial due to the specified risk factors and exclusion criteria. The presence of risk factors like hypertension, diabetes, and other conditions among the trial participants might raise concerns about potential interactions or adverse effects of Aggrenox, contributing to the historical failure rate. 3. Efficacy of Aggrenox Capsule (Efficacy Agent): The Aggrenox capsule, combining aspirin and dipyridamole, is beneficial in preventing strokes by its antiplatelet and vasodilatory effects. Given the criteria provided for the clinical trial specifically targeting patients with a history of cerebral infarction who meet certain age and risk factor requirements, Aggrenox capsule’s mechanism of action aligns with the needs of this patient population. Draw Conclusion. Therefore, considering the historical failure rate, safety concerns, and the moderate enrollment difficulty level associated with the clinical trial design and the use of the Aggrenox capsule in patients with cerebral infarction, the predicted success rate of the trial is low, at 0.0. (groundtruth is 0)

4.3. Implementation Details

In this section, we provide detailed descriptions of the implementation processes to enhance the reproducibility of our study.

Role Assignment to Agents Each agent within the CT-Agent framework is designated a specific role, which is integrated directly into the LLM’s system prompt for clarity and focus. For instance, the role of the Efficacy Agent is defined as follows:

”As an efficacy expert, you have the capability to assess a drug’s efficacy against diseases by examining its effectiveness on the disease.”

This role definition is crucial as it guides the LLM to prioritize responses based on the assigned expert domain, leveraging the model’s inherent capability to focus more acutely on instructed tasks than on general information.

Defining External Tools External tools are defined in a structured format to facilitate their integration and usage within the LLM environment. These definitions are crafted in JSON format, specifying the function name, description, and necessary parameters. Examples can be found in the [A.2](#).

This structured approach allows for the direct transmission of function calls to the LLM, which in turn provides detailed responses including the function name and arguments. These responses enable the execution of functions locally and the retrieval of results in a structured manner.

Enhanced Few-Shot Reasoning To improve the model’s reasoning capabilities, we incorporate examples of sub-problems and corresponding labels within the system prompt. This method, known as few-shot learning, aids the LLM in understanding the context and methodology required to solve complex problems by referencing similar, previously solved problems. This approach not only enhances the accuracy of the model’s outputs but also its ability to generalize from limited examples to new, unseen scenarios.

These implementation strategies collectively ensure that each component of CT-Agent operates effectively and that the integration between different agents and external tools is seamless, fostering an environment conducive to robust, reproducible research.

4.4. Quantitative Results

Table 2 presents the performance of various methods. Specifically, CT-Agent obtain the highest ROC-AUC score at 0.8347 among all the compared methods. We observe that CT-Agent achieve competitive performance among all the well-established methods, e.g., GBDT. Also, compared with the standard prompt (basic GPT) model, our method consistently improves all six evaluation metrics. We also compare the performance of GPT-3.5 and GPT-4 for the standard prompting method in clinical trial outcome prediction. As observed in Table 2, GPT-4 demonstrates a superior performance across most metrics compared to GPT-3.5, suggesting that newer versions of large language models may offer incremental improvements in predicting clinical trial outcomes using the standard prompting method.

Table 2: Predictive performance of various methods.

Method	Accuracy	ROC-AUC	PR-AUC	Precision	Recall	F1
GBDT	0.6250	0.8000	0.8669	0.6250	1.0000	0.7692
HAtten	0.7500	0.7573	0.8718	0.8947	0.6800	0.7727
GPT-3.5	0.5250	0.4853	0.4419	0.4000	0.5333	0.4571
GPT-4	0.6500	0.6800	0.4582	0.5385	0.4666	0.5000
CT-Agent	0.7000	0.8347	0.7908	0.5714	0.8000	0.6667

4.5. Case Study

We analyze a realistic case study in Table 1. The NCTID is NCT00311402. The trial focused on evaluating the treatment effect of the Aggrenox capsule on cerebrovascular accidents. Aggrenox is a popular drug and contains a combination of aspirin and dipyridamole. First, the user describes the problem using natural language: “I have designed a clinical trial and hope you can help me predict whether this trial can pass” and attach the drug name, disease name, and inclusion/exclusion criteria as features. Then, the Planning Agent decomposes the whole problem into three subproblems based on clinical knowledge. It was enhanced by few-shot learning, which gives some representative examples to the GPT as prompt. These three subproblems are clinical trial enrollment, Aggrenox’s safety to human bodies, and Aggrenox’s efficacy in treating cerebrovascular accidents, which are handled by enrollment agents, safety agents, and efficacy agents, respectively. These agents of subproblems can be solved with the help of external tools. For example, we train an enrollment success prediction model and the estimated enrollment failure probability is 0.359. Enrollment Agent automatically recognizes and call the predictive AI model and insert the results (0.359) into the text. Similarly, Safety Agent (drug safety) identifies that the historical failure rate of Aggrenox capsules is 100%, indicating its high risk. Combining this information, the Reasoning Agent will make the final decision that the trial is highly likely to fail, which is correctly forecasted.

4.6. Ablation Study

4.6.1. DIFFERENT VERSIONS OF GPT

In this ablation study, we assess the effectiveness of two iterations of the Generative Pre-trained Transformer: GPT-3.5 and GPT-4. Our objective is to explore their performance in the context of predicting outcomes in clinical trials using a standard prompting approach. The comparative analysis focuses on key performance metrics, as detailed in the subsequent table.

The results, as summarized in Table 2, distinctly illustrate that GPT-4 outperforms GPT-3.5 across the majority of evaluated metrics. This enhancement in performance with GPT-4 underscores the potential benefits of integrating more advanced versions of large language models in the domain of clinical trial outcome prediction.

4.6.2. IMPACT OF FEW-SHOT LEARNING

This segment of the ablation study investigates the influence of few-shot learning techniques when integrated into CT-Agent, a multi-agent framework employing large language models (LLMs). Our comparative analysis pits the version of the model that incorporates few-shot learning against a baseline version that does not utilize these adaptations. The comparative results are encapsulated in the table below.

Table 3: Impact of few-shot learning on CT-Agent performance.

Method	Accuracy	ROC-AUC	PR-AUC	Precision	Recall	F1
CT-Agent w few-shot	0.7	0.8347	0.7908	0.5714	0.8	0.6667
CT-Agent w/o few-shot	0.75	0.824	0.6793	0.647	0.7333	0.6875

Table 3 reveals that while the accuracy and F1 score marginally favor the model without few-shot learning, the few-shot adapted model (CT-Agent) exhibits superior performance in terms of ROC-AUC and PR-AUC. This indicates that the incorporation of few-shot learning significantly bolsters the model’s proficiency in accurately classifying positive instances, despite a minor trade-off in overall accuracy and precision-recall balance.

5. Discussion

This study introduces the groundbreaking Multi-Agent Clinical Trial Helper (CT-Agent), a multi-agent framework that synergizes the advanced capabilities of GPT-4 with sophisticated agent architectures and cutting-edge reasoning technologies like LEAST-TO-MOST and ReAct. Our system significantly enhances the performance of large language models (LLMs) in clinical settings, managing complex trial processes and introducing novel functionalities such as predictive analytics, comprehensive failure analysis, and precise trial duration estimations.

Our evaluations, which include computational benchmarks and expert feedback, underscore the efficiency and effectiveness of CT-Agent in improving clinical trial outcomes. This integration of LLMs with multi-agent systems not only manages the complexities inherent in clinical trials but also bridges the gap between conversational AI and actionable intelligence in healthcare. CT-Agent establishes a new benchmark in the application of LLMs to clinical trials, promising a future where advanced AI tools play a crucial role in advancing medical research and patient care.

Limitations While the CT-Agent demonstrates the capability to automatically recognize and decompose user issues, directing them to specialized agents for resolution, it still relies significantly on human intervention for its design and configuration. This dependency on manual input for agent creation limits the system’s scalability and adaptability, particularly in dynamic environments where user requirements and contexts evolve rapidly. Further development could focus on integrating machine learning techniques to enable the CT-Agent to learn from interactions and autonomously update its problem-solving strategies, thereby reducing the need for frequent human oversight and redesign. These points underscore the

need for continuous research and development to fully realize the potential of AI-driven clinical trials, addressing both technical and clinical implications.

References

- Yi-Tan Chang, Eric P Hoffman, Guoqiang Yu, David M Herrington, Robert Clarke, Chiung-Ting Wu, Lulu Chen, and Yue Wang. Integrated identification of disease specific pathways using multi-omics data. *bioRxiv*, page 666065, 2019.
- Lulu Chen, Chiung-Ting Wu, Robert Clarke, Guoqiang Yu, Jennifer E Van Eyk, David M Herrington, and Yue Wang. Data-driven detection of subtype-specific differentially expressed genes. *Scientific reports*, 11(1):332, 2021.
- Tianyi Chen, Nan Hao, Yingzhou Lu, and Capucine Van Rechem. Uncertainty quantification on clinical trial outcome prediction. *arXiv preprint arXiv:2401.03482*, 2024.
- Tianfan Fu, Kexin Huang, Cao Xiao, Lucas M Glass, and Jimeng Sun. HINT: Hierarchical interaction network for clinical-trial-outcome predictions. *Patterns*, 3(4):100445, 2022.
- Junyi Gao, Cao Xiao, Lucas M Glass, and Jimeng Sun. COMPOSE: Cross-modal pseudo-siamese network for patient trial matching. In *Proceedings of the 26th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pages 803–812, 2020.
- Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. Measuring massive multitask language understanding. In *9th International Conference on Learning Representations, ICLR 2021, Virtual Event, Austria, May 3-7, 2021*. OpenReview.net, 2021. URL <https://openreview.net/forum?id=d7KBjmI3GmQ>.
- Daniel Scott Himmelstein, Antoine Lizée, Christine Hessler, Leo Brueggeman, Sabrina L Chen, Dexter Hadley, Ari Green, Pouya Khankhanian, and Sergio E Baranzini. Systematic integration of biomedical knowledge prioritizes drugs for repurposing. *Elife*, 6:e26726, 2017.
- Di Jin, Eileen Pan, Nassim Oufattole, Wei-Hung Weng, Hanyi Fang, and Peter Szolovits. What disease does this patient have? A large-scale open domain question answering dataset from medical exams. *CoRR*, abs/2009.13081, 2020. URL <https://arxiv.org/abs/2009.13081>.
- Qiao Jin, Bhuwan Dhingra, Zhengping Liu, William W. Cohen, and Xinghua Lu. Pubmedqa: A dataset for biomedical research question answering. In Kentaro Inui, Jing Jiang, Vincent Ng, and Xiaojun Wan, editors, *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing, EMNLP-IJCNLP 2019, Hong Kong, China, November 3-7, 2019*, pages 2567–2577. Association for Computational Linguistics, 2019. doi: 10.18653/V1/D19-1259. URL <https://doi.org/10.18653/v1/D19-1259>.
- Tiffany H Kung, Morgan Cheatham, Arielle Medenilla, Czarina Sillos, Lorie De Leon, Camille Elepaño, Maria Madriaga, Rimel Aggabao, Giezel Diaz-Candido, James Maningo, et al. Performance of chatgpt on usml: Potential for ai-assisted medical education using large language models. *PLoS digital health*, 2(2):e0000198, 2023.
- Jinhyuk Lee, Wonjin Yoon, Sungdong Kim, Donghyeon Kim, Sunkyu Kim, Chan Ho So, and Jaewoo Kang. BioBERT: a pre-trained biomedical language representation model for biomedical text mining. *Bioinformatics*, 36(4):1234–1240, 2020.

- Jianning Li, Amin Dada, Behrus Puladi, Jens Kleesiek, and Jan Egger. Chatgpt in healthcare: a taxonomy and systematic review. *Computer Methods and Programs in Biomedicine*, page 108013, 2024.
- Jialin Liu, Changyu Wang, and Siru Liu. Utility of chatgpt in clinical practice. *Journal of Medical Internet Research*, 25:e48568, 2023.
- Valentin Liévin, Christoffer Egeberg Hother, and Ole Winther. Can large language models reason about medical questions?, 2023.
- Cai Long, Kayle Lowe, Jessica Zhang, André dos Santos, Alaa Alanazi, Daniel O’Brien, Erin Wright, and David Cote. A novel evaluation model for assessing chatgpt on otolaryngology–head and neck surgery certification examinations: Performance study. *JMIR Medical Education*, 10, 2023. URL <https://api.semanticscholar.org/CorpusID:265110947>.
- Yingzhou Lu, Chiung-Ting Wu, Sarah J Parker, Zuolin Cheng, Georgia Saylor, Jennifer E Van Eyk, Guoqiang Yu, Robert Clarke, David M Herrington, and Yue Wang. COT: an efficient and accurate method for detecting marker genes among many subtypes. *Bioinformatics Advances*, 2(1):vbac037, 2022.
- Yingzhou Lu, Tianyi Chen, Nan Hao, Capucine Van Rechem, Jintai Chen, and Tianfan Fu. Uncertainty quantification and interpretability for clinical trial approval prediction. *Health Data Science*, 2024.
- Renqian Luo, Liai Sun, Yingce Xia, Tao Qin, Sheng Zhang, Hoifung Poon, and Tie-Yan Liu. Biogpt: generative pre-trained transformer for biomedical text generation and mining. *Briefings in bioinformatics*, 23(6):bbac409, 2022.
- OpenAI. Introducing chatgpt, 2022. URL <https://openai.com/blog/chatgpt>. Accessed: 2023-05-11.
- OpenAI. Gpt-4 technical report, 2023.
- Karan Singhal, Shekoofeh Azizi, Tao Tu, S Sara Mahdavi, Jason Wei, Hyung Won Chung, Nathan Scales, Ajay Tanwani, Heather Cole-Lewis, Stephen Pfohl, et al. Large language models encode clinical knowledge. *Nature*, 620(7972):172–180, 2023a.
- Karan Singhal, Tao Tu, JuraJ Gottweis, Rory Sayres, Ellery Wulczyn, Le Hou, Kevin Clark, Stephen Pfohl, Heather Cole-Lewis, Darlene Neal, et al. Towards expert-level medical question answering with large language models. *arXiv preprint arXiv:2305.09617*, 2023b.
- David Vilares and Carlos Gómez-Rodríguez. HEAD-QA: A healthcare dataset for complex reasoning. In Anna Korhonen, David R. Traum, and Lluís Màrquez, editors, *Proceedings of the 57th Conference of the Association for Computational Linguistics, ACL 2019, Florence, Italy, July 28- August 2, 2019, Volume 1: Long Papers*, pages 960–966. Association for Computational Linguistics, 2019. doi: 10.18653/V1/P19-1092. URL <https://doi.org/10.18653/v1/p19-1092>.
- Yue Wang, Yingzhou Lu, Yinlong Xu, Zihan Ma, Hongxia Xu, Bang Du, Honghao Gao, and Jian Wu. Twin-gpt: Digital twins for clinical trials via large language model. *arXiv preprint arXiv:2404.01273*, 2024.
- David S Wishart, Yannick D Feunang, An C Guo, Elvis J Lo, Ana Marcu, Jason R Grant, Tanvir Sajed, Daniel Johnson, Carin Li, Zinat Sayeeda, et al. Drugbank 5.0: a major update to the drugbank database for 2018. *Nucleic acids research*, 46(D1):D1074–D1082, 2018.

- Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. React: Synergizing reasoning and acting in language models. *arXiv preprint arXiv:2210.03629*, 2022.
- Ling Yue, Jonathan Li, Md Zabirul Islam, Bolun Xia, Tianfan Fu, and Jintai Chen. Trialdura: Hierarchical attention transformer for interpretable clinical trial duration prediction. *arxiv*, 2024.
- Li Yunxiang, Li Zihan, Zhang Kai, Dan Ruilong, and Zhang You. Chatdoctor: A medical chat model fine-tuned on llama model using medical domain knowledge. *arXiv preprint arXiv:2303.14070*, 2023.
- Bai Zhang, Yi Fu, Zhen Zhang, Robert Clarke, Jennifer E Van Eyk, David M Herrington, and Yue Wang. DDN2.0: R and python packages for differential dependency network analysis of biological systems. *bioRxiv*, pages 2021–04, 2021.
- Xingyao Zhang, Cao Xiao, Lucas M Glass, and Jimeng Sun. Deepenroll: Patient-trial matching with deep embedding and entailment prediction. In *Proceedings of The Web Conference 2020*, pages 1029–1037, 2020.
- Denny Zhou, Nathanael Schärli, Le Hou, Jason Wei, Nathan Scales, Xuezhi Wang, Dale Schuurmans, Claire Cui, Olivier Bousquet, Quoc Le, et al. Least-to-most prompting enables complex reasoning in large language models. *arXiv preprint arXiv:2205.10625*, 2022.

Appendix A.

A.1. Enrollment Model

Model Architecture The architecture of the model consists of the following components

- A transformer encoder layer that processes embeddings of inclusion and exclusion criteria, drugs, and diseases. This layer is designed to capture interactions across these different types of information, crucial for understanding the complexity of trial eligibility.
- A fully connected layer that maps the high-dimensional features from the transformer encoder to a single output, indicating the probability of successful enrollment.
- A sigmoid activation function applied to the output of the fully connected layer, converting it into a probability measure.

Data Processing and Feature Extraction For each trial, the Enrollment Agent performs the following steps:

- **Criteria Segmentation:** The agent first segments the eligibility criteria into inclusion and exclusion categories. This segmentation allows for targeted analysis of factors that either qualify or disqualify potential participants.
- **Embedding Generation:** Using the pretrained BioBERT model, the agent converts text data from the criteria, drugs, and diseases into dense vector representations. These embeddings capture deep semantic features that are essential for accurate model predictions.
- **Feature Aggregation:** The embeddings are then aggregated and fed into the transformer encoder, which processes them to capture the complex dependencies among the criteria, drugs, and diseases.

Model Training and Evaluation The model is trained on a dataset comprising historical trial data, where each record includes the eligibility criteria, associated drugs and diseases, and the trial outcome regarding enrollment success. The training process involves:

- Balancing the dataset to handle disparities in the number of successful and unsuccessful enrollments.
- Employing a cross-entropy loss function adjusted for class imbalance, ensuring that the model accurately learns from both positive (successful enrollment) and negative (unsuccessful enrollment) examples.
- Evaluation on the test dataset yielded an ROC-AUC score of 0.7037 and an accuracy of 0.7689, indicating effective class differentiation and prediction accuracy. This model exemplifies the integration of advanced NLP and neural architectures to improve clinical trial design and efficiency.

A.2. Function Calling Definitions

The function for retrieving drug information from the DrugBank database is defined as follows:

```
{
  "type": "function",
  "function": {
    "name": "retrieval_drugbank",
    "description": "Retrieves information about a drug from the DrugBank database
                  using the drug's name as input.",
    "parameters": {
      "type": "object",
      "properties": {
        "drug_name": {
          "type": "string",
          "description": "The name of the drug."
        }
      }
    },
    "required": ["drug_name"]
  }
}
```

The function for retrieving the path connecting the drug to the disease from the Hetionet Knowledge Graph is defined as follows:

```
{
  "type": "function",
  "function": {
    "name": "retrieval_hetionet",
    "description": "
      Given the names of a drug and a disease, the model retrieves
      the path connecting the drug to the disease from
      the Hetionet Knowledge Graph.
      Hetionet is a comprehensive knowledge graph that integrates diverse
      biological information by connecting genes, diseases, compounds,
      and more into an interoperable framework.
    "
  }
}
```

It structures real-world biomedical data into a network, facilitating advanced analysis and discovery of new insights into disease mechanisms, drug repurposing, and the genetic underpinnings of health and disease."

```
"parameters": {
  "type": "object",
  "properties": {
    "drug_name": {
      "type": "string",
      "description": "The drug name",
    },
    "disease_name": {
      "type": "string",
      "description": "The disease name",
    }
  },
  "required": ["drug_name", "disease_name"],
},
}
```