
BattleAgent: Multi-modal Dynamic Emulation on Historical Battles to Complement Historical Analysis

Shuhang Lin*
Rutgers University

Wenyue Hua*
Rutgers University

Lingyao Li
University of Michigan

Che-Jui Chang
Rutgers University

Lizhou Fan
University of Michigan

Jianchao Ji
Rutgers University

Hang Hua
University of Rochester

Mingyu Jin
Rutgers University

Jiebo Luo
University of Rochester

Yongfeng Zhang
Rutgers University

Abstract

This paper presents **BattleAgent**, a detailed emulation demonstration system that combines the Large Vision-Language Model (VLM) and Multi-Agent System (MAS). This novel system aims to simulate complex dynamic interactions among multiple agents, as well as between agents and their environments, over a period of time. It emulates both the decision-making processes of leaders and the viewpoints of ordinary participants, such as soldiers. The emulation showcases the current capabilities of agents, featuring fine-grained multi-modal interactions between agents and landscapes. It develops customizable agent structures to meet specific situational requirements, for example, a variety of battle-related activities like scouting and trench digging. These components collaborate to recreate historical events in a lively and comprehensive manner while offering insights into the thoughts and feelings of individuals from diverse viewpoints. The technological foundations of BattleAgent establish detailed and immersive settings for historical battles, enabling individual agents to partake in, observe, and dynamically respond to evolving battle scenarios. This methodology holds the potential to substantially deepen our understanding of historical events, particularly through individual accounts. Such initiatives can also aid historical research, as conventional historical narratives often lack documentation and prioritize the perspectives of decision-makers, thereby overlooking the experiences of ordinary individuals. This biased documentation results in a considerable gap in our historical understanding, as many stories remain untold. BattleAgent leverages the current advancements in Artificial Intelligence (AI) to provide some insights to bridge this gap. It illustrates AI's potential to revitalize the human aspect in crucial social events, thereby fostering a more nuanced collective understanding and driving the progressive development of human society. Quantitative evaluations are computed on the final emulation result, showing reasonable performance and effectiveness of the approach. The data and code for this project are accessible at <https://github.com/agiresearch/battleagent> and the demo will be released in one month.

*Shuhang Lin and Wenyue Hua have equal contributions. **Author Emails:** {shuhang.lin,wenyue.hua,yongfeng.zhang}@rutgers.edu

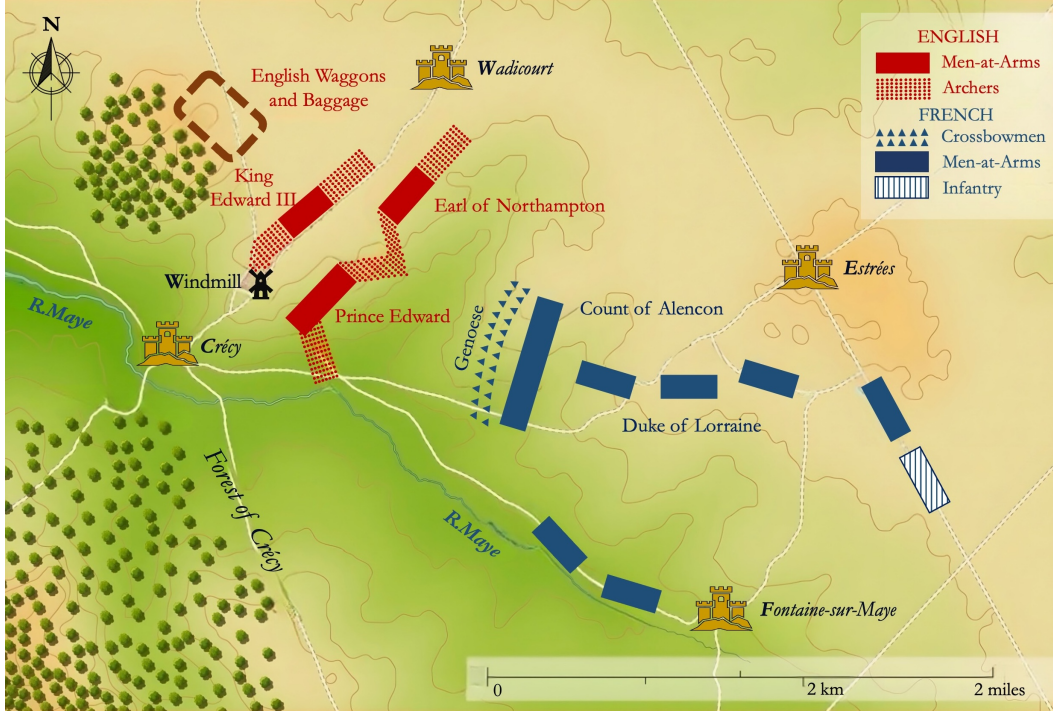


Figure 1: Demonstration of the emulated Battle of Crécy, 1346: Troop formations and movements depicting the positions of the English and French forces during the historical engagement, with key locations and leaders marked.

Image adjusted from <https://the-past.com/feature/the-battle-of-crecy-26-august-1346/>

1 Introduction

An agent is defined as a system that has the ability to perceive its surroundings and make informed decisions based on these perceptions to accomplish specific objectives [1]. The recent progress in Large Language Models (LLMs) [2, 3] has demonstrated impressive reasoning capabilities [4, 5], indicating their potential to serve as the foundation for agents. These models have shown remarkable proficiency in following instructions [6, 7], interpreting commands, and emulating human reasoning and learning processes [8, 9, 10]. Additionally, the development of large Vision Language Models (VLM) [11] has facilitated the creation of various agent applications that support multi-modal information interaction [12, 13]. When combined with external tools, either physical or virtual, these agents employ LLMs or VLM as their reasoning backbone to determine how tasks should be addressed, how tools should be utilized, and what information should be retained in memory. This enhancement equips agents to manage an array of natural language processing tasks and engage with their environment using language.

A multitude of agent applications have been created using LLM and VLM, with a focus on enhancing reasoning [14, 15, 16, 17], production capabilities [18, 19, 20, 21, 22, 23], gaming [24, 25, 26, 27], and social simulation [28, 29, 30, 31, 32], among others. WarAgent [32] is the pioneering LLM-based MAS simulation of historical events, examining the behaviors of systems at the macro level, such as nations and governments, rather than the micro-level simulation of detailed and dynamic events occurring during battles or individual experiences in such dynamic time periods. Therefore, BattleAgent, building on the foundation laid by WarAgent in historical event simulation, investigates the potential of LLM and VLM for detailed historical situation recovery and the exploration of individual experiences within the simulation.

The study of history has long been a pursuit to understand the human experience through the lens of past events. Traditional historical narratives often focus on the perspectives of leaders and decision-makers, leaving the experiences of ordinary individuals in the shadows. This selective approach to history has created a significant gap in our understanding, as the stories and experiences of the

common participants, such as soldiers, are frequently overlooked. The motivation behind this research is to address this imbalance and provide a more comprehensive view of historical events by leveraging advancements in AI. Oral history [33, 34] has been one method used to capture the experiences of individuals, offering a more personal account of historical events. However, this approach is limited to recent history and is constrained by the availability of witnesses, often leaving many details undiscovered. As we move further away from events in time, the voices of those who lived through them fade, and with them, the richness of history’s tapestry.

In response to these challenges, our study introduces BattleAgent, a novel emulation framework that utilizes LMM-based MAS for the detailed reconstruction of historical events, with an emphasis on delineating the experiences of ordinary individuals, notably soldiers. BattleAgent emulates historical combats within complex terrains and hierarchical command structures, incorporating sophisticated military logistics and strategic planning. Central to our model, we meticulously craft 30 individual soldier agents, each endowed with a richly detailed background and a distinct personality, thereby infusing them with vibrancy and depth. BattleAgent is designed to emulate and document the experiences of these agents, capturing their actions, injuries, emotional responses, and psychological states throughout the course of the battles. By analyzing these elements, we generate personalized narratives that reflect the multifaceted experiences of ordinary individuals engaged in warfare. This approach not only immortalizes the actions and sentiments of these agents but also furnishes a nuanced, individualized perspective on the common soldier’s experience within the broader tapestry of historical conflict.

To emulate such a complex scenario, our emulation incorporates the following three key features:

- **Enhanced 2-D Realism Features:** BattleAgent simulates detailed interactions within environments, including terrain engagement, temporal progression, and interactions between agents.
- **Immersive Multi-agent Interactions:** It integrates MAS to facilitate dynamic interactions among agents in battle emulations, accurately reflecting the historical milieu and the intricacies of military engagements, from strategic maneuvers to logistical considerations and communication dynamics.
- **Dynamic Agent Structure:** The framework introduces adaptable agent configurations and multi-modal interactions. The system can “self improvise” its structure to fork, merge, and prune agents to continuously maintain the emulation effectiveness. It boasts the capability to autonomously adjust its architecture to optimize emulation fidelity.

The contributions of our study to historical analysis and society can be summarized as follows:

- **Emphasis on individual perspectives and granularity:** Providing a platform for the voices of common people to be heard and understood in historical events. This platform aims to enhance the accuracy of historical reconstructions by incorporating individual perspectives.
- **Connection and resonance with the past:** Helping to prevent future conflicts by learning from the detailed analysis of past mistakes and human costs. This platform fosters empathy and a deeper connection to the past by humanizing the experiences of those involved in historical battles.
- **Educational tool for understanding history:** Providing an educational tool to help people understand the intricacies of history and the harsh realities of historical events. Its immersive and interactive platform can foster empathy and a more nuanced perspective on the past, making it a valuable resource for students and history enthusiasts.
- **Potential as a next-generation game engine:** Providing a fully automated process to create immersive and dynamic historical emulations, making it a potential next-generation game engine. By using LLM-based agents and VLM-based agents, it can generate detailed and realistic environments, characters, and events, offering a unique and engaging gaming experience.

2 Related Work

2.1 Multi-Agent System

MAS has revolutionized the landscape of AI, offering a platform for simulating complex interactions and scenarios [32]. With the development of the reasoning intelligence of LLMs, especially their outstanding reasoning ability in complex scenarios [35, 36, 37], the integration of MAS into AI systems

has demonstrated their versatility and efficacy. The initial classification of MAS into reasoning-enhancement, non-player character (NPC) multi-agent players, and production-enhancement systems has been foundational in understanding their diverse applications. Notable developments such as LLM-Debate [14], ChatEval [15], and MAD [17] have significantly advanced reasoning-enhancement systems. Similarly, in NPC multi-agent systems, the emergence of Generative Agents [38] and GPT-Bargaining [39] has paved the way for more human-like agent behaviors. In the production-enhancement domain, innovations like MetaGPT [18] and OpenAGI [20] have streamlined and enhanced collaborative efforts in software development. Many works have also explored agents’ potential in scientific experiments [40, 41, 42] setting.

In the context of humanities and historical research, the WarAgent [32] initiative has exemplified the application of LLM-based MAS for simulating international conflicts, with each agent representing a different country to explore the dynamics of international relations and conflicts. Building on the humanitarian insights gleaned from WarAgent, our study seeks to refine this approach by concentrating on the granular emulation of historical scenarios from the vantage point of ordinary individuals. This shift towards focusing on the micro-level experiences within historical events aims to provide a more detailed and empathetic understanding of the past, leveraging the advancements in MAS and LLM technologies to capture the nuanced perspectives of common people in historical narratives. By doing so, we aspire to enrich the tapestry of historical understanding with a deeper, more inclusive examination of human experiences during pivotal moments in history.

Recent advancements in multi-modal multi-agent AI systems have further expanded the capabilities of MAS. AppAgent [43] demonstrates the use of multimodal agents as smartphone users, enhancing our understanding of human-computer interactions. The integration of generative AI and multi-modal agents in AWS [44] has unlocked new potentials in financial markets. LLaVA-Plus’s contribution [45] in teaching agents to use various tools has opened up new avenues for agent adaptability and functionality. Furthermore, the implementation of “multimodal chain-of-action agents” [46] has provided a novel perspective on agent interaction with digital interfaces, contributing to more intuitive user interface designs and realistic simulations in digital realms.

BattleAgent emulation is the first large multimodal model-based multi-agent application that introduces a novel quantitative dimension to historical and humanities studies and underscores the broader impact of AI in understanding human history and shaping future scenarios. By exploring alternative historical pathways and key determinants, our work demonstrates the significant contributions of LMM and MAS in enhancing our comprehension of the past and potentially guiding a more informed and peaceful future.

2.2 Challenges in Granular Historical Analysis

The pursuit of simulating historical events using computational methods has evolved significantly over the years. Beginning with human simulations, transitioning to human-program hybrid systems, and finally arriving at fully computerized simulations, each stage has brought unique insights and challenges [32]. Human simulations, as outlined by Dickson [47], provided a foundational approach. In educational scenarios, such simulations involved role-playing exercises, enabling students to delve into the complexities of historical events, such as the United States’ entry into WWI. The advent of human-program hybrid systems, exemplified by the Inter-Nation Simulation model [48] and its various applications [49]. These systems merged human decision-making with computational processes, creating a more dynamic and interactive environment for simulating international conflicts. However, the reliance on human input still presented limitations in terms of scalability and the possible depth of analysis.

In the past decade, there has been a significant shift in leveraging computing power to create more sophisticated simulations. The OneSAF Objective System (OOS) [50] and the JAVA-based simulation of the Bay of Biscay submarine war [51] are prime examples. These simulations used detailed models of military operations and game theory, enhancing the accuracy and depth of historical analysis. More recently, through the development of generative AI methods, high-level simulation of social system dynamics becomes a reality [52]. Despite these advancements, the complexity of human behavior and the vastness of historical data remained challenging to fully encapsulate in these models. Moreover, while fully computerized simulations can achieve the most detailed and accurate simulation among all of the three stages, they still are focuses on the panoramic and high-level simulation of historical

analysis, often unable to delve deeper into the individual witnesses’ reflection and granular analysis beyond countries’ or famous leaders’ perspectives.

The ideas of “agent-based computational models” and “generative social science” have been well-known theories before the start of the 21st century [53]. Rule-based agents, while can reconstruct complex social behaviors, to some extent, to investigate “backward to the future” [52], are often unable to maintain human-like intelligence, which is the key to simulating and understanding human society. Granular historical analysis with generative methods seems to be unrealistic even with the most cutting-edge Computational Social Science (CSS) methodologies. The current landscape of CSS, particularly in fields like sentiment analysis [54, 55, 56], primarily operates on contemporary data sources. This presents a significant challenge in historical analysis, as historical data often lacks the granularity and digital format required for computational analysis.

Our research confronts this challenge by employing an LLM within a MAS framework. This approach represents a novel step in historical simulation, blending the comprehensive data processing capabilities of modern AI with the intricate modeling of MAS. This integration marks a significant departure from traditional methods, as it attempts to overcome the limitations of data scarcity and quality in historical research. By utilizing advanced language models, we can infer, reconstruct, and simulate historical narratives and events with a level of depth and accuracy previously unattainable. Thus, we refer to this granular simulation approach as “history emulation”.

Our MAS framework not only models individual agents and their interactions but also incorporates broader socio-political and economic contexts derived from limited historical data. This approach allows for a more nuanced and granular exploration of historical events, shedding light on the complex interplay of factors that shaped these occurrences. Our work, therefore, stands at the forefront of historical emulation, or as we redefined “history emulation,” offering a unique blend of AI-powered analysis and traditional historical scholarship. This synergy aims to provide new perspectives on historical events, facilitating a granular, diverse, and deeper understanding of the past and its implications for the future.

3 Emulation Setting

This section outlines the emulation framework and setting for our research demonstration. We commence with an exposition of the historical context of the four significant European battles that our emulation seeks to emulate: the Battle of Crécy, the Battle of Agincourt, the Battle of Poitiers, and the Battle of Falkirk. Each battle has been selected for its notable use of cold weapons and the strategic bipartite confrontations that characterized warfare during their respective periods. Building upon the historical context, we elaborate on the configuration of agents and their designated roles within our emulation framework. Our model incorporates two distinct categories of agents to capture the complexity of the battlefield: commanding agents and soldier agents. This dual-layered agent structure enhances the emulation’s fidelity, offering nuanced insights into the command and control hierarchies, as well as the individual soldier experiences of historical warfare. To accurately emulate the historical engagements, each agent type has a specific set of actions available.

3.1 Battle Histories

Battle of Crécy [57] a pivotal clash during the Hundred Years’ War, occurred on 26 August 1346 in northern France. The English forces under King Edward III confronted the French army led by King Philip VI. The English army is believed to have consisted of *around 10,000 to 15,000* men, while the French forces are estimated to have been *between 20,000 and 35,000* strong. As the English army marched through northern France, they were assaulted by the French, leading to a decisive English victory marked by substantial French casualties. As for casualties, the English losses were relatively light, with estimates ranging from *a few hundred to around 2,000* men. On the other hand, the French suffered heavy losses, with estimates suggesting that *between 10,000 and 15,000* French soldiers were killed, including many high-ranking nobles. The Battle of Crécy marked a significant turning point in the Hundred Years’ War and demonstrated the effectiveness of the longbow against traditional knightly charges.

Battle of Agincourt [58] fought on 25 October 1415 near Azincourt in northern France, stands as a landmark English triumph in the Hundred Years’ War. The English army consisted of *around 6,000*



Figure 2: Battle of Crécy, from an illuminated manuscript of Jean Froissart's Chronicles.

men, primarily made up of archers armed with longbows. The French, on the other hand, had a force of *approximately 12,000 to 36,000 men*, composed of knights, men-at-arms, and infantry. Estimates of casualties vary, but it is generally agreed that the English suffered relatively light losses, with *up to 600 men killed*. In contrast, the French suffered heavy casualties, with estimates suggesting that *between 4,000 and 10,000 French soldiers* were killed, including many high-ranking nobles. The confrontation took place on Saint Crispin's Day and ended with a surprising victory for the outnumbered English forces. This victory significantly elevated English morale and status, inflicted severe damage on France, and initiated a phase of English ascendancy in the war that persisted for 14 years. This period ended with England's defeat at the hands of France in 1429 during the Siege of Orléans.

Battle of Poitiers [59] On 19 September 1356, the Battle of Poitiers was waged between the French army, commanded by King John II, and an Anglo-Gascon contingent led by Edward, the Black Prince. The battle unfolded in western France, near Poitiers, where a French force numbering *between 14,000 to 16,000 men* launched an attack against a strategically fortified position held by a *6,000-strong* Anglo-Gascon army. It is generally agreed that the English suffered relatively light losses, with *around 40 men-at-arms* killed. In contrast, the French suffered heavy casualties, with estimates suggesting that *around 4,500 French soldiers* were killed, including many high-ranking nobles.

Battle of Falkirk [60] a major engagement in the First War of Scottish Independence, took place on 22 July 1298. The English army, commanded by King Edward I, achieved a significant victory over the Scots, who were led by William Wallace. The English army consisted of *around 12,000 to 15,000 men*, primarily made up of heavy infantry and archers. The Scottish forces were believed to have numbered *around 5,000 to 8,000 men*, composed mainly of spearmen and some cavalry. English had *around 2,000 men* killed. In contrast, the Scottish suffered heavy casualties, with estimates suggesting that *between 2,000 Scottish soldiers* were killed, including many high-ranking nobles. The defeat at Falkirk led to Wallace's resignation as Guardian of Scotland.

3.2 Agent Definition

Our basic emulation setting contains the profile definition of agents and their action space. Our emulation encompasses two types of agents:

1. **Commanding Agents:** These agents symbolize the individuals or collective entities responsible for strategic decision-making in the theater of war. They are programmed to emulate the tactical acumen and leadership styles of historical commanders, thereby influencing the broader outcome of the simulated conflict.
2. **Soldier Agents:** These agents personify the rank-and-file soldiers who executed the battle plans on the ground. Each soldier agent is equipped with a comprehensive profile that includes personal history, psychological traits, and potential responses to combat stimuli.

Commanding Agents These agent profiles include general information about an army (a group of soldiers). Decisions and Strategies of the commanding agent will be made based on the whole army information, which includes the following aspects:

1. **ID:** The ID of a commanding agent is represented by a hash code that is generated to uniquely identify each commanding agent within the emulation sandbox. This is necessary due to the dynamic agent structure employed in our emulation, which allows for the creation of additional agents beyond the initial (two) commanding agents as the emulation progresses. The use of a hash code ensures that each agent can be accurately identified and tracked throughout the course of the emulation.
2. **Military Command Structure:** This involves the hierarchical organization and leadership dynamics within each military faction.
3. **Morale and Discipline:** An assessment of the troops' psychological readiness, their discipline levels, and overall morale.
4. **Military Strategy:** The overarching tactical approaches and plans employed by each side in the conflict.
5. **Military Capability:** An inventory of the weapons and defense tools at each side's disposal.
6. **Force size and composition:** This aspect includes the total number of soldiers and their composition including information about the types of troops, their roles, and their proportions in the overall force.
7. **Location:** The current location of the agent is represented by its coordinates. These coordinates provide a precise indication of the agent's position within the simulated environment, allowing for accurate tracking and analysis of its movements and interactions with other agents and the environment.

Below is an example of an English army agent which is initialized at the very start of the emulation for the Battle of Crécy:

```

1 1. ID: ARMY-13e370a8
2
3 2. Command Structure: experienced commanders with significant autonomy
4
5 3. Morale and Discipline:
6 (1) High morale and strict discipline
7 (2) Enhanced by tactical innovations and effective use of the longbow
8
9 4. Military Strategy:
10 (1) Aggressively defensive posture, emphasizing strategic high ground
    for offensive strikes. Vigorously exploits terrain advantages for
    combative engagements.
11 (2) Forceful application of longbows, enabling a confrontational yet
    adaptive offense, underscored by ambitious tactical innovations.
12
13 5. Military Capability:
14 (1) Longbow: Exhibits a rapid rate of firing and extensive range, with
    the capability to pierce armor.
15 (2) Gunpowder Weapons: Encompasses a range of types, incorporating
    dismounted men-at-arms and selective deployment of cannons.
16
17 6. Force Size and Composition:
18 Size: 10,000. Includes men-at-arms, longbowmen, hobelars, and spearmen
    .

```

```

19
20 6. Armament and Protection:
21 (1) Armor: Men-at-arms wore quilted gambeson under mail armor,
    supplemented by plate armor, bascinets with movable visors, and
    mail for throat, neck, and shoulders.
22 (2) Shields: Heater shields made from thin wood overlaid with leather.
23 (3) Weapons: included lances used as pikes, swords, and battle axes.
24 (4) Special Weapon: Longbow.
25
26 7. Location: [0, 0]

```

Listing 1: An example of commanding agent profile

Soldier Agents Soldier agents are characterized by a wide range of attributes, including name, age, familial ties, occupation, personality traits, social standing, potential health conditions, physical fitness, hobbies and interests, conversational style, unique idiosyncrasies, and personal secrets or controversies. These attributes are designed to provide a comprehensive and nuanced representation of the soldier agents.

3.3 Anonymization of Battle Details

To prevent providing explicit hints to the LLM and VLM regarding the specific battle being emulated, in case it has memorized certain battle information, we anonymize various battle details. This includes the names of countries, leaders, specific year and date of the battle, and location names. By doing so, we ensure that the LLM’s decisions are based solely on the information provided in the prompt and not influenced by any prior knowledge it may have.

3.4 Action Space

Action Space of Commanding Agents Our emulation framework is characterized by its comprehensive detail, featuring an action space that encompasses 51 distinct actions. Agents within the emulation have the flexibility to select any combination of these actions at each decision point. The actions available in the action space are organized into seven categorically distinct groups:

1. Reposition. This category includes actions that involve the movement of an army or a subsection thereof to a different location: *Reposition Forces, Create Decoy Units*
2. Preparation. Actions in this group are geared towards readying forces for an impending attack: *Deploy Longbows, Rally Troops, Employ Artillery, Use of Gunpowder Weapons, Resupply Archers, Destroy Enemy Morale, Deploy Archers in Flanking Positions, Organize Night Raids, Organize Raiding Parties, Digging trenches*
3. Attack. This group encapsulates a variety of common attack strategies, such as skirmishing, ambushing, besieging, cavalry charges, and direct firing, among others: *Initiate Skirmish, Charge Cavalry, Ambush Enemy, Launch Full Assault, Archery Duel, Siege Tactics, Hand-to-Hand Combat, Counterattack, Conduct Reconnaissance, Direct Artillery Fire, Engage in Siege Warfare, Execute Flanking Maneuvers, Use Cavalry for Shock Tactics, Employ Archers Strategically*
4. Defense. Encompasses actions such as shielding, fortification, and the creation of obstacles: *Construct Defenses, Prepare Defenses, Develop Counter-Siege Measures, Form Defensive Shields, Establish Defensive Fortifications, Fortify Rear Guards, Fortify Position, Create Obstacles for Enemy Cavalry, Form Defensive Pike Formations, Set Traps*
5. Observation. Focused on gathering information about the surrounding area and the current situation of the enemy: *Scout Enemy Position, Gather Intelligence, Intercept Enemy Supplies, Establish Communication Lines*
6. Retreat. Actions related to strategic withdrawal in the face of adverse conditions: *Retreat and Regroup, Tactical Retreat, Plan Feigned Retreat*

Each action requires 3 inputs: initiator, location, and recipient. For each action, the agent is required to (1) identify the specific agent or entity responsible for emulation of the action. It defines who is carrying out the action within the emulation. (2) the geographical point or area within the emulation

where the action takes place. It is crucial for contextualizing the action within the broader landscape of the battle scenario. (3) the target or beneficiary of the action. It specifies towards whom or what the action is directed, whether it is an opposing force, a specific unit, or another entity within the emulation.

Action Space of Soldier Agents These soldier agents follow the orders given by their commanding agents and are affected by combat outcomes based on the general fatality rate, which is calculated according to the actions taken, the size of the armies, and the types of weapons involved. By going through this process, the individual agents provide a more granular view of the battlefield dynamics, allowing for a better understanding of the experiences and feelings of soldiers on the ground.

4 Emulation Sandbox

In our emulation framework, we concentrate on a relatively straightforward scenario: a bipartite battle. The process begins with the establishment of the geographical context for the entire scenario, both textual description as well as a visual map. Subsequently, we define the two initial opposing commanding agents, considering seven profile aspects to ensure a comprehensive and realistic representation. Notice that in this section, we simply use “agent” to refer to “commanding agent”.

4.1 Time and Space in Sandbox

In order to accurately simulate the dynamics of historical battles, it is crucial to effectively manage the time and space within the sandbox environment. In this section, we will discuss our approach to time and space management in the sandbox.

Quantized Time Management The battlefield environment is characterized by continuous dynamic changes as shown in both Figure 3 and Figure 6, with agents frequently altering their actions and positions. To accurately emulate these dynamics while preserving the discrete decision-making process in our agent-based emulation, we employ a time quantization approach. Specifically, we discretize the continuous flow of time [61, 62] into 15-minute intervals. For each quantized time block, agents have the flexibility to either maintain their current action or adapt their actions based on the unfolding situation.

Coordinate Generation based on Map We obtain the initial map of the battlefield from historical documents¹. To generate the coordinates from the original image, we use one army position as the reference point, designated as the (0,0) position. We then use a scale of 10 yards as one unit of the coordinate system. The coordinates of key landscapes on the map are estimated and expressed with a range in the textual description.

For example, in the Battle of Crécy, the coordinates of a river are generated as follows:

```

1 RiverB: {
2   path:
3     {
4       start: [-100, -200],
5       end: [100, 50],
6       description: Flow path of the river begins from the southwest
          corner towards the northeast.
7     },
8   properties:
9     {
10      type: River
11    },
12   description: Meandering waterway providing natural boundaries and
          obstacles.
13 }
```

Listing 2: An example of coordinates of a river in the battlefield of Battle of Crécy

¹https://en.wikipedia.org/wiki/Battle_of_Crecy, https://en.wikipedia.org/wiki/Battle_of_Agincourt, https://en.wikipedia.org/wiki/Battle_of_Falkirk, https://en.wikipedia.org/wiki/Battle_of_Poitier

The coordinates of these key positions, including both key landscapes and existing agents on the battlefield, serve as reference points for agent movement decisions. When making decisions on agent movement, the agent will refer to the coordinates of these key positions and choose the target position to move to. This approach enables agents to navigate the battlefield in a realistic and contextually appropriate manner, taking into account the presence of natural boundaries and obstacles.

4.2 General Sandbox Process

Here we provide a very simple and crude overview of the emulation sandbox, as presented in Figure 3. We initiate the emulation based on historical map which contain information about geography as well as the position of the armies. The following represents a high-level overview of the steps involved in the emulation process:

- Step 1: Each agent starts by observing its surroundings and gathering information. This observation process involves text-based reasoning from the agent’s prompt as well as obtaining visual information from the map.
- Step 2: Based on the gathered information, each agent decides on its actions, such as preparing for battle (e.g., digging trenches, reinforcing troops), collecting further information, or making organizational agent structure changes like forking, merging, or being pruned to dynamically create new agents or diminish existent agents.
- Step 3: For every 15-minute interval in the emulation sandbox, the general information, including agent locations and properties, is updated according to the actions taken by all agents.
- The process then loops back to Step 1, with agents continuing to observe, make decisions, and act based on the updated information and evolving battlefield situation.

By following this iterative process, our agent-based emulation can effectively capture the complex dynamics of the battlefield while maintaining the discrete decision-making process inherent to the emulation.

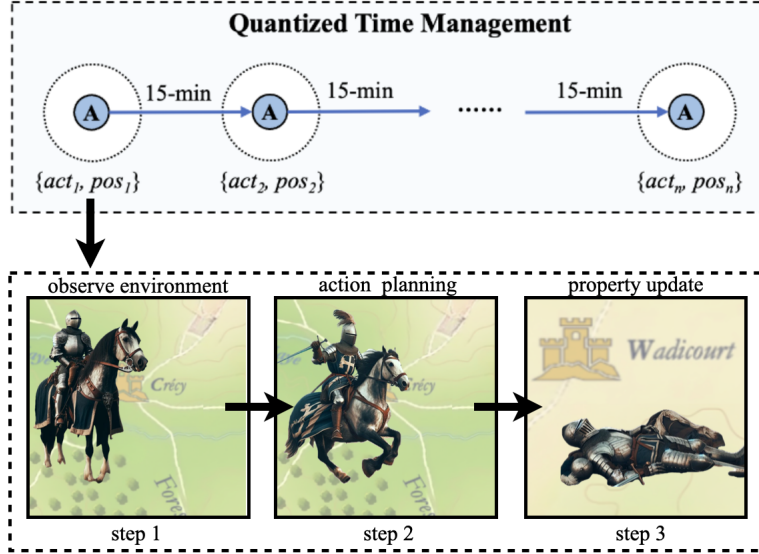


Figure 3: General Sandbox Process: The agent actions and properties are updated every 15 minutes, and for each 15-minute interval, each agent observes the environment, plans its actions, and acts based on its updated properties. This process is repeated for the duration of the simulation to emulate the dynamics of historical battles.

4.3 Detailed Emulation Process

In this subsequent section, we delve into a more comprehensive explication of the general sandbox process previously introduced, elaborating on several key aspects. These include the management of

time in a quantized manner, the process of making observations grounded in the map, the planning and execution of actions by the agents, the evaluation of casualty for each agent at every quantized time interval and the methodology behind updating the properties of each agent.

4.3.1 Observation based on Map Information



Figure 4: Observation based on Map Information.

In line with the behavior of all agents, the initial step prior to planning or executing any actions involves observing the overall environment information. To achieve this, we adopt a multi-modal approach that amalgamates both textual and visual representations of the environment. A visual map is presented to the agents, depicting their starting locations and specific places of interest, such as villages. In addition, the agents are provided with corresponding textual descriptions of the coordinates in their prompts. Subsequently, each agent proceeds to collect data in two primary aspects:

1. Geographic landscapes: Geographic information offers a macroscopic view, providing agents with critical insights for strategic planning and navigation.
2. Nearby agents: Agents information including affiliation (friend or foe), current actions, precise positioning (distance and bearing), and detailed agent profiles including the remaining number of soldiers and their composition.

Together, these observation streams form a comprehensive situational awareness framework, crucial for the agents' adaptive responses and operational effectiveness in diverse settings. An example is presented in Figure 4.

To simulate the agents' limited field of view, each agent's prompt only includes the coordinates of agents and places within a 10k-meter range. All agents then calculate the relative distance of other agents and places within their 10k-meter sight range using both textual and visual information. It is important to note that if obstacles like forests or villages come into an agent's sight, the agent will not be able to see through these obstacles. By combining textual and visual descriptions, our multi-modal approach enables agents to interact effectively with their surroundings while maintaining a realistic representation of their limited field of view.

If sub-agents collect information and communicate with the parent agent, the parent agent's knowledge can extend beyond its immediate sight range. This information exchange allows the parent agent to

make more informed decisions based on the broader context provided by its sub-agents, enabling more effective coordination and strategic planning. In such cases, the parent agent’s awareness of the battlefield is not strictly limited to its own 1,000-meter sight range but can be expanded through the information gathered by its sub-agents. This extended knowledge can include details about enemy positions, terrain features, or other relevant factors that might influence the parent agent’s decision-making process. As a result, the dynamic multi-agent system can better adapt to the complex and evolving battlefield environment by leveraging the combined knowledge of its constituent agents.

4.3.2 Action Planning

At each discrete time point, an agent has the ability to choose from a multitude of potential actions. In this part, we will outline four common actions that agents typically engage in: location movement, dynamic agent structure, interaction with the landscape, and interaction with other agents. These actions encompass a range of strategic and tactical considerations that agents must take into account when making decisions in the context of the battlefield.

Location Movement In the context of location movement, an agent possesses the capability to traverse to a different location for strategic purposes. This may involve moving closer to enemy agents to initiate an attack, or distancing itself from potential threats. In terms of the mechanics of location movement, the agent will generate the coordinates of its intended final destination, which it aims to reach within the subsequent 15-minute timeframe.

The following represents an illustrative output of a location movement action for an agent belonging to England with the ID ARMY-A606969B. In this scenario, the agent moves from its current location, represented by the coordinates [95, -55], to a new location with coordinates [100, -50]. During the course of this movement, the agent sustains losses of 30 soldiers.

```

1 {
2   identity: England,
3   id: ARMY-a606969b,
4   action_description: Reposition Forces Reposition to [100, -50] to
                        utilize ForestF for cover and longbow deployment, supporting
                        nearby friendly units and executing flanking maneuvers,
5   current_location: [95, -55],
6   target_location: [100, -50],
7   remaining_number: 170,
8   original_number: 200,
9   lost_number: 30
10 }
```

Listing 3: An example output of location movement action

Dynamic Agent Structure The battlefield environment is highly dynamic and fluid, with a multitude of situations arising unpredictably. To address this complexity, we propose a dynamic agent structure [63, 64] that enables agents to adapt their organizational configurations according to the current situation. Our proposed dynamic agent structure supports several adaptive mechanisms, as shown in Figure 5:

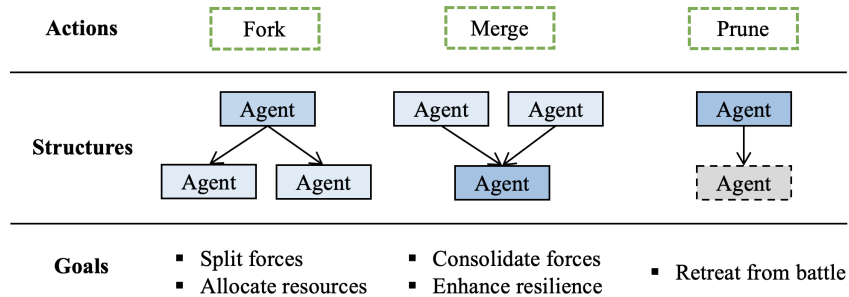


Figure 5: Dynamic agent structure.

1. Fork: An agent may decide to fork another autonomous agent for a specific task, effectively splitting its forces and allocating resources to address multiple objectives simultaneously.
2. Merge: In scenarios where an agent is under significant pressure but chooses to continue fighting, it may merge with the closest allied agent to consolidate forces and enhance its resilience.
3. Prune: In cases where an agent is overwhelmed or retreats from the battlefield, the dynamic agent structure accommodates this change by pruning the agent from the active force.

This dynamic agent structure allows for flexible and adaptive decision-making in response to the ever-changing conditions of the battlefield. For instance, an agent may fork a sub-agent by allocating a portion of its soldiers to create an autonomous agent responsible for collecting information about the surroundings and communicating it back. Additionally, the agent may dispatch autonomous sub-agents to focus on handling specific segments of the incoming enemy forces based on the situation's requirements. In cases where agents lack the necessary strength to address the current situation, they may choose to retreat from the map or merge with other agents to bolster their combined power. Therefore, our project introduces a dynamic multi-agent system in which the structure and number of agents adapt according to the needs of the given situation. This adaptive approach allows the system to effectively respond to the diverse and unpredictable challenges that arise in the battlefield environment.

The following represents an example of a forking action, where the current agent with the ID ARMY-4C9B5218 decides to split its army into two distinct agents. The first new agent consists of 500 soldiers, all of whom are light cavalry, and is assigned the mission of initiating a skirmish at the location of [25, -75] with an enemy agent with the ID ARMY-AC57A3F0. This new agent is assigned the ID ARMY-AC57A3F0. The second new agent consists of 4,000 soldiers, all of whom are heavy cavalry, and is assigned the mission of rallying troops at the location of [25, -75]. This new agent is assigned the ID ARMY-70B7FA20. This example demonstrates the dynamic nature of agent structure within the emulation sandbox, where agents can split and form new agents with distinct properties and missions.

```

1 Current agent:
2 {
3   identity: France,
4   agent_id: ARMY-4c9b5218,
5   troop_number: 6000,
6   current_location: [20, -70]
7 }
8
9 Given the current battlefield situation, where the enemy is deploying
   longbows and fortifying positions to counter our cavalry charge,
   we need to adapt our strategy. Our initial mission was to charge
   cavalry, but the enemy's longbows pose a significant threat to our
   heavily armored knights, especially if we charge without dealing
   with them first.
10
11 ## Thought for Action:
12 1. We need to disrupt the enemy's longbowmen, who are currently
   thebiggest threat to our cavalry.
13 2. We should avoid charging directly into fortified positions or where
   longbowmen are deployed.
14 3. Coordinating with friendly forces is crucial to create a unified
   front and increase the effectiveness of our actions.
15 4. Considering the terrain and the positioning of our forces, we
   should aim to flank the enemy where possible.
16 5. We must maintain high morale among our troops to ensure the
   effectiveness of our charge.
17
18 ## The Actions to Perform:
19 - **Reposition Forces**: To avoid a head-on charge into the enemy's
   longbows, we will reposition our forces to prepare for a flanking
   maneuver.
20
21 fork agents:
22 [

```

```

23 {
24   subAgent_NextActionType: Initiate Skirmish,
25   troopType: Light Cavalry,
26   deployedNum: 500,
27   target_position: [25, -75],
28   target_agent_id: ARMY-75b59d12,
29   agent_id: ARMY-ac57a3f0,
30   remarks: Engaging enemy longbowmen to disrupt their formation and
31   reduce their effectiveness.
32 },
33
34 {
35   subAgent_NextActionType: Rally Troops,
36   troopType: Heavy Cavalry,
37   deployedNum: 4000,
38   target_position: [20, -70],
39   target_agent_id: None,
40   agent_id: ARMY-70b7fa20,
41   remarks: Rallying the main force to maintain high morale and prepare
         for the charge.
42 }
43 ]

```

Listing 4: An example of forking of agents

In Section 3.2, the general profile of commanding agents has been delineated. However, we can see from the dynamic agent structure here that agents are dynamic entities in our sandbox, and within a single country's army, there may be numerous distinct agents, each engaged in different tasks. Therefore, in addition to the general information inherited from the overall commanding agent profile as defined in Section 3.2, each agent possesses more granular and unique information, as defined by the following dynamic properties:

1. Initial mission assigned when being created
2. Current location represented by coordinates
3. The number of soldiers at its disposal
4. The type of soldiers under its command

These properties are subject to evolution over time. For instance, the number of soldiers associated with an agent may fluctuate as a result of soldiers joining the agent, thereby increasing its forces, or from soldiers being killed or wounded in battle, leading to a decrease in its forces. The current location of the agent may also change as it navigates the battlefield, and its initial mission may adapt in response to shifting circumstances and strategic considerations.

Interaction with Landscape Environment To accurately emulate battle dynamics, it is crucial for agents to be able to interact with the physical surroundings as shown in Figure 6 (c), such as rivers, forests, villages, and other features. For example, when encountering a river, agents may build a bridge to cross it; when encountering a forest, agents might choose to hide within it to ambush enemies; and when encountering a village, agents could decide to circumvent it. To facilitate these interactions, it is essential to maintain a relative distance between agents and specific locations on the map, as well as between agents themselves.

The following represents an example of an agent utilizing the natural cover provided by a forest. In this scenario, the agent strategically positions itself within the forest to gain a tactical advantage, such as concealment from enemy agents or protection from ranged attacks.

```

1 {
2   identity: England,
3   id: ARMY-53c7a137,
4   action_description: Fortify Position Constructing defenses in ForestF
         to utilize the natural cover against cavalry charges and to
         bolster the position of our longbowmen,
5   location: [100, -50],

```

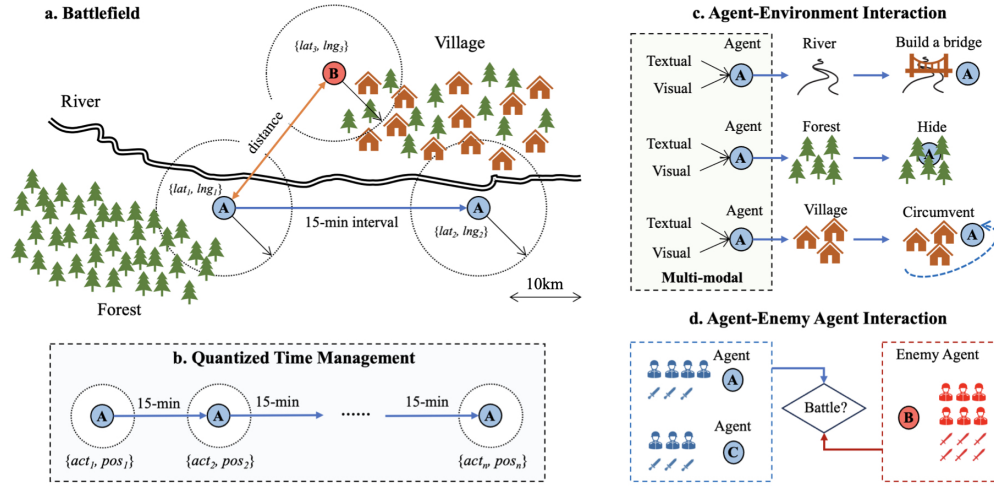


Figure 6: Battlefield interaction (a) Battlefield environment, (2) Quantized time management, (c) Agent-environment interaction, and (d) Agent and enemy agent interaction.

```

6 remaining_number: 300,
7 original_number: 300,
8 lost_number: 0
9 }

```

Listing 5: An example of an agent interacting with the landscape of forest

Interaction with Other Agents Given the observation agents make about their surrounding situations, agents will make decisions regarding whether and when to engage in interactions with other agents, particularly those identified as enemies, as depicted in Figure 6 (d). The specific nature and timing of these interactions are not predetermined; rather, they are initiated by the agents themselves. For instance, when an enemy agent is within close proximity, an agent may opt to engage in combat or launch an attack. The outcome of these interactions between agents is contingent upon various factors, such as the number of soldiers at their disposal and the types of weapons they possess.

The following represents two examples of offensive interactions between two enemy agents. In the first example, an agent executes a flanking maneuver against another agent, as shown in the code snippet below:

```

1 {
2 identity: France,
3 id: ARMY-d2ff280c,
4 action_description: Execute Flanking Maneuvers Flanking enemy unit '
  ARMY- 9418a275' in the midst of a Tactical Retreat to destabilize
  their fortification efforts,
5 location: [10, -100],
6 remaining_number: 8200,
7 original_number: 12000,
8 lost_number: 3800
9 }

```

Listing 6: An example of an agent interacting with a target agent with attacks

In this scenario, the agent with the ID ARMY-D2FF280C, representing France, executes a flanking maneuver against the enemy unit with the ID ARMY-9418A275. The maneuver is intended to destabilize the enemy's fortification efforts while they are in the midst of a tactical retreat. The agent's

current location is represented by the coordinates [10, -100], and it has a remaining force of 8,200 soldiers out of an original force of 12,000, having lost 3,800 soldiers.

The second example involves an agent ambushing an enemy cavalry unit, as shown in the code snippet below:

```
1 {  
2 identity: England,  
3 id: ARMY-2508af97,  
4 action_description: Ambush Enemy Proactively engaging enemy cavalry at  
   [83.8, -17.4] to disrupt their maneuvers and prevent them from  
   supporting their troops,  
5 location: [20, -25],  
6 remaining_number: 600,  
7 original_number: 600,  
8 lost_number: 0  
9 }
```

Listing 7: Another example of an agent interacting with a target agent with attacks

Notice that there are numerous actions that agents can undertake even when they are not directly interacting with each other. These actions may include fortifying their position, rallying troops, and other similar activities that contribute to their overall strategic advantage on the battlefield.

4.3.3 Casualty Evaluation by Observer

In the event that one agent initiates an aggressive action towards another, hereafter referred to as the target agent, both parties may sustain corresponding casualty losses. The loss is evaluated by an objective evaluator supported by GPT-4, which can be seen as an observer. The observer determines the casualties based on several factors:

1. Current profile information from the agents, including their force size, force composition, command architecture, and location.
2. The actions undertaken by the agents, including the action name and a more detailed description of the action generated alongside the action name by the agent. For example, “Deploy Longbows: Deploying longbows in coordination with nearby friendly forces to initiate a skirmish against the nearest enemy cavalry unit and disrupt their advance.”
3. The location and relative distance between the agents, as well as relevant landscape information surrounding them. This information is used to assess the tactical advantages or disadvantages of the agents’ positions.
4. Objective information about the specific weapon utilized, including performance metrics such as the range and accuracy of the weapon. This information is obtained from reputable sources such as Wikipedia.

In order to improve the accuracy and reliability of casualty assessments, it is recommended that future iterations of the emulation sandbox incorporate an *expert system* with a more comprehensive understanding of the weapons involved. Such a system would be able to provide more nuanced and accurate evaluations of casualties based on a deeper understanding of the capabilities and limitations of different weapons, as well as the tactics and strategies employed by the agents.

4.3.4 Agent Profile Update

The profile information of each agent, encompassing force size, force composition, command architecture, and location, is dynamically updated based on the actions undertaken by the agent at the current time. The profile information is updated for every quantized time period.

The factors taken into account to update the agent profile include the actions undertaken by the agent, their interactions with other agents, and any movement that occurs. Specifically, the force size of an agent is determined by three key factors: the change in casualty numbers, the forking of the agent, or the merging of other agents. An agent can decide whether to fork more agents or to merge with other agents, which will result in an increase or decrease in force size, respectively. The location of an agent is contingent upon their movement. If a movement action is executed, the agent’s location is

updated accordingly. The updates to the agent profile ensure that the emulation sandbox accurately reflects the current state of the battlefield and the evolving dynamics of the conflict.

4.3.5 Historical Action Trajectory

For all agents, once decisions regarding actions have been made at each quantized time period, these actions are subsequently recorded into a historical action trajectory, which is then incorporated into the general prompt. As a result, all future decision-making processes will be informed not only by the current environmental information but also by the historical trajectory of actions that have been previously undertaken. This approach enables agents to make informed decisions that are grounded in both the current context and the historical record of actions on the battlefield.

4.4 Emulation of Single Soldier

Our sandbox also simulates the experiences of individual soldiers and reflects their personal perspectives. Throughout the emulation, these soldier agents document their experiences and emotions based on their current actions, previous action trajectories, and any wounds they may have sustained from enemy attacks on the agent to which they belong. This process provides valuable insights into the personal aspects of warfare that complement the higher-level strategic decision-making processes carried out by the commanding agents.

By adopting a multi-layered approach, our emulation is able to capture both the macro-level strategic dynamics and the micro-level personal experiences of the battlefield. This results in a more comprehensive representation of the complexities of war, encompassing both the broader strategic considerations and the individual experiences of soldiers on the ground.

5 Experiment

The primary objective of these experiments is to investigate the extent to which agents based on LLMs and VLMs can reasonably simulate historical battles, which are characterized by a high degree of complexity and dynamism. Specifically, we aim to evaluate the ability of agents to effectively navigate and adapt to the rapidly evolving and unpredictable situations that typically arise during battles. By doing so, we hope to gain insights into the potential of LLMs and VLMs as tools for simulating and analyzing historical conflicts.

We conduct experiments on 4 distinct historical scenarios, namely the Battle of Crécy, the Battle of Agincourt, the Battle of Falkirk, and the Battle of Poitiers. The experiments are performed using 3 strong language models and vision-language models: Claude-3-opus [65], GPT-4-1106-preview [66], and GPT-4-vision [67]. For each scenario and each language model, we execute the emulation 5 times using the same setting within a sandbox environment to account to randomness, continuing until the casualty figures for both armies converge, or in other words, reach a state of stability.

5.1 Evaluation Metrics

As historical battles often lack comprehensive records and documentation, and are typically characterized by unpredictable events that are challenging to replicate in an emulation, our evaluation methodology is divided into three distinct components.

Evaluation aspect	Description
Final battle casualty	Comparison with historical data, focusing on the final casualty figures for both armies
Human analysis on location movement	Assessment of the dynamic structure of agents and their movement on the battlefield as a whole
Human analysis of agent action	Evaluation of the reasonableness of the actions conducted by the agents.

Table 1: Three aspects of evaluation and demonstration.

Battle Final Casualty The first aspect focuses on the final battle casualty, where we compare the simulated casualty figures with historical data. Specifically, we evaluate the casualty of each army for every quantized time period, recording the casualty for each time for both armies. We then compute the mean and variance at each time and compare the final result with historical data to determine the degree of alignment between the emulation results and the historical record. This approach allows us to assess the accuracy and reliability of the emulation sandbox in replicating historical battles.

Human Analysis on Location Movement The second part of the evaluation involves a human analysis of the agents’ behavior on the battlefield, specifically their location movements and dynamic structure. To facilitate an intuitive assessment of whether the agents can rationally interact with the environment and make decisions about where to go, we present images of the action dynamics and the positions of the agents across time. This allows us to evaluate the agents’ ability to adapt to changing circumstances and execute historically plausible strategies, providing insights into the effectiveness of the emulation sandbox in simulating realistic battle scenarios.

Human Analysis on Agent Action The third aspect of the evaluation involves checking the actions taken by the agents. To provide a clear and understandable representation, we present a sequence of diagrams that track the action trajectories taken by specifically two agents.

Although the latter two aspects are challenging to quantitatively evaluate, hopefully by presenting the agents’ behavior and actions in a visual format, we can offer insights into the agents’ decision-making processes and their ability to execute historically plausible strategies.

5.2 First Aspect: Battle Result Performance

In this section, we present the results of the first aspect of the evaluation, which focuses on comparing the simulated casualty figures with historical data. Figures 7, Figures 8, 9, and 10 illustrate the results of simulating the four battles using the three LLM-based and VLM-based agents.

For each battle emulation, we utilize 3 different backbone models to generate results. The results are presented in 3 separate images, with the leftmost figure showing the results from the Claude-3 simulations, the middle figure showing the results from the GPT-4 simulations, and the rightmost figure showing the results from the GPT-4-vision simulations. This allows us to compare the performance of the different language model models in replicating historical battles and to evaluate the impact of the choice of language model on the accuracy and reliability of the simulation results. For each image, the x-axis represents the process of time in emulation, recorded in minutes, while the y-axis represents the casualty figures for armies belonging to the two countries. The casualty of each army is represented by a mean line and a variance band, where the mean and variance are computed based on the 5 simulations run in the sandbox in the current setting. This allows us to assess the accuracy and reliability of the emulation sandbox in replicating historical battles and to evaluate the impact of the choice of language model on the simulation results.

Overall, the results indicate that Claude-3 predicts a significantly higher casualty rate compared to GPT-4 and GPT-4-vision. This discrepancy suggests that the choice of backbone language model can have a significant impact on the accuracy and reliability of the emulation sandbox in replicating historical battles.

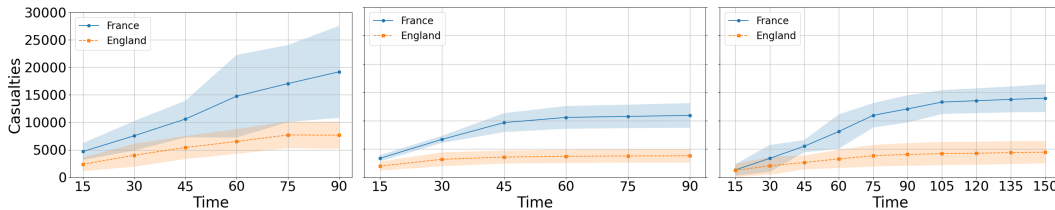


Figure 7: Emulated casualty results for Battle of Crécy.

Battle of Crécy In the emulated Battle of Crécy, the results indicate that all three models predict a higher casualty rate for the French soldiers compared to the English soldiers, which aligns with historical records. Specifically, Claude-3’s prediction suggests that the final casualty figure for

the French army is approximately 2.4 times that of the English army. However, the prediction is characterized by high variance, which increases as the emulation progresses, suggesting a high degree of randomness in the emulation process. In contrast, GPT-4’s prediction suggests that the final casualty figure for the French army is approximately 2.75 times that of the English army, with a relatively low variance, indicating a stable emulation process. Similarly, GPT-4-vision’s prediction suggests that the final casualty figure for the French army is approximately 3.5 times that of the English army, with an acceptable level of variance.

Overall, the results suggest that all three agents are able to simulate the casualty figures for the Battle of Crécy with varying degrees of accuracy and stability. While Claude-3’s prediction is characterized by high variance, GPT-4 and GPT-4-vision’s predictions are more stable and align more closely with historical records.

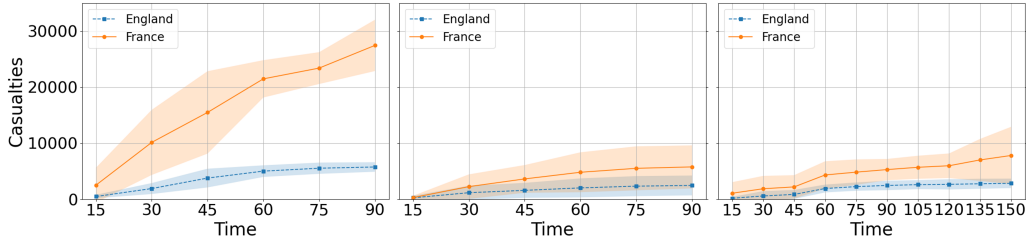


Figure 8: Emulated casualty results for Battle of Agincourt.

Battle of Agincourt In the emulated Battle of Agincourt, the results indicate that all three models predict a higher casualty rate for the French soldiers compared to the English soldiers, which aligns with historical records. Based on historical documentation, only a few hundred died in the English army and about 4,000 to 10,000 died in the French army. Among the three models, GPT-4 and GPT-4-vision predict a closer casualty result, though the casualty of the English is still much higher than historical fact, but the casualty in French are relatively similar. The variance in the predictions is within an acceptable range. However, Claude-3 predicts a much higher casualty for both English and French armies, with the French army almost dying out, which is not very reasonable.

Overall, the results suggest that GPT-4 and GPT-4-vision can provide a relatively reasonable result for the Battle of Agincourt, while Claude-3’s predictions are less accurate. This highlights the importance of selecting an appropriate language model for the emulation sandbox to achieve historically plausible results.

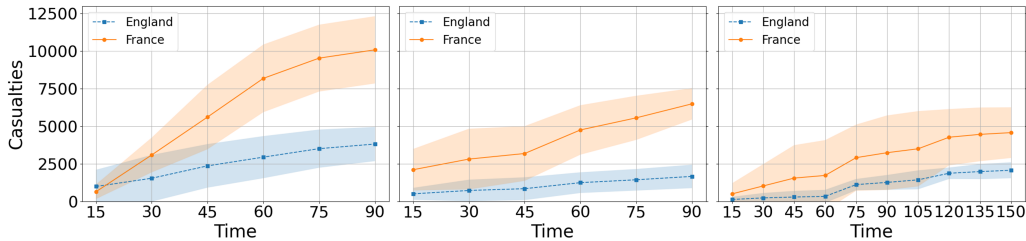


Figure 9: Emulated casualty results for Battle of Poitiers.

Battle of Poitiers In the emulated Battle of Poitiers, the results show a much higher casualty for the French than the English army, conforming to historical documentation: only around 40 soldiers died in this battle in the English army, but more than 4,500 men-at-arms were killed or captured. GPT-4 and GPT-4-vision predict an acceptable result for the French army, which is around 4,000 to 6,000, but too high for the English army, which is close to 2,000 for both models. Claude-3 again predicts too high casualty for both English and French armies.

Overall, the results suggest that GPT-4 and GPT-4-vision can provide a relatively reasonable result for the Battle of Poitiers, while Claude-3’s predictions are less accurate.

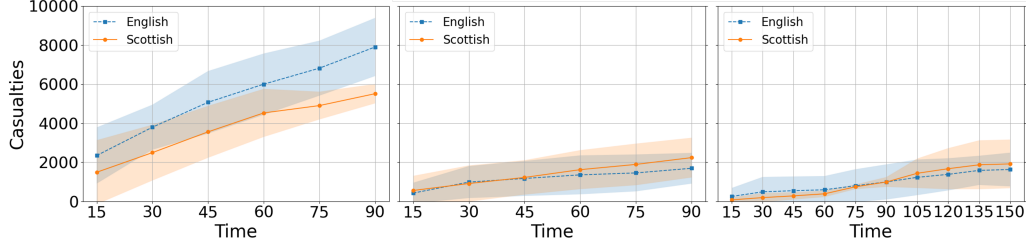


Figure 10: Emulated casualty results for Battle of Falkirk.

Battle of Falkirk In the emulated Battle of Falkirk, the results show a similar casualty for the English army and the Scottish army, conforming to historical documentation: around 2,000 soldiers died in this battle in both the English and Scottish armies. GPT-4 and GPT-4-vision predict a very close result, with the casualty rate for both English and Scottish armies around 2,000. Again, Claude-3 predicts too high casualty for both armies, with English around 8,000 and Scottish around 6,000.

Overall, the results suggest that GPT-4 and GPT-4-vision can provide a relatively reasonable result for the Battle of Falkirk, while Claude-3’s predictions are less accurate.

5.3 Spatial Movement Result Performance

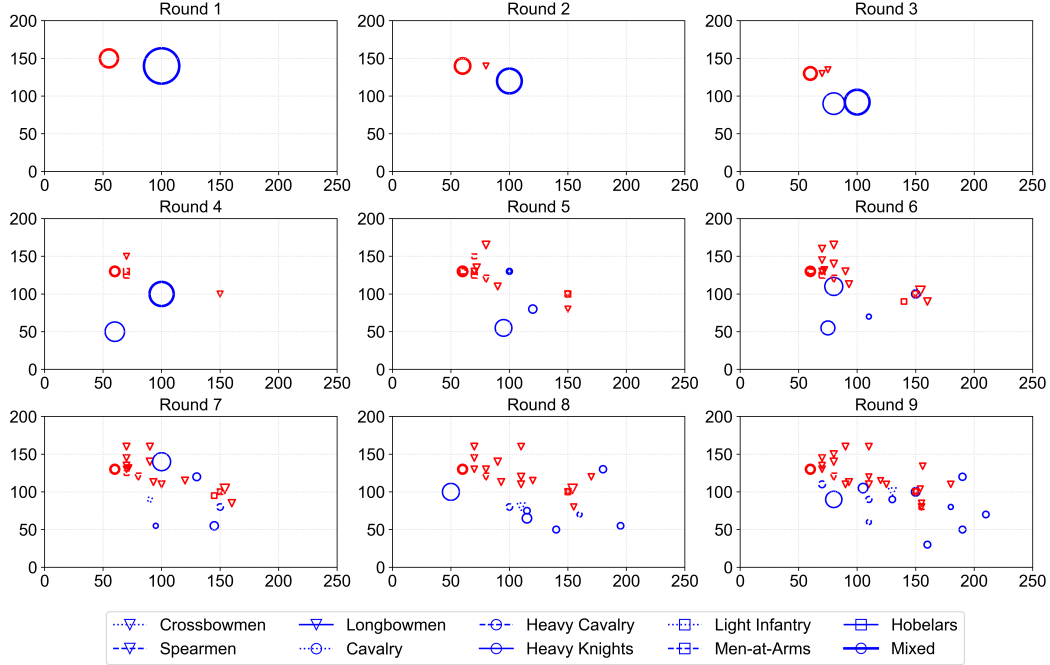


Figure 11: All agent movement and dynamic agent structure on battlefield.

Figure. 11 illustrates the general agent location dynamics of a single emulation of the Battle of Crécy using GPT-4. The English army is represented by red symbols, while the French army is represented by blue symbols. The sizes of the symbols are normalized to correspond to the number of soldiers contained in each agent. Different line type

At a glance, we can observe that as the emulation progresses, both armies are gradually split into smaller teams, especially the English army. Notably, some of the longbowmen tend to maintain a safe distance from the enemy for extended periods, using their longbows to inflict casualties from afar. As time progresses, the advantage of the French army’s larger number of soldiers is diminishing over time, particularly in the case of the heavy cavalry and heavy knights. This is likely due to the

effectiveness of the English longbowmen in inflicting casualties from a safe distance, as well as the challenging terrain of the battlefield, which made it difficult for the heavily armored French knights to maneuver effectively.

To further evaluate the performance of the LLMs and VLMs in simulating historical battles, we can examine the paths taken by individual agents over time. This can provide insights into whether these models have a good sense of distance and can make reasonable decisions based on the overall environment.

5.4 Agent Action throughout Emulation

Figure 12 provides an illustrative example of the actions undertaken by two agents, one representing a part of the army belonging to England and the other representing a part of the army belonging to France, throughout the entire emulation time. The English agent’s cautious approach is reflected in its movements and actions, while the French agent’s aggressive strategy is evident in its frequent attacks and resulting losses. This example provides a reasonable representation of how historical battles may have unfolded.

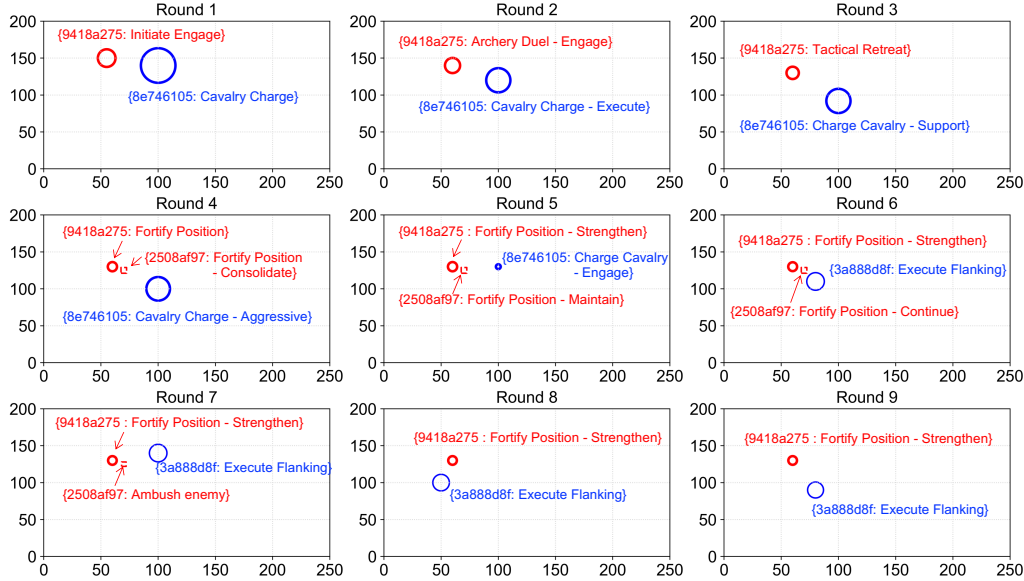


Figure 12: Agent action tracker over time.

5.5 Experiment Summary

The evaluation results of the emulation sandbox for historical battles indicate that the choice of language model can have a significant impact on the accuracy and reliability of the simulation. Specifically, GPT-4 and GPT-4-vision have shown to provide relatively reasonable results in terms of casualty figures for both armies in the emulated battles of Agincourt, Poitiers, and Falkirk, while Claude-3 has consistently predicted much higher casualty rates than historical records.

In terms of the strategies and tactical maneuvers employed by the agents, the results suggest that the agents can adapt to changing circumstances and execute historically plausible strategies. However, further analysis is needed to assess the agents’ ability to rationally interact with the environment and make decisions about where to go.

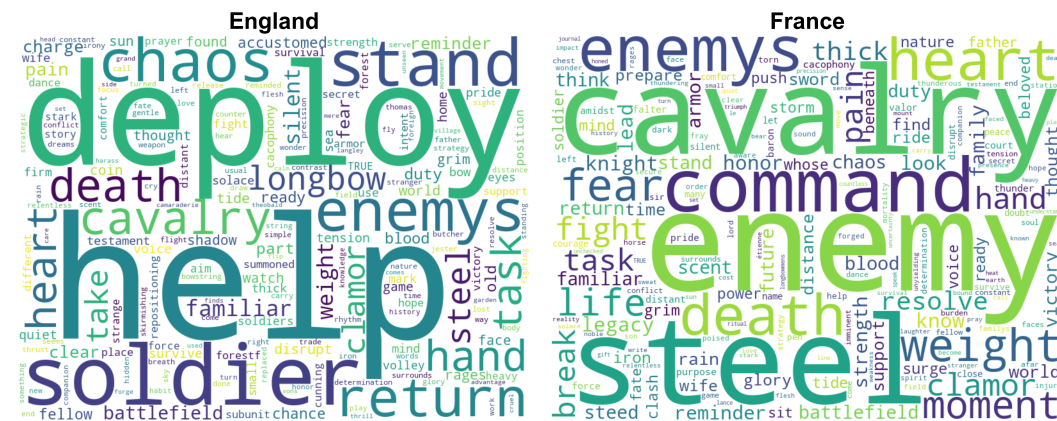
Overall, the evaluation methodology employed in this study, which combines a comparison with historical data and a human analysis of the agents’ behavior, has proven to be effective in assessing the accuracy and reliability of the emulation sandbox. However, further refinement is needed to improve the accuracy of the predictions and the agents’ decision-making capabilities.

6 Soldier Agent Experience

The experiences of individual soldiers are recorded for each quantized time interval during the simulation. Each soldier agent generates a document based on their unique background and experiences. Therefore, for a simulation involving 30 soldier agents, a total of 30 documents are produced, each containing multiple episodes.

To provide an overview of the content of these documents, we utilize word frequency analysis and present the results using a word cloud. This involves merging all documents generated by soldier agents from the same country, regardless of the quantized time interval. The resulting word cloud offers a visual representation of the most frequently occurring words and phrases in the documents, providing insights into the experiences and perspectives of the soldier agents.

To generate the word cloud, we first preprocess the text data by removing words with specific part-of-speech tags, such as those belonging to the tag group [RB, MD, IN, CD, DT, NNS, PRP, NNS, FW]. We also remove functional words that do not convey meaningful information. Additionally, we filter out high-frequency words that do not provide significant insights, such as “think”, “feel”, “battle”, and “war”. This preprocessing step helps to ensure that the resulting word cloud accurately reflects the most salient themes and concepts present in the soldier agents’ documents.



The above two figures present word clouds generated from the documents written by soldier agents in the English and French armies during the Battle of Crécy. Despite being the winning army, the English army's word cloud reveals a high frequency of words such as "death", "chaos", and "return", indicating the intense and chaotic nature of the battle. On the other hand, the French army's word cloud shows a high frequency of words such as "death" and "fear", reflecting the fear and uncertainty experienced by the soldiers. These word clouds provide a visual representation of the experiences and emotions of the soldier agents in both armies during the battle.

7 Conclusion and Future Work

Conclusion In this study, we have demonstrated the potential of LLM and VLM to support highly complex and dynamic simulations of historical battles. Our emulation sandbox provides a comprehensive evaluation of the simulated battles, including a comparison of casualty figures with historical data and a human analysis of the strategies and tactical maneuvers employed by both armies. Our approach also presents the individual experiences of soldiers on the battlefield using soldier agents, providing valuable insights into the personal aspects of warfare that complement the higher-level strategic decision-making processes carried out by the commanding agents.

We believe that our work can also provide new pedagogical methods for students and researchers interested in historical analysis. By simulating historical battles and presenting the results in an interactive and intuitive way, students can gain a deeper understanding of the complexities and dynamics of warfare. Moreover, our approach can be used to support research in various fields, such as military history, AI, and game theory.

Future Work Our study has demonstrated the potential of LLM and VLM to support complex and dynamic simulations of historical battles. However, there is still room for improvement and expansion of our approach.

Firstly, we aim to develop **additional evaluation metrics** for such dynamic simulations to establish their effectiveness more comprehensively. This will enable us to better assess the accuracy and reliability of the simulation results and identify areas for improvement.

Secondly, we plan to extend our approach to **simulate different types of battles beyond barpitite medieval battles**. This will allow us to evaluate the versatility of our approach and its applicability to a wider range of historical battles.

Thirdly, we aim to **incorporate expert systems** for various parts of the simulation, such as information collection for observation and casualty estimation. This will enable us to improve the accuracy and realism of the simulation results, while LLM remains solely responsible for decision-making.

Furthermore, we are interested in developing **more realistic simulations of individual soldiers** beyond just adopting prompting with personalized information. This will enable us to capture the personal experiences of soldiers on the battlefield more accurately and comprehensively.

Finally, we plan to **explore the interaction between commanding agents and soldier agents**, enabling soldier agents to not only follow commands but also actively affect the decision-making process. This will provide insights into the dynamics of command and control in historical battles and enhance the realism of the simulation.

In summary, our future work aims to extend and enhance our approach to provide even more realistic and comprehensive simulations of historical battles, capturing the complexities and dynamics of warfare and providing valuable insights into the strategies and tactics employed by both armies.

References

- [1] Zhiheng Xi, Wenxiang Chen, Xin Guo, Wei He, Yiwen Ding, Boyang Hong, Ming Zhang, Junzhe Wang, Senjie Jin, Enyu Zhou, et al. The rise and potential of large language model based agents: A survey. *arXiv preprint arXiv:2309.07864*, 2023.
- [2] Wayne Xin Zhao, Kun Zhou, Junyi Li, Tianyi Tang, Xiaolei Wang, Yupeng Hou, Yingqian Min, Beichen Zhang, Junjie Zhang, Zican Dong, et al. A survey of large language models. *arXiv preprint arXiv:2303.18223*, 2023.
- [3] Lizhou Fan, Lingyao Li, Zihui Ma, Sanggyu Lee, Huizi Yu, and Libby Hemphill. A bibliometric review of large language models research from 2017 to 2023. *arXiv preprint arXiv:2304.02020*, 2023.
- [4] Jie Huang and Kevin Chen-Chuan Chang. Towards reasoning in large language models: A survey. *arXiv preprint arXiv:2212.10403*, 2022.
- [5] Mingyu Jin, Qinkai Yu, Haiyan Zhao, Wenye Hua, Yanda Meng, Yongfeng Zhang, Mengnan Du, et al. The impact of reasoning step length on large language models. *arXiv preprint arXiv:2401.04925*, 2024.
- [6] Jeffrey Zhou, Tianjian Lu, Swaroop Mishra, Siddhartha Brahma, Sujoy Basu, Yi Luan, Denny Zhou, and Le Hou. Instruction-following evaluation for large language models. *arXiv preprint arXiv:2311.07911*, 2023.
- [7] Zhiyuan Zeng, Jiatong Yu, Tianyu Gao, Yu Meng, Tanya Goyal, and Danqi Chen. Evaluating large language models at evaluating instruction following. *arXiv preprint arXiv:2310.07641*, 2023.
- [8] Boshi Wang, Hao Fang, Jason Eisner, Benjamin Van Durme, and Yu Su. Llms in the imagination: Tool learning through simulated trial and error. *arXiv preprint arXiv:2403.04746*, 2024.
- [9] Chenyu Wang, Weixin Luo, Qianyu Chen, Haonan Mai, Jindi Guo, Sixun Dong, Zhengxin Li, Lin Ma, Shenghua Gao, et al. Tool-lmm: A large multi-modal model for tool agent learning. *arXiv preprint arXiv:2401.10727*, 2024.

- [10] Weizhou Shen, Chenliang Li, Hongzhan Chen, Ming Yan, Xiaojun Quan, Hehong Chen, Ji Zhang, and Fei Huang. Small llms are weak tool learners: A multi-llm agent. *arXiv preprint arXiv:2401.07324*, 2024.
- [11] Jingyi Zhang, Jiaying Huang, Sheng Jin, and Shijian Lu. Vision-language models for vision tasks: A survey. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2024.
- [12] Zane Durante, Qiuyuan Huang, Naoki Wake, Ran Gong, Jae Sung Park, Bidipta Sarkar, Rohan Taori, Yusuke Noda, Demetri Terzopoulos, Yejin Choi, et al. Agent ai: Surveying the horizons of multimodal interaction. *arXiv preprint arXiv:2401.03568*, 2024.
- [13] Junlin Xie, Zhihong Chen, Ruifei Zhang, Xiang Wan, and Guanbin Li. Large multimodal agents: A survey. *arXiv preprint arXiv:2402.15116*, 2024.
- [14] Yilun Du, Shuang Li, Antonio Torralba, Joshua B. Tenenbaum, and Igor Mordatch. Improving factuality and reasoning in language models through multiagent debate, 2023.
- [15] Chi-Min Chan, Weize Chen, Yusheng Su, Jianxuan Yu, Wei Xue, Shanghang Zhang, Jie Fu, and Zhiyuan Liu. Chateval: Towards better llm-based evaluators through multi-agent debate, 2023.
- [16] Qiushi Sun, Zhangyue Yin, Xiang Li, Zhiyong Wu, Xipeng Qiu, and Lingpeng Kong. Corex: Pushing the boundaries of complex reasoning through multi-model collaboration, 2023.
- [17] Tian Liang, Zhiwei He, Wenxiang Jiao, Xing Wang, Yan Wang, Rui Wang, Yujiu Yang, Zhaopeng Tu, and Shuming Shi. Encouraging divergent thinking in large language models through multi-agent debate, 2023.
- [18] Sirui Hong, Xiawu Zheng, Jonathan Chen, Yuheng Cheng, Ceyao Zhang, Zili Wang, Steven Ka Shing Yau, Zijuan Lin, Liyang Zhou, Chenyu Ran, et al. Metagpt: Meta programming for multi-agent collaborative framework. *arXiv preprint arXiv:2308.00352*, 2023.
- [19] Zhiwei Liu, Weiran Yao, Jianguo Zhang, Le Xue, Shelby Heinecke, Rithesh Murthy, Yihao Feng, Zeyuan Chen, Juan Carlos Niebles, Devansh Arpit, et al. Bolaa: Benchmarking and orchestrating llm-augmented autonomous agents. *arXiv preprint arXiv:2308.05960*, 2023.
- [20] Yingqiang Ge, Wenyue Hua, Kai Mei, Jianchao Ji, Juntao Tan, Shuyuan Xu, Zelong Li, and Yongfeng Zhang. OpenAGI: When LLM meets domain experts. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [21] Zhao Yang, Jiayuan Liu, Yucheng Han, Xin Chen, Zebiao Huang, Bin Fu, and Gang Yu. Appagent: Multimodal agents as smartphone users. *arXiv preprint arXiv:2312.13771*, 2023.
- [22] Kai Mei, Zelong Li, Shuyuan Xu, Ruosong Ye, Yingqiang Ge, and Yongfeng Zhang. Llm agent operating system. *arXiv preprint arXiv:2403.16971*, 2024.
- [23] Yingqiang Ge, Yujie Ren, Wenyue Hua, Shuyuan Xu, Juntao Tan, and Yongfeng Zhang. Llm as os, agents as apps: Envisioning aios, agents and the aios-agent ecosystem. *arXiv e-prints*, pages arXiv-2312, 2023.
- [24] Ran Gong, Qiuyuan Huang, Xiaojian Ma, Hoi Vo, Zane Durante, Yusuke Noda, Zilong Zheng, Song-Chun Zhu, Demetri Terzopoulos, Li Fei-Fei, et al. Mindagent: Emergent gaming interaction. *arXiv preprint arXiv:2309.09971*, 2023.
- [25] Yuzhuang Xu, Shuo Wang, Peng Li, Fuwen Luo, Xiaolong Wang, Weidong Liu, and Yang Liu. Exploring large language models for communication games: An empirical study on werewolf. *arXiv preprint arXiv:2309.04658*, 2023.
- [26] Yihuai Lan, Zhiqiang Hu, Lei Wang, Yang Wang, Deheng Ye, Peilin Zhao, Ee-Peng Lim, Hui Xiong, and Hao Wang. Llm-based agent society investigation: Collaboration and confrontation in avalon gameplay. *arXiv preprint arXiv:2310.14985*, 2023.
- [27] Sihao Hu, Tiansheng Huang, Fatih Ilhan, Selim Tekin, Gaowen Liu, Ramana Kompella, and Ling Liu. A survey on large language model-based game agents. *arXiv preprint arXiv:2404.02039*, 2024.
- [28] Xianghe Pang, Shuo Tang, Rui Ye, Yuxin Xiong, Bolun Zhang, Yanfeng Wang, and Siheng Chen. Self-alignment of large language models via multi-agent social simulation. In *ICLR 2024 Workshop on Large Language Model (LLM) Agents*.
- [29] Xuhui Zhou, Zhe Su, Tiwalayo Eisape, Hyunwoo Kim, and Maarten Sap. Is this the real life? is this just fantasy? the misleading success of simulating social interactions with llms. *arXiv preprint arXiv:2403.05020*, 2024.

- [30] Karthik Sreedhar and Lydia Chilton. Simulating human strategic behavior: Comparing single and multi-agent llms. *arXiv preprint arXiv:2402.08189*, 2024.
- [31] Chengxing Xie, Canyu Chen, Feiran Jia, Ziyu Ye, Kai Shu, Adel Bibi, Ziniu Hu, Philip Torr, Bernard Ghanem, and Guohao Li. Can large language model agents simulate human trust behaviors? *arXiv preprint arXiv:2402.04559*, 2024.
- [32] Wenyue Hua, Lizhou Fan, Lingyao Li, Kai Mei, Jianchao Ji, Yingqiang Ge, Libby Hemphill, and Yongfeng Zhang. War and peace (waragent): Large language model-based multi-agent simulation of world wars. *arXiv preprint arXiv:2311.17227*, 2023.
- [33] Linda Shopes. Oral history. *The SAGE handbook of qualitative research*, pages 451–465, 2011.
- [34] Alessandro Portelli. What makes oral history different. In *The oral history reader*, pages 77–88. Routledge, 2002.
- [35] Lizhou Fan, Wenyue Hua, Lingyao Li, Haoyang Ling, Yongfeng Zhang, and Libby Hemphill. Nphardeval: Dynamic benchmark on reasoning ability of large language models via complexity classes. *arXiv preprint arXiv:2312.14890*, 2023.
- [36] Lizhou Fan, Wenyue Hua, Xiang Li, Kaijie Zhu, Mingyu Jin, Lingyao Li, Haoyang Ling, Jinkui Chi, Jindong Wang, Xin Ma, et al. Nphardeval4v: A dynamic reasoning benchmark of multimodal large language models. *arXiv preprint arXiv:2403.01777*, 2024.
- [37] Yadong Zhang, Shaoguang Mao, Tao Ge, Xun Wang, Adrian de Wynter, Yan Xia, Wenshan Wu, Ting Song, Man Lan, and Furu Wei. Llm as a mastermind: A survey of strategic reasoning with large language models. *arXiv preprint arXiv:2404.01230*, 2024.
- [38] Joon Sung Park, Joseph C O’Brien, Carrie J Cai, Meredith Ringel Morris, Percy Liang, and Michael S Bernstein. Generative agents: Interactive simulacra of human behavior. *arXiv preprint arXiv:2304.03442*, 2023.
- [39] Yao Fu, Hao Peng, Tushar Khot, and Mirella Lapata. Improving language model negotiation with self-play and in-context learning from ai feedback. *arXiv preprint arXiv:2305.10142*, 2023.
- [40] Andres M Bran, Sam Cox, Oliver Schilter, Carlo Baldassari, Andrew White, and Philippe Schwaller. Augmenting large language models with chemistry tools. In *NeurIPS 2023 AI for Science Workshop*, 2023.
- [41] Wenyue Hua, Xianjun Yang, Zelong Li, Cheng Wei, and Yongfeng Zhang. Trustagent: Towards safe and trustworthy llm-based agents through agent constitution. *arXiv preprint arXiv:2402.01586*, 2024.
- [42] Kexin Chen, Junyou Li, Kunyi Wang, Yuyang Du, Jiahui Yu, Jiamin Lu, Guangyong Chen, Lanqing Li, Jiezhong Qiu, Qun Fang, et al. Towards an automatic ai agent for reaction condition recommendation in chemical synthesis. *arXiv preprint arXiv:2311.10776*, 2023.
- [43] Zhao Yang, Jiakuan Liu, Yucheng Han, Xin Chen, Zebiao Huang, Bin Fu, and Gang Yu. Appagent: Multimodal agents as smartphone users. *arXiv preprint arXiv:2312.13771*, 2023.
- [44] Amazon Web Service. Generative ai and multi-modal agents in aws: The key to unlocking new value in financial markets, 1 2024.
- [45] Shilong Liu, Hao Cheng, Haotian Liu, Hao Zhang, Feng Li, Tianhe Ren, Xueyan Zou, Jianwei Yang, Hang Su, Jun Zhu, et al. Llava-plus: Learning to use tools for creating multimodal agents. *arXiv preprint arXiv:2311.05437*, 2023.
- [46] Zhuosheng Zhan and Aston Zhang. You only look at screens: Multimodal chain-of-action agents. *arXiv preprint arXiv:2309.11436*, 2023.
- [47] Ted Dickson. The road to united states involvement in world war i: A simulation. *OAH Magazine of History*, 17(1):48–56, 2002.
- [48] Harold Steere Guetzkow, Chadwick F Alger, and Richard A Brody. Simulation in international relations: Developments for research and teaching. (*No Title*), 1963.
- [49] Charles F Hermann and Margaret G Hermann. An attempt to simulate the outbreak of world war i. *American Political Science Review*, 61(2):400–416, 1967.
- [50] Eric Tollefson, M Martin, Andrew Fletcher, and ARMY TRADOC ANALYSIS CENTER MONTEREY CA. Onesaf objective system (oos) behavior model verification. *US Army TRADOC Analysis Center—Monterey, Monterey, CA*, 2008.

- [51] Raymond R Hill, Lance E Champagne, and Joseph C Price. Using agent-based simulation and game theory to examine the wwii bay of biscay u-boat campaign. *The Journal of Defense Modeling and Simulation*, 1(2):99–109, 2004.
- [52] Navid Ghaffarzadegan, Aritra Majumdar, Ross Williams, and Niyousha Hosseinichimeh. Generative agent-based modeling: Unveiling social system dynamics through coupling mechanistic models with generative artificial intelligence. *arXiv preprint arXiv:2309.11456*, 2023.
- [53] Joshua M Epstein. Agent-based computational models and generative social science. *Complexity*, 1999.
- [54] Lizhou Fan, Huizi Yu, and Zhanyuan Yin. Stigmatization in social media: Documenting and analyzing hate speech for covid-19 on twitter. *Proceedings of the Association for Information Science and Technology*, 57(1):e313, 2020.
- [55] Zhanyuan Yin, Lizhou Fan, Huizi Yu, and Anne J Gilliland. Using a three-step social media similarity (tsms) mapping method to analyze controversial speech relating to covid-19 in twitter collections. In *2020 IEEE International Conference on Big Data (Big Data)*, pages 1949–1953. IEEE, 2020.
- [56] Lingyao Li, Zihui Ma, Lizhou Fan, Sanggyu Lee, Huizi Yu, and Libby Hemphill. Chatgpt in education: A discourse analysis of worries and concerns on social media. *arXiv preprint arXiv:2305.02201*, 2023.
- [57] Alfred H Burne. *The Crecy War: A Military History of the Hundred Years War from 1337 to the Peace of Bretigny in 1360*. Casemate Publishers, 2016.
- [58] Anne Curry. *The battle of Agincourt: sources and interpretations*. Boydell Press, 2000.
- [59] Chris Given-Wilson and Françoise Bériac. Edward iii’s prisoners of war: the battle of poitiers and its context. *The English Historical Review*, 116(468):802–833, 2001.
- [60] André Geraque Kiffer. *Battle Of Falkirk, July 22, 1298*. Clube de Autores, 2019.
- [61] Toshifumi Matsuoka, Takahiro Hasegawa, Yasuhiro Yamada, Tetsuya Tamagawa, and Yuzuru Ashida. Computer simulation for sandbox experiments. In *SEG International Exposition and Annual Meeting*, pages SEG–2001. SEG, 2001.
- [62] Ahmed A Al Rowaei, Arnold H Buss, and Stephen Lieberman. The effects of time advance mechanism on simple agent behaviors in combat simulations. In *Proceedings of the 2011 Winter Simulation Conference (WSC)*, pages 2426–2437. IEEE, 2011.
- [63] Zijun Liu, Yanzhe Zhang, Peng Li, Yang Liu, and Diyi Yang. Dynamic llm-agent network: An llm-agent collaboration framework with agent team optimization. *arXiv preprint arXiv:2310.02170*, 2023.
- [64] Shanshan Han, Qifan Zhang, Yuhang Yao, Weizhao Jin, Zhaozhuo Xu, and Chaoyang He. Llm multi-agent systems: Challenges and open problems. *arXiv preprint arXiv:2402.03578*, 2024.
- [65] Anthropic. The claude 3 model family: Opus, sonnet, haiku. https://www-cdn.anthropic.com/de8ba9b01c9ab7cbabf5c33b80b7bbc618857627/Model_Card_Claude_3.pdf, 2024.
- [66] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- [67] OpenAI. Gpt-4v(ision) system card. https://cdn.openai.com/papers/GPTV_System_Card.pdf, 2023.