

A Deep Reinforcement Learning Methods based on Deterministic Policy Gradient for Multi-Agent Cooperative Competition

Xuan Zuo* Hui-Feng Xue** Xiao-Yin Wang** Fu-Kai Jia** Wan-Ru Du**
Tao Tian** Shan Gao* Pu Zhang*

* School of Automation, Northwestern Polytechnical University, Xi'an 710072
P.R.China (e-mail: 1069114233@qq.com).

** China Aerospace Academy of Systems Science and Engineering, Beijing 100048
P.R.China (e-mail: xhf0616@163.com)}

Abstract: Deep reinforcement learning in multi-agent scenarios is important for real-world applications but presents challenges beyond those seen in single agent settings. This paper considers deterministic policy gradient algorithms for multi-agent control and expands the Actor-Critic method that consider action policies of other agents to cooperative competition mission. Our method is applicable not only to cooperative settings with shared rewards, but also adversarial settings including attacking and defending. In the computer experiments, agents are divided into attacking agents and defending agents. The results show that attacking agents which play the roles of deceivers can attract most of defending agents and help the other attacking agents to reach their targets successfully. Choosing appropriate length of training could help agents learn better action policy. The experiments results reveal that the number of agents has an effect on the performance of our proposed method. Increasing the number of deceivers in attacking agents can significantly increase the mission success of attacking party, but the computational complexity will rise and more episodes are needed to train agents.

Keywords: Machine learning, reinforcement learning, multi-agent, cooperative competition, artificial intelligence.

1. INTRODUCTION

Much progress towards artificial intelligence has been made using supervised learning systems that are trained to replicate the decisions of human experts (Silver et al., 2017; Hastie et al., 2009; LeCun et al., 2015; Krizhevsky et al., 2012). Deep reinforcement learning (DRL) represents a step towards building autonomous systems with a higher level understanding of the visual world. In DRL (Arulkumaran et al., 2017; Francois-Lavet et al., 2018), deep neural networks are trained to approximate the optimal policy or the value function. In this way the deep neural networks, serving as function approximator, enables powerful generalization. One of the key advantages of DRL is that it enables reinforcement learning (RL) to scale to problems with high-dimensional state and action spaces (Hernandez-Leal, Kartal and Taylor, 2019).

Recently, there has been rapid progress using deep neural networks trained by reinforcement learning. These systems have outperformed humans in games, such as AlphaGo (Silver, Huang, and Maddison, 2016), AlphaGo Zero (Silver et al., 2017) and Alpha Zero (Silver, Hubert and Schrittwieser, 2017, 2018). Otherwise, in the study of antagonistic video games based on local visual information, Deepmind's AlphaStar (Vinyals, Babuschkin and Czarnecki, 2019) in the real-time strategy game starcraft II, and OpenAI Five (OpenAI, 2018) in the multiplayer game DOTA 2 show outstanding performance beyond top human players. The

outstanding ability of deep reinforcement learning in solving single individual policy optimization problem urges researchers to try to apply related methods to solve multi-agent cooperation or competition problems.

The nature of interaction between agents can either be cooperative, competitive, or both and many algorithms are designed only for a particular nature of interaction. These algorithms are generally not applicable in competitive or mixed settings. These algorithms are generally not applicable in competitive or mixed settings (Lowe, Wu and Tamar, 2017). As multi-agent learning needs to search for the optimal policy in the high-dimensional state space, traditional deep reinforcement learning algorithms will run into limitations in multi-agent environments. The high-dimensional state space will bring about an unstable training environment, if every agent learns policy from its own perspective without cooperative communication or setting of sharing global rewards. In recent years, the actor-critic algorithm architecture which could optimize multi-agent's policy model with centralized training and test the model using trained policy and local observation of every agent with decentralized execution, has been introduced into multi-agent deep reinforcement learning to share agents' rewards in various ways.

The multi-agent cooperative competition scenario in this paper includes both competition and cooperation between agents. Different kinds of agents take on different missions in

cooperative competition, and cooperation between agents' policy could be realized through the centralized Critic. In the cooperative competition scenario, the agents in attacking party which undertake the attacking mission can arrive their targets synchronously by minimizing the standard deviation of the distance between agents and targets. The other agents in attacking party which undertake deception mission can induce adversaries at close range by maximizing the value of a Gauss-Quadratic mixing function about the distance between them and defending agents. The defending agents intercept adversaries by minimizing the distance to the attacking agents and giving him a constant reward on contact. In this paper, multi-agent deep deterministic strategy gradient algorithm which combines Actor-Critic framework and DDPG algorithm is used to train agents learning policy, so that the multi-agent multi-mission multi-target co-control could be realized. This paper also analyzes the effects of the number of training episodes and the number of agents undertook deception mission on the training results.

2. RELATE WORKS

In the face of some complex decision-making problems in real scenarios, the decision-making ability of single-agent system is not enough to meet the challenges. In most of the multi-agent reinforcement learning algorithms, the training mechanism is assigned to each agent separately. For example, independent Q learning algorithm is adopted to train each agent. The distributed learning architecture above reduces the difficulty of implementing learning and the complexity of calculation. The distributed learning architecture above reduces the difficulty of implementing learning and the complexity of calculation. For the DRL problem of large-scale state space, a simple multi-agent DRL system can be constructed by using DQN algorithm instead of Q learning algorithm to train each agent individually (Liu et al., 2018). Tampuu et al. (Tampuu, Mätiisen and Kodelja, 2017) used the above ideas to expand the framework of deep Q learning, dynamically adjusted the reward mode according to different goals, and proposed a DRL model in which multiple agents could cooperate and compete with each other. In the face of a class of reasoning missions that require multiple agents to communicate with each other, the DQN model cannot usually learn effective strategies. To solve this problem, Foerster et al. (Foerster, Assael and de Freitas, 2016) proposed a model called distributed deep loop Q network (DDRQN), which solved the problem of multi-agent communication and cooperation that can be observed in the state part.

The machine learning method based on Q learning is challenged by an inherent non-stationarity of the environment, while policy gradient suffers from a variance that increases as the number of agents grows. In recent years, some new approaches have been proposed, including deep reinforcement learning based on attention mechanism and deep reinforcement learning based on Actor-Critic framework.

In 2017, Choi et al. (Choi, Lee and Zhang, 2017) proposed a multi-focus attention network (MANet) method that can simulate human spatial extraction ability. This method first divides the low-level input into several segments representing

local states. After this segmentation, parallel attention layers attend to the partial states relevant to solving the mission. The network model estimates state-action values using these attended partial states, and attains better learning effect with significantly less experience samples. In the environment of large-scale multi-agent cooperation, it is difficult for agents to differentiate valuable information that helps cooperative decision making from globally shared information. The predefined communication architectures, on the other hand, restrict communication among agents and thus restrain potential cooperation. To tackle these difficulties, in 2018, Jiang et al. (Jiang and Lu, 2018) proposed an attentional communication model that learns when communication is needed and how to integrate shared information for cooperative decision making.

In 2017, Foerster et al. (Foerster, Farquhar and Afouras, 2018) proposed a new multi-agent actor-critic method called counterfactual multi-agent (COMA) policy gradients based on the Actor-Critic. COMA uses a centralized critic to estimate the Q-function and decentralized actors to optimize the agents' policies. To address the challenges of multi-agent credit assignment, it uses a counterfactual baseline that marginalizes out a single agent's action, while keeping the other agents' actions fixed. Lowe R et al. (Lowe, Wu and Tamar, 2017) present an adaptation of actor-critic methods named MADDPG (Multi-agent Deep Deterministic Policy Gradient) that considers action policies of other agents and is able to successfully learn policies that require complex multiagent coordination to maximize the global reward. They also introduce a training regimen utilizing an ensemble of policies for each agent that weakened overfitting and leads to more robust multi-agent policies. They show the success of their approach compared to existing methods in cooperative as well as competitive scenarios.

3. MULTI-AGENT COOPERATIVE COMPETITION AND DECEPTION TACTION

3.1 Multi-agent centralized training framework

In this paper, the Multi-Agent Deep Deterministic Policy Gradient (MADDPG) algorithm is adopted to train all agents in the competitive scenarios, so that both attacking and defending party can gradually learn cooperative tactics in the training process. Based on the Deep Deterministic Policy Gradient (DDPG) algorithm (Lillicrap, Hunt and Pritzel, 2015), MADDPG algorithm adds the policy information of other agents, including the state, action and reward value of each agent. Similar to DDPG, the MADDPG algorithm also uses the Actor-Critic algorithm framework (Sutton and Barto, 2018), as shown in Figure 1.

During training, the policy network in the Actor selects an action according to the observation information of the current agent, and the action is executed in the environment to get the new state and return value. After that, each agent's state, action, reward value, and new state form a sample to be saved to the experience replay buffer.

This paper focuses on the application of deep reinforcement learning algorithm and functionalized return value. Through centralized training, the agents that undertakes the decoy mission and the agents that undertakes the attack mission realize the multi-agent, multi-mission, multi-target collaborative tactics. In the battlefield environment, there are often many mobile defense units near important targets. When a target encounters an attack or is about to encounter an attack, the defending units of the defender can move in time and intercept the incoming units. In order to safely break through the defense of the defending party and reach the targets, the attacking units of the attacking party need to have good tactical coordination. In this paper and experiments, we assume that the defending party has two types of agents, one is the commander with global insight, and the other is the ordinary guardian who can only observe the barrier-free area. The attacking party also has two types of agents, one is the attacker who undertakes the collaborative attack mission, and the other is the deceiver who actively approaches and attracts the defending units to cover attackers.

In this paper, the reward function of the defending agents consists of two parts: one is the punishment of the distance between the defending party and the nearest unit of the attacking party; the other part is the reward when the defending unit intercepts the attacking unit with a constant value. Thus, the reward function of the defending agents can be expressed as:

$$R_i^{defense} = \sum_{t=1}^T r_{i1}^{defense}(t) + r_{i2}^{defense}(t) \quad (5)$$

where $i = 1, 2, \dots, m$ indicates the i^{th} defending agent, t indicates the t^{th} step in each episode. The first part of the above equation that represents punishment can be written as:

$$r_{i1}^{defense}(t) = a * \min\{d_{i,j}^{defense-attack}\}$$

where a is a negative constant, $j = 1, 2, \dots, n$ indicates the j^{th} attacking agent, and $d_{i,j}^{defense-attack}$ is the distance between the attacking agent and the defending agent. The second part of the formula (5) that represents reward can be written as:

$$r_{i2}^{defense}(t) = \sum_i \sum_j g_{i,j}$$

With

$$g_{i,j} = \begin{cases} c, & \text{if } d_{i,j}^{defense-attack} < l + l' \\ 0, & \text{if } d_{i,j}^{defense-attack} \geq l + l' \end{cases}$$

where c is a positive constant, l and l' represent the radius of the attacking agents and defending agents respectively.

As the attackers and the deceivers in attacking side have different tactical missions, the settings of their reward

functions are also different. The attackers' reward function consists of five parts: one is the distance from the attacker to the nearest target as punishment; the second is a constant value as the reward for the attacker to reach the target; the third is the standard deviation of the distance between each attacker and its nearest target that is used as punishment to train agents to attack synchronously; the fourth is to use a constant value as punishment when the attacker is intercepted by the defending agents; the fifth is the distance between the attacker and its nearest defending agent. Therefore, the attackers' reward function can be written as follows:

$$R_j^{attack} = \sum_{t=1}^T r_{j1}^{attack}(t) + r_{j2}^{attack}(t) + r_{j3}^{attack}(t) + r_{j4}^{attack}(t) + r_{j5}^{attack}(t) \quad (6)$$

The first part of the above equation can be written as:

$$r_{j1}^{attack}(t) = b * \min_k \{d_{k,j}^{target-attack}\}$$

where b is a negative constant, $k = 1, 2, \dots, o$ indicates the k^{th} target, and $d_{k,j}^{target-attack}$ is the distance between the target and the attacker. The second part of formula (6) can be written as:

$$r_{j2}^{attack}(t) = \begin{cases} c', & \text{if } d_j^{attack-target} < l' + l'' \\ 0, & \text{if } d_j^{attack-target} \geq l' + l'' \end{cases}$$

where c' is a positive constant, l'' is the radius of the target. The third part of formula (6) can be written as:

$$r_{j3}^{attack}(t) = b' * std. \left\{ \min_k \{d_{j,k}^{attack-target}\} \right\}_j$$

where b' is a negative constant, $d_{j,k}^{attack-target}$ is the distance between the attacker and the target, and $std.$ represents the standard deviation of elements in the collection. The fourth part of formula (6) can be written as:

$$r_{j4}^{attack}(t) = \sum_{i=1}^m g'_{j,i}$$

with

$$g'_{j,i} = \begin{cases} c'', & \text{if } d_{i,j}^{defense-attack} < l + l' \\ 0, & \text{if } d_{i,j}^{defense-attack} \geq l + l' \end{cases}$$

where c'' is a negative constant. The fifth part of formula (6) can be written as:

$$r_{j5}^{attack}(t) = b'' * \min_i \{d_{i,j}^{defense-attack}\}$$

where b'' is a positive constant.

The deceivers' reward function consists of three parts: one is a Gauss-Quadratic mixing function about the distance between the deceiver and defending agent as a reward to train

the deceiver to learn to approach defending agent actively and keep the distance from defending agent; the second is a constant as punishment when the deceiver is contacted by defending agent; the third is the distance between the deceiver and attacker as a reward to train the deceiver to learn to evade attackers in the mission. Therefore, deceivers' reward function can be written as follows:

$$R_{j'}^{deception} = \sum_{t=1}^T r_{j',1}^{deception}(t) + r_{j',2}^{deception}(t) + r_{j',3}^{deception}(t) \quad (7)$$

the first part of above equation can be written as:

$$r_{j',1}^{deception}(t) = \sum_{i=1}^m h_1 * e^{-\alpha_1 * (d_{i,j'}^{defense-deception} - f_1)^2} + \alpha_2 * (h_2 - (d_{i,j'}^{defense-deception} - f_2)^2)$$

where $d_{i,j'}^{defense-deception}$ is the distance between the defending agent and the deceiver, $f_1, f_2, h_1, h_2, \alpha_1, \alpha_2$ are all positive constant parameters. The second part of formula (7) can be written as:

$$r_{j',2}^{deception}(t) = \sum_{i=1}^m g''_{j',i}$$

with

$$g''_{j',i} = \begin{cases} c''', & \text{if } d_{i,j'}^{defense-deception} < l + l' \\ 0, & \text{if } d_{i,j'}^{defense-deception} \geq l + l' \end{cases}$$

where c''' is a negative constant. The third part of formula (7) can be written as:

$$r_{j',3}^{deception}(t) = b''' * \min\{d_{j',j}^{deception-attack}\}$$

where b''' is a positive constant, $d_{j',j}^{deception-attack}$ is the distance between the deceiver and the attacker.

In addition to the targets and agents in both attacking and defending party, obstacles and forests are also set up in the mission scenario in this paper. The obstacles are impassable, agents can only detour after touching. The forest can cover the observation ability of the agent. When agents are outside the forest, they cannot observe the agents in the forest. Similarly, when agents are in the forest, they cannot observe other agents outside the forest. The commander in defending agents is the only agent with global observation capability, and all the agents in the scenario can be observed regardless of whether the commander is in forests or outside forests.

The setting of obstacles and forests can be used to simulate some dangerous obstacles and low-detection areas in complex battlefield environment, respectively. Agents are

trained to avoid obstacles and move purposefully into and out of forests. The attacking agents can learn to use the forests intelligently to avoid being intercepted, and the defending agents can learn to cooperate with each other to deal with the attacking agents in forests.

5. EXPERIMENTS

5.1 The experiment of cooperative competition and deception tactics

In the computer experiments, the two types of attacking agents can be trained to learn policy to maximize the global reward without any particular communication method. Thus, the agents can cooperate on the attacking missions.

In order to measure the performance of the multi-agent cooperative tactics quantitatively, three indicators are used to describe the process of attackers and deceivers learning cooperative tactics, including the total number of times attacking agents were intercepted by defending agents, the number of times the attackers were intercepted, and the number of times the attackers reached the targets for each 1000 episodes during the training process. The strength of each episode is set to 25 steps in the experiments. The experimental environments will be reset at the beginning of the new episode, randomly generating new locations for targets, agents, obstacles and forests. There are three defending agents, one of which is the commander and others are ordinary guardians, and there are five attacking agents, two of which are attackers and three are deceivers. In addition, there are 2 targets, 1 obstacle, and 2 forests in the experimental environments, as shown in Figure 2.

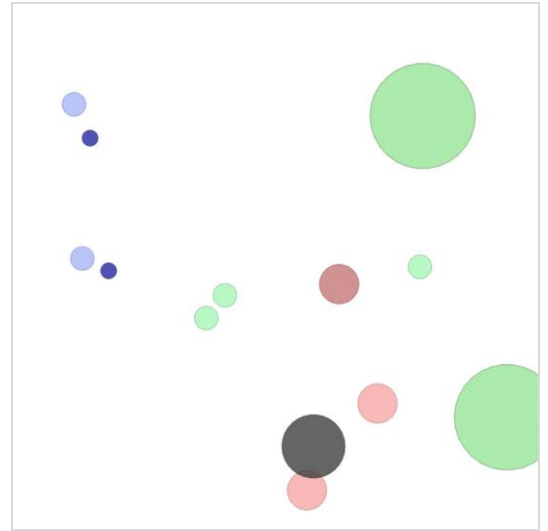


Figure 2. Multi-agent cooperative attack and deception mission scenarios. The larger red circles indicate the defending agents, among which the dark red one is the commander; the smaller blue and green circles are the attackers and deceivers in attacking agents respectively. The dark blue dots indicate the targets, black circles indicate obstacle, and green large circles indicate forests.

The experiment results show that multi-agent cooperative attack and deception mission in the above experimental environments can achieve satisfied results after about 80,000 times of training. The experiment results after 81,000 episodes of training are given below. It can be observed that the three defending agents are completely confused by the deceivers in the attacking party, neither defending the attackers nor defending and staying near the targets. The attackers successfully avoid the defending agents and approach the targets in the mission of this episode.

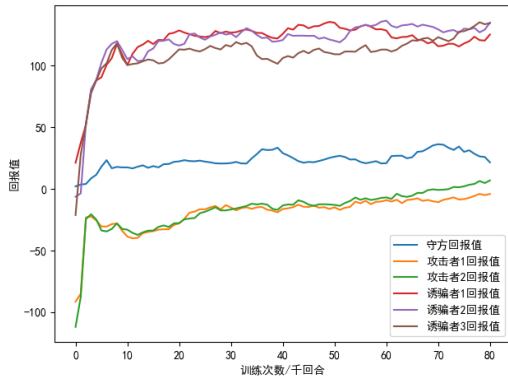


Figure 3. The reward values of agents for the cooperative attack and deception mission during 81,000 training episodes. The ordinate indicates the sum of the reward values per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

Figure 3 shows the reward value curves of all agents in the training process. In the first 10,000 episodes, with the various agents moving from blind action to gradually learning their own tasks, the reward values rise rapidly. Then, as the opponents' tactical level improved, their own reward values decrease significantly. In the subsequent training process, the cooperative tactics of both the attack and defense sides have gradually matured, and the reward values of all types of agents have increased. After training up to 70,000 episodes, the attacking agents gain a comparative advantage and their reward values continue to rise, while the return values of the defending agent decrease significantly.

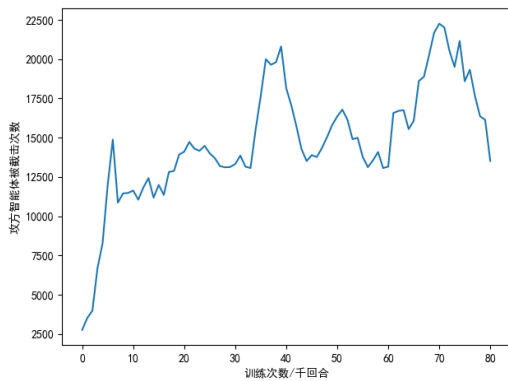


Figure 4. The number of times attacking agents are intercepted by defending agents during the 81,000 training episodes. The ordinate indicates the sum of the number of

interceptions per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

Figure 4 and Figure 5 show the number of times all the attacking agents are intercepted by the defending agents and the number of times attackers are intercepted per 1,000 episodes during the training process respectively. It can be observed that as the tactical ability of the defending agents continue to enhance after the start of training, the number of times attacking agents are intercepted increases rapidly, reaching a maximum of 14,724 times per thousand episodes. However, after only a few thousand episodes of training, the attacking agents soon learn to avoid being intercepted by the defending agents, and the number of interceptions decreases significantly. After about 30,000 episodes of training, the number of interceptions fluctuates between 12,000 and 22,500 per thousand episodes as the competition between the defending and attacking party becomes more intense. The number of times attackers are intercepted increases rapidly during the initial stage of training, and then decreases significantly. Between 29,000 and 35,000 episodes, the number of times attackers are intercepted is only about one-tenth of the number of times all attacking agents are intercepted, indicating that most of the defending agents are induced by the deceivers. However, as can be seen from Figure 6, the attacker can only reach the target about once per episode on average, which means that the attackers have not learned to attack the targets cooperatively.

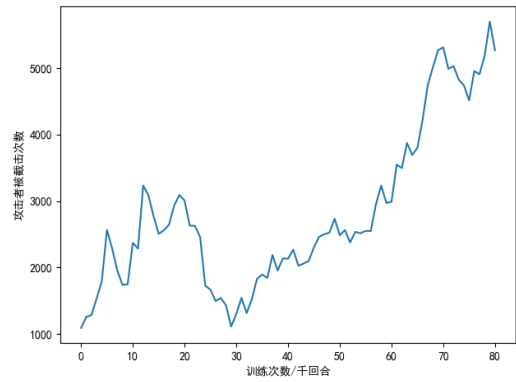


Figure 5. The number of times attackers are intercepted by defending agents during the 81,000 training episodes. The ordinate indicates the sum of the number of interceptions per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

As shown in Figure 5 and Figure 6, the number of times the attackers reach the targets is significantly more than the number of interceptions after about 70,000 episodes of training, indicating that the cooperative attack tactics of the attacking agents are effective at this time. When training to 80,001 to 81,000 episodes, the number of times attackers reach targets peaks at 16,040 per 1,000 episodes. As the tactical ability of the defending agents continue to enhance at this time, the number of times attackers are intercepted also rises rapidly. If training continues, the attackers' tactics will tend to avoid the risk of being intercepted and even give up

many opportunities to get close to targets. Therefore, the training process is terminated at 81,000th episodes. Fig. 1 meant 16000 A/m or 0.016 A/m. Figure labels should be legible, approximately 8 to 12 point type.

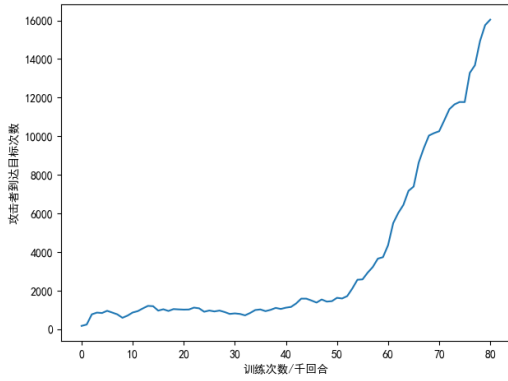


Figure 6. The number of times the attackers reach the targets during 81,000 episodes of training. The ordinate indicates the sum of the number of times attackers reach targets per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

5.2 The controlled experiment

In order to reveal the role of deceivers more directly and verify the effect of deception tactics, the controlled experiment has been carried out. In the controlled experiment, the mission and the experimental environment are exactly the same as the previous experiment, the only difference is the setting of the agents' reward function. The deceivers' action policy is simply to avoid being intercepted by the defending agents, instead of actively approaching defending agents at appropriate distance. In the following controlled experiment, agents are also trained 81,000 episodes. The results are shown below.

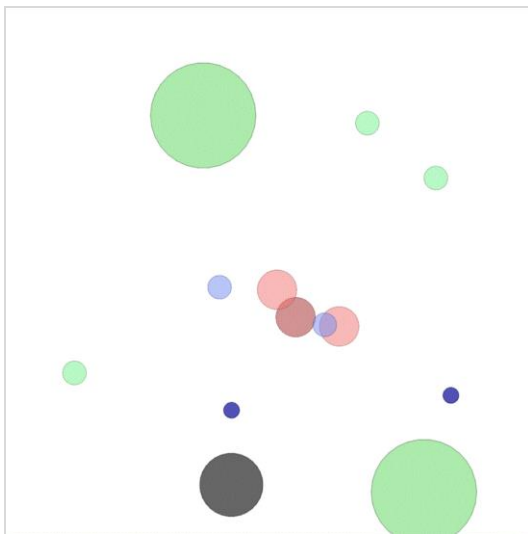


Figure 7. The controlled experimental scenario without deception tactics. The larger red circles indicate the defending agents, among which the dark red one is the

commander; the smaller blue circles are the attackers, and the small green circles are deceivers without deception tactics. The dark blue dots indicate the targets, black circles indicate obstacle, and green large circles indicate forests.

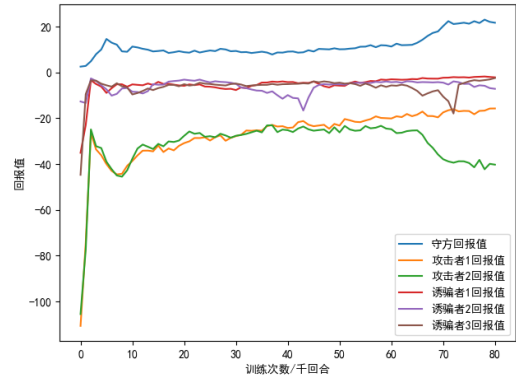


Figure 8. The reward values of agents in the controlled experiment. The ordinate indicates the sum of the reward values per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

In Figure 7, it can be observed that the attackers are surrounded by defending agents, while the attacking agents that originally played the roles of deceivers can only elude the defending agents and can no longer effectively deceive defending agents. Figure 8 shows the reward value curves of agents in the training process. It can be observed that after training to 65,000 episodes, the reward values of defending agents are rapidly increased, while the reward value of an attacker in attacking agents decreases significantly. This shows that the tactical mature defending agents can distinguish different types of attacking agents. After the attackers lose the cover of the deceivers, they are more likely to be intercepted by defending agents.

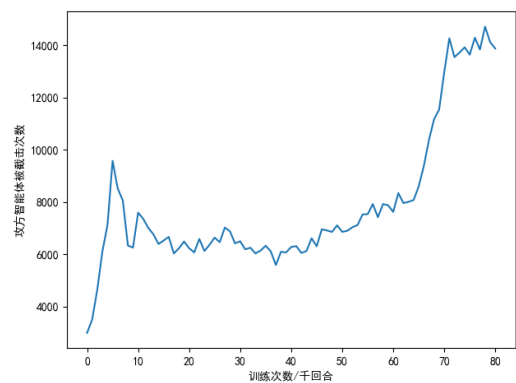


Figure 9. The number of times attacking agents are intercepted by defending agents during the 81,000 training episodes. The ordinate indicates the sum of the number of interceptions per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

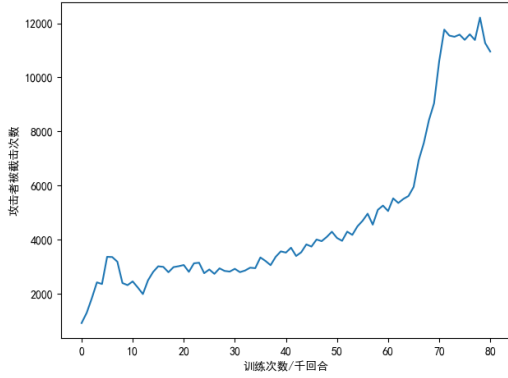


Figure 10. The number of times attackers are intercepted by defending agents during the 81,000 training episodes. The ordinate indicates the sum of the number of interceptions per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

Figure 9 and Figure 10 respectively show the number of times all attacking agents are intercepted and the number of times attackers are intercepted during the training of the controlled experiment. It is easy to see that the proportion of the number of times attackers are intercepted accounts for the number of times all attacking agents are intercepted is significantly higher than that of the above situation with deception tactics. When training to around 70,000 episodes, the number of times attackers are intercepted quickly rises to more than 10,000 times. The attacking agents intercepted by defending agents at this time are almost all attackers. In other words, the three attacking agents without deception policy in the controlled experiment can hardly cover the attackers.

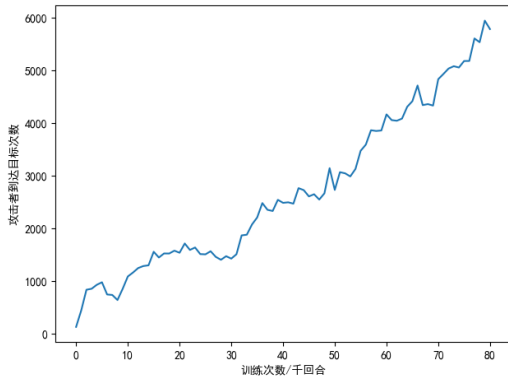


Figure 11. The number of times the attackers reach the targets during 81,000 episodes of training. The ordinate indicates the sum of the number of times attackers reach targets per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

Figure 11 shows the number of times the attackers reach targets during the training process in the controlled experiment, which shows a nearly linear increase with the number of training episodes. When train to 79,001 to 80,000 episodes, the attackers reach the targets 5,946 times per thousand episodes. Compared with the above scenarios with deception tactics, the number of times of reaching targets in

the controlled experiment decreased by 62.5%. Since the number of times attacker are intercepted has far exceeded the number of times of reaching targets, most attacks are actually failed.

The results of the above controlled experiments reflect that the deception tactics play an important role in multi-agent cooperative competition. By actively approaching the defending agents and keeping an appropriate distance, the deceivers can confuse and induce opponents, thus covering and helping the attackers to safely complete their attacking mission.

5.3 The effect of training episodes on multi-agent policy

In the training process, the competition situation and agents' tactical level will directly affect the opponents' policy optimization process. To achieve the goal of the cooperative mission in this paper, the two types of attacking agents should learn to cooperate with each other to perform their tasks. On the other hand, the two types of defending agents should learn to work together to besiege and intercept attacking agents, so as to maximize the reward value of all of them.

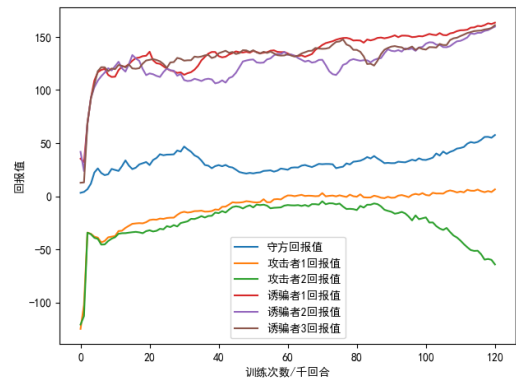


Figure 12. The reward values of agents for the cooperative attack and deception mission during 121,000 training episodes. The ordinate indicates the sum of the reward values per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

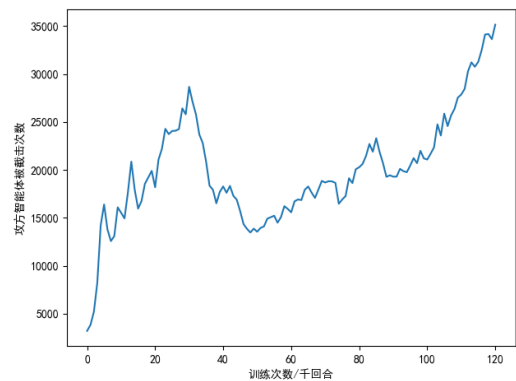


Figure 13. The number of times attacking agents are intercepted by defending agents during the 121,000 training episodes. The ordinate indicates the sum of the number of interceptions per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

Figure 12 to Figure 15 show the results of three defending agents against two attackers and three deceivers during 121,000 episodes of training. As the number of training episodes increases, it can be observed that the tactics of the defending agents are more and more tend to work together to besiege a single attacking agent, and the total number of interceptions and the number of times the attackers are intercepted are rapidly increased, as shown in Figures 12, 13 and 14. After training to about 90,000 episodes, this trend will force the attackers to adopt a more conservative policy of avoiding being intercepted rather than getting close to targets, as shown in figure 15.

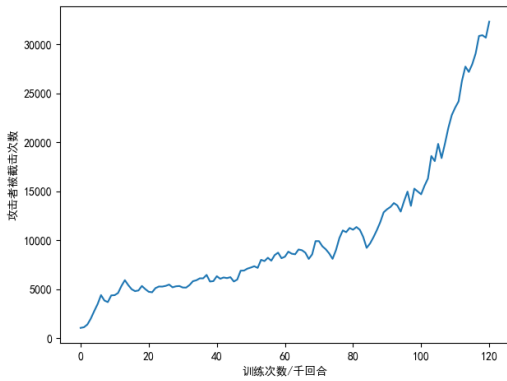


Figure 14. The number of times attackers are intercepted by defending agents during the 121,000 training episodes. The ordinate indicates the sum of the number of interceptions per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

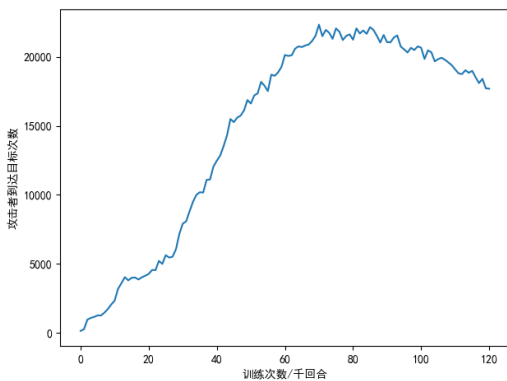


Figure 15. The number of times the attackers reach the targets during 121,000 episodes of training. The ordinate indicates the sum of the number of times attackers reach targets per 1,000 training episodes; the abscissa indicates the number of training episodes in thousands of episodes.

The above results show that if the agents encounter powerful opponents and accumulates too much failure experience in the training process, the agents' policy will be more and more conservative. As the training episodes increases, the attacking agents will gradually give up the opportunity to take the risk to complete their mission, which will eventually lead to the agents learning passive escape policy. This is obviously inconsistent with the expected goal of the cooperative competitive mission. Therefore, the number of training episodes should be reasonably set according to the requirements of the mission and the data of experiments.

By comparing the data of above experiments where three defending agents (including one commander and two guardians) and five attacking agents (including two attackers and three deceivers) are set up, it can be observed that the attacking agents can learn relatively satisfied tactical policy after about 80,000 episodes of training.

5.4 The effect of the number of deceivers on cooperative attacking

In the cooperative competition mission in this paper, the deceivers play important roles in inducing defending agents and covering attackers to reach targets safely. This paper researches the computer experiment results of cooperative attacking and deception tactics with variable number of deceivers by tuning the number of deceivers. It is found that the number of deceivers has a significant effect on the performance of deception tactics. The length of the training is set to 81,000 episodes in the experiments, the number of deceivers ranges from 0 to 4, and other settings in the experimental environment are consistent with the previous experiments. Table 1 statistics the experiments data of the last 1,000 episodes of the training process, i.e., the total number of interceptions, the number of times attackers are intercepted and the number of times attackers reach targets in the 80,001-81,000 episodes. The proportion of the number of times attackers are intercepted accounts for the total number of interceptions is also listed in Table 1.

Table 1. Statistical results of training data for the 80001-81000 episodes. N_d is the number of deceivers, T is the total number of interceptions, T_a is the number of times attackers are intercepted, T_r is the number of times attackers reach targets.

N_d	T	T_a	T_a/T	T_r
0	9666	9666	100%	1446
1	10938	7818	71.48%	3950
2	17679	8970	50.74%	10065
3	13512	5271	39.00%	16040
4	17793	4929	27.70%	5717

From the data in Table 1, it can be observed that the number of times attackers are intercepted shows a roughly downward trend with the increase of the number of deceivers. But when the number of deceivers increases to two, the number of times attackers are intercepted exceeded that of only one deceiver, which is caused by the inherent randomness of reinforcement learning. This random disturbance is reflected in the statistical data, which inevitably caused the training results to deviate from the overall trend to some extent. However, the proportion of the number of times attackers are intercepted accounts for the total number of interceptions decreases monotonously as the number of deceivers increases. This demonstrates that the more deceivers there are, the more defending agents they can induce, and the better the attackers' chances of surviving.

This paper tests the policy network model obtained after 81,000 episodes of training. The testing environment is set as the same as the training environment. The length of test is set to 1,000 episodes, and the test data are listed in Table 2.

Table 2. Statistical results of test data of 1000 episodes. N_d is the number of deceivers, T is the total number of interceptions, T_a is the number of times attackers are intercepted, T_r is the number of times attackers reach targets.

N_d	T	T_a	T_a / T	T_r
0	6030	6030	100%	1287
1	7515	5415	72.06%	3067
2	11436	5259	45.99%	5856
3	10365	4254	41.04%	10810
4	14865	3999	26.90%	3744

In Table 2, as the number of deceivers increases, the number of times attackers are intercepted monotonically decreases, and the proportion of the number of times attackers are intercepted accounts for the total number of interceptions also monotonically decreases. As shown in Table 1 and Table 2, when the number of deceivers increase from 0 to 4 in training and testing, the number of times attackers reach targets increases as the number of deceivers increases. This shows that, as more deceivers join, the attacking agents get a higher winning percentage. However, when the number of deceivers increases to more than 4, the previous training times are not enough for attackers to learn better policy in the face of more complex situations, so that the number of times attackers reach targets decreases compare to previous situation with less agents.

6. CONCLUSIONS

The decision-making ability of a single-agent system is far from enough for many complex decision-making problems in real-world scenarios. MADDPG provides an effective

method to solve the problem of multi-agent cooperative control. This paper realizes multi-agent cooperative competition and deception tactics by designing functional reward values for multi-agents based on MADDPG algorithm, and verifies the role of deception tactics through the controlled experiments. We set up multiple types of agents to share various missions and obtain the cooperative attacking and deception policy modes through computer experiments to realize the cooperation of multi-mission among multiple agents. This paper also researches the effect of training episodes on multi-agent learning policy in the experiments. The experiments data shows that too little training is not enough for agents to learn satisfied policy, while too much training may make agents learn overly aggressive or conservative policy, so the number of training episodes should be reasonably set in combination with the mission requirements and experimental results. Otherwise, we analyze the effect of the number of deceivers on the performance of deception tactics by tuning the number of deceivers in the experiments. The training and testing results show that increasing the number of deceivers in the mission can significantly improve the performance of deception tactics by covering attackers to reach targets safely from the interception of defending agents. But the computational complexity of the multi-agent environments will rise with the increase of the number of agents, and more training episodes are needed to ensure a satisfied learning effect.

Although the deep reinforcement learning method has made a significant breakthrough in solving cooperative mission of multi-agent, there are still some problems that have not been solved well. First, the number of training episodes relies on human experience to set in the existing multi-agent algorithms based on deep reinforcement learning. It is difficult to define how many episodes of training can obtain a model that best conforms to the expected result before a number of training and testing. Second, the existing deep reinforcement learning algorithms continuously optimize the agents' policy in the direction of maximizing the Q values in training process by updating network parameters. Training a type of agents in this way can only perform one type of mission, unable to balance multi-mission goals and flexibly adjust policy according to the dynamic environment, and the problem of multi-agent multi-mission cooperation can only be solved by assigning different missions to different types of agents. At last, as the deep reinforcement learning method has a strong inherent randomness, the experiments results are difficult to be reproduced. The random multi-agent dynamic environment and experience sampling process will cause the policy learning process to be unstable and make the experimental data fluctuate in a large range. The above three challenges restrict the application of deep reinforcement learning in practice. It is a practical problem in the cluster control of unmanned intelligent system, which is worth further exploring.

REFERENCES

- Arulkumaran K., Deisenroth M., Brundage M. and Others (2017). A brief survey of deep reinforcement learning. IEEE Signal Processing Magazine, 34, 26–38.

- Choi J., Lee B., and Zhang B. (2017). Multi-focus attention network for efficient deep reinforcement learning. In Workshops at the Thirty-First AAAI Conference on Artificial Intelligence.
- Foerster J., Assael I., de Freitas N. and Others (2016). Learning to communicate with deep multi-agent reinforcement learning. In Advances in Neural Information Processing Systems, pp. 2137–2145.
- Foerster J., Farquhar G., Afouras T. and Others (2018). Counterfactual multi-agent policy gradients. In Thirty-Second AAAI Conference on Artificial Intelligence.
- François-Lavet V., Henderson P., Islam R. and Others (2018). An introduction to deep reinforcement learning. Foundations and Trends $\mathbb{R}$ in Machine Learning, 11 (3-4), 219–354.
- Hastie T., Tibshirani R., and Friedman J. (2009). The elements of statistical learning: data mining, inference, and prediction. Springer.
- Hernandez-Leal P., Kartal B., and Taylor M. E. (2019). A survey and critique of multiagent deep reinforcement learning. Autonomous Agents and Multi-Agent Systems, 33 (6), 750–797.
- Jiang J., and Lu Z. (2018). Learning attentional communication for multi-agent cooperation. In Advances in Neural Information Processing Systems, pp. 7254–7264.
- Krizhevsky A., Sutskever I., and Hinton G. E. (2012). Imagenet classification with deep convolutional neural networks. In Advances in Neural Information Processing Systems, pp. 1097–1105.
- LeCun Y., Bengio Y., and Hinton G. (2015). Deep learning. nature, 521 (7553), 436–444.
- Lillicrap T., Hunt J., Pritzel A. and Others (2015). Continuous control with deep reinforcement learning. In arXiv preprint, Vol. 1509.02971.
- Liu Q., Zhai J., Zhang Z. and Others (2018). A survey on deep reinforcement learning. Chinese Journal of Computers, 40, 1–27.
- Lowe R., Wu Y., Tamar A. and Others (2017). Multi-agent actor-critic for mixed cooperative competitive environments. In Advances in Neural Information Processing Systems, pp. 6379–6390.
- OpenAI (2018). Openai five. In <https://blog.openai.com/openai-five/>.
- Silver D., Huang A., Maddison C. and Others (2016). Mastering the game of go with deep neural networks and tree search. Nature, 529, 484–489.
- Silver D., Hubert T., Schrittwieser J. and Others (2017). Mastering chess and shogi by self-play with a general reinforcement learning algorithm. In arXiv preprint, Vol. 1712.01815.
- Silver D., Hubert T., Schrittwieser J. and Others (2018). A general reinforcement learning algorithm that masters chess, shogi, and go through self-play. Science, 362, 1140–1144.
- Silver D., Lever G., Heess N. and Others (2014). Deterministic policy gradient algorithms. In International Conference on Machine Learning, pp. 387–395.
- Silver D., Schrittwieser J., Simonyan K. and Others (2017). Mastering the game of go without human knowledge. Nature, 550, 354–359.
- Sutton R., and Barto A. (2018). Reinforcement learning: An introduction. MIT press.
- Tampuu A., Matiisen T., Kodelja D. and Others (2017). Multiagent cooperation and competition with deep reinforcement learning. Plos One, 12(4), e0172395.
- Vinyals O., Babuschkin I., Czarnecki W. and Others (2019). Grandmaster level in starcraft ii using multi-agent reinforcement learning. Nature, 575, 350–354.