

Title	仲間の立場から人間の行動を誘導するゲーム AI に向けての考察
Author(s)	山田, 直央
Citation	
Issue Date	2022-03
Type	Thesis or Dissertation
Text version	author
URL	http://hdl.handle.net/10119/17640
Rights	
Description	Supervisor:池田 心, 先端科学技術研究科, 修士(情報科学)

In recent years, as the improvements in the communication techniques, the delay problem has been reduced. This enables people living remotely and a large number of people to play games even that lags of several frames may cause discomfort. Among games that have benefited from this fact, multi-player shooting games are a popular genre worldwide. It is also a popular genre for e-sports.

Many people want to start playing games belonging to this genre. However, when beginners with poor skills play these games, they are likely to be beaten easily by other players due to the difference in skills. Beginners may have no much opportunity to learn through trial and error while enjoying the games. Also, in each game, specific gestures are used to communicate intentions. It is often that beginners overlook such gestures, making them difficult to exchange information that leads to clearing the games. In this sense, players new to these games do not have time to learn how to play and may not enjoy the games because they cannot communicate smoothly with their teammates.

If machines can communicate in a way that is easy to understand and human-like, they can behave in such a way and serve as a starting point for humans to learn strategies for playing games. These machines will provide more opportunities for learning strategies. In order to provide humans with an appropriate environment for practice and learning, we consider it necessary to model intention communications that are conveyed without words.

To model intention communications that are conveyed without words, we need to analyze and reproduce the following topics: *how humans communicate intentions*, *how to utilize human-like intention communications*, and *how intention communications influence the opponents/teammates*. In this research, we aim at the second topic, analyzing and reproducing how to utilize human-like intention communications.

For this purpose, we discussed 3 conditions that human-like intention communications are likely to occur during gameplay: *human players have different information*, *it is easy to observe teammates' actions*, and *players have some rooms in terms of action space and some spare time*. Based on these conditions, we implemented such an experiment environment that intention communications are likely to occur. In more detail, each intention communication in this experiment environment is performed by a button push, making us easy to collect labels of the timings when the intention communications took place. In this research, we had 6 kinds of intention

communications, coming from 3 circumstances: during stage selection, during item emergence, and during battle, where each has 2 choices (2 buttons).

We collected about 1,351 minutes of human behavior data from the implemented experiment environment. We then used these data to do supervised learning to predict what kinds of intention communications will be done at which timings during gameplay. The learning network contained the architecture called C-LSTM, a combination of convolutional layers and LSTM (long short-term memory) layers. For the input data, we used map data of 86×43 and 56 data not in the same form as maps. Map data were input to the convolutional layers, and the obtained outputs were combined with the remaining input data, which were then input to LSTM layers. We set the output data of the network to 3 parameters, representing *whether the players are communicating intentions*, *whether the players press button 1 for intention communication*, and *whether the players press button 2 for intention communication*.

For the supervised learning, we tried different settings for comparison, including masking some information from the input data and varying the time length of the input data. We compared the prediction accuracy of these networks and confirmed the following results: the prediction on timings were strongly influenced by *which human players were put in the same team*, and the accuracy was higher when the input data contained 100 frames (10 seconds) than 500 frames (50 seconds). Among the predictions by the networks, the highest F-score was 0.306, which was not a high value. However, since it is known that there are differences between humans in *whether to press the buttons or not and how fast or slow they press the buttons*, the result alone does not mean poor performance.

In addition, we conducted an experiment to confirm whether the timings of doing intention communications by the proposed method were natural to human players. We made two sets of videos based on the collected human behavior data: *videos showing the timings that intention communications were actually done* and *videos showing the timings that intention communications were predicted to happen by the supervised learning model*. Participants of this experiment were asked to watch the videos and answer whether they would do intention communications at the next moment. 62.7% of the answers matched the data for making the former set of videos (those showing the timings that intention communications were actually done), and 55.0% matched the latter set (videos showing the timings that intention communications were predicted to happen by the supervised learning model). From the results, the naturalness of the timings predicted by the proposed method was slightly worse than the timings from human data. However, compared to the bad impression received from an F-score of 0.306, the difference from

human data in this experiment was much more natural. From this point, we believe that the goal of this research has been achieved to some extent.