



ARL-MR-1120 • FEB 2025



Generalization between Two Human–Automation Collaborative Tasks via Data-Based Transfer Learning

by Torin Adamson, Liz DiGioia, Yazied Hasan, Ava Naffah,
Matthias Mehl, Lydia Tapia, and Evan C. Carter

DISTRIBUTION STATEMENT A. Approved for public release: distribution is unlimited.

NOTICES

Disclaimers

The findings in this report are not to be construed as an official Department of the Army position unless so designated by other authorized documents.

Citation of manufacturer's or trade names does not constitute an official endorsement or approval of the use thereof.

Destroy this report when it is no longer needed. Do not return it to the originator.



Generalization between Two Human–Automation Collaborative Tasks via Data-Based Transfer Learning

Liz DiGioia, Yazied Hasan, and Lydia Tapia
University of New Mexico

Ava Naffah and Matthias Mehl
University of Arizona

Torin Adamson and Evan C. Carter
DEVCOM Army Research Laboratory

REPORT DOCUMENTATION PAGE					
1. REPORT DATE		2. REPORT TYPE		3. DATES COVERED	
February 2025		Memorandum Report		START DATE 1 October 2023 END DATE 30 September 2024	
4. TITLE AND SUBTITLE Generalization between Two Human–Automation Collaborative Tasks Via Data-based Transfer Learning					
5a. CONTRACT NUMBER		5b. GRANT NUMBER		5c. PROGRAM ELEMENT NUMBER	
5d. PROJECT NUMBER		5e. TASK NUMBER		5f. WORK UNIT NUMBER	
6. AUTHOR(S) Torin Adamson, Liz DiGioia, Yazied Hansan, Ava Naffah, Matthias Mehl, Lydia Tapia, and Evan C. Carter					
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) DEVCOM Army Research Laboratory ATTN: FCDD-RLA-FA Aberdeen Proving Ground, MD 21005				8. PERFORMING ORGANIZATION REPORT NUMBER ARL-MR-1120	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)			10. SPONSOR/MONITOR'S ACRONYM(S)		11. SPONSOR/MONITOR'S REPORT NUMBER(S)
12. DISTRIBUTION/AVAILABILITY STATEMENT DISTRIBUTION STATEMENT A. Approved for public release: distribution unlimited.					
13. SUPPLEMENTARY NOTES ORCID: Evan C. Carter, 0000-0001-7471-8769					
14. ABSTRACT When collaborating with AI, human behavior can evolve unpredictably outside a controlled lab environment. Improving our understanding of these variations in behavior requires a larger volume of data. Our approach gamifies human–AI collaboration with a mobile video game, making it feasible to collect sufficient data by repeating experiments daily over a long period of time. This data can then be analyzed to generalize knowledge of human–AI collaboration across different tasks. To demonstrate this, we presented 83 participants with one of two dynamic obstacle-avoidance tasks over the course of a 180-day study. The two tasks differed in the role played by the human. They consisted of a supervisory role (where the human takes control from the AI when they believe it is necessary) and a defense role (where the human prevents objects from reaching the AI-controlled payload). To evaluate whether collaboration behavior could be generalized, a transfer learning model was trained on participant data from the supervisor task to predict behavior in the defensive task. The capability to generalize using transfer learning was demonstrated when models showed improved prediction of the defensive task after data collected from both experimental tasks was combined in training sets in unequal measure.					
15. SUBJECT TERMS human–autonomy teaming, mobile sensing, mobile video games, ambulatory assessment, transfer learning, machine learning, Humans in Complex Systems					
16. SECURITY CLASSIFICATION OF:				17. LIMITATION OF ABSTRACT	
a. REPORT UNCLASSIFIED		b. ABSTRACT UNCLASSIFIED		c. THIS PAGE UNCLASSIFIED	
				UU	
18. NUMBER OF PAGES 23					
19a. NAME OF RESPONSIBLE PERSON Evan C. Carter				19b. PHONE NUMBER (Include area code) 520-271-4966	

STANDARD FORM 298 (REV. 5/2020)
Prescribed by ANSI Std. Z39.18

Contents

List of Figures	iv
List of Tables	v
1. Introduction	1
2. Background	2
3. Methods	3
3.1 Participant Study with Busy Beeway	3
3.2 Two Series of Experiments	4
3.3 Data Preparation	4
3.4 Data Analysis and Player Preference Transfer	7
4. Results and Discussion	9
5. Conclusions	11
6. References	12
List of Symbols, Abbreviations, and Acronyms	15
Distribution List	16

List of Figures

- Fig. 1 Performance of each trained model on predicting preference in the target task. Each model is trained with different ratios of source:target data: the blue points represent samples from AI 3, the orange points represent samples from AI 4, and the red dotted line represents a simple line fitted to the data. 11

List of Tables

Table 1	Mean and standard deviation prediction performance (root-mean squared error) of the 10 trained models on the holdout fold of the k-fold cross validation, the test data of the source task, and the test data of the target task. Each model is trained on a different source:target ratio of training data.	7
Table 2	The amount of time (given as a ratio) each participant spent in the defender role is listed and separated by the AI agent that was assigned. These ratios are of the combined total time spent with each AI throughout the entire study period.....	10

1. Introduction

Existing human–automation interaction studies typically focus on specific aspects in the design of systems under tight experimental control in laboratory settings. Although this yields robust data useful for improving these specific designs, these kinds of studies do not capture realistic variation in human–automation behavior. Outside of the laboratory, the human’s interactions with the AI agent could be affected by changing internal (e.g., stress) or external (e.g., environmental) factors, which we refer to as the “human context.” In addition to restricted context, typical human–AI interaction studies tend to fix the roles of the humans and agents given the specific application under study (e.g., the human supervising an autonomous vehicle). More flexible research methods that place humans and automation in various roles offer the potential to produce results that could be generalized to other applications.

In this report, we detail analyses of an ongoing study that combines two research platforms, Busy Beeway and StudentLife, to gather data. The former conducts human–automation studies in a flexible manner by supporting remote configurable experimental parameters,¹ and the latter gathers human context data through wearable devices and mobile phones.² Each participant is prompted to complete daily motion planning tasks in Busy Beeway for 180 days while StudentLife records context continuously. Two separate series of experiments are included in this study that place the participant in supervisory and defending roles. At the time of writing, 58 participants have completed the first study, and 30 have completed the second. The total daily sessions of Busy Beeway recorded from both studies are 7,546 and 3,653, respectively.

The similarities between the two roles allow us to use the knowledge gained about the behavior of humans in the supervisory role to help understand their behavior in the defending role. A direct application of this approach would be through transfer learning, which is the concept of reusing parts of an ML model to improve the training on a new, unseen task. By showing that a learning model trained with knowledge from the supervisory role can be more accurate than a model that does not leverage previous knowledge, we demonstrate the generalizability and utility of the knowledge gained from such specialized tasks.

2. Background

Existing research on systems that pair humans with AI agents focuses on important factors that affect human behavior and overall performance of the team. These include studies on improving the legibility of AI behavior so the human can better understand its intentions,^{3,4} how to manage a level of trust appropriate to the AI’s capabilities,^{5,6} and how the two can communicate through interfaces that do not place too large a cognitive load on the human.⁷ A large amount of this research is conducted in controlled lab environments to investigate how humans act in supervisory roles with autonomous vehicles, often in realistic simulators that recreate the driver’s seat physically.^{6,8–10} While this approach is immersive and yields robust data for important factors in designing human–AI teams, less is known about human behavior in everyday life outside the lab or when the human is placed into different roles. For example, in multiagent systems where a valuable payload must be escorted by defender units, many existing studies focus on solutions that rely on automation alone.^{11–13} The behavior of humans may not be consistent over longer periods of time as they become accustomed to the system or are influenced by external factors. This work seeks to expand on the existing body of research by conducting research over long periods of time that place human subjects in different roles as they collaborate with AI to solve obstacle-avoidance tasks.

The difference in experimental designs required for the different human roles and long study periods make conventional in-lab methods infeasible. Busy Beeway allows this flexibility by adapting experimental tasks in a video game format.¹ This is a process known as “gamification” that has also been applied to make protein-folding tasks,¹⁴ image-labeling tasks,¹⁵ and robotic swarm-control tasks¹⁶ more familiar and enjoyable to participants. By targeting mobile devices in particular, this also allows for each subject to participate in their own environment and submit collected data remotely. This simplifies the logistics of longitudinal studies by using daily gameplay sessions of Busy Beeway.

To understand the cooperative behavior of humans with their AI partners, we operationalize human preference as the way in which the human modifies overall system behavior as compared to what would have happened without human input. This concept is inspired by previous work on the linearly-solvable Markov decision process (LMDP) framework,^{17,18} which considers control as deviations from so-called

passive dynamics that would occur without any input. Our approach is aimed at learning how human preferences change as a function of different AI partners and different role requirements in the human–AI team.

To show that the knowledge extracted about the players’ preferences can extend across different roles the human may assume, we perform transfer learning on the player preferences across the different player roles. Transfer learning is the practice of using the knowledge gained from learning one task, known as the source, to help the learning process on a different task, known as the target.^{19,20} The transfer of knowledge can be in the form of transferring the data collected to train a model, transferring the learned weights of a network, or reusing a model entirely to predict a new task. This has been shown to help improve performance, reduce training time, and produce robust models. For the purposes of this work, we focus on the performance benefits of transfer learning as an indicator of the generalizability and utility of the knowledge gained from the source task.

3. Methods

3.1 Participant Study with Busy Beeway

Busy Beeway is a video game designed for mobile devices that collects data from human interactions with AI as they cooperate to solve motion planning tasks.¹ It can deploy different experimental conditions directly to individual participants and collect the resulting gameplay and self-report survey data. In this research, participants and AI agents controlled a payload (represented by the bee character under camera focus) to reach seven target locations in order while avoiding obstacles moving with stochastic dynamics. Participants completed three of levels of this task in a session, with difficulty increased at each level through the addition of obstacles.

The data used in this report is sourced from an ongoing study. The participants were recruited at the University of Arizona. During the onboarding process, Busy Beeway was installed on the participant’s phone, a demographic survey was administered, and the game mechanics were explained with a tutorial. Over the course of a 180-day study period, participants were prompted to play daily sessions of Busy Beeway with a different experimental condition each day. Details of the onboarding process and study protocol can be found in Adamson et al.²¹

3.2 Two Series of Experiments

The ongoing study currently involves two series of experiments where the same obstacle-avoidance task is presented, but with the humans and AI agents in different roles. The first series of experiments observed patterns of human intervention in a shared control system where the AI agent was in control unless the participant chose to assume control by interacting with the virtual on-screen joystick. Once they took control, the participant could release control back to the AI by releasing the joystick. In each daily session, the participant was given one of four different AI agent configurations which varied in the degree to which they prioritized reaching the goal swiftly or avoiding the obstacles. More details on this study and its gameplay can be found in Adamson et al.²¹

The second series of experiments also required human intervention; however, the human and AI did not share control of one entity. Rather, there were two entities performing separate roles (a defender and a payload). The defender could repel obstacles by colliding with them while the payload was required to avoid the obstacles or risk crashing. Each entity was controlled either by the human or AI, but not both. If the human chose to stop controlling their entity, that entity took no further action until control was resumed by the human. In each daily session, the human was assigned to play as either the payload or the defender, or was given the ability to choose which entity they would rather control. In this latter experiment, the human could swap entities at will. When the human was assigned to the defender role or was allowed to choose, they were assigned to collaborate with either the Evasive AI or Hasty AI (see Section 3.3). More details on this study can be found in DiGioia et al.²²

3.3 Data Preparation

Two AI agents between both series of experiments are of the same algorithm and configuration in each. They will serve as a reference to understand changes in participant behavior as the roles change. One agent was designed to prioritize collision avoidance and thus is referred to as the evasive AI in this report. The other AI was designed to prioritize reaching the goal as quickly as possible and is referred to as the hasty AI. Because different AI agents had different priorities, the participants were required to modify their own strategies to maintain a desired balance of speed and safety.

Only data from participants who successfully completed their 180-day study period are considered in this report. In the first series, only sessions where participants shared control with the hasty and evasive agents are included in the analysis, ignoring the two other agents. For the second series, the instances where the participant was in the defender role were included. To further compare the actions of the participant against the AI agents, control data with the participant in the payload role is also analyzed.

To prepare the data for the learning tasks, we mapped the categorical features to one-hot encoding, and we normalized the scale of the numerical features to the range of [0–1]. The resulting features we used are as follows:

- *DayInStudy*: Which day in the 180-day study the session was held.
- *AI*: 3 distinct one-hot encoded features where one indicates the assigned AI of the session. Options include AI1, AI2, and None.
- *Upset*: Daily self-report of mood as selected on a 5-point Likert scale.
- *Interested*: Daily self-report of mood as selected on a 5-point Likert scale.
- *Determined*: Daily self-report of mood as selected on a 5-point Likert scale.
- *Irritable*: Daily self-report of mood as selected on a 5-point Likert scale.
- *ElapsedDays*: The number days of active participation.
- *Behavioral Inhibition System*: This measures concerns regarding the possible occurrence of negative events and the sensitivity to such events when they do occur. Integer value up to 8.
- *Behavioral Approach System (BAS)-R*: Reward responsiveness, comprising items that focus on positive responses to the occurrence or anticipation of reward. Integer value up to 8.
- *BAS-F*: Fun-seeking, comprising items that reflect both a desire for new rewards and a willingness to approach a potentially rewarding event on the spur of the moment. Integer value up to 8.
- *BAS-D*: Drive, comprising items that reflect a concern for the possible occurrence of negative events and sensitivity in the response to their occurrence. Integer value up to 8.

- *Gender*: Three distinct features where one indicates the selected gender. Options are Male, Female, and Other.
- *Ethnicity*: Five distinct features where one indicates the selected ethnicity. Options are American Native, Asian or Pacific Islander, Black or African American, White/Caucasian, Other.
- *Education*: Five distinct one-hot encoded features where one indicates the selected education level. Options are Below High School, High School or Equivalent (GED), 2-Year Associates Degree, 4-Year Bachelors Degree, Masters Degree, Doctorate Degree.
- *Occupation*: Twenty-four distinct features where one indicates the possible occupation. Options are Management; Business and Financial Operations; Computer and Mathematical; Architectural and Engineering; Life, Physical, and Social Science; Community and Social Service; Legal; Education, Training, and Library; Arts, Design, Entertainment, Sports and Media; Healthcare Practitioners; Healthcare Support and Technical; Protective Service; Food Preparation and Serving; Building, Grounds Cleaning and Maintenance; Personal Care and Service; Sales and Retail; Office and Administrative Support; Farming, Fishing and Forestry; Construction and Extraction; Installation, Maintenance and Repair; Production; Transportation and Materials Moving; Student; (Not Listed Here); Unemployed.
- *GameSelf*: Demographic survey on single-player gaming habits, as selected on a 10-point frequency scale. Options are Never, Once a Month, Several Times a Month, Once a Week, Several Times a Week, Once a Day, Several Times a Day, Once an Hour, Several Times an Hour, All The Time.
- *GameRoom*: Demographic survey on local multiplayer gaming habits, as selected on a 10-point frequency scale. Options are Never, Once a Month, Several Times a Month, Once a Week, Several Times a Week, Once a Day, Several Times a Day, Once an Hour, Several Times an Hour, All The Time.
- *GameOnline*: Demographic survey on online multiplayer gaming habits, as selected on a 10-point frequency scale. Options are Never, Once a Month, Several Times a Month, Once a Week, Several Times a Week, Once a Day, Several Times a Day, Once an Hour, Several Times an Hour, All The Time.

Additionally, the data was segmented to allow for testing and training. The segmen-

tation process happened as follows: first, the data was separated by series, resulting in a source and target datasets that correspond to the two tasks. Each dataset was then shuffled independently. Second, we reserved 30% of each dataset as a test dataset, leaving the remaining 70% to be sampled for training. Third, we created training datasets by sampling data from the unreserved source and target data from the previous step based on the source:target ratios, specified in Section 4. Finally, for the training process, we performed 10-fold cross validation on the desired training set. This process involved splitting the training data set into 10 subsets, or folds. Then, one of the folds was held out as a validation set and a model was trained on the remaining nine folds and assessed on the holdout. This process was repeated for each of the folds, resulting in 10 models. After training was completed, we tested the 10 models on the test sets of the source and target tasks that were withheld in the second step described above. Means and standard deviations of prediction performance from the 10 models on the validation sets, source test data, and target test data are given in Table 1. Prediction performance is defined as root-mean squared error.

Table 1. Mean and standard deviation prediction performance (root-mean squared error) of the 10 trained models on the holdout fold of the k-fold cross validation, the test data of the source task, and the test data of the target task. Each model is trained on a different source:target ratio of training data.

Source:target ratio	Validation		Source test		Target test	
	Mean	STD	Mean	STD	Mean	STD
0:1	0.010	± 0.001	0.313	± 0.001	0.095	± 0.002
1:0	0.006	$\pm <0.001$	0.076	± 0.002	0.270	± 0.008
1:1	0.011	$\pm <0.001$	0.074	± 0.002	0.118	$\pm <0.001$
1:3	0.010	± 0.001	0.097	± 0.010	0.091	± 0.001

3.4 Data Analysis and Player Preference Transfer

When humans are teamed up with AI agents in uncertain environments (like those filled with stochastic dynamic obstacles in Busy Beeway), trust becomes an important factor. Using Lee and See’s definition of trust as the human’s belief that the agent will help them achieve their own goals,²³ we assume a divergence from the AI agent’s decision-making as a measure of the human’s preference. For this study, if a participant does not agree or trust an AI agent, we expect they will attempt to alter the AI behavior in some way using their own control input. This is directly

permitted in the first series of experiments, as the participant can take control away from the AI at any time. In the second, they can only influence the AI with whatever strategy they use to clear obstacles away from its path.

Strategies in the obstacle-avoidance task can prioritize either reaching the next goal quickly or avoiding obstacles. To measure which of these strategies are dominant in participant or AI decisions, the change in distance to the next goal can be taken among individual actions. If there were no obstacles, the optimal solution would always choose actions directly moving to the next goal, reducing this distance by the greatest amount. With obstacles, it is sometimes necessary to move to avoid a collision, so a smaller decrease or increase in goal distance would indicate avoidance. In this research, actions occur every game frame where the payload is moved in a direction. Comparing the actions of the payload under full AI control against the participant in various roles will reveal this difference in preference.

To quantify this difference, a ratio of rank-sum tests comparing the change in goal distances between two distributions of actions is taken. Given a distribution of actions taken by AI alone x , and a distribution of actions taken while the AI is teamed with a participant y , the ratio $d = \frac{U(x,y)}{U(y,x)}$ is calculated where U is the number of times a sample in the second distribution comes before (has a lower change in goal distance) a sample in the first (as calculated by MATLAB’s ranksum function). If this ratio is less than 1, the actions in y bring the payload further away from the goal, on average, than AI alone. If greater than 1, then the payload is brought closer to the next goal, on average, than AI alone.

The regression model of choice is the boosted ensemble trees²⁴ method (as implemented in the *fitresensemble()* function from the MATLAB statistics and ML package²⁵), which trains an ensemble of 100 regression trees. This method was chosen for its accommodation for the types of features in the data and the distribution of labels. The regression models are trained on the task of estimating the preference ratio d , based on the player attributes from both demographic and daily mood surveys, and the day of study and assigned AI, as described in Section 3. In addition to withholding 30% of the data for testing, we also used 10-fold cross-validation during the training.

4. Results and Discussion

In addition to the individual decisions made every frame, participant preference can manifest in a way specific to the two series of experiments. In the first study, the participant decides when (and if) to intervene and take control from the AI. Across all gameplay sessions, participants asserted control 10.19% of the time with the evasive AI and 21.49% with the hasty AI. This is expected, as the hasty AI is more prone to colliding with obstacles if left alone. In the second study, two of the five experimental conditions allowed the participant to choose their role between controlling the payload (as they would in the first series when asserting control) or the defender agent to prevent obstacles from reaching the AI-controlled payload. Table 2 lists the ratio of time each participant spent in the defender role when given the choice. Most participants clearly preferred to be in this role, with 22 of 25 having ratios above 0.5 when paired with the evasive AI, and 18 when paired with the hasty AI. Though the roles are different in nature, any alteration in the AI behavior of both roles that was influenced by the participant’s actions would be reflected in the frame-by-frame actions.

To evaluate whether or not it is possible to generalize the change in human preferences to different AI agents and roles, data from the two series of experiments was used in an example transfer learning task. Table 1 shows baseline cases where data from within each series was used to predict preferences in the second series (using a training and test set) and cases with varying ratios of data from either series. Figure 1 illustrates the fitness for each of these ratios. A ratio of 0:1 represents training only using second series data and 1:0 represents transfer learning with just the first series data. A 10-fold cross-validation was used to evaluate the accuracy of the predictions for the second series of experiments. With no data transferred (0:1), the training yields a model with an average fold-wise error of 0.095. Transferring a model trained fully on source data (1:0) or an evenly combined data set (1:1) performed worse at the target task, yielding average fold-wise errors of 0.27 and 0.12, respectively. However, adding some source data to the target data for training (1:3) allowed models to predict player preference in the target task with an average fold-wise error of 0.091, lower (with 0.001 deviation) than the model that was trained purely on target task data. This shows that while a simple approach to transfer learning may not be feasible, it is possible to form a more general model between these different tasks by adjusting the ratio of trained data. What this ratio is, and the de-

gree to which it may depend on the target task configuration or not, remains to be seen.

Participant	Defender control ratio	
	(Evasive AI)	(Hasty AI)
A	0.98	0.88
B	0.97	0.97
C	0.98	0.87
D	0.77	0.31
E	0.96	0.95
F	0.33	0.13
G	0.87	0.55
H	0.09	0.05
I	0.73	0.31
J	0.71	0.70
K	0.93	0.68
L	0.62	0.36
M	0.84	0.79
N	0.80	0.71
O	0.98	0.86
P	0.92	0.95
Q	0.90	0.63
R	0.74	0.51
S	0.82	0.33
T	0.88	0.72
U	0.91	0.66
V	0.83	0.70
W	0.22	0.12
X	0.99	0.74
Y	0.94	0.79

Table 2. The amount of time (given as a ratio) each participant spent in the defender role is listed and separated by the AI agent that was assigned. These ratios are of the combined total time spent with each AI throughout the entire study period.

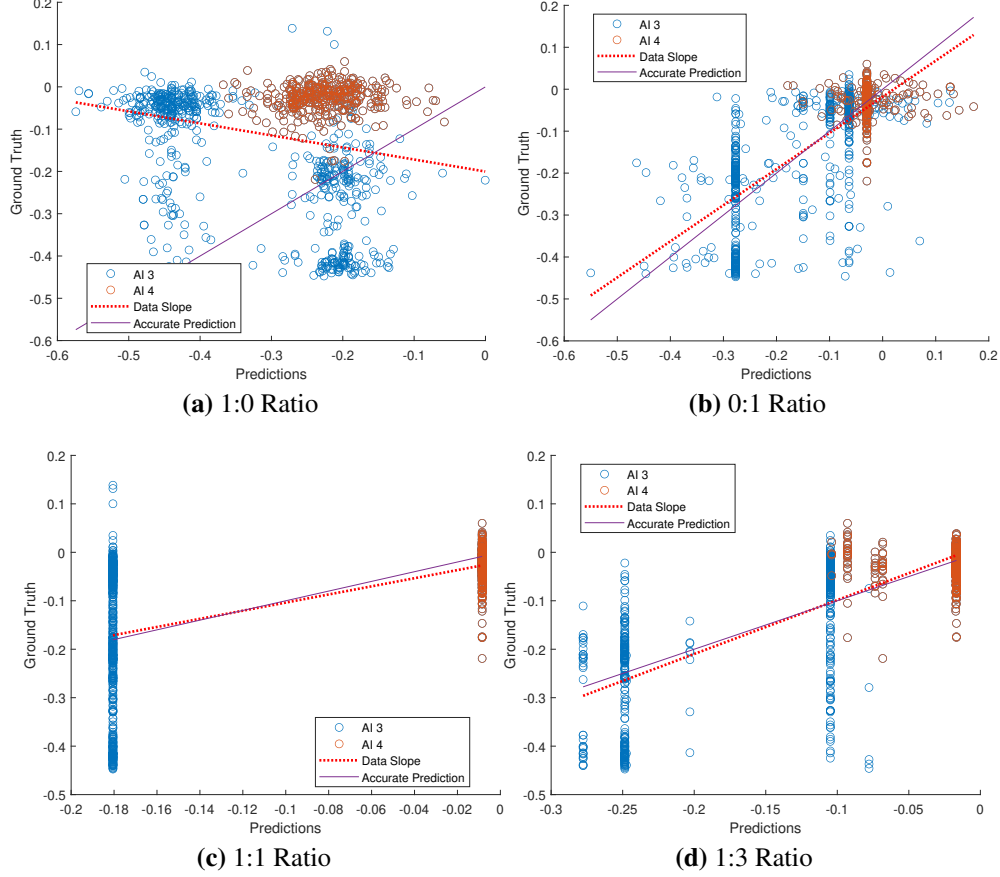


Figure 1. Performance of each trained model on predicting preference in the target task. Each model is trained with different ratios of source:target data: the blue points represent samples from AI 3, the orange points represent samples from AI 4, and the red dotted line represents a simple line fitted to the data.

5. Conclusions

Generalization of human–AI collaboration across different roles using transfer learning is possible with the right combination of training data. In this preliminary study, certain ratios of data among the roles to be used in a boosted ensemble trees learning method resulted in lower prediction error than with using the target role data alone, and interestingly, less than combining the data sets in equal measure. However, it is not clear what an optimal ratio would be, if it is specific to the role or not, or if it is a cause of the specific learning method used here. Exploring each of these possibilities to learn which are the cause of this result will help improve generalization through ML, which would allow more informed decisions on which kinds of AI to pair with particular humans.

6. References

1. Adamson T, Oishi M, Chiang HTL, Tapia L. Busy Beeway: a game for testing human-automation collaboration for navigation. Proceedings of the 10th International Conference on Motion in Games (MIG '17); 2017 Nov 8–10; Barcelona, Spain. 2017. p. 1–6. <https://doi.org/10.1145/3136457.3136471>
2. Wang R et al. StudentLife: using smartphones to assess mental health and academic performance of college students. In: Rehg JM, Murphy SA, Kumar S, editors. Mobile health. Springer; c2017. p. 7–33.
3. Hetherington NJ, Croft EA, Van der Loos HM. Hey robot, which way are you going? nonverbal motion legibility cues for human-robot spatial interaction. IEEE Rob Auto Letters. 2021;6(3):5010–5015. <https://doi.org/10.1109/LRA.2021.3068708>
4. Dragan AD, Lee KC, Srinivasa SS. Legibility and predictability of robot motion. Proceedings of the 8th ACM/IEEE International Conference on Human-Robot Interaction (HRI); 2013 March 3–6; Tokyo, Japan. 2013. p. 301–308. <https://doi.org/10.1109/HRI.2013.6483603>
5. Xu A, Dudek G. Maintaining efficient collaboration with trust-seeking robots. Proceedings of the IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS); 2016 Oct 9–14; Daejeon, Korea. 2016. p. 3312–3319. <https://doi.org/10.1109/IROS.2016.7759510>
6. Yin W et al. Effects of trust in human-automation shared control: a human-in-the-loop driving simulation study. Proceedings of the IEEE International Intelligent Transportation Systems Conference (ITSC); 2021 Sep 19–22; Indianapolis, IN. 2021. p. 1147–1154. <https://doi.org/10.1109/ITSC48978.2021.9565080>
7. Ma Y, Zhang Y, Wang C, Xiao Y. Influencing factors and variation of subjective workload in the process of takeover during autonomous driving. Proceedings of the 6th International Conference on Transportation Information and Safety; 2021 Oct 22–24; Wuhan, China. 2021. p. 1496–1502.
8. Cutlip S, Wan Y, Sarter N, Gillespie RB. The effects of haptic feedback and transition type on transfer of control between drivers and vehicle automation. IEEE Trans Hum Mach Sys. 2021;51(6):613–621.

9. Li R et al. Indirect shared control for cooperative driving between driver and automation in steer-by-wire vehicles. *IEEE Trans Intl Transp Sys.* 2021;22(12):7826–7836.
10. Huang C, Lv C, Hang P, Hu Z, Xing Y. Human–machine adaptive shared control for safe driving under automation degradation. *IEEE Intelligent Transportation Systems Magazine.* March–April 2022;14(2):53–66.
11. Sheikh HU, Bölöni L. Multi-agent reinforcement learning for problems with combined individual and team reward. *Proceedings of the International Joint Conference on Neural Networks (IJCNN)*; 2020 July 19–24; Glasgow, UK. 2020. p. 1–8. <https://doi.org/10.1109/IJCNN48605.2020.9206879>
12. Lee ES, Zhou L, Ribeiro A, Kumar V. Learning decentralized strategies for a perimeter defense game with graph neural networks. *arXiv*; 2022 Sep 23. [arXiv:2211.01757](https://doi.org/10.48550/arXiv.2211.01757). <https://doi.org/10.48550/arXiv.2211.01757>
13. Lee ES, Zhou L, Ribeiro A, Kumar V. Graph neural networks for decentralized multi-agent perimeter defense. *Front Control Eng.* 2023;4. <https://doi.org/10.3389/fcteg.2023.1104745>
14. Cooper S et al. Predicting protein structures with a multiplayer online game. *Nature.* 2010;466:756–760. <https://doi.org/10.1038/nature09304>
15. von Ahn L, Dabbish L. Labeling images with a computer game. *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems (CHI '04)*; 2004 Apr 24–29; Vienna, Austria. 2004. p. 319–326. <https://doi.org/10.1145/985692.98573>
16. Becker A, Ertel C, McLurkin J. Crowdsourcing swarm manipulation experiments: a massive online user study with large swarms of simple robots. *Proceedings of the IEEE International Conference on Robotics and Automation (ICRA)*; 2014 May 31–June 7; Hong Kong, China. 2014. p. 2825–2830. <https://doi.org/10.1109/ICRA.2014.6907264>
17. Dvijotham K, Todorov E. Inverse optimal control with linearly-solvable MDPs. *Proceedings of the 27th International Conference on Machine Learning*; 2010 June 21–24; Haifa, Israel. 2010. p. 335–342.

18. Uchibe E. Model-free deep inverse reinforcement learning by logistic regression. *Process Lett.* 2018;47:891–905.
19. Pan SJ, Yang Q. A survey on transfer learning. *IEEE Trans Knowl Data Eng.* 2010;22(10):1345–1359.
20. Zhuang F et al. A comprehensive survey on transfer learning. *Proceedings of the IEEE.* 2021;109(1):43–76.
21. Adamson T et al. Long-term analysis of a human-AI collaboration study using a mobile game. DEVCOM Army Research Laboratory (US); 2023 Feb. Report No.: ARL-TR-9628.
22. DiGioia L et al. Navigating with a defensive agent: role switching for human automation collaboration. *Proceedings of the 16th ACM SIGGRAPH Conference on Motion, Interaction and Games (MIG '23)*; 2023 Nov 15–17; Rennes, France. 2023. p. 1–7. <https://doi.org/10.1145/3623264.362445>
23. Lee JD, See KA. Trust in automation: designing for appropriate reliance. *Hum Factors.* 2004;46(1):50–80.
24. Mienye ID, Sun Y. A survey of ensemble learning: concepts, algorithms, applications, and prospects. *IEEE Access.* 2022;10:99129–99149.
25. The MathWorks Inc. Statistics and machine learning toolbox. 2024.

List of Symbols, Abbreviations, and Acronyms

AI	artificial intelligence
BAS	Behavioral Approach System
ML	machine learning

1 DEFENSE TECHNICAL
(PDF) INFORMATION CTR
DTIC OCA

1 DEVCOM ARL
(PDF) FCDD RLB CI
TECH LIB