

REPORT DOCUMENTATION PAGE			2	Form Approved OMB NO. 0704-0188	
The public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA, 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to any penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number. PLEASE DO NOT RETURN YOUR FORM TO THE ABOVE ADDRESS.					
1. REPORT DATE (DD-MM-YYYY) 09-07-2014		2. REPORT TYPE Conference Proceeding		3. DATES COVERED (From - To) -	
4. TITLE AND SUBTITLE Accelerating Imitation Learning in Relational Domains via Transfer by Initialization			5a. CONTRACT NUMBER W911NF-13-1-0432		
			5b. GRANT NUMBER		
			5c. PROGRAM ELEMENT NUMBER 611102		
6. AUTHORS Sriram Natarajan, Phillip Odom, Saket Joshi, Tushar Khot, Kristian Kersting, Prasad Tadepalli			5d. PROJECT NUMBER		
			5e. TASK NUMBER		
			5f. WORK UNIT NUMBER		
7. PERFORMING ORGANIZATION NAMES AND ADDRESSES Indiana University at Bloomington 509 E 3RD ST Bloomington, IN 47401 -3654			8. PERFORMING ORGANIZATION REPORT NUMBER		
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS (ES) U.S. Army Research Office P.O. Box 12211 Research Triangle Park, NC 27709-2211			10. SPONSOR/MONITOR'S ACRONYM(S) ARO		
			11. SPONSOR/MONITOR'S REPORT NUMBER(S) 64577-NS-YIP.1		
12. DISTRIBUTION AVAILABILITY STATEMENT Approved for public release; distribution is unlimited.					
13. SUPPLEMENTARY NOTES The views, opinions and/or findings contained in this report are those of the author(s) and should not be construed as an official Department of the Army position, policy or decision, unless so designated by other documentation.					
14. ABSTRACT The problem of learning to mimic a human expert/teacher from training trajectories is called imitation learning. To make the process of teaching easier in this setting, we propose to employ transfer learning (where one learns on a source problem and transfers the knowledge to potentially more complex target problems). We consider multi-relational environments such as real-time strategy games and use functional-gradient boosting to capture and transfer the models learned in these environments. Our experiments demonstrate that our learner learns a very good initial model from the simple scenario and effectively transfers the knowledge to the more complex scenario thus					
15. SUBJECT TERMS Statistical Relational Learning, Transfer Learning, Imitation Learning					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT UU	15. NUMBER OF PAGES	19a. NAME OF RESPONSIBLE PERSON Sriram Natarajan
a. REPORT UU	b. ABSTRACT UU	c. THIS PAGE UU			19b. TELEPHONE NUMBER 812-856-7615

Report Title

Accelerating Imitation Learning in Relational Domains via Transfer by Initialization

ABSTRACT

The problem of learning to mimic a human expert/teacher from training trajectories is called imitation learning. To make the process of teaching easier in this setting, we propose to employ transfer learning (where one learns on a source problem and transfers the knowledge to potentially more complex target problems). We consider multi-relational environments such as real-time strategy games and use functional-gradient boosting to capture and transfer the models learned in these environments. Our experiments demonstrate that our learner learns a very good initial model from the simple scenario and effectively transfers the knowledge to the more complex scenario thus achieving a jump start, a steeper learning curve and a higher convergence in performance.

Conference Name: Inductive Logic Programming

Conference Date: August 28, 2013

Accelerating Imitation Learning in Relational Domains via Transfer by Initialization

Sriraam Natarajan¹, Phillip Odom¹, Saket Joshi², Tushar Khot³, Kristian Kersting⁴, and Prasad Tadepalli⁵

¹ Indiana University Bloomington, USA

² Cycorp Inc, USA

³ University of Wisconsin-Madison, USA

⁴ Fraunhofer IAIS, Germany

⁵ Oregon State University, USA

Abstract. The problem of learning to mimic a human expert/teacher from training trajectories is called imitation learning. To make the process of teaching easier in this setting, we propose to employ transfer learning (where one learns on a source problem and transfers the knowledge to potentially more complex target problems). We consider multi-relational environments such as real-time strategy games and use functional-gradient boosting to capture and transfer the models learned in these environments. Our experiments demonstrate that our learner learns a very good initial model from the simple scenario and effectively transfers the knowledge to the more complex scenario thus achieving a jump start, a steeper learning curve and a higher convergence in performance.

1 Introduction

It is common knowledge that both humans and animals learn new skills by observing others. This problem, which is called *imitation learning*, can be formulated as learning a representation of a policy – a mapping from states to actions – from examples of that policy. Imitation learning has a long history in machine learning and has been studied under a variety of names including learning by observation [1], learning from demonstrations [2], programming by demonstrations [3], programming by example [4], apprenticeship learning [5], behavioral cloning [6], and some others. Techniques used from supervised learning have been successful for imitation learning [7]. We follow this tradition and investigate the use of supervised learning methods to learn behavioral policies.

Our focus is on relational domains where states are naturally described by relations among an indefinite number of objects. Examples include real time strategy games such as Warcraft, regulation of traffic lights, logistics, and a variety of planning domains. A supervised learning method for imitation learning was recently proposed [8]. This approach assumes an efficient hypothesis space



Fig. 1. Wargus Scenarios (*left*) The two tower scenario where providing examples is easier. *right* The three tower scenario which is significantly more complicated and requires more training trajectories.

for the policy function, and learns only policies in this space that are closest to the training trajectories [9, 10]. This approach is based on functional gradient boosting [11] where a set of relational regression trees [12] are used to compactly represent a complex relational policy. This approach was demonstrated to be successful in many problems.

One of the key assumptions in the proposed approach is that the policies can be generalized across the objects in the domain. While one of the advantages of a logical representation is the generalization capability, it is also quite possible that in several large problems, the optimal policies can vary greatly as the number of the objects in the domains can increase. In such cases, the learner has to be provided with new example trajectories to learn the policies. Since the complexity of the domain has increased, the number of trajectories required for learning can also increase significantly. For instance, consider the two scenarios presented in Figure 1 where the goal is to defend the towers from being destroyed by the enemy units. In the left figure, there are two towers and two enemy and friendly footmen and archers. In the right figure, all the numbers increase by one. As we show empirically, the optimal policies for the two scenarios can be very different. More importantly, the number of trajectories required to converge to the optimal policy is higher in the case of the more complex scenario.

In order to train on such complex scenarios, we propose to employ transfer learning [13, 14] for learning in a (simpler) source problem and then transferring the learned knowledge to a (more complex) target task. More precisely, we aim to employ transfer by initialization [15] where the models learned from the source task are used to initialize the models in the learning task. Following prior work [8], we perform search through the space of policies using functional gradient boosting but initialize the gradients with the models learned in the source task. Our hypothesis is that this initialization will allow the learner to explore more complex policy spaces that might not have been accessed easily if the search started out with uniform policies. We verify this claim empirically.

In summary, we consider the problem of imitation learning in relational domains where the optimal policies can be significantly different as the number of objects in the domain increases. Generalization of policies is still a desired property as the properties of the objects themselves can change across situations with the same number of objects. When the number of objects change, we propose to employ transfer learning by initialization to initialize the gradients in the target task. We evaluate the hypothesis in a real time strategy game and show that we are able to achieve a jump start, faster convergence to a more optimal policy.

The rest of the paper is organized as follows: we introduce the background and the prior work on relational imitation learning next. We then present our transfer algorithm for initialization and evaluate the algorithm on a complex RTS game and conclude the paper by outlining some challenges for future work.

2 Background

An MDP is described by a set of discrete states $\mathbf{S}$, a set of actions $\mathbf{A}$, a reward function $r_s(a)$ that describes the expected immediate reward of action a in state s , and a state transition function $p_{ss'}^a$ that describes the transition probability from state s to state s' under action a . A policy, π , is defined as a mapping from states to actions, and specifies what action to execute in each state. In the imitation learning, we assume that the reward function is not directly obtained from the environment. Our input consists of S , A and supervised trajectories generated by a Markov policy. We try to match it using a parameterized policy.

3 Relational Imitation Learning

Following Ratliff et al. [16], we assume that the discount factor are absorbed into the transition probabilities and policies are described by $\mu \in \mathbf{G}$ where $\mathbf{G}$ is the space of all state-action frequency counts. We assume a set of features $\mathbf{F}$ that describe the state space of the MDP and the expert chooses the action a_i at any time step i based on the set of feature values $\langle f_i \rangle$ according to some function. For simplicity, we denote the set of features at any particular time step i of the j th trajectory as $\mathbf{f}_i^j$ and we drop j whenever it is fairly clear from the context.

The goal of our algorithm is to learn a policy that suitably mimics the expert. More formally, we assume a set of training instances $\{\langle \mathbf{f}_i^j, a_i \rangle_{i=1}^{m^j}\}_{j=1}^n$ that is provided by the expert. Given these training instances, the goal is to learn a policy μ that is a mapping from $\mathbf{f}_i^j$ to a_i^j for each set of features $\mathbf{f}_i^j$. The key aspect of our setting is that the individual features are relational i.e., objects and relationships over these objects. The features are denoted in standard logic notation where $p(X)$ denotes the predicate p whose argument is X . The problem of imitation learning given these relational features and expert trajectories can now be posed as a regression problem or a supervised learning problem over these trajectories.

In our previous work [8], we employed Functional-Gradient Boosting for learning relational policies. The goal is to find a policy μ that is captured using

the trajectories (i.e., features $\mathbf{f}_i^j$ and actions a_i^j) provided by the expert, i.e., the goal is to determine a policy $\mu = P(a_i | \mathbf{f}_i; \psi) \forall a, i$ where the features are relational. These features could define the objects in the domain (squares in a gridworld, players in robocup, blocks in blocksworld, archers or footmen in a real-time strategy game etc.), their relationships (type of objects, teammates in robocup etc.), or temporal relationships (between current state and previous state) or some information about the world (traffic density at a signal, distance to the goal etc.).

We assume a functional parametrization over the policy and consider the conditional distribution over actions a_i given the features to be,

$$P(a_i | \mathbf{f}_i; \psi) = e^{\psi(a_i; \mathbf{f}_i)} / \sum_{a'_i} e^{\psi(a'_i; \mathbf{f}_i)}, \forall a_i \in \mathbf{A} \quad (1)$$

where $\psi(a_i; \mathbf{f}_i)$ is the potential function of a_i given the grounding $\mathbf{f}_i$ of the feature predicates at state s_i and the normalization is over all the admissible actions in the current state. Formally, functional gradient ascent starts with an initial potential ψ_0 and iteratively adds gradients Δ_i . Here, Δ_m is the functional gradient at episode m and is

$$\Delta_m = \eta_m \times E_{x,y} [\partial / \partial \psi_{m-1} \log P(y|x; \psi_{m-1})] \quad (2)$$

where η_m is the learning rate. Note that in Equation 2, the expectation $E_{x,y}[\dots]$ cannot be computed as the joint distribution $P(\mathbf{x}, \mathbf{y})$ is unknown (in our case, y 's are the actions while x 's are the features). Instead of computing the gradients over the potential function, the gradients are computed for each training example:

$$\Delta_m(a_i^j; \mathbf{f}_i^j) = \nabla_{\psi} \sum_j \sum_i \log(P(a_i^j | \mathbf{f}_i^j; \psi))|_{\psi_{m-1}} \quad (3)$$

These are point-wise gradients for examples $\langle \mathbf{f}_i^j, a_i^j \rangle$ on each state i in each trajectory j conditioned on the potential from the previous iteration (shown as $|_{\psi_{m-1}}$). Now this set of local gradients form a set of training examples for the gradient at stage m . The main idea in the gradient-tree boosting is to fit a regression-tree on the training examples at each gradient step [17]. The idea of functional gradient boosting is presented in Figure 2.

The functional-gradient w.r.t $\psi(a_i^j; \mathbf{f}_i^j)$ of the likelihood for each example $\langle \mathbf{f}_i^j, a_i^j \rangle$ is given by:

$$\frac{\partial \log P(a_i^j | \mathbf{f}_i^j; \psi)}{\partial \psi(\hat{a}_i^j; \mathbf{f}_i^j)} = I(a_i^j = \hat{a}_i^j | \mathbf{f}_i^j) - P(a_i^j | \mathbf{f}_i^j; \psi) \quad (4)$$

where $\hat{a}_i^j$ is the action observed from the trajectory and I is the indicator function that is 1 if $a_i^j = \hat{a}_i^j$ and 0 otherwise. The key feature of the above expression is that the functional-gradient at each state of the trajectory is dependent on the observed action $\hat{a}$. If the example is positive (i.e., it is an action executed by the expert), the gradient $(I - P)$ is positive indicating that the policy should

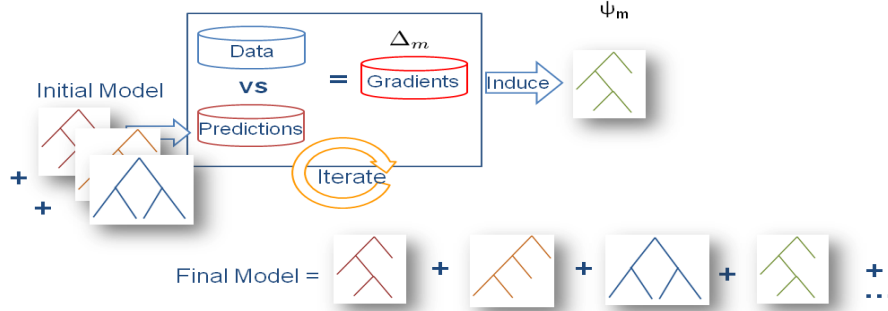


Fig. 2. Relational FGB. This is similar to the standard FGB where trees are induced in stage-wise manner; the key difference being that the trees are relational regression trees. To compute the predictions, a query, x is applied to each tree in turn, and the numerical values at the leaf reached in each tree are summed to obtain $\psi(x)$.

increase the probability of choosing the action. On the contrary if the example is a negative example (i.e., for all other actions), the gradient is negative implying that it will push the probability of choosing the action towards 0.

Following prior work [18–20], we used *Relational Regression Trees* (RRTs)[12] to fit the gradient function at every feature in the training example [8]. Hence the distribution over each action is represented as a set of RRTs on the features. These trees are learned such that at each iteration the new set of RRTs aim to maximize the likelihood. Hence, when computing $P(a(X)|f(x))$ for a particular value of state variable X (say x), each branch in each tree is considered to determine the branches that are satisfied for that particular grounding (x) and their corresponding regression values are added to the potential. For example, X could be a particular unit in Wargus, or a certain block in the blockworld.

4 Relational Transfer

In this work, we extend the previous work in imitation learning by employing the ideas for *inductive transfer* [13]. While the previous approach was able to achieve generalization in a imitation learning setting, the generalized policies might not be sufficient in some other variations of the problems. For instance, in the scenarios considered in Figure 1, the optimal policies for the three tower defense scenario can be significantly different from the easier task of two tower scenario. Also, since the three tower case is a harder task, as we show empirically, learning in this setting might require more example trajectories from the expert. In such cases, it is easier to transfer the knowledge gained from the two tower scenario to the three tower case for initialization and then improve upon the knowledge by learning in the three tower case. This will enable the learner to: (a) learn a better policy than the one generalized from the two tower scenario

Table 1. Transfer Learning Algorithm

```

1: function TRANSFER( $T_{source}, T_{target}$ )
2:    $\Lambda_s = \text{TIL}(\{\}, T_{source})$  ▷ Learn with source Trajectories
3:    $\Lambda_t = \text{TIL}(\Lambda_s, T_{target})$  ▷ Use learned models and learn on Target Trajectories
4:   return  $\Lambda_t$ 
5: end function
6: function TIL( $\Lambda, \text{Trajectories } T$ )
7:    $\Lambda_0 = \Lambda$ 
8:   for  $1 \leq k \leq |\Lambda|$  do ▷ Iterate through each action
9:     for  $1 \leq m \leq M$  do ▷ M gradient steps
10:       $S_k := \text{GenExamples}(k; T; \Lambda_{m-1}^k)$ 
11:       $\Delta_m(k) := \text{FitRRT}(S_k; L)$  ▷ Gradient
12:       $\Lambda_m^k := \Lambda_{m-1}^k + \Delta_m(k)$  ▷ Update models
13:    end for
14:     $P(A = k | \mathbf{f}) \propto \psi^k$ 
15:  end for
16:  return  $\Lambda$ 
17: end function
18: function GENEXAMPLES( $k, T, \Lambda$ )
19:    $S := \emptyset$ 
20:   for  $1 \leq j \leq |T|$  do ▷ Trajectories
21:     for  $1 \leq i \leq |S^j|$  do ▷ States of trajectory
22:       Compute  $P(\hat{a}_i^j = k | \mathbf{f}_i^j)$  ▷ Probability of user action being the current
23:       action
24:        $\Delta_m(k; \mathbf{f}_i^j) = I(\hat{a}_i^j = k) - P(\hat{a}_i^j = k | \mathbf{f}_i^j)$ 
25:        $S := S \cup [(\hat{a}_i^j, \mathbf{f}_i^j), \Delta(\hat{a}_i^j; \mathbf{f}_i^j)]$  ▷ Update relational regression examples
26:     end for
27:   end for
28:   return  $S$  ▷ Return regression examples
29: end function

```

and (b) converge to the optimal policy faster in the three tower scenario i.e., from fewer trajectories when compared to learning with no knowledge.

The form of the functional gradients facilitate easy transfer. Since they perform gradient descent in function space, we can initialize the models (ψ_0) for the three tower scenario with some of the trees learned in the two tower scenario. Conceptually, this is essentially the same as using the result of the first few gradient steps in the source problem while learning in the target problem. After initializing the gradient ascent with the initial set of trees, we propose to learn new set of trees in the target task that build upon the initial model. As mentioned earlier, this initial set of trees for the ψ_0 .

To perform learning in the target task, the trajectories must be weighted given the initial model. Similar to Equation 4, we compute the value of $I(a_i^j = \hat{a}_i^j | \mathbf{f}_i^j) - P(a_i^j | \mathbf{f}_i^j; \psi_0)$ for each action of each trajectory, i.e., the weight of each observed action is the difference between the indicator function of that action and the marginal probability of that action given the initial potential function.

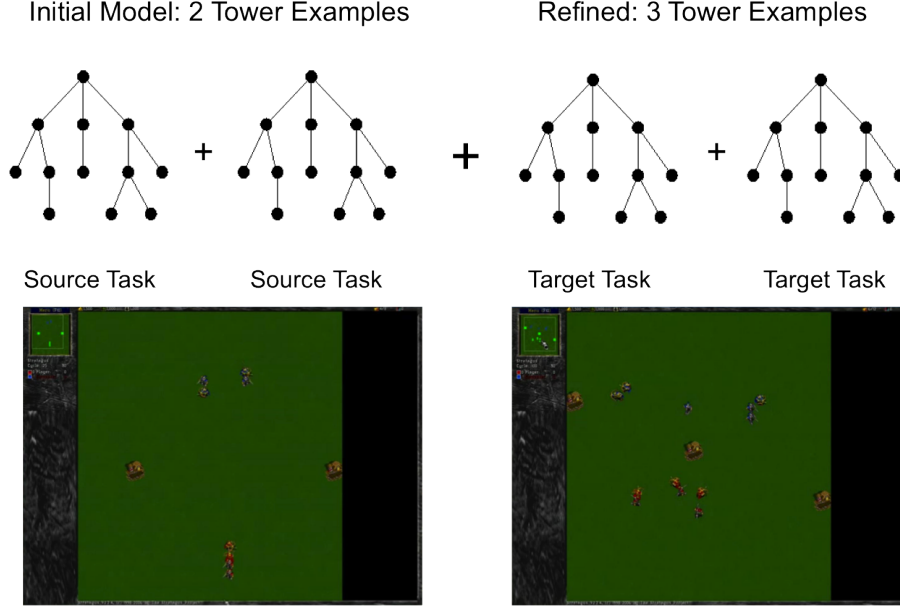


Fig. 3. Proposed transfer approach. The key idea is to learn a small set of trees from the source task and use them to initialize the RFGB algorithm for the harder task.

Once these weights are computed for the given trajectory, they serve as the examples for learning new set of trees. This idea is presented in Figure 3. First we learn a few set of trees in the source task and then use them to initialize the models in the target task. We then learn a new set of trees in the target task.

Our proposed transfer learning algorithm is presented in Algorithm 1. The function *Transfer* is the main algorithm that takes as input trajectories from the source and the target tasks. The algorithm first calls the *TIL* function (which stands for *Tree – based Imitation Learning*) with an empty potential function (empty set of trees) and the source trajectories. The *TIL* function then learns a set of trees in the source task. The *TIL* function is the same one presented in [8] with the modification that it can use an initial set of trees. For each action (k), it generates the examples for our regression tree learner (called using function *FitRRT*) to get the new regression tree and updates its model (Λ_m^k). This is repeated up to a pre-set number of iterations M (typically, $M = 20$). We found empirically that increasing M has no effect on the performance as the example weights nearly become 0 and the regression values in the leaves are close to 0 as well. Note that the after m steps, the current model Λ_m^k will have m regression trees each of which approximates the corresponding gradient for the action k . These regression trees serve as the individual components ($\Delta_m(k)$) of the final potential function.

Once the set of trees have been learned in the source task, a subset of those trees (typically we use 20 trees in our experiments), is then used as the initial model for the target task and the *TIL* function is called with this initial set and the trajectories. The function then returns a new set of trees which are then used for evaluating in the target task. It must be mentioned that when choosing to act in the target task, inference over the actions is performed using *all* the trees (the initial set of source trees plus the target trees). It is easy to see that we cannot ignore the transferred trees since they form the first step of the gradient ascent when learning the policy in the target domain.

The proposed approach is closely related to the idea of modular policies of Driessens [21]. He observed that the use of functional gradients to represent policies allows us to separate the gradient updates to different subtasks of the agent’s task. Thus, we can create separate potential functions for each part of the task and the natural addition operator of functional gradients allows then to obtain the final policy which is essentially a sum of different regression trees. We extend the above to transfer learning where we consider an initial set of trees for a different task (or potentially a subtask) and a new set of trees are then learned for the new task. Hence, combining these two ideas, it is possible to learn a higher level policy in a hierarchy by transferring from the lower level subtasks.

5 Experiments

We present the empirical evaluation of our proposed algorithm on a real-time strategy game. We are particularly interested in the following questions:

Q1: How do the transferred models compare against the models that are generalized using the relational imitation learning algorithm?

Q2: How do the transferred models compare against the models that are learned directly on the target task with no prior models from source task?

Experimental Setup: Stratagus is an open-source real-time strategy (RTS) game engine written in C based off the Warcraft series of games. Like all RTS games, it allows multiple agents to be controlled simultaneously in a fully observed setting, making an ideal test bed for imitation learning. A java client was written, revised at Oregon State University⁶, to connect to the Stratagus game engine via a socket connection. The client collects all of the game information from the game engine and can issue detailed commands to all units of a player in the game. This client allows for the learned policies to be executed directly in the game environment as opposed to simulation creating more realistic performance metrics.

The setting in which transfer is being tested is the tower defense scenario shown in Figure 1. The map used for the experiments consisted of 6 x 6 grid world. Our scenarios consist of two opposing teams-one attacking, one defending-each with two kinds of units. Footman have more health but must be close to an enemy to attack them while ranged archers are easily killed but can attack from a distance. Towers exist on the map in set locations. The defending team must

⁶ <http://beaversource.oregonstate.edu/projects/stratagusai>

prevent the attacking team from destroying the towers while the attacking team must destroy as many towers as they can. The defending team must divide its units among the various towers to prevent one tower from falling while another is being saved. This dynamic creates complex policies.

Predicates	Description
friendlyobject	The type of defending unit
enemyobject	The type of attacking unit
dead	Enemy unit that is dead
locationId	Location of friendly unit
strength	Strength (hit points) of friendly unit
distance	Distance of a friendly unit to a tower
enemyattower	Tower that enemy unit is attacking
attacking	Enemy unit that a friendly unit was attacking in the previous state

Fig. 4. Features that describe the state in the two scenarios. We omit the arguments of the predicates for brevity.

We used the following features to describe the state: the strength (high, medium, low) and location of all friendly units, the type (footman, archer) of all units in the game, which tower each enemy unit is currently attacking, and the enemy unit that all friendly units were attacking the previous state. Friendly units are unaware of the strength of enemy units or their exact location. The full set of information given at every state is included in Table 4. The actions available to the friendly units are to move to a location and attack

a particular unit. The nature of the objects and the relationships between the objects in this game naturally allow for a relational representation. Each type of unit (footman or archers) shares traits such as their attacking range and their total health so certain policy rules will naturally apply to all units of that type.

As mentioned earlier, the goal of this experiment is to learn to protect two towers in a source scenario and transfer the learned knowledge to a target scenario. The attacking team’s strategy is as follows: At the start of each game, each member of the attacking team randomly selects a tower to attack. However, if approached by an enemy archer or footmen, they will change their target to eliminate the opposing player’s offensive units. After destroying one tower, they will randomly select another tower to attack until there are no towers left and the game ends. The goal of the game is to defend the towers. The game ends if either the enemy team manages to destroy all of the towers or the friendly team kills all of the enemy units. The number of friendly and attacking units vary between the source and target scenarios. In the source scenario, there are two footmen and two archers while in the target scenario, there are three footmen and three archers. In both the scenarios, the towers cannot defend themselves.

Results: We used two performance metrics that are based on the number of towers saved. The 3-tower win percentage is the percentage of games in which all three towers were saved. This is a difficult task because it requires the friendly team to defend all three towers simultaneously; else one tower may fall while they are all defending the others. The 1-tower win percentage is the easier metric,

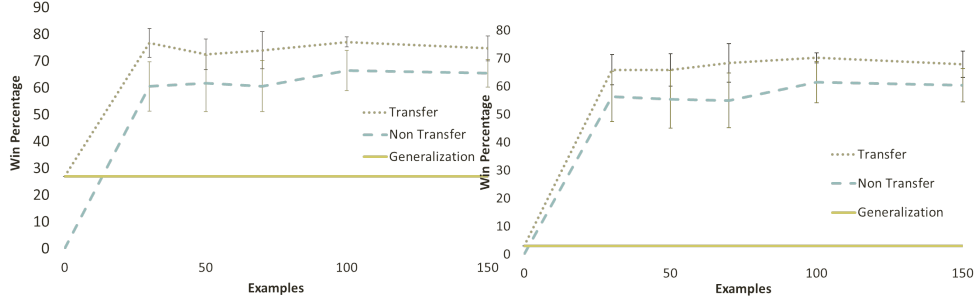


Fig. 5. Results of saving at least one tower (left) and saving all three towers(right). The transferred models dominate the one learned from two towers (generalization) and the one learned on three towers without an initial policy (non-transfer)

only requiring 1 of the 3 towers to be saved. We used randomly selected samples of 30, 50, 70, 100, and 150 3-tower games from a pool of approximately 1000 expert games for training in the target scenario. For the source scenario, we used 100 games for training. This experiment was repeated 10 times and the average percentage of winnings games were computed.

We used three models for evaluation – the transfer model, the non-transferred model which learns only on the target task and the generalization model which learns only on the source and not the target task. We learned 20 trees on the source task and learned a further 20 on the target task. Hence, the final transfer model had 40 trees while the other two models had 20 trees each.

The results are presented in Figures 5 corresponding to saving at least one tower and three towers respectively. As can be seen from the figures, the transfer model is superior to both the generalization and the pure imitation learner in the twin tasks. Also, the results follow the original transfer learning goals of jump start, steeper learning curve and better convergence. It is fairly clear that in defending at least one tower, the use of the initial models from the source scenario provides a bigger jump start than in defending all towers. On the other hand, when saving all towers, the jump start is not fairly high but the learning curve is steeper. It appears that the use of initial models allow for the learner in the target scenario to explore a space of policies that might not have been otherwise reached from an uniform policy. Similarly, the difference between transfer and non-transfer models is higher in the one tower case than saving all the towers though the difference is statistically significant in both the cases. Also, it must be mentioned that even when the non-transferred models were provided with a lot more trajectories, it is not able to match the performance of the transferred model. This suggests that using initial policies in some cases is more useful than obtaining more expert trajectories. It would be interesting to evaluate these results in other domains as well.

In summary, our experiments answer both the questions affirmatively in that the transfer models dominate both the original imitation learning models.

6 Discussion and Conclusion

We address the issue of sample complexity in imitation learning settings. In scenarios where the expert’s time is expensive/valuable and we have access to only a few training examples from the expert our approach is to divide the expert’s time between simple (smaller domain size) and harder (larger domain size) problems. Although policies induced from the simpler problem training instances can be employed to solve the larger domain via relational generalization, in scenarios we provide (such as Wargus) this does not translate to better performance. We have presented transfer learning by using the simpler policy as our initial models and building an updatable relational model by learning from the harder examples. We observe not only a superior performance to generalization but also a drastic reduction in the sample complexity as compared with the naive method of directly inducing a model on the complex examples.

Imitation learning encounters two major problems when dealing with large state spaces. First, assuming a tabular representation of the policy to be learned is likely to exceed memory due to the large state and action space. Second, it can only make use of a limited amount of expert traces compared to the excessive amount of possible traces. The implicit feedback gained by the expert’s traces on the best action to take in a state might be so sparse that a well-generalizing policy will only be discovered slowly.

The first problem can be solved using relational imitation learning (RIL) for structural domains. However, the problem of sparse feedback has not been addressed by RIL yet. For relational reinforcement learning, there is a compelling and simple solution to this problem: inject traces of execution of a reasonable policy for the task at hand [22]. Unfortunately, this does not work for imitation learning. The input consists already of traces of execution of a reasonable policy, namely the policy of the expert. Thus, we do not gain anything despite enlarging the training set as can be seen from our results. The non-transfer model seemed to have converged to a inferior policy. To overcome this problem, we intuitively propose to inject traces of a policy of a reasonably well related task. Specifically, we directly inject the complete "related" policy into a functional gradient boosting approach to RIL. This appears to be an interesting result in that sometimes prior policies have a better impact on the performance compared to more trajectories. Our immediate challenge is to validate this hypothesis on other more complex domains. Another interesting direction is the possibility of employing active learning methods for extracting the best complex examples given the initial model thereby further improving on the performance of transfer.

Acknowledgements

SN and PO thank Army Research Office grant number W911NF-13-1-0432 under the Young Investigator Program. SN and TK gratefully acknowledge the support of the DARPA DEFT Program under the Air Force Research Laboratory (AFRL) prime contract no. FA8750-13-2-0039. Any opinions, findings, and conclusion or recommendations expressed in this material are those of the authors and do not necessarily reflect the view of the DARPA, AFRL, or the US government. SJ was supported by a Computing Innovations Postdoctoral Fellowship. Kristian Kersting was supported by the

Fraunhofer ATTRACT fellowship STREAM and by the European Commission under contract number FP7-248258-First-MM. Prasad Tadepalli acknowledges the support of ONR grant N000141110106.

References

1. A. Segre and G. DeJong. Explanation-based manipulator learning: Acquisition of planning ability through observation. In *Conf on Robotics and Automation*, 1985.
2. B. Argall, S. Chernova, M. Veloso, and B. Browning. A survey of robot learning from demonstration. *Robotics and Autonomous Systems*, 57:469–483, 2009.
3. S. Calinon. *Robot Programming By Demonstration: A probabilistic approach*. EPFL Press, 2009.
4. H. Lieberman. Programming by example (introduction). *Communications of the ACM*, 43:72–74, 2000.
5. A. Ng and S. Russell. Algorithms for inverse reinforcement learning. In *ICML*, 2000.
6. C. Sammut, S. Hurst, D. Kedzier, and D. Michie. Learning to fly. In *ICML*, 1992.
7. N. Ratliff, A. Bagnell, and M. Zinkevich. Maximum margin planning. In *ICML*, 2006.
8. S. Natarajan, S. Joshi, P. Tadepalli, K. Kersting, and J. Shavlik. Imitation learning in relational domains: A functional-gradient boosting approach. In *IJCAI*, 2011.
9. R. Khordon. Learning action strategies for planning domains. *Artificial Intelligence*, 113:125–148, 1999.
10. S. Yoon, A. Fern, and R. Givan. Inductive policy selection for first-order mdps. In *UAI*, 2002.
11. J.H. Friedman. Greedy function approximation: A gradient boosting machine. *Annals of Statistics*, 29, 2001.
12. H. Blockeel. Top-down induction of first order logical decision trees. *AI Commun.*, 12(1-2), 1999.
13. S. Pan and Q. Yang. A survey on transfer learning. *IEEE Transactions on Knowledge and Data Engineering*, 22:1345–1359, 2010.
14. Stephan Al-Zubi and Gerald Sommer. Imitation learning and transferring of human movement and hand grasping to adapt to environment changes. In *Human Motion*, volume 36 of *Computational Imaging and Vision*, pages 435–452. 2008.
15. N. Mehta, S. Natarajan, P. Tadepalli, and A. Fern. Transfer in variable-reward hierarchical reinforcement learning. *Machine Learning*, 73(3):289–312, 2008.
16. N. Ratliff, D. Silver, and A. Bagnell. Learning to search: Functional gradient techniques for imitation learning. *Autonomous Robots*, pages 25–53, 2009.
17. T.G. Dietterich, A. Ashenfelter, and Y. Bulatov. Training conditional random fields via gradient tree boosting. In *ICML*, 2004.
18. B. Gutmann and K. Kersting. TildeCRF: Conditional random fields for logical sequences. In *ECML*, 2006.
19. S. Natarajan, T. Khot, K. Kersting, B. Guttmann, and J. Shavlik. Gradient-based boosting for statistical relational learning: The relational dependency network case. *Machine Learning*, 2012.
20. K. Kersting and K. Driessens. Non-parametric policy gradients: A unified treatment of propositional and relational domains. In *ICML*, 2008.
21. K. Driessens. Non-disjoint modularity in reinforcement learning through boosted policies. In *Multi-disciplinary symposium on Reinforcement Learning*, 2009.
22. K. Driessens and S. Dzeroski. Integrating guidance into relational reinforcement learning. *Machine Learning*, 57(3):271–304, 2004.