

thesis.pdf

by Warwick Masson

FILE	WARWICK_MASSON_388188_THESIS.PDF (2.41M)		
TIME SUBMITTED	18-MAR-2016 12:42PM	WORD COUNT	11311
SUBMISSION ID	647041006	CHARACTER COUNT	54865



University of the Witwatersrand
Computer Science Masters
Reinforcement Learning with Parameterized Actions

Warwick Masson
388188
Supervised by Pravesh Ranchod and George Konidakis

March 18, 2016

1

Declaration

I declare that this thesis is my own, unaided work. It is being submitted for the Degree of Master of Science at the University of the Witwatersrand, Johannesburg. It has not been submitted before for any degree or examination at any other University.

(Signature of candidate)

____ day of ____ 20 ____ in ____

Abstract

In order to complete real-world tasks, autonomous robots require a mix of fine-grained control and high-level skills. A robot requires a wide range of skills to handle a variety of different situations, but must also be able to adapt its skills to handle a specific situation. Reinforcement learning is a machine learning paradigm for learning to solve tasks by interacting with an environment. Current methods in reinforcement learning focus on agents with either a fixed number of discrete actions, or a continuous set of actions.

We consider the problem of reinforcement learning with parameterized actions—discrete actions with continuous parameters. At each step the agent must select both which action to use and which parameters to use with that action. By representing actions in this way, we have the high level skills given by discrete actions and adaptability given by the parameters for each action.

We introduce the Q-PAMDP algorithm for model-free learning in parameterized action Markov decision processes. Q-PAMDP alternates learning which discrete actions to use in each state and then which parameters to use in those states. We show that under mild assumptions, Q-PAMDP converges to a local optimum. We compare Q-PAMDP with a direct policy search approach in the goal-scoring and Platform domains. Q-PAMDP out-performs direct policy search in both domains.

Acknowledgements

First, thanks to my parents for their support throughout this work. Although it was trying at times, they were always there for me when I needed them.

This thesis would not have been possible without all the work done by my supervisor George Konidaris. We worked closely together to turn what started as a broad idea into a focused piece of research. He encouraged me to submit this research to two conferences, which gave the work a vital stress-test.

Pravesh Ranchod, my other supervisor, deserves special thanks for helping me negotiate the processes towards finishing this Masters.

I am grateful to my fellow students for providing input and for being a sounding board for my ideas. It was extremely useful to be able to explain concepts to someone who was unfamiliar with my work. Thanks in particular to Dean Wookey for his help with editing, theory, and with the phrasing and formalization of the Q-PAMDP algorithm.

Contents

1	Introduction	6
2	Background	8
2.1	Reinforcement Learning	8
2.2	Value Functions	11
2.3	Policy Search	12
2.4	Summary	16
3	Reinforcement Learning with Parameterized Actions	17
3.1	Naive Approaches	19
3.2	Related Work	20
3.3	Summary	21
4	The Q-PAMDP Algorithm	22
4.1	Convergence Properties of Q-PAMDP(1)	23
4.2	Gradient of $H(\theta)$	25
4.3	Convergence properties of Q-PAMDP(∞)	26
4.4	Summary	26
5	The Robot Soccer Goal Domain	27
5.1	Domain	28
5.2	Implementation Details	29
5.3	Results	30
6	The Platform Domain	33
6.1	Specification	33
6.2	Approach	34
6.3	Results	9
6.4	Summary	36
7	Conclusions and Future Work	37
7.1	Future Work	37
7.2	Conclusion	37
	References	40
	Appendices	41
A	Theoretical Definitions	42
B	Disconnected Action Spaces	43

C	Discontinuous Policies	44
D	RSS Simulator	45

List of Figures

1.1	Parameterized actions for a soccer player robot.	6
2.1	Interaction between an agent and the environment.	8
2.2	A simple maze problem.	9
2.3	Discrete and continuous state representations	9
2.4	Policies for discrete and continuous state spaces.	10
2.5	Discrete and continuous action spaces.	10
2.6	Examples of 2-dimensional Fourier basis functions.	12
2.7	Gradient ascent on a surface.	14
3.1	A quad-rotor delivery system. There are two actions: $\text{move-to}(x, y, z)$, which moves the quadrotor to position (x, y, z) , and $\text{drop-payload}(x, y)$, which drops the payload at position (x, y)	17
3.2	Discrete, continuous, and parameterized action spaces.	18
3.3	Action selection with a two-tier policy.	19
3.4	The WAM robot arm.	19
5.1	The robot soccer 2D simulator.	27
5.2	The robot soccer goal domain.	28
5.3	Average return in the goal domain.	31
5.4	Average goal scoring probability.	32
5.5	A converged robot soccer goal episode.	32
6.1	A screenshot from the Platform domain.	33
6.2	The components of the Platform domain.	34
6.3	The Platform domain's actions: run, hop, and leap.	34
6.4	Average percentage distance covered in the Platform domain.	35
6.5	A successful episode of the Platform domain.	35
B.1	The ball sorting problem: balls must be dropped into one of the boxes, but not in between them.	43
C.1	A discontinuous policy domain.	44
C.2	Optimal policy $\pi^*(x)$ for the discontinuous policy domain.	44

Chapter 1

Introduction

As research in robotics and autonomous control advances, efforts are being made to tackle complex, real-world problems. These problems are too difficult to solve with a fixed, manually specified algorithm, and require agents that learn from experience. Reinforcement learning is a field that is interested in learning solutions for control problems.

40

In reinforcement learning, we typically consider either a discrete or a continuous action space [Sutton and Barto 1998]. With a discrete action space, we make decisions about which distinct action we wish to perform from a finite action set. With a continuous action space, we express the selected action as a real-valued vector. If we use a continuous action space, we lose the ability to consider differences in kind: all actions must be expressible as a single vector. However, if we only use discrete actions, we either lose the ability to apply the same action with different parameters, or suffer a blow-up in the number of actions as we must represent variations of the same action as distinct.

A parameterized action is a discrete action parameterized by a real-valued vector. Modeling actions this way introduces structure into the action space by treating different kinds of continuous actions as distinct. At each step an agent must choose both which action to execute based on the state and what parameters to execute it with.

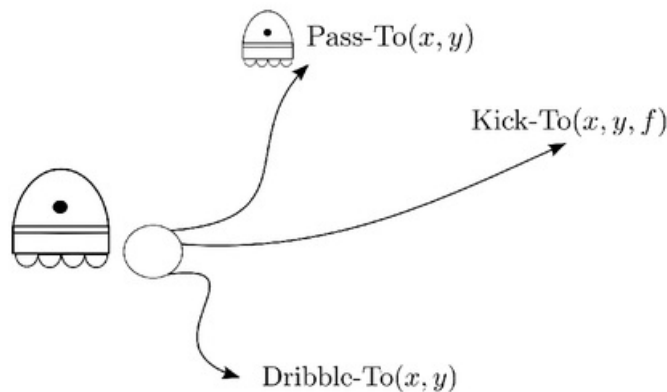


Figure 1.1: Parameterized actions for a soccer playing robot: dribble, kick, and pass.

For example, consider a soccer playing robot that can kick, pass, or run. We can associate a continuous parameter vector to each of these actions: we can kick the ball to a target with a given force, pass to another player at a specific position, or dribble to another position. This is illustrated in figure 1.1. Each

of these actions has its own distinct parameter vector.

This thesis introduces parameterized action Markov decision processes (PAMDPs), which model reinforcement learning problems with either distinct actions which require parameters to adjust the action to different situations, or where there are multiple mutually incompatible continuous actions. Each of these actions is parameterized in its own way, so expressing an action as a single vector containing all the parameters for each of these movements would be redundant, as only a small subset of the parameters are used for any action.

Although prior research has included parameterized and hybrid discrete-continuous actions, it has only done so from a planning perspective. Additionally, such work has largely assumed the parameter set is the same for all actions [Hoe *et al.* 2013; Rachelson 2009; Guestrin *et al.* 2004]. We consider the problem from a reinforcement learning perspective where we have no information about the model or reward function, and allow for each action to have its own set of parameters.

We focus on how to learn an action-selection policy given pre-defined parameterized actions. We introduce the Q-PAMDP algorithm, which alternates learning action-selection and parameter-selection policies and compare it to a direct policy search method. We show that with appropriate update rules Q-PAMDP converges to a local optimum. These methods are compared empirically in the goal and Platform domains, where Q-PAMDP out-performed direct policy search and fixed parameter SARSA.

39 This thesis is organized as follows. In chapter 2 we explain Markov decision processes, reinforcement learning, and policy search. In chapter 3 we introduce the idea of parameterized actions and define the parameterized action MDP. Chapter 4 details the Q-PAMDP algorithm and provides theoretical results pertaining to it. Chapter 5 and 6 give the experimental results for the goal domain and Platform domain, respectively. Chapter 7 concludes the thesis and covers related research, and discusses potential future work on this topic.

Chapter 2

Background

In this chapter we survey reinforcement learning and policy search. We consider the agent and environment model, Markov decision processes, value functions and Q-learning. We then consider policy search, looking at the episodic natural actor-critic (eNAC) algorithm in particular.

2.1 Reinforcement Learning

In reinforcement learning, we are concerned with solving control problems with a reward signal. The control problem involves an agent taking actions based on its state, after which it receives a new state and a reward. This process is illustrated in figure 2.1.

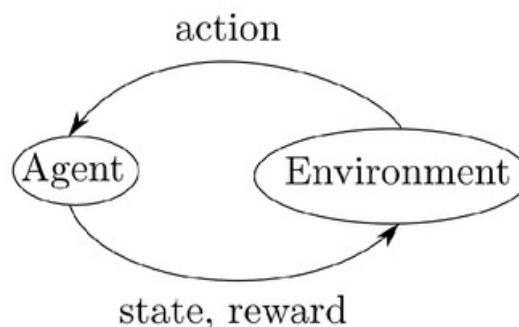


Figure 2.1: Interaction between an agent and the environment. At each step, the agent selects an action and executes it, receiving a new state and a reward.

As an example of a problem, we might have a robot (the agent) with a location (the state), that is navigating a maze (the environment). This is depicted in figure 2.2. The robot can move forward, left, right, or back (the actions) and receives a number indicating the cost of taking a single step, or a large reward for reaching the goal. Our task is to design an algorithm that will learn a policy that maps from states to actions, essentially determining how the agent acts in each state. To return to the maze example, a successful policy would tell the robot which direction to move in each state so as to finish the maze.

We can formalize this problem as a Markov decision process. A Markov decision process (MDP) is a tuple $\langle S, A, P, R, \gamma \rangle$, where S is a set of states, A is a set of actions, $P(s, a, s')$ is the probability of

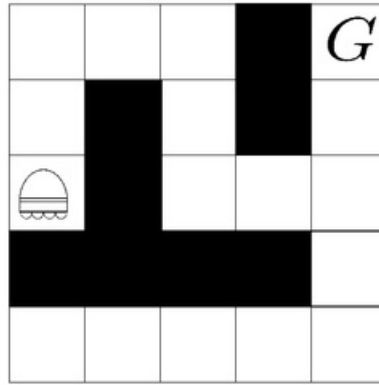
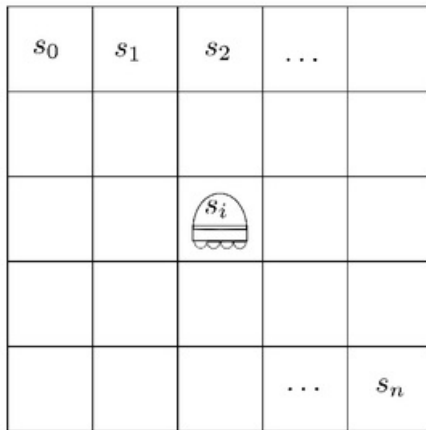


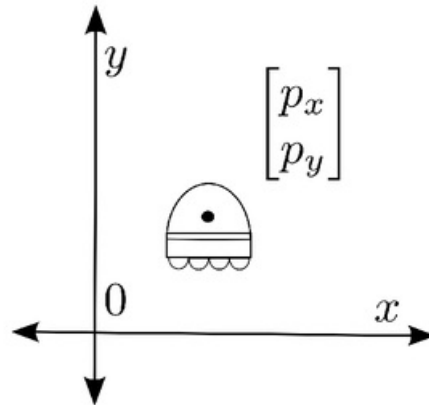
Figure 2.2: A simple maze problem. The robot (middle left) must navigate the maze to find the goal (marked G).

12
 38
 transitioning to state s' from state s after taking action a , $R(s, a, r)$ is the probability of receiving reward r for taking action a in state s , and γ is a discount factor [Sutton and Barto 1998].

The solution to an MDP is a policy π . A policy can be either stochastic or deterministic. With a deterministic policy, the policy π mapping $\pi : S \rightarrow A$ which selects an action for each state. A stochastic policy distribution $\pi(a|s)$ which determines the probability of selecting action a in state s . We wish to find a policy which maximizes the expected sum of discounted rewards (the return). An agent interacts with an MDP by taking an action and receiving a state and a reward.



(a) Discrete state space.



(b) Continuous state space.

Figure 2.3: Discrete and continuous state representations for the position of a robot. The discrete state space is a finite set of distinct positions. The continuous state space is a subset of all real-valued positions.

60
 35
 A state space can be either discrete or continuous. A discrete state space has a finite set of states $S = \{s_1, \dots, s_n\}$. A continuous state space has an infinite set of states, $S \subseteq \mathbb{R}^n$. Consider a two-dimensional robotics problem, where the robot has an x and a y position. With a discrete representation there are a finite number of completely distinct possible positions, as depicted in figure 2.3a. In a continuous representation the state is given by a 2-dimensional value, as depicted in figure 2.3b.

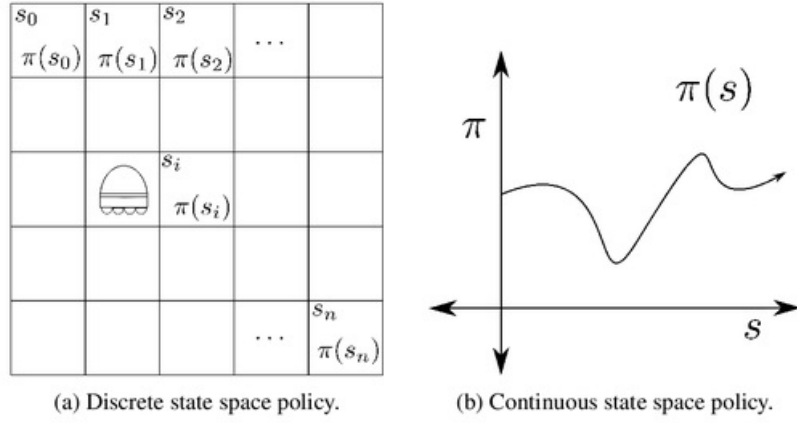


Figure 2.4: Policies for discrete and continuous state spaces.

How we represent a policy depends on the state space. A discrete state space requires a different choice for each state (figure 2.4a) A continuous state space assigns values based on a function (figure 2.4b).

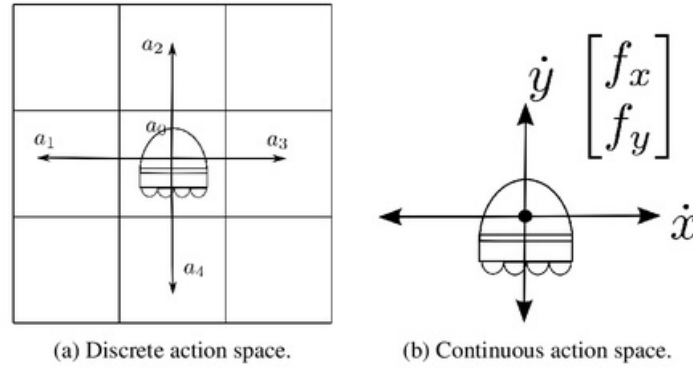


Figure 2.5: Discrete and continuous action spaces. The discrete action space is a finite set of distinct directions, while the continuous action space is a real-valued powers.

Similarly, the action space can be either discrete or continuous. A discrete action space is a finite set of actions $A = \{a_1, \dots, a_n\}$. To continue with our robotics example the robot may the move in different cardinal directions, as depicted in figure 2.5a. With a continuous action space the robot selects a real-value vector from an action space $A \subseteq \mathbb{R}^m$. The robot can move with force (f_x, f_y) , as depicted in figure 2.5b. Note that a continuous state space can have a discrete or a continuous action space.

A model of an MDP contains information about its transition and reward functions. In a planning problem we have an accurate model of the MDP, and we must compute the policy using this knowledge. In reinforcement [20](#) learning we do not have access [59](#) to the MDP but we can interact with it. Our approach can either be model-based or model-free. With model-based reinforcement learning, the agent constructs a model of the MDP using its experience of the MDP. We are concerned with model-free learning, where we learn without constructing a model.

Algorithm 1 SARSA [Sutton and Barto 1998]

Input:**Algorithm:** $s \leftarrow$ initial state $a \sim \pi(a|s)$ **repeat** Take action a , receive reward r , state s' $a' \leftarrow \pi(a'|s')$ $Q(s, a) \leftarrow Q(s, a) + \alpha [r + \gamma Q(s', a') - Q(s, a)]$ $s \leftarrow s'$ $a \leftarrow a'$ **until** Q converges

2.2 Value Functions

The value function $V^\pi(s)$ is defined as the expected discounted return achieved by policy π starting at state s :

$$V^\pi(s) = \mathbb{E}_\pi \left[\sum_{t=0}^{\infty} \gamma^t r_t \right].$$

Similarly, we define the action-value function

$$Q^\pi(s, a) = \mathbb{E}_\pi [r_0 + \gamma V^\pi(s)],$$

as the expected return obtained by taking action a in state s , and then following policy π thereafter. If we want to use a value function for control, we require a model because we need to know which state we'll be in after each action. We would prefer to avoid this requirement, therefore we learn the values of each action independently. Knowing the action-value function allows us to select the action which maximizes $Q^\pi(s, a)$. We can learn Q for an optimal policy using a method such as SARSA [Sutton and Barto 1998].

With a discrete state space, we can represent Q with a table that contains the value for each action-value pair. SARSA (algorithm 1) can learn the values for such a tabular representation. SARSA is an on-policy algorithm, which means that it bases policy on the current action-value function while learning it. The algorithm uses a temporal difference update rule

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \alpha [r_{t+1} + \gamma Q(s_{t+1}, a_{t+1}) - Q(s_t, a_t)].$$

which is analogous to the temporal difference update rule for $V(s)$.

In domains with a continuous state space, we can represent $Q(s, a)$ using parametric function approximation with a set of parameters ω . The simplest form of parametric function approximation is linear function approximation:

$$Q_\omega(s, a) = \omega^T \phi_a(s),$$

where ω is a parameter vector and ϕ_a is a feature vector for action a .

One basis scheme (which we adopt in this thesis) is the Fourier basis function, which uses features:

$$\phi_j(s) = \cos(\pi c_j \cdot s),$$

where c_j is a frequency vector with each element an integer between 0 and d , and d is the order of the Fourier basis. Figure 2.6 depicts several example Fourier basis functions. The basis is obtained by

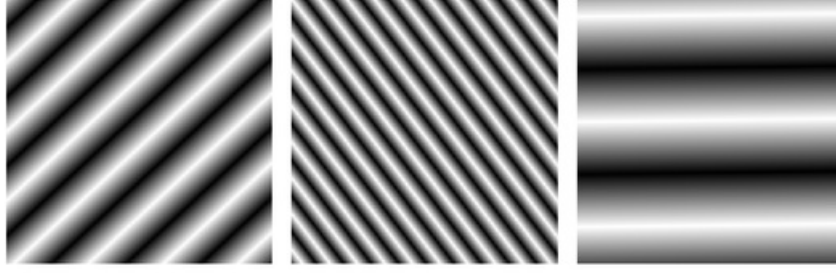


Figure 2.6: Examples of 2-dimensional Fourier basis functions.

enumerating all such frequency vectors. This results in $(d + 1)^n$ basis functions. We can reduce this number by placing restrictions on the coupling of variables used in the vectors c_j . The idea behind this basis is Fourier series approximation: for a function f with period T , it can be approximated with the series

$$\bar{f}(x) = \frac{a_0}{2} + \sum_{k=1}^n \left[a_k \cos\left(k \frac{2\pi}{T} x\right) + b_k \sin\left(k \frac{2\pi}{T} x\right) \right]. \quad 26$$

As the function we are approximating is unknown, we must determine the variables a_k . If we scale the variables between 0 and 1, we can assume that the function is periodic over $[-1, 1]$ (period 2), and is an even function i.e. $b_k = 0$. This leaves us with an approximation

$$\bar{f}(x) = \frac{a_0}{2} + \sum_{k=1}^n a_k \cos(2\pi kx),$$

which gives us the basis functions used above [Konidaris *et al.* 2011].

Algorithms such as gradient descent SARSA [Sutton and Barto 1998] learn this representation of Q . The idea is that we can perform a gradient update, which takes the form

$$\omega_{t+1} = \omega_t + \alpha \left[r_{t+1} + \gamma Q_{\omega_t}(s_{t+1}, a_{t+1}) - Q_{\omega_t}(s_t, a_t) \right] \nabla_{\omega_t} Q_{\omega_t}(s_t, a_t). \quad 32$$

If we use linear function approximation, the gradient is simply

$$\nabla_{\omega_t} Q_{\omega_t}(s_t, a_t) = \phi(s_t, a_t).$$

57 To speed up the learning process, we can use traces. The idea behind traces is that the reward can be attributed not just to the current state and action but also previous ones. By doing so, we can propagate rewards back to an action that led to a useful state. SARSA(λ) uses traces with a trace factor λ . Algorithm 2 gives the full description of gradient descent SARSA(λ).

2.3 Policy Search

For problems with a continuous action space ($A \subseteq \mathbb{R}^m$) we encounter the action selection problem: selecting the optimal action with respect to $Q(s, a)$ is non-trivial, as it requires finding a global maximum for a function in a continuous space. We can avoid this problem by using a policy search algorithm, given a class of policies parameterized by a set of parameters θ . A parameterized policy π_θ can directly return a continuous action and so avoids searching the entire action space at each step. This transforms

Algorithm 2 Gradient Descent SARSA(λ) [Sutton and Barto 1998]**Input:**Action-Value representation $Q_\omega(s, a)$ Policy π dependent on Q_ω **Algorithm:**Initialize ω randomly $s \leftarrow$ initial state $a \sim \pi(a|s)$ **repeat**Take action a , receive reward r , state s' **if** s' is terminal **then** $\delta \leftarrow r - Q_\omega(s, a)$ **else** $a' \sim \pi(a'|s')$ $\delta \leftarrow r + \gamma Q_\omega(s', a') - Q_\omega(s, a)$ **end if** $e \leftarrow \gamma \lambda e + \nabla_\omega Q_\omega(s, a)$ $\omega \leftarrow \omega + \alpha \delta e$ $s \leftarrow s'$ $a \leftarrow a'$ **until** θ converges

the problem into one of direct optimization over θ . Several approaches to policy search exist, including policy gradient methods, entropy-based approaches, path integral approaches, and sample-based approaches [Deisenroth *et al.* 2013].

Policy gradient methods use gradient ascent to optimize the parameters. Gradient ascent works by updating the parameters of a function in the direction of fastest increase of that function. This direction is given by the gradient of the function with respect to its parameters. The idea is that if $J(\theta)$ represents the expected discounted return for policy π_θ , and J is differentiable with respect to θ , then we can update θ with

$$\theta_{t+1} = \theta_t + \alpha_t \nabla_\theta J(\theta), \quad (2.1)$$

where α_t is a decreasing update rate. This works because $\nabla_\theta J(\theta)$ is the direction of the fastest increase in $J(\theta)$ with respect to θ . As long as α_t is sufficiently small, $J(\theta_{t+1}) \geq J(\theta_t)$. This is a local optimization method, so it will converge to some local optima θ^* with respect to $J(\theta)$.

The difficulty lies in estimating $\nabla_\theta J(\theta)$ [Deisenroth *et al.* 2013]. The policy gradient is given by

$$\nabla_\theta J(\theta) = \int_{\tau} \nabla_\theta p_\theta(\tau) R(\tau) d\tau, \quad (2.2)$$

where τ is a trajectory, and $p_\theta(\tau)$ is the probability of encountering trajectory τ under policy π_θ [Deisenroth *et al.* 2013]. Using the identities

$$\begin{aligned} \nabla_\theta p_\theta(\tau) &= p_\theta(\tau) \nabla_\theta \log p_\theta(\tau) \\ p_\theta(\tau) &= p(s_0) \prod_{t=1}^T p(s_{t+1}|s_t, a_t) \pi_\theta(a_t|s_t), \end{aligned}$$

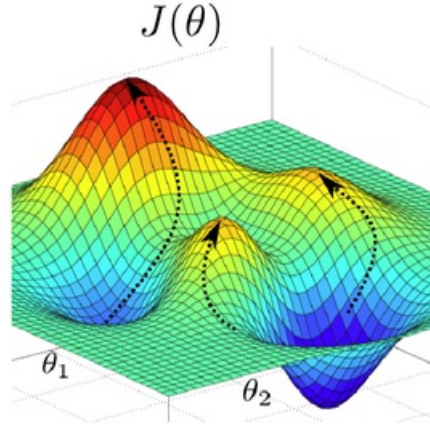


Figure 2.7: Gradient ascent on surface $J(\theta)$ with respect to $\theta = (\theta_1, \theta_2)$. As gradient ascent is a local optimization method, the maximum it converges to depends on the initial parameters.

we can write the gradient as the expectation

$$\nabla_{\theta} J(\theta) = \mathbb{E}_{p_{\theta}(\tau)} [\nabla_{\theta} \log p_{\theta}(\tau) R(\tau)] \quad (2.3)$$

$$= \mathbb{E}_{p_{\theta}(\tau)} \left[\sum_{t=0}^{T-1} \nabla_{\theta} \log \pi_{\theta}(a_t | s_t) R(\tau) \right], \quad (2.4)$$

as taking the log allows us to turn the product into a sum, and remove the dependence on the transition probabilities $p(s_{t+1} | s_t, a_t)$ [Deisenroth *et al.* 2013].

With a standard gradient, we assume that all parameters have a similar effect on the objective function. If we do not make this assumption, then differences between $p_{\theta}(\tau)$ and $p_{\theta+\Delta\theta}(\tau)$ can be large [Deisenroth *et al.* 2013]. The natural gradient accounts for this difference using the Fisher information matrix, given by

$$F_{\theta} = \mathbb{E}_{p(\tau)} [\nabla_{\theta} \log p_{\theta}(\tau) \nabla_{\theta} \log p_{\theta}(\tau)^T].$$

The natural gradient is given by

$$\nabla_{\theta}^{NG} J(\theta) = F_{\theta}^{-1} \nabla_{\theta} J(\theta),$$

and works by inverting the difference between successive θ so that each parameter has a similar effect on J [Deisenroth *et al.* 2013].

The episodic natural actor critic (eNAC) algorithm, given in algorithm 3, computes a natural gradient using n sample paths τ [Peters and Schaal 2008]. The policy gradient is given as

$$\begin{aligned} \nabla_{\theta}^{\text{eNAC}} J(\theta) &= F_{\theta}^{-1} \nabla_{\theta} J(\theta) \\ &= \mathbb{E}_{p_{\theta}(\tau)} [\psi \psi^T] \mathbb{E}_{p_{\theta}(\tau)} [\psi R(\tau)]. \end{aligned}$$

where

$$\psi = \sum_{t=0}^{T-1} \nabla_{\theta} \log \pi_{\theta}(a_t | s_t).$$

We can compute these expectations by taking the average over runs. This kind of evaluation tends to be very noisy, which results in our estimate of the gradient having a large variance. By introducing the baseline b for reward:

$$\nabla_{\theta} J_{\theta} = \mathbb{E}_{p_{\theta}(\tau)} [\nabla_{\theta} \log p_{\theta}(\tau) (R(\tau) - b)],$$

Algorithm 3 Episodic Natural Actor Critic (eNAC)

Input:Parameterized policy π_θ Initial parameters θ Initial feature function ϕ **Algorithm:****for** $i = 1$ **to** N **do**Collect samples $s_{1:T}^{[i]}, a_{1:T-1}^{[i]}, r_{1:T}^{[i]}$

$$R^{[i]} = \sum_{t=1}^T r_t^{[i]}$$

$$\psi^{[i]} = \begin{bmatrix} \sum_{t=0}^{T-1} \nabla_\theta \log \pi_\theta(a_t^{[i]} | s_t^{[i]}) \\ \phi(s_0^{[i]}) \end{bmatrix}$$

end for

$$R = \begin{bmatrix} R^{[1]}, \dots, R^{[N]} \end{bmatrix}^T$$

$$\Psi = \begin{bmatrix} \psi^{[1]}, \dots, \psi^{[N]} \end{bmatrix}^T$$

$$\begin{bmatrix} w \\ v \end{bmatrix} = (\Psi^T \Psi)^{-1} \Psi^T R$$

return $\nabla_\theta^{\text{eNAC}} J(\theta) = w$

we can reduce the variance by selecting the appropriate baseline [Deisenroth *et al.* 2013]. The baseline does not bias the gradient estimate, as

$$\begin{aligned}
 \mathbb{E}_{p_\theta(\tau)} [\nabla_\theta \log p_\theta(\tau) b] &= b \int_\tau \nabla_\theta p_\theta(\tau) d\tau \\
 &= b \nabla_\theta \int_\tau p_\theta(\tau) d\tau \\
 &= b \nabla_\theta 1 \\
 &= 0.
 \end{aligned}$$

We choose b so as to minimize the variance of the gradient estimate. The optimal b depends on the particular method of estimating the gradient. The eNAC algorithm uses a vector $\phi(s_0)$ to compute the baseline b [Deisenroth *et al.* 2013].

2.4 Summary

We have covered reinforcement learning, value functions, Q-learning and function approximation, and policy search. All these methods assume the agent is equipped with discrete or continuous actions. In the next chapter we introduce a formalization of the reinforcement learning problem that includes both.

Chapter 3

Reinforcement Learning with Parameterized Actions

A parameterized action is an discrete action that takes continuous set of parameters. Such an action consists of two parts: the discrete action and the continuous parameters. Intuitively, the discrete part determines the kind of action (the “what”), while the parameters determine its application (the “how”).

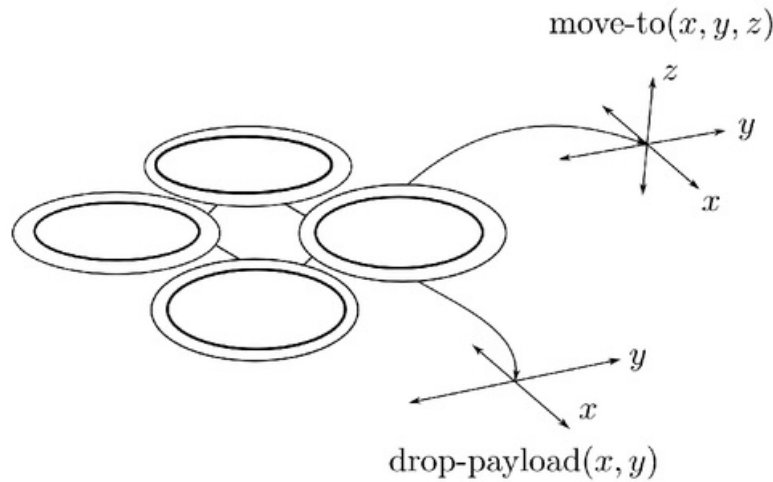
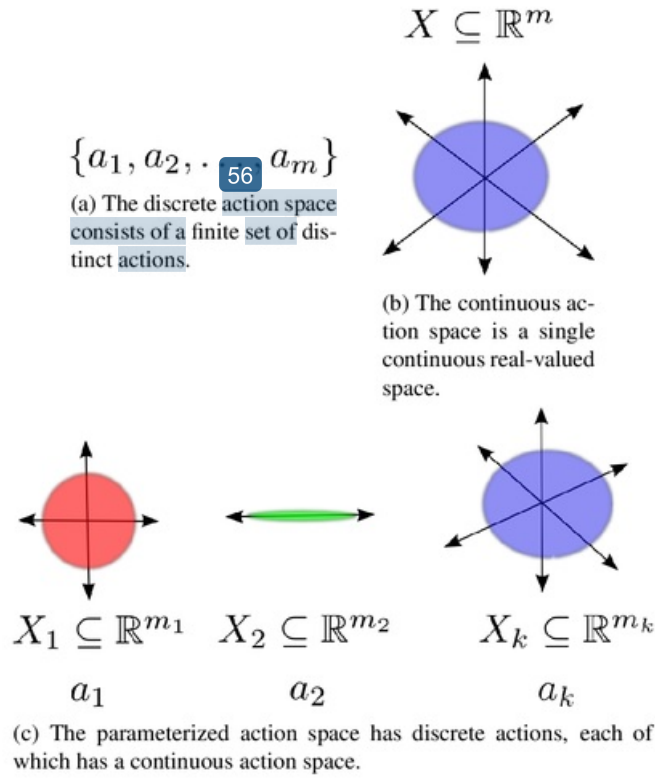


Figure 3.1: A quad-rotor delivery system. There are two actions: $\text{move-to}(x, y, z)$, which moves the quadrotor to position (x, y, z) , and $\text{drop-payload}(x, y)$, which drops the payload at position (x, y) .

For example, a quad-rotor delivery system might have a move action, $\text{move-to}(x, y, z)$, and an action $\text{drop-payload}(x, y)$. Figure 3.1 depicts this domain. Each of these actions is totally distinct from the other, and we would not want to combine these actions into a single continuous action lest we experiment with dropping payloads while moving.

We consider MDPs where the state space is continuous ($S \subseteq \mathbb{R}^n$) and the actions are parameterized. By this, we mean there is a finite set of actions $A_d = \{a_1, a_2, \dots, a_k\}$, and for each action $a \in A_d$, there is a set of continuous parameters $X_a \subseteq \mathbb{R}^{m_a}$. The action space is then given by

$$A = \bigcup_{a \in A_d} \{(a, x) \mid x \in X_a\}.$$



55
Figure 3.2: The three types of action spaces: discrete, continuous, and parameterized.

We refer to such MDPs as parameterized action MDPs (PAMDPs). Figure 3.2 depicts the different action spaces: discrete, continuous, and parameterized.

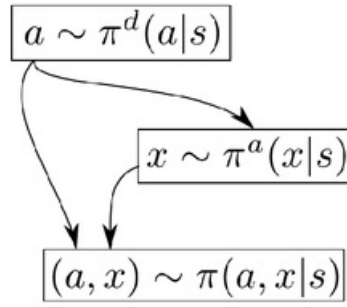
We apply a two-tiered approach for domains with parameterized actions: selecting the parameterized action, and selecting the parameters for that action. We denote the action policy by $\pi^d(a|s)$. To select the parameters for the action, we define the action-parameter policy for action a as $\pi^a(x|s)$. The policy is then given by $\pi(a, x|s) = \pi^d(a|s)\pi^a(x|s)$. In other words, to select a complete action (a, x) , we first sample a discrete action a from $\pi^d(a|s)$ and then sample a parameter x from $\pi^a(x|s)$. Figure 3.3 depicts this selection process.

Consider a robotics domain where we can apply a torque at each joint in a seven jointed robot arm. as depicted in figure 3.4. An action can be specified as a torque to each joint so that

$$a = (\tau_1, \tau_2, \tau_3, \tau_4, \tau_5, \tau_6, \tau_7)^T.$$

This gives us an action space $A \subseteq \mathbb{R}^7$. Working with such a low-level high-dimensional action space is difficult, so learning high-level low-dimensional skills is preferable. For example, we may learn a skill for positioning the end-effector at a position (x, y, z) , or a skill for moving an object to a position (x, y) . Each skill would translate the control into the appropriate low-level torques. As such, we could not combine skills as they would each indicate different torques. Therefore we could treat each skill as a parameterized action.

PAMDPs can also be used to represent disconnected action spaces (appendix B) or discontinuous action spaces (appendix C).



53

Figure 3.3: Action selection with a two-tier policy. At each step the agent first selects action a using the discrete policy π^d , then selects parameters x using parameter policy π^a . The complete action is then (a, x) .

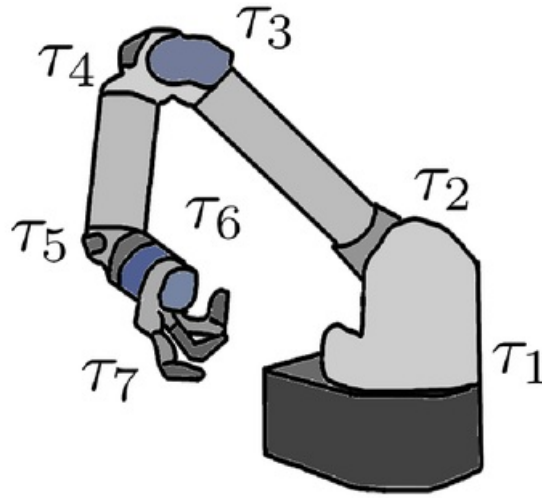


Figure 3.4: The WAM robot arm, with 7 torques.

3.1 Naïve Approaches

One approach for handling a continuous action space is to select a finite number of vectors from that space and use the set of these vectors as a discrete action space. We refer to this approach as discretization. For a continuous action space $A \subseteq \mathbb{R}^m$, we select actions $a_1, a_2, \dots, a_k \in A$. Our discretized action space becomes $A_d = \{a_1, a_2, \dots, a_k\}$. There are different schemes for selecting these actions.

The simplest is to evenly space actions over the action space. As an example, suppose we had a continuous action space $A = [-1, 1]$. One discretization would be $A_d = \{-1, -0.5, 0, 0.5, 1\}$. We can increase the granularity by using more actions, so that

$$A_d = \{-1, -0.75, -0.5, -0.25, 0, 0.25, 0.5, 0.75, 1\},$$

and so on. If we have a multi-dimensional action space, the number of actions needed to cover it will increase exponentially with the number of dimensions.

We can apply this approach by discretizing each parameter space individually. If we have discrete actions $a_1, \dots, a_k$ with parameter spaces $X_1, \dots, X_k$, then for each X_i select a set of parameters

$$X_i^d = x_{i1}, x_{i2}, \dots, x_{il_i} \subseteq X_i.$$

Then the complete discretized action space is given by

$$A = \{(a_1, x_{11}), \dots, (a_1, x_{1l_1}), (a_2, x_{21}), \dots, (a_2, x_{2l_2}), \dots, (a_k, x_{k1}), \dots, (a_k, x_{kl_k})\}.$$

The drawback of using discretization is adequately representing the action space may require many discrete actions. This can make learning prohibitively expensive. Discretization is an inappropriate approach for such action spaces.

The next approach we consider is one of direct policy search. The action policy is defined by the parameters ω and is denoted by $\pi_\omega^d(a|s)$. The action-parameter policy for action a is determined by a set of parameters θ_a , and is denoted $\pi_{\theta_a}^a(x|s)$. Let $\Theta = (\theta, \omega)$. We use a direct policy search method to optimize

$$J(\Theta) = \mathbb{E}_{s_0 \sim D} \left[V^{\pi_\Theta}(s_0) \right],$$

with respect to (Θ) , where s_0 is an initial state sampled according to the initial state distribution D . J is the expected return for a given policy starting at an initial state.

Policy gradient algorithms (eNAC for example) may require computation of the log policy gradient features, given by $\nabla_\Theta \log \pi_\Theta(a, x|s)$. To compute these features for a two-tier policy, note that

$$\begin{aligned} \nabla_\theta \log \pi_\Theta(a, x|s) &= \nabla_\theta \log \pi_\omega^d(a|s) \pi_{\theta_a}^a(x|s) \\ &= \nabla_\theta \log \pi_{\theta_a}^a(x|s), \\ \nabla_\omega \log \pi_\Theta(a, x|s) &= \nabla_\omega \log \pi_\omega^d(a|s) \pi_{\theta_a}^a(x|s) \\ &= \nabla_\omega \log \pi_\omega^d(a|s), \\ \nabla_\Theta \log \pi_\Theta(a, x|s) &= \begin{pmatrix} \nabla_\theta \log \pi_{\theta_a}^a(x|s) \\ \nabla_\omega \log \pi_\omega^d(a|s) \end{pmatrix}. \end{aligned}$$

3.2 Related Work

A parameterized task's problem depends on task parameters τ that is fixed through out the episode. The goal is to maximize

$$\int P(\tau) J(\pi_\tau, \tau) d\tau$$

where π_τ is a task dependent policy, J is the objective function, and $P(\tau)$ is the probability of task parameters τ .

Kober *et al.* [2012] developed algorithms to adjust dynamic motor primitives to different task parameters. They apply this to learn table-tennis and darts with different starting positions and targets. da Silva *et al.* [2012] introduced parameterized skills as a task dependent parameterized policy, and devised a method for learning different skills as distinct charts in policy space. The system works by sampling a set of tasks, learning their associated parameters, and trying to find a mapping from tasks to parameters [da Silva *et al.* 2012].

Either of these methods could be used to create parameterized actions. For example in the robot soccer domain, we could train a kick action with task parameters such as the target and power and then learn another skill for running at a particular speed. The task parameters for these skills become the parameters for our parameterized actions, and we are faced with the problem of selecting which parameterized actions to use, and how to select its parameters.

Guestrin *et al.* [2004] introduced an algorithm for solving factored MDPs with a hybrid discrete-continuous action space. However, their formalism has an action space with a mixed set of discrete and continuous components, whereas our domain has distinct actions with a different number of continuous components for each action. Furthermore, they assume the domain has a compact factored representation, and only consider planning.

Rachelson [2009] encountered parameterized actions in the form of an action to wait for a given period of time in his research on time dependent, continuous time MDPs (TMDPs). He developed XMDPs, which are TMDPs with a parameterized action space [Rachelson 2009]. He developed a Bellman operator for this domain, and in a later paper mentions that the TiMDP_{poly} algorithm can work with parameterized actions, although this specifically refers to the parameterized wait action [Rachelson *et al.* 2009]. This research also takes a planning perspective, and only considers a time dependent domain. Additionally, the size of the parameter space for the parameterized actions is the same for all actions.

Hoey *et al.* [2013] considered mixed discrete-continuous actions in their work on Bayesian affect control theory. They model affect control theory, a formalization of interpersonal interaction, as a POMDP with hybrid actions. To approach this problem they use a form of POMCP, a Monte Carlo sampling algorithm, using domain specific adjustments to compute the continuous action components [Silver and Veness 2010]. They note that the discrete and continuous components of the action space reflect different control aspects: the discrete control provides the “what”, while the continuous control describes the “how” [Hoey *et al.* 2013].

In their research on symbolic dynamic programming (SDP) algorithms, Zamani *et al.* [2012] considered domains with a set of discrete parameterized actions. Each of these actions has a different parameter space. Symbolic dynamic programming is a form of planning for relational or first-order MDPs, where the MDP has a set of logical relationships defining its dynamics and reward function. Their algorithms represent the value function as an extended algebraic decision diagram (XADD). As such, this work is limited to MDPs with predefined logical relations.

An option is a closed-loop policy that can be followed for multiple steps [Sutton *et al.* 1999]. Options generally represent a high-level skill. We could view each parameterized action as an option with a continuous policy. This would be useful if we want to allow for skills that are extended over a number of steps.

3.3 Summary

Parameterized actions are actions with both a discrete and a continuous component. We have described parameterized action MDPs (PAMDPs), and explained how a policy might be designed for such a domain. Such a policy consists of a discrete policy and a continuous policy. We have discussed two naive approaches for learning in PAMDPs: discretization and direct policy search. In the next chapter we introduce the Q-PAMDP algorithm, which alternates learning the discrete and continuous policies.

Chapter 4

The Q-PAMDP Algorithm

The Q-PAMDP algorithm alternates updating the parameter-policy and learning an action-value function. We can approach learning the action-value function by fixing the parameter-policy and treating it as a discrete-action MDP. Then we can update the parameter-policy by fixing the action-value function and treating it as an optimization problem.

To fix the parameter-policy, we define the corresponding discrete-action MDP as follows. For any PAMDP $M = \langle S, A, P, R, \gamma \rangle$ with a fixed parameter-policy π_θ , there exists a corresponding discrete action MDP, M_θ . Definitions used in this paper are given in appendix A.

Definition 4.1 (M_θ). M_θ is a tuple given by

$$M_\theta = \langle S, A_d, P_\theta, R_\theta, \gamma \rangle$$

where A_d is the discrete action set and

$$P_\theta(s, a, s') = \int_{x \in X_a} \pi_\theta^a(x|s) P(s, (a, x), s') dx$$

$$R_\theta(s, a, r) = \int_{x \in X_a} \pi_\theta^a(x|s) R(s, (a, x), r) dx.$$

The objective function for M_θ is the same as M for a fixed θ . We represent the action-value function for M_θ using function approximation with parameters ω . For M_θ , there exists an optimal set of representation weights ω_θ^* which maximizes $J(\theta, \omega)$ with respect to ω .

Let $W(\theta) = \omega_\theta^*$ be a function which selects an optimal ω for θ . We can compute $W(\theta)$ using a Q-learning algorithm such as gradient-descent SARSA(λ) [Sutton and Barto 1998].

Definition 4.2 ($W(\theta)$). The function W is given by

$$W(\theta) = \arg \max_{\omega} J(\theta, \omega).$$

Definition 4.3 ($H(\theta)$). The function H is given by

$$H(\theta) = J(\theta, W(\theta)).$$

Algorithm 4 Q-PAMDP(k)

Input:Initial parameters θ_0, ω_0

Parameter update method P-UPDATE

Q-learning algorithm Q-LEARN

Algorithm: $\omega \leftarrow \text{Q-LEARN}^{(\infty)}(M_\theta, \omega_0)$ **repeat** $\theta \leftarrow \text{P-UPDATE}^{(k)}(J_\omega, \theta)$ $\omega \leftarrow \text{Q-LEARN}^{(\infty)}(M_\theta, \omega)$ **until** θ converges

Algorithm 8 describes a method for alternating updating θ and ω . P-UPDATE should optimize θ with respect to $H(\theta) = J(\theta, W(\theta)) = J(\theta, \omega)$. The algorithm takes two methods, P-UPDATE and Q-LEARN. Q-LEARN can be any algorithm for Q-learning with function approximation, provided $\text{Q-LEARN}(M_\theta, \omega) = W(\theta)$, and is left as a design choice. P-UPDATE performs a single update of θ , after which ω is learned for the new MDP M_θ .

4.1 Convergence Properties of Q-PAMDP(1)

We now show that Q-PAMDP(1) converges to a local or bal optimum with mild assumptions. We assume that iterating P-UPDATE converges to some θ^* with respect to a given objective function f . As the P-UPDATE step is a design choice, it can be any algorithm with the appropriate convergence property such as eNAC. Q-PAMDP(1) is equivalent to the sequence

$$\begin{aligned}\omega_{t+1} &= W(\theta_t) \\ \theta_{t+1} &= \text{P-UPDATE}(J_{\omega_{t+1}}, \theta_t),\end{aligned}$$

if Q-LEARN converges to $W(\theta)$ for each given θ .

Theorem 4.1.1 (Convergence to a Local Optimum). *For any θ_0 , if the sequence*

$$\theta_{t+1} = \text{P-UPDATE}(H, \theta_t), \tag{4.1}$$

converges to a local optimum with respect to H , then Q-PAMDP(1) converges to a local optimum with respect to J .

Proof. By definition of the sequence above $\omega_t = W(\theta_t)$, so it follows that

$$\begin{aligned}J_{\omega_t} &= J(\theta, W(\theta)) \\ &= H(\theta).\end{aligned}$$

In other words, the objective function J equals H if $\omega = W(\theta)$. If we replace J with H in our update for θ , to obtain the update rule

$$\theta_{t+1} = \text{P-UPDATE}(H, \theta_t).$$

Therefore by equation 68 the sequence θ_t converges to a local optimum θ^* with respect to H . Let $\omega^* = W(\theta^*)$. As θ^* is a local optimum with respect to H , by definition there exists $\epsilon > 0$, s.t.

$$\|\theta^* - \theta\|_2 < \epsilon \implies H(\theta) \leq H(\theta^*).$$

Therefore for any ω ,

$$\begin{aligned} \left\| \begin{pmatrix} \theta^* \\ \omega^* \end{pmatrix} - \begin{pmatrix} \theta \\ \omega \end{pmatrix} \right\|_2 < \epsilon &\implies \|\theta^* - \theta\|_2 < \epsilon \\ &\implies H(\theta) \leq H(\theta^*) \\ &\implies J(\theta, \omega) \leq J(\theta^*, \omega^*). \end{aligned}$$

Therefore (θ^*, ω^*) is a local optimum with respect to J . $\square$

In summary, if we can locally optimize θ ; and $\omega = W(\theta)$ at each step, then we will find a local optimum for $J(\theta, \omega)$. The conditions for the previous theorem can be met by assuming that P-UPDATE is a local optimization method such as a gradient based policy search. A similar argument shows that if the sequence θ_t converges to a global optimum with respect to H , then Q-PAMDP(1) converges to a global optimum (θ^*, ω^*) .

One problem is that at each step we must re-learn $W(\theta)$ for the updated value of θ . We now show that if updates to θ are bounded and $W(\theta)$ is a continuous function, then the required updates to ω will also be bounded. Intuitively, we are supposing that a small update to θ results in a small change in the weights specifying which discrete action to choose. The assumption that $W(\theta)$ is continuous is strong, and may not be satisfied by all PAMDPs. It is not necessary for the operation of Q-PAMDP(1), but when it is satisfied we do not need to completely re-learn ω after each update to θ . We show that by selecting an appropriate α we can shrink the differences in ω as desired.

Theorem 4.1.2 (Bounded Updates to ω). *If W is continuous with respect to θ , and updates to θ are of the form*

$$\theta_{t+1} = \theta_t + \alpha_t P\text{-UPDATE}(\theta_t, \omega_t),$$

with the norm of each P-UPDATE bounded by

$$0 < \|P\text{-UPDATE}(\theta_t, \omega_t)\|_2 < \delta,$$

for some $\delta > 0$, then for any difference in ω $\epsilon > 0$, there is an initial update rate $\alpha_0 > 0$ such that

$$\alpha_t < \alpha_0 \implies \|\omega_{t+1} - \omega_t\|_2 < \epsilon.$$

Proof. Let $\epsilon > 0$ and

$$\alpha_0 = \frac{\delta}{\|P\text{-UPDATE}(\theta_t, \omega_t)\|_2}.$$

As $\alpha_t < \alpha_0$, it follows that

$$\begin{aligned} \delta &> \alpha_t \|P\text{-UPDATE}(\theta_t, \omega_t)\|_2 \\ &= \|P\text{-UPDATE}(\theta_t, \omega_t)\|_2 \\ &= \|\omega_{t+1} - \theta_t\|_2. \end{aligned}$$

So we have

$$\|\omega_{t+1} - \theta_t\|_2 < \delta.$$

As W is continuous, this means that

$$\|W(\theta_{t+1}) - W(\theta_t)\|_2 = \|\omega_{t+1} - \omega_t\|_2 < \epsilon.$$

□

In other words, if our update to θ is bounded and W is continuous, we can always adjust the learning rate α so that the difference between ω_t and ω_{t+1} is bounded.

4.2 Gradient of $H(\theta)$

With Q-PAMDP(1) we want P-UPDATE to optimize $H(\theta)$. One logical choice would be to use a gradient update. The next theorem shows that gradient of H is equal to the gradient of J $\omega = W(\theta)$. This is useful as we can apply existing gradient-based policy search methods to compute the gradient of J with respect to θ . The proof follows from the fact that we are at a global optimum of J with respect to ω , and so the gradient $\nabla_{\omega} J$ is zero. This theorem requires that W is differentiable (and therefore also continuous).

Theorem 4.2.1 (Gradient of $H(\theta)$). If $J(\theta, \omega)$ is differentiable with respect to θ and ω and $W(\theta)$ is differentiable with respect to θ , then the gradient of H is given by $\nabla_{\theta} H(\theta) = \nabla_{\theta} J(\theta, \omega^*)$, where $\omega^* = W(\theta)$.

Proof. If $\theta \in \mathbb{R}^n$ and $\omega \in \mathbb{R}^m$, then we can compute the gradient of H by the chain rule:

$$\begin{aligned} \frac{\partial H(\theta)}{\partial \theta_i} &= \frac{\partial J(\theta, W(\theta))}{\partial \theta_i} \\ &= \sum_{j=1}^n \frac{\partial J(\theta, \omega^*)}{\partial \theta_j} \frac{\partial \omega_j}{\partial \theta_i} + \sum_{k=1}^m \frac{\partial J(\theta, \omega^*)}{\partial \omega_k^*} \frac{\partial \omega_k^*}{\partial \theta_i}, \\ &= \frac{\partial J(\theta, \omega^*)}{\partial \theta_i} + \sum_{k=1}^m \frac{\partial J(\theta, \omega^*)}{\partial \omega_k^*} \frac{\partial \omega_k^*}{\partial \theta_i}, \end{aligned}$$

where $\omega^* = W(\theta)$. Note that by definitions of W ,

$$\begin{aligned} \omega^* &= W(\theta) \\ &= \arg \max_{\omega} J(\theta, \omega), \end{aligned}$$

we have that the gradient of J with respect to ω is zero

$$\frac{\partial J(\theta, \omega^*)}{\partial \omega_k^*} = 0,$$

as ω is a global maximum with respect to J for fixed θ . It follows that

$$\nabla_{\theta} H(\theta) = \nabla_{\theta} J(\theta, \omega^*).$$

□

To summarize, if $W(\theta)$ is continuous and P-UPDATE converges to a global or local optimum, then Q-PAMDP(1) will converge to a global or local optimum, respectively, and the Q-LEARN step will be bounded if the update rate of the P-UPDATE step is bounded. As such, if P-UPDATE is a policy gradient update step then Q-PAMDP by Theorem 4.3 will converge to a local optimum and by Theorem 4.4 the Q-LEARN step will require a fixed number of updates. This policy gradient step can use the gradient of J with respect to θ .

4.3 Convergence properties of Q-PAMDP(∞)

With Q-PAMDP(∞) each step performs a full optimization on θ and then a full optimization of ω . The θ step would optimize $J(\theta, \omega)$, not $H(\theta)$, as we do update ω while we update θ . Q-PAMDP(∞) has the disadvantage of requiring global convergence properties for the P-UPDATE method.

Theorem 4.3.1 (Local Convergence of Q-PAMDP(∞)). *If at each step of Q-PAMDP(∞) for some bounded set Θ :*

$$\begin{aligned}\theta_{t+1} &= \arg \max_{\theta \in \Theta} J(\theta, \omega_t), \\ \omega_{t+1} &= W(\theta_{t+1}),\end{aligned}$$

then Q-PAMDP(∞) converges to a local optimum.

Proof. By definition of W , $\omega_{t+1} = \arg \max_{\omega} J(\theta_{t+1}, \omega)$. This algorithm takes the form of direct alternating optimization. As such, it converges to a local optimum [Bezdek and Hathaway 2002]. $\square$

Q-PAMDP(∞) has weaker convergence properties than Q-PAMDP(1), as it requires a globally convergent P-UPDATE. However, it has the potential to bypass nearby local optima [Bezdek and Hathaway 2002].

4.4 Summary

We have examined the convergence properties of Q-PAMDP(1) and Q-PAMDP(∞). By making weak assumptions both methods convergence to a local optimum. Q-PAMDP(1) requires a locally convergent P-UPDATE method and Q-PAMDP has the stronger requirement that the P-UPDATE method be globally convergent.

Chapter 5

The Robot Soccer Goal Domain

Robot soccer is a challenge domain for designing autonomous soccer-playing robots. The difficulty of this problem allows for the exploration of a number of different problems in artificial intelligence such as planning, co-operation, skill-development, vision, communication, and state inference. We consider a two-dimensional simulated robot soccer problem, depicted in 5.1.

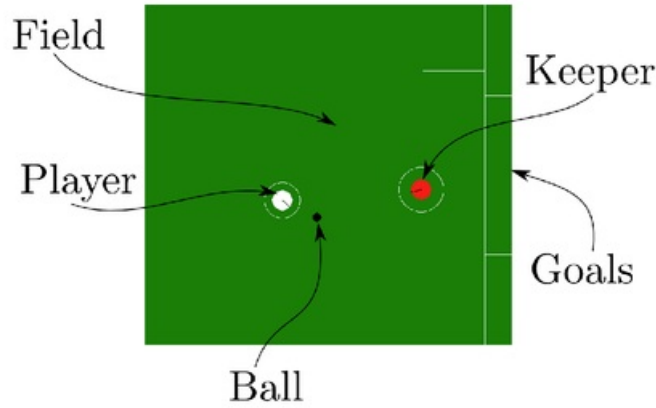


Figure 5.1: The robot soccer 2D simulator.

Each robot is represented by a circular entity with a position (x, y) , velocity $(\dot{x}, \dot{y})$, and an orientation θ . There are three primitive actions:

- $\text{turn}(\theta)$: turns the player θ degrees.
- $\text{dash}(\rho)$: accelerates the player forward with power ρ .
- $\text{kick}(\rho, \theta)$: if the ball is in the player's possession, accelerates the ball with power ρ in the direction θ .

Each action adds some noise to the control. The simulator uses a simple physics system: if two entities overlap they are shifted apart, and the positions are updated with the rules:

$$\begin{aligned}\hat{v}_{t+1} &= d\hat{v}_t + \hat{a}_{t+1} \\ \hat{p}_{t+1} &= \hat{p}_t + \hat{v}_{t+1},\end{aligned}$$

for position $\hat{p}$, velocity $\hat{v}$, and acceleration $\hat{a}$ and where d is the velocity decay factor. We describe the implementation details of this domain in appendix D.

5.1 Domain



Figure 5.2: The robot soccer goal domain. The player (white) attempts to kick the ball (black) past the keeper (red) and into the goal. The radius around players depicts their area of interaction with the ball.

Although the primitive actions turn, dash, and kick represent all the variety of actions we require to play the game, we wish to use higher-level actions. We therefore develop the skills to-ball, kick-to, and shoot-goal.

The to-ball() skill is used to move the player towards the ball until it gains possession of the ball. If the player is facing the ball, the player simply dashes at power 10, otherwise the player turns to face the ball.

The kick-to(x, y) is a skill that kicks the ball towards position $\hat{p} = (x, y)$. This requires computing the required power and direction for such a kick. First, note that starts at position $\hat{b}_0$ and velocity $\hat{v}_0$, and we assume no other force acts on it, then the system is given by

$$\begin{aligned} \hat{v}_{t+1} &= d\hat{v}_t \\ \hat{b}_{t+1} &= \hat{b}_t + \hat{v}_{t+1}. \end{aligned}$$

Expanding this equation, we have that

$$\begin{aligned} \hat{b}_t &= \hat{b}_0 + \sum_{i=0}^{t-1} d^i \hat{v}_0 \\ &= \hat{b}_0 + \frac{1 - d^t}{1 - d} \hat{v}_0. \end{aligned}$$

Note that the velocity diminishes exponentially, but is always positive provided $\hat{v}_0$ is non-zero. As such, to kick the ball to a position p , we must have that

$$\begin{aligned} \hat{p} &= \hat{b}_\infty \\ &= \lim_{t \rightarrow \infty} \hat{b}_t \\ &= \lim_{t \rightarrow \infty} \hat{b}_0 + \frac{1 - d^t}{1 - d} \hat{v}_0 \\ &= \hat{b}_0 + \frac{1}{1 - d} \hat{v}_0. \end{aligned}$$

Therefore, if we want to kick the ball to position (x, y) , it must start at velocity $\hat{v}_0$, and then

$$\hat{v}_0 = (1 - d)(\hat{p} - \hat{b}_0).$$

If the ball is already moving at velocity $\hat{v}$ we must then impart an acceleration $\hat{a}$ so that

$$\hat{a} + \hat{v} = \hat{v}_0.$$

Therefore

$$\begin{aligned}\hat{a} &= \hat{v}_0 - \hat{v} \\ &= (1 - d)(\hat{p} - \hat{b}_0) - \hat{v}.\end{aligned}$$

The kick-to command is reduced to $\text{kick}(\rho, \theta)$, where

$$\begin{aligned}\rho &= \|\hat{a}\|_2 \\ \theta &= \arctan 2(\hat{a}_1, \hat{a}_0).\end{aligned}$$

The $\text{shoot-goal}(h)$ action kicks the ball into the net at position h along the goal line. This is simply reduced to a kick-to action with target $(\text{GOAL-LINE} + \text{BALL-WIDTH}, h)$. We kick to a target inside the net by one BALL-WIDTH.

Each episode starts with the player at a random position along the left bound of the field. The player starts with the ball in its possession, and the keeper is positioned between the ball and the goal. The game takes place in a 2D environment where the player and the keeper have a position, velocity and orientation and the ball has a position and velocity resulting in 14 continuous state variables.

An episode ends when the keeper possesses the ball, the player scores a goal, or the ball leaves the field. The reward for an action is 0 for non-terminal state, 50 for a terminal goal state, and $-d$ for a terminal non-goal state, where d is the distance of the ball to the goal. The player has two parameterized actions: $\text{kick-to}(x, y)$, and $\text{shoot-goal}(h)$. If the player is not in possession of the ball, it moves towards it using $\text{to-ball}()$. The keeper has a fixed policy: it moves towards the ball, and if the player shoots at the goal, the keeper moves to intercept the ball.

To score a goal, the player must shoot around the keeper. This means that at some positions we must shoot left past the keeper, and at others shoot to the right past the keeper. However at no point do we shoot at the keeper, so an optimal policy would be discontinuous. This policy would be difficult to represent in a purely continuous action space, but is simpler in a parameterized action domain. We split the action into two parameterized actions: $\text{shoot-goal-left}(h)$, $\text{shoot-goal-right}(h)$. This allows us to use a simple action selection policy instead of complex continuous action policy.

5.2 Implementation Details

We represent the action-value function for the discrete action a using linear function approximation: $Q_\omega(s, a) = \omega_a^T \phi_a(s)$, where ω_a is a vector of weights, and $\phi_a(s)$ gives the features for state s . For this domain, we use Fourier basis features [Konidaris *et al.* 2011]. As we have 14 state variables, we must be selective in which basis functions to use. We only use basis functions with two non-zero elements and exclude all velocity state variables. We use the soft-max discrete action policy [Sutton and Barto 1998]

$$\pi_\omega^d(a|s) = \frac{\exp(Q_\omega(s, a)/\tau)}{\sum_{b \in A_d} \exp(Q_\omega(s, b)/\tau)},$$

where τ is the action selection temperature.

We represent the action-parameter policy π_θ^a as a normal distribution around a weighted sum of features

$$\pi_\theta^a(x|s) = \mathcal{N}(\theta_a^T \psi_a(s), \Sigma),$$

where θ_a is a matrix of weights, and $\psi_a(s)$ gives the features for state s , and Σ is a fixed covariance matrix. We use specialized features for each action. For the shoot-goal actions we use using a simple linear basis

$$\psi_{\text{sg}} = \begin{pmatrix} 1 \\ g \end{pmatrix},$$

where g is the projection of the keeper onto the goal line. For kick-to we use linear features

$$\psi_{\text{kt}}(s) = \left(1, bx, by, bx^2, by^2, \frac{bx - kx}{\|b - k\|_2}, \frac{by - ky}{\|b - k\|_2} \right)^T,$$

where (bx, by) is the position of the ball and (kx, ky) is the position of the keeper. These linear features allow for a policy that directs the ball around the keeper and towards the goal.

The policy is differentiable, allowing us to use a policy gradient method. To represent the discrete policy we use a polynomial basis which considers only the position variables. This is because using the Fourier basis described previously would require too many parameters to optimize. For the direct policy search approach, we use the episodic natural actor critic (eNAC) algorithm [Peters and Schaal 2008], optimizing $J(\omega, \theta)$ with respect to (ω, θ) . This approach shows the performance of currently available methods.

For the Q-PAMDP(1) and Q-PAMDP(∞) approaches we use the gradient-descent SARSA(λ) algorithm for Q-learning, and the eNAC algorithm for policy search [Peters and Schaal 2008]. Each eNAC update uses 50 episodes to estimate the gradient of J with respect to θ . For the direct policy search approach, we use the episodic natural actor critic (eNAC) algorithm [Peters and Schaal 2008], computing the gradient of $J(\omega, \theta)$ with respect to (ω, θ) . For the Q-PAMDP approach we use the gradient-descent SARSA(λ) algorithm for Q-learning, and the eNAC algorithm for policy search. At each step we perform one eNAC update based on 50 episodes and then refit Q_ω using 50 gradient descent SARSA(λ) episodes.

5.3 Results

We plot return in figure 5.3. As it is easier to interpret, we plot goal scoring probability in figure 5.4. Return is directly correlated with goal scoring probability, so their graphs are close to identical. We can see that direct eNAC is outperformed by Q-PAMDP(1) and Q-PAMDP(∞). This is likely due to the difficulty of optimizing the action selection parameters directly, rather than with Q-learning.

For both methods, the goal probability is greatly increased: while the initial policy rarely scores a goal, both Q-PAMDP(1) and Q-PAMDP(∞) increase the probability of a goal to roughly 35%. Direct eNAC converged to a local maximum of 15%. Finally, we include the performance of SARSA(λ) where the action parameters are fixed at the initial θ_0 . This achieves roughly 20% scoring probability. Both Q-PAMDP(1) and Q-PAMDP(∞) strongly out-perform fixed parameter SARSA, but eNAC does not. Figure 5.5 depicts a single episode using a converged Q-PAMDP(1) policy—the player draws the keeper out and strikes when the goal is open.

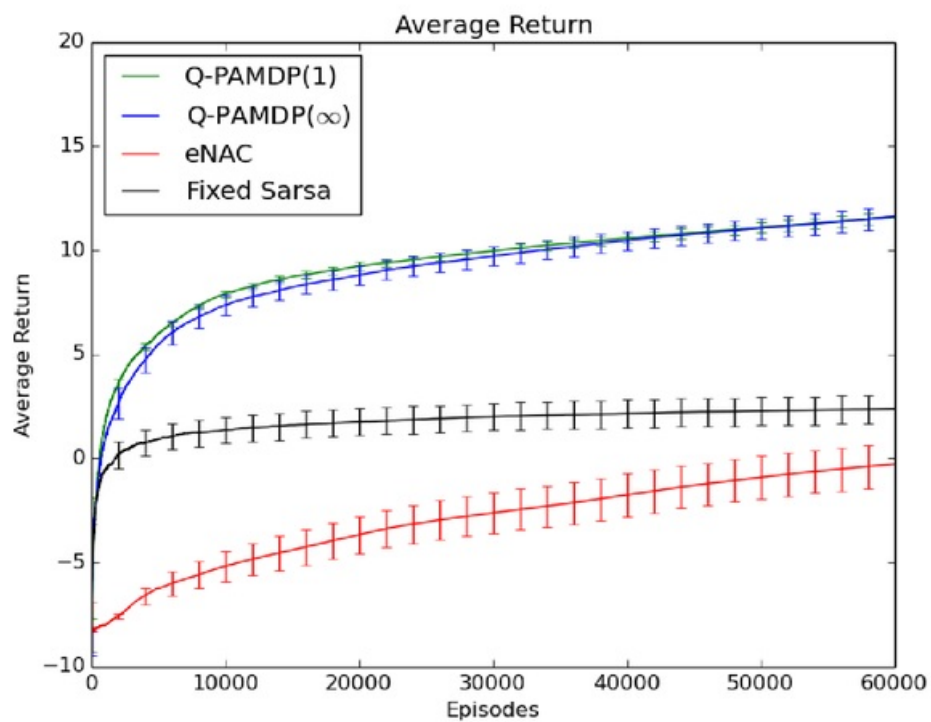


Figure 5.3: Average return, averaged over 20 runs for Q-PAMDP(1), Q-PAMDP(∞), fixed parameter Sarsa, and eNAC in the goal domain. Error bars show standard error.

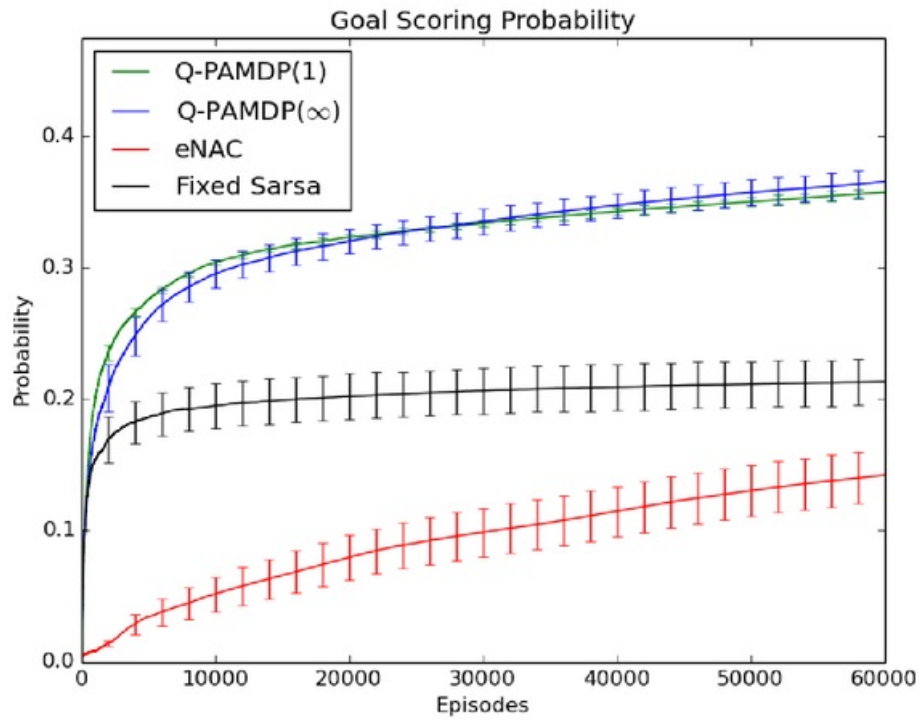


Figure 5.4: Average goal scoring probability, averaged over 20 runs for Q-PAMDP(1), Q-PAMDP(∞), fixed parameter Sarsa, and eNAC in the goal domain. Intervals show standard error.

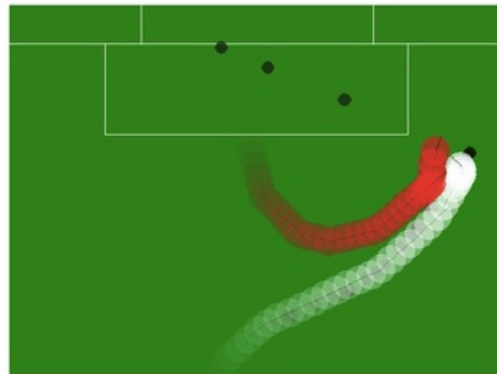


Figure 5.5: A robot soccer goal episode using a converged Q-PAMDP(1) policy. The player runs to one side, then shoots immediately upon overtaking the keeper.

Chapter 6

The Platform Domain

6

Next we consider the Platform domain, where the agent starts on a platform and must reach a goal while avoiding enemies. If the agent reaches the goal platform, touches an enemy, or falls into a gap between platforms, the episode ends. This domain is depicted in figure 6.1.

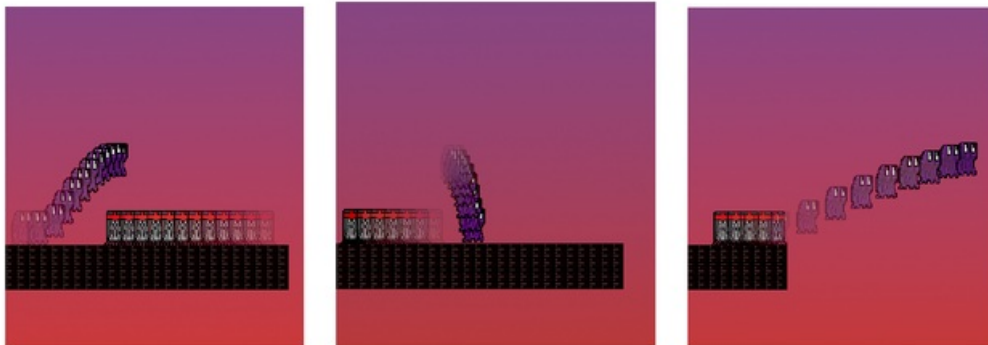


Figure 6.1: A screenshot from the Platform domain. The player (purple) hops over an enemy (grey and red), and then leaps over a gap.

6.1 Specification

The Platform domain has three types of objects: players, enemies, and platforms. The domain is 2-dimensional, and we specify positions by two variables x, y . Platforms are stationary and can be specified by a position and a width. Players and enemies have positions as well as velocities which are denoted $\dot{x}, \dot{y}$.

Each enemy is placed on a platform and “patrols” the platform, moving from one end to the other at a constant speed. An enemy only moves when the player is on or above its platform, so we only need to consider one enemy at the time. As an enemy’s y is fixed and its velocity $\dot{y} = 0$, we can specify an enemy’s position by its position x and velocity $\dot{x}$.

23

The reward for a step is the change in x value for that step, divided by the total length of all the platforms and gaps. The agent has two primitive actions: run or jump, which continue for a fixed period or until the agent lands again respectively. There are two different kinds of jumps: a high jump to get

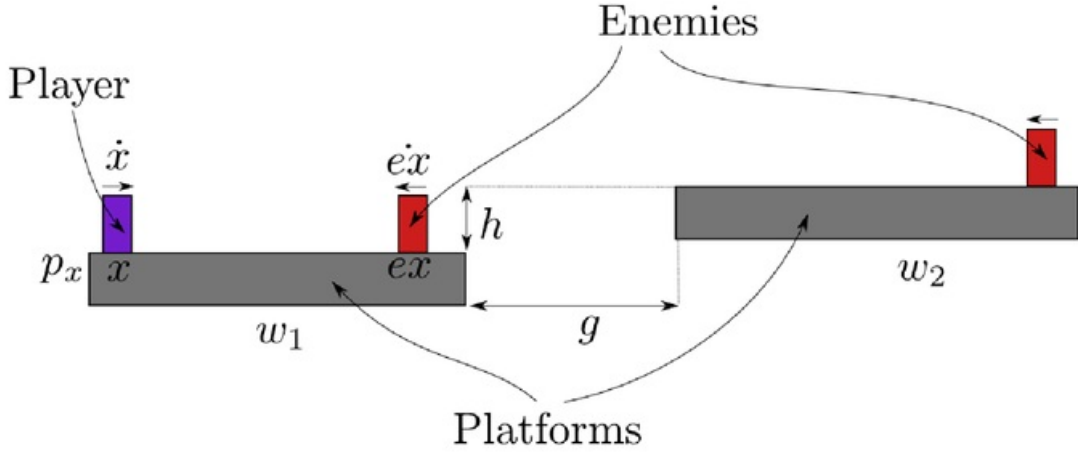


Figure 6.2: The components of the Platform domain. x , $\dot{x}$, ex , and $\dot{ex}$ represent the positions and velocities of the player and an enemy respectively. p_x is the position of the first platform, w_1 and w_2 are the width of the platforms, while g and h are the horizontal and vertical gaps between platforms.

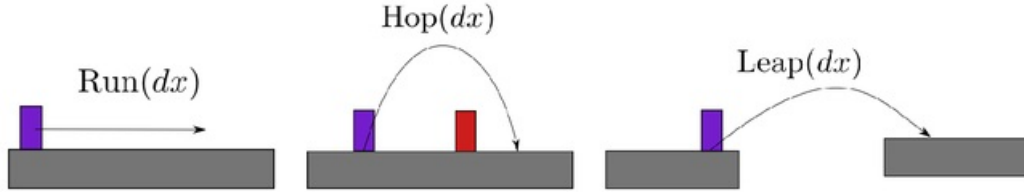


Figure 6.3: The Platform domain's actions: run, hop, and leap.

over enemies, and a long jump to get over gaps between platforms. The domain therefore has three parameterized actions: $\text{run}(dx)$, $\text{hop}(dx)$, and $\text{leap}(dx)$ (21). These actions are depicted in figure 6.3. The agent only takes actions while on the ground. The source code for this domain is available at <http://Github.com/WarwickMasson/aaai-platformer>.

6.2 Approach

The state space consists of four variables $(x, \dot{x}, ex, \dot{ex})$, representing the agent position, agent speed, enemy position, and enemy speed respectively. For learning Q_ω , as in the previous domain, we use linear function approximation with the Fourier basis. We apply a softmax discrete action policy with Q_ω .

We use a Gaussian parameter policy with parameter features

$$\psi_a(s) = [1, x, \dot{x}, ex, \dot{ex}, w_1, w_2, g, p_x, h]^T,$$

where w_1 is the width of the current platform, w_2 is the width of the next platform, g is the gap between platforms, p_x is the x position of the current platform and h is the difference in height between platforms. We have labeled these parameters in figure 6.2. All parameters are scaled to between $[-1, 1]$ so they have a similar effect on the parameter policy.

6.3 Results

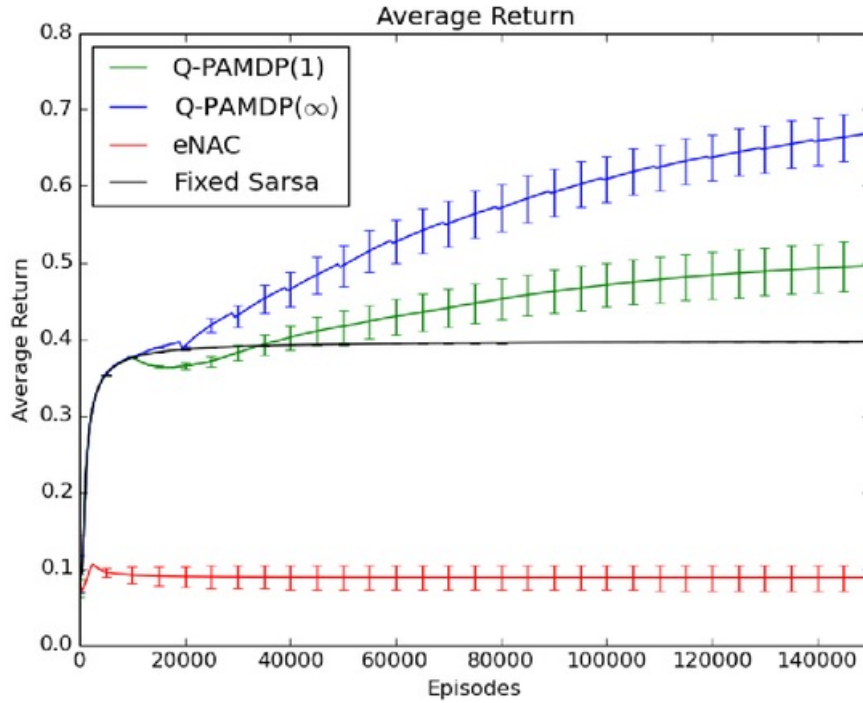


Figure 6.4: Average percentage distance covered, averaged over 20 runs for Q-PAMDP(1), Q-PAMDP(∞), and eNAC in the Platform domain. Intervals show standard error.

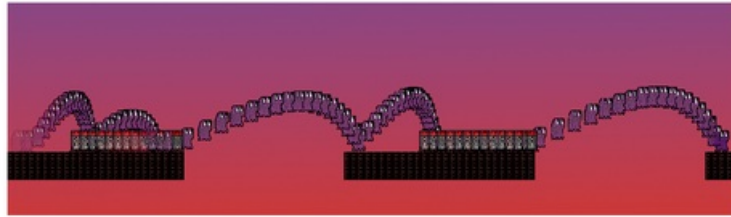


Figure 6.5: A successful episode of the Platform domain. The agent hops over the enemies, leaps over the gaps, and reaches the last platform.

Figure 6.4 shows the performance of eNAC, Q-PAMDP(1), Q-PAMDP(∞), and SARSA with fixed parameters. Both Q-PAMDP(1) and Q-PAMDP(∞) outperformed the fixed parameter SARSA method, reaching on average 50% and 65% of the total distance respectively. This is better than the fixed SARSA baseline of 40%, and much better than direct optimization using eNAC which reached 10%. Figure 6.5 shows a successfully completed episode of the Platform domain.

6.4 Summary

The Platform domain is a simulated platforming game. We demonstrated that both versions of Q-PAMDP outperformed fixed parameter SARSA and eNAC. Overall, Q-PAMDP(∞) performed better than Q-PAMDP(1).

Chapter 7

Conclusions and Future Work

7.1 Future Work

One aspect of parameterized actions explored in this paper is using multiple parameterizations of a single action. This would be useful in a domain which has a complicated optimal policy or for a domain where the optimal policy is discontinuous, and therefore difficult to represent with a single parameterization.

In a high dimensional continuous action space, we can learn a number of low dimensional parameterized skills which become the parameterized actions for a domain. The work of da Silva *et al.* [2012] and of Kober *et al.* [2012] would be particularly useful for the development of such skills.

One consideration would be combining the learning of θ and ω into a single step. With an appropriate learning rate, one might be able to update θ and ω together. This would require that an approximation of $H(\theta)$ be made using $J(\theta, \omega)$ where ω is sufficiently close to $W(\theta)$.

Another avenue of research would be to connect the computation of $\nabla_{\theta} J(\theta, W(\theta))$ and knowledge of $Q(s, a)$. This might lead to an actor-critic style algorithm, where we use a TD-error computed using $Q(s, a)$ to compute the gradient [Bhatnagar *et al.* 2009].

Further work is also needed to establish the theoretical standing of Q-PAMDP(∞), and to clarify its theoretical relationship with Q-PAMDP(1). It would also be useful to determine in what domains we would expect Q-PAMDP(1) to out-perform Q-PAMDP(∞) or vice versa.

7.2 Conclusion

The PAMDP formalism models reinforcement learning domains with parameterized actions. Parameterized actions give us the adaptability of continuous domains and the ability to use distinct kinds of actions. They also allow for simple representation of discontinuous policies without complex parameterizations. We have presented three approaches for model-free learning in PAMDPs: direct optimization and two variants of the Q-PAMDP algorithm. We have shown that Q-PAMDP(1), with an appropriate P-UPDATE method, converges to a local or global optimum. Q-PAMDP(∞) with a global optimization step converges to a local optimum.

We have examined the performance of these approaches in the goal scoring domain and the Platformer domain. The robot soccer goal domain models the situation where a striker must out-manuever a keeper to score a goal. Of these, Q-PAMDP(1) and Q-PAMDP(∞) outperformed eNAC and fixed parameter

SARSA. Q-PAMDP(1) and Q-PAMDP(∞) performed similarly well in terms of goal scoring, learning policies that score goals roughly 35% of the time. In the Platform domain we found that both Q-PAMDP(1) and Q-PAMDP(∞) outperformed eNAC and fixed SARSA.

By exploring the idea of reinforcement learning with parameterized actions, we can expand the range of potential applications of reinforcement learning. In particular, we can deal with more complex action spaces. The ultimate goal is to apply these advances to handle real-world control problems.

References

- [Bezdek and Hathaway 2002] JC Bezdek and RJ Hathaway. Some notes on alternating optimization. In *Advances in Soft Computing*, pages 288–300. Springer, 2002.
- [Bhatnagar *et al.* 2009] S Bhatnagar, RS Sutton, M Ghavamzadeh, and M Lee. Natural Actor–Critic Algorithms. *Automatica*, 45(11):2471–2482, 2009.
- [da Silva *et al.* 2012] BC da Silva, G Konidaris, and AG Barto. Learning parameterized skills. In *Proceedings of the Twenty-Ninth International Conference on Machine Learning*, pages 1679–1686, 2012.
- [Deisenroth *et al.* 2013] MP Deisenroth, G Neumann, and J Peters. *A Survey on Policy Search for Robotics*, volume 2. Now Publishers, 2013.
- [Guestrin *et al.* 2004] C Guestrin, M Hauskrecht, and B Kveton. Solving factored MDPs with continuous and discrete variables. In *Proceedings of the Twentieth Conference on Uncertainty in Artificial Intelligence*, pages 235–242, 2004.
- [Hoey *et al.* 2013] J Hoey, T Schroder, and A Alhothali. Bayesian affect control theory. In *Proceedings of the Fifth International Conference on Affective Computing and Intelligent Interaction*, pages 166–172. IEEE, 2013.
- [Kober *et al.* 2012] J Kober, A Wilhelm, E Oztop, and J Peters. Reinforcement learning to adjust parametrized motor primitives to new situations. *Autonomous Robots*, 33(4):361–379, 2012.
- [Konidaris *et al.* 2011] G Konidaris, S Osentoski, and PS Thomas. Value function approximation in reinforcement learning using the Fourier basis. In *Proceedings of the Twenty-Fifth International Conference on Artificial Intelligence*, pages 380–385, 2011.
- [Peters and Schaal 2008] J Peters and S Schaal. Natural actor-critic. *Neurocomputing*, 71(7):1180–1190, 2008.
- [Rachelson *et al.* 2009] E Rachelson, P Fabiani, and F Garcia. TiMDP-poly: an improved method for solving time-dependent MDPs. In *Proceedings of the Twenty First International Conference on Tools with Artificial Intelligence*, pages 796–799. IEEE, 2009.
- [Rachelson 2009] E Rachelson. *Temporal Markov Decision Problems: Formalization and Resolution*. PhD thesis, University of Toulouse, France, March 2009.
- [Silver and Veness 2010] D Silver and J Veness. Monte-Carlo planning in large POMDPs. In *Advances in Neural Information Processing Systems*, volume 23, pages 2164–2172, 2010.
- [Sutton and Barto 1998] RS Sutton and AG Barto. *Introduction to Reinforcement Learning*. MIT Press, Cambridge, MA, USA, 1998.

- [Sutton *et al.* 1999] RS Sutton, D Precup, and S Singh. Between mdps and semi-mdps: A framework for temporal abstraction in reinforcement learning. *Artificial Intelligence*, 112(1):181–211, 1999.
- [Zamani *et al.* 2012] Z Zamani, S Sanner, and C Fang. Symbolic dynamic programming for continuous state and action MDPs. In *Proceedings of the Twenty-Sixth AAAI Conference on Artificial Intelligence*, 2012.

Appendices

Appendix A

Theoretical Definitions

The following are the exact definitions used in this paper.

Definition A.1 (Euclidean Norm). For a vector $x \in \mathbb{R}^n$,

$$\|x\|_2 = \sqrt{\sum_{i=1}^n x_i^2}.$$

Definition A.2 (Continuity). A function $f(x)$, where $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$ is said to be continuous with respect to x if for any $\epsilon > 0$, $\exists \delta > 0$ such that $\forall x, y \in \mathbb{R}^n$,

$$\|x - y\|_2 < \delta \implies \|f(x) - f(y)\|_2 < \epsilon.$$

Definition A.3 (Convergence of a Sequence). We say that a sequence $\{a_n\}$ converges to a if for any $\epsilon > 0$, $\exists N \in \mathbb{N}$ such that

$$n \geq N \implies \|a_n - a\|_2 < \epsilon.$$

Definition A.4 (Global Optima). For a function $f : \mathbb{R}^n \rightarrow \mathbb{R}$, we say that x^* is a global optima for f if for any $x \in \mathbb{R}^n$,

$$f(x) \leq f(x^*).$$

Definition A.5 (Local Optima). For a function $f : \mathbb{R}^n \rightarrow \mathbb{R}$, we say that x^* is a local optima for f if there exists $\epsilon > 0$ such that for any x

$$\|x - x^*\|_2 < \epsilon \implies f(x) \leq f(x^*).$$

Appendix B

Disconnected Action Spaces

One use of the PAMDPs formalism is for handling disconnected action spaces. By disconnected action spaces, we mean that there are sets $X_1, X_2, \dots, X_k$, such that

$$A = \bigcup_{i=1}^k X_i,$$

but where for each X_i, X_j there is no path between one element of X_i to an element of X_j .

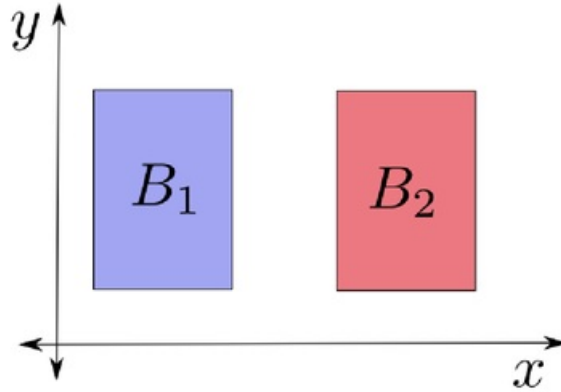


Figure B.1: The ball sorting problem: balls must be dropped into one of the boxes, but not in between them.

Consider a robotics problem involving a robot arm sorting balls into two boxes. At each step, the robot is given a ball of a given colour and size and must place it into the correct box. The colour of the ball determines which box it should be sorted into, while the size determines where it should be placed in the box. We have an action $\text{drop-ball-at}(x, y)$, which drops the ball at position (x, y) . We can only drop a ball into a box, any other position is invalid. The action space for this problem would look like figure B.1: two disconnected areas B_1 and B_2 , with the rest of the space $\mathbb{R}^2$ an invalid target. While we could treat the space $A = B_1 \cup B_2$ as a single continuous space, it makes more sense to treat these as two distinct actions with specific parameters. This is where parameterized actions are helpful: we can have two parameterized actions $(a_1, (x, y))$ and $(a_2, (x, y))$ corresponding to spaces B_1 and B_2 .

Appendix C

Discontinuous Policies

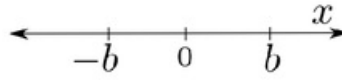


Figure C.1: Discontinuous policy domain. The agent must move to either $-b$ or b , whichever is closer.

In continuous parameterized policies, we may encounter a situation where the optimal policy is discontinuous. As an example, consider a simple domain where we have a single state variable $x \in \mathbb{R}$, with $x_0 = 0$, where we can apply action $a \in \mathbb{R}$ such that $x_{t+1} = x_t + a_t$, and we must move to either $x = -b$ or $x = b$, whichever is closer. This domain is depicted in figure C.1.

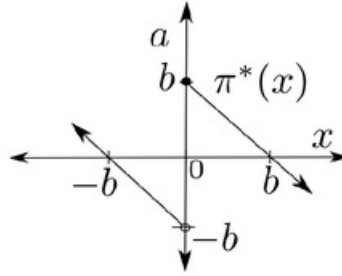


Figure C.2: Optimal policy $\pi^*(x)$ for the discontinuous policy domain.

The optimal policy for this domain is given by a piecewise continuous function

$$\pi^*(x) = \begin{cases} b - x & \text{if } x \geq 0 \\ x - b & \text{if } x < 0 \end{cases}.$$

This policy is depicted in figure C.2. Now consider how to represent this using a parameterized policy. A standard parameterized policy uses a linear form $\pi(s) = \theta^T \phi(s)$, where the features $\phi(s)$ tend to be linear with respect to s . Such a policy cannot represent the optimal policy π^* . This means we must specially design these features to cope with the discontinuity. If we plan to use a gradient update the policy must also be differentiable, which can complicate matters. We could alternatively treat this as a parameterized action problem: create two actions a_+ and a_- , and essentially create two policies are linear with respect to x . This particular strategy will be used in the goal domain for the shoot-goal action.

Appendix D

RSS Simulator

The 2D Robo-cup soccer server (RCSS) is a simplified simulated version of the robot soccer problem. Different sub-problems of soccer can be considered in the simulator. The soccer server is designed as a single server that agents connect to via sockets. It involves a complex message passing system, and systems for monitoring, player vision, and communication between players. These features are unnecessary and at times create obstacles for its use in this domain: the requirement of the server to be running independently means the server runs at its own schedule, which can complicate the use of learning agents and increase running times. More-over, the server uses message passing through UDP, which means that messages can be lost. The requirement that the messages be parsed from S-expressions adds additional computational load to the problem.

For all these reasons, we created a limited form of the simulator that implements the simulation rules. We ported these update rules to a python package. The simulator updates only when called with a new action, and as such there is no possibility of messages being lost. We implemented a visualizer for the game using the pygame package.

21

The source code for this domain is available at <http://github.com/WarwickMasson/aaai-goal>. This code is divided into modules for running the simulator, constructing the interface, running learning experiments, and drawing graphs.

ORIGINALITY REPORT

9%

SIMILARITY INDEX

6%

INTERNET SOURCES

7%

PUBLICATIONS

2%

STUDENT PAPERS

PRIMARY SOURCES

1

Submitted to University of Witwatersrand

Student Paper

<1%

2

www.ausy.tu-darmstadt.de

Internet Source

<1%

3

S. Schaal. "Reinforcement Learning for Parameterized Motor Primitives", The 2006 IEEE International Joint Conference on Neural Network Proceedings, 2006

Publication

<1%

4

J. Zico Kolter. "Regularization and feature selection in least-squares temporal difference learning", Proceedings of the 26th Annual International Conference on Machine Learning - ICML 09 ICML 09, 2009

Publication

<1%

5

Lecture Notes in Computer Science, 2008.

Publication

<1%

6

leenissen.dk

Internet Source

<1%

7

Bartana, A.. "Laser cooling of molecules by dynamically trapped states", Chemical

<1%

8	www.eea-esem.com Internet Source	<1 %
---	---	------

9	web.it.kth.se Internet Source	<1 %
---	---	------

10	www.ausy.informatik.tu-darmstadt.de Internet Source	<1 %
----	---	------

11	Jens Kober. "Policy search for motor primitives in robotics", Machine Learning, 11/06/2010 Publication	<1 %
----	---	------

12	Bei Li, , Siddharth Gangadhar, Samuel Cheng, and Pramode K. Verma. "Maximize user rewards in distributed generation environments using reinforcement learning", IEEE 2011 EnergyTech, 2011. Publication	<1 %
----	--	------

13	www.cs.cmu.edu Internet Source	<1 %
----	---	------

14	Lecture Notes in Computer Science, 2010. Publication	<1 %
----	---	------

15	A.M. Tehrani. "FUZZY REINFORCEMENT LEARNING FOR EMBEDDED SOCCER AGENTS IN A MULTI-AGENT CONTEXT", International Journal of Robotics and Automation, 2006 Publication	<1 %
----	---	------

16	Natural Language Dialog Systems and Intelligent Assistants, 2015. Publication	<1 %
17	Lecture Notes in Computer Science, 2007. Publication	<1 %
18	www.ijcai.org Internet Source	<1 %
19	citeseerx.ist.psu.edu Internet Source	<1 %
20	Adaptation Learning and Optimization, 2012. Publication	<1 %
21	www.lri.fr Internet Source	<1 %
22	ijcai.org Internet Source	<1 %
23	www.bio-nica.info Internet Source	<1 %
24	Submitted to University of Edinburgh Student Paper	<1 %
25	R.E. Skelton. "A convexifying algorithm for the design of structured linear controllers", Proceedings of the 39th IEEE Conference on Decision and Control (Cat No 00CH37187) CDC-00, 2000 Publication	<1 %
26	Submitted to University of Florida	

<1 %

27

www.cs.duke.edu

Internet Source

<1 %

28

Lecture Notes in Computer Science, 2005.

Publication

<1 %

29

www.db-thueringen.de

Internet Source

<1 %

30

Lecture Notes in Computer Science, 2003.

Publication

<1 %

31

db.s2.chalmers.se

Internet Source

<1 %

32

H. Tsujino. "Connectionist Reinforcement Learning with Cursory Intrinsic Motivations and Linear Dependencies to Multiple Representations", The 2006 IEEE International Joint Conference on Neural Network Proceedings, 2006

Publication

<1 %

33

Roberto Battiti. "Learning While Optimizing an Unknown Fitness Surface", Lecture Notes in Computer Science, 2008

Publication

<1 %

34

Cheng, Yuhu, Huanting Feng, and Xuesong Wang. "Efficient data use in incremental actor–critic algorithms", Neurocomputing, 2013.

<1 %

-
- | | | |
|----|---|------|
| 35 | www.dtic.upf.edu
Internet Source | <1 % |
|----|---|------|
-
- | | | |
|----|---|------|
| 36 | rldm.org
Internet Source | <1 % |
|----|---|------|
-
- | | | |
|----|---|------|
| 37 | roboticsproceedings.org
Internet Source | <1 % |
|----|---|------|
-
- | | | |
|----|---|------|
| 38 | users.soe.ucsc.edu
Internet Source | <1 % |
|----|---|------|
-
- | | | |
|----|---|------|
| 39 | dsl.serc.iisc.ernet.in
Internet Source | <1 % |
|----|---|------|
-
- | | | |
|----|---|------|
| 40 | Yamaguchi, Akihiko, Jun Takamatsu, and Tsukasa Ogasawara. "DCOB: Action space for reinforcement learning of high DoF robots", Autonomous Robots, 2013.
Publication | <1 % |
|----|---|------|
-
- | | | |
|----|--|------|
| 41 | Nguyen, Sao Mai, and Pierre-Yves Oudeyer. "Socially guided intrinsic motivation for robot learning of motor skills", Autonomous Robots, 2014.
Publication | <1 % |
|----|--|------|
-
- | | | |
|----|---|------|
| 42 | Levine, Sergey, Nolan Wagener, and Pieter Abbeel. "Learning contact-rich manipulation skills with guided policy search", 2015 IEEE International Conference on Robotics and Automation (ICRA), 2015.
Publication | <1 % |
|----|---|------|
-

43	www.irp.oist.jp Internet Source	<1 %
44	people.cs.ubc.ca Internet Source	<1 %
45	Apelian, . "The Derivative", Pure and Applied Mathematics, 2009. Publication	<1 %
46	Magdi S. Mahmoud. "Mathematical Foundations", Switched Time-Delay Systems, 2010 Publication	<1 %
47	Lecture Notes in Computer Science, 2006. Publication	<1 %
48	busoniu.net Internet Source	<1 %
49	hto-b.usc.edu Internet Source	<1 %
50	Lassaigne, Richard, and Sylvain Peyronnet. "Approximate planning and verification for large Markov decision processes", International Journal on Software Tools for Technology Transfer, 2015. Publication	<1 %
51	Reinhart, René Felix, and Jochen Jakob Steil. "Efficient policy search in low-dimensional embedding spaces by generalizing motion primitives with a parameterized skill	<1 %

52

www.ri.cmu.edu

Internet Source

<1 %

53

da Motta Salles Barreto, A.. "Restricted gradient-descent algorithm for value-function approximation in reinforcement learning", Artificial Intelligence, 200803

Publication

<1 %

54

refoteka.ru

Internet Source

<1 %

55

Dijst, Martin. "ICTs and Accessibility: An Action Space Perspective on the Impact of New Information and Communication Technologies", Advances in Spatial Science, 2004.

Publication

<1 %

56

Wawrzyński, Paweł, and Ajay Kumar Tanwani. "Autonomous reinforcement learning with experience replay", Neural Networks, 2012.

Publication

<1 %

57

allserv.kahosl.be

Internet Source

<1 %

58

www.elec.qmul.ac.uk

Internet Source

<1 %

59

www.davidandre.com

Internet Source

<1 %

60	148.204.65.169 Internet Source	<1 %
61	www.cs.mcgill.ca Internet Source	<1 %
62	apps.cs.utexas.edu Internet Source	<1 %
63	Undergraduate Texts in Mathematics, 1996. Publication	<1 %
64	Xiaoming Chen. "Stability analysis and control design for 2-D fuzzy systems via basis-dependent Lyapunov functions", Multidimensional Systems and Signal Processing, 11/17/2011 Publication	<1 %
65	"Asymptotic Optimality", Springer Texts in Statistics, 1998 Publication	<1 %
66	Lecture Notes in Computer Science, 2013. Publication	<1 %
67	Shun-ichi Amari. "Natural Gradient Works Efficiently in Learning", Neural Computation, 02/1998 Publication	<1 %
68	Rosa, Paulo, Gary J. Balas, Carlos Silvestre, and Michael Athans. "A Synthesis Method of LTI MIMO Robust Controllers for Uncertain LPV Plants", IEEE Transactions on Automatic	<1 %

69

Wookey, Dean S., and George D. Konidakis. "Regularized feature selection in reinforcement learning", Machine Learning, 2015.
Publication

<1 %
