

Reinforcement learning

In many domains it is difficult to give a precise evaluation function, which would allow an agent to evaluate the effectiveness or correctness of her actions, except after she has reached a terminal state. In the terminal state the agent by default obtains an objective evaluation of her actions which is termed a **reward**, or a **reinforcement**. It is convenient to state the agent's task so that it would have to act in such a way to maximize this reward or reinforcement. This is called **reinforcement learning**.

In a general case the agent may not have full information about her environment, as well as a precise, or any, description of her actions. The situation of this agent is similar to one of the possible statement of a full task of artificial intelligence — the agent is placed in an environment she does not know, and cannot act in it. She has to learn to act in this environment effectively, to maximize some criterion, available as reinforcements.

We shall consider a probabilistic model of the agent's action outcomes. In fact, we shall assume that we are dealing with an unknown Markov decision problem (MDP).

Passive reinforcement learning

Let us first consider **passive reinforcement learning**, where we assume that the agent's policy $\pi(s)$ is fixed. Agent is therefore bound to do what the policy dictates, although the outcomes of her actions are probabilistic. The agent may watch what is happening, so she knows what states she is reaching and what rewards she gets there.

The agent's job is to learn the utilities of the states $U^\pi(s)$, computed according to the equation:

$$U^\pi(s) = E \left[\sum_{t=0}^{\infty} \gamma^t R(s_t) \right]$$

For the 4x3 example environment used here for illustration we will assume $\gamma = 1$.

Direct utility determination

The objective of the agent is to compute the state utilities $U^\pi(s)$ generated by the current policy $\pi(s)$. The state utilities are defined as expected values of the sums of the (discounted) reinforcements received by the agent, who started from a given state, and is acting according to her policy:

$$U^\pi(s) = E \left[\sum_{t=0}^{\infty} \gamma^t R(s_t) \right]$$

The agent may approximate the above quantity using the trials by computing the **reward-to-go** for each state visited along the way. At the end of each trial the agent takes the reward obtained in that state to be its utility, and then, backtracking along the trial, computes the reward-to-go for subsequent states, by summing up the rewards obtained in the last section of the trial.¹

For example, for the trial:

$(1, 1) \rightarrow (1, 2) \rightarrow (1, 3) \rightarrow (1, 2) \rightarrow (1, 3) \rightarrow (2, 3) \rightarrow (3, 3) \rightarrow (4, 3) \rightarrow +1$

we get $R_{tg}(4, 3) = 1$, $R_{tg}(3, 3) = 0.96$, $R_{tg}(2, 3) = 0.92$, $R_{tg}(1, 3) = 0.88$, $R_{tg}(1, 2) = 0.84$, $R_{tg}(1, 3) = 0.80$, $R_{tg}(1, 2) = 0.76$, $R_{tg}(1, 1) = 0.72$

¹In case of discounting coefficient $\gamma \neq 1.0$ the rewards should be properly discounted.

Trials

Recall the 4×3 environment we studied earlier:

3	→	→	→	+1
2	↑		↑	-1
1	↑	→	↑	←
	1	2	3	4

3	0.812	0.868	0.918	+1
2	0.762		0.660	-1
1	0.705	0.655	0.611	0.388
	1	2	3	4

The agent executes the **trials** where she performs actions according to the policy, until reaching a terminal state, and receives percepts indicating both the current state and the reinforcement. Example trials:

$(1, 1) \rightarrow (1, 2) \rightarrow (1, 3) \rightarrow (1, 2) \rightarrow (1, 3) \rightarrow (2, 3) \rightarrow (3, 3) \rightarrow (4, 3) \rightarrow +1$
 $(1, 1) \rightarrow (1, 2) \rightarrow (1, 3) \rightarrow (2, 3) \rightarrow (3, 3) \rightarrow (3, 2) \rightarrow (3, 3) \rightarrow (4, 3) \rightarrow +1$
 $(1, 1) \rightarrow (2, 1) \rightarrow (3, 1) \rightarrow (3, 2) \rightarrow (4, 2) \rightarrow -1$

By averaging over a large number of samples she can determine the subsequent approximations of the expected value of state utilities, which converge in the infinity to the real expected values. This way the reinforcement learning task is reduced to a simple inductive learning.

This approach works, but is not very efficient, since it requires a large number of trials. The problem is, that the algorithm, by using simple averaging, ignores important information contained in the trials, namely, that state utilities in neighboring states are related.

$(1, 1) \rightarrow (1, 2) \rightarrow (1, 3) \rightarrow (1, 2) \rightarrow (1, 3) \rightarrow (2, 3) \rightarrow (3, 3) \rightarrow (4, 3) \rightarrow +1$
 $(1, 1) \rightarrow (1, 2) \rightarrow (1, 3) \rightarrow (2, 3) \rightarrow (3, 3) \rightarrow (3, 2) \rightarrow (3, 3) \rightarrow (4, 3) \rightarrow +1$
 $(1, 1) \rightarrow (2, 1) \rightarrow (3, 1) \rightarrow (3, 2) \rightarrow (4, 2) \rightarrow -1$

For example, in the second trial of the previous example, the algorithm evaluates the utility of state (3,2) as the reward-to-go only from this trial, but ignores the fact, that the successor state in this state is (3,3), which has an already known, significantly higher utility. The Bellman equation relates the utilities of successor states, but this approach cannot take advantage of this.

Adaptive dynamic programming

The **adaptive dynamic programming** (ADP) is a process similar to the dynamic programming, combined with learning the model of the environment, ie. the state transition probability distribution and the reward function. It works by counting the transitions from the state-action pairs to the next states. The trials provide the training data of such transitions. The agent can compute the probability of the transitions as frequencies of their occurrences in the trials.

For example, in the presented trials, the action $\rightarrow$ (Right) was executed three times in the state (1,3). Two of these times the successor state was (2,3), so the agent should compute $P((2,3)|(1,3), \text{Right}) = \frac{2}{3}$.

(1,1)_{-0.04} $\leadsto$ (1,2)_{-0.04} $\leadsto$ (1,3)_{-0.04} $\leadsto$ (1,2)_{-0.04} $\leadsto$ (1,3)_{-0.04} $\leadsto$ (2,3)_{-0.04} $\leadsto$ (3,3)_{-0.04} $\leadsto$ (4,3)₊₁
 (1,1)_{-0.04} $\leadsto$ (1,2)_{-0.04} $\leadsto$ (1,3)_{-0.04} $\leadsto$ (2,3)_{-0.04} $\leadsto$ (3,3)_{-0.04} $\leadsto$ (3,2)_{-0.04} $\leadsto$ (3,3)_{-0.04} $\leadsto$ (4,3)₊₁
 (1,1)_{-0.04} $\leadsto$ (2,1)_{-0.04} $\leadsto$ (3,1)_{-0.04} $\leadsto$ (3,2)_{-0.04} $\leadsto$ (4,2)₋₁

After executing every single action the agent updates the state utilities by solving the (simplified) Bellman equation using one of the appropriate methods. The equation is simplified because we only know the distribution of the effects of actions belonging to the policy, and we cannot compute the best action for each state. Since we are computing the U^π we take just these actions.

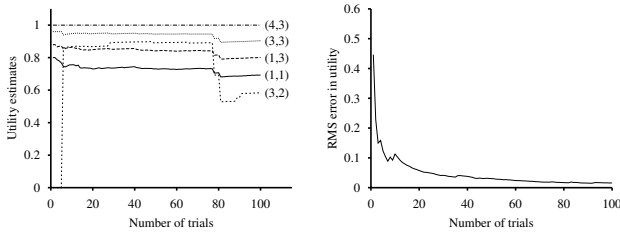
Adaptive dynamic programming — the algorithm

function PASSIVE-ADP-AGENT(*percept*) **returns** an action
inputs: *percept*, a percept indicating the current state s' and reward signal r'
persistent: π , a fixed policy
 mdp , an MDP with model P , rewards R , discount γ
 U , a table of utilities, initially empty
 N_{sa} , a table of frequencies for state-action pairs, initially zero
 $N_{s'|sa}$, a table of outcome frequencies given state-action pairs, initially zero
 s, a , the previous state and action, initially null

if s' is new **then** $U[s'] \leftarrow r'$; $R[s'] \leftarrow r'$
if s is not null **then**
 increment $N_{sa}[s, a]$ and $N_{s'|sa}[s', s, a]$
for each t such that $N_{s'|sa}[t, s, a]$ is nonzero **do**
 $P(t | s, a) \leftarrow N_{s'|sa}[t, s, a] / N_{sa}[s, a]$
 $U \leftarrow \text{POLICY-EVALUATION}(\pi, U, mdp)$
if $s'.\text{TERMINAL?}$ **then** $s, a \leftarrow \text{null}$ **else** $s, a \leftarrow s', \pi[s']$
return a

Adaptive dynamic programming — efficiency

The ADP algorithm updates the utility values as best as it is possible, and in this respect it is a standard that the other algorithms can be compared to. The policy computation procedure, which solves a system of linear equations, can be intractable for problems with large state spaces (eg. for the backgammon game we get 10^{50} equations with 10^{50} unknowns).



The above charts show the convergence for an example learning experiment for the 4×3 environment. In this experiment the first trial ending in the “bad” terminal state first occurs around the 78th trial, which causes large updates to some utility values.

Temporal difference learning

Instead of solving the full equation system for each trial, it is possible to update the utilities using the currently observed reinforcements. Such algorithm is called the **temporal difference** (TD) method:

$$U^\pi(s) \leftarrow U^\pi(s) + \alpha(R(s) + \gamma U^\pi(s') - U^\pi(s))$$

In this case the single utility value is updated from a single observed state transition, instead of the expected value of all transitions. This is why we take this correction — the difference between the utility of the move and the utility of the state — reduced by a factor $\alpha < 1$. This introduces small corrections after each move. The correction converges to zero when the state utility becomes equal to the discounted utility of a move.

Note that this method does not require having a model of the environment $P(s'|s, a)$, nor does it compute one.

Temporal difference learning — the algorithm

function PASSIVE-TD-AGENT(*percept*) **returns** an action
inputs: *percept*, a percept indicating the current state s' and reward signal r'
persistent: π , a fixed policy
 U , a table of utilities, initially empty
 N_s , a table of frequencies for states, initially zero
 s, a, r , the previous state, action, and reward, initially null

if s' is new **then** $U[s'] \leftarrow r'$
if s is not null **then**
 increment $N_s[s]$
 $U[s] \leftarrow U[s] + \alpha(N_s[s])(r + \gamma U[s'] - U[s])$
if $s'.\text{TERMINAL?}$ **then** $s, a, r \leftarrow \text{null}$ **else** $s, a, r \leftarrow s', \pi[s'], r'$
return a

Convergence of the temporal difference method

There is a close relationship and similarity between the ADP and TD algorithms. While the latter makes only local updates in utility values, their average values converge to the same values as for the ADP algorithm.

In case of learning with many transition examples, the frequencies of state occurrences agrees with their probability distributions, and it can be proved that the utilities converge to the correct values.

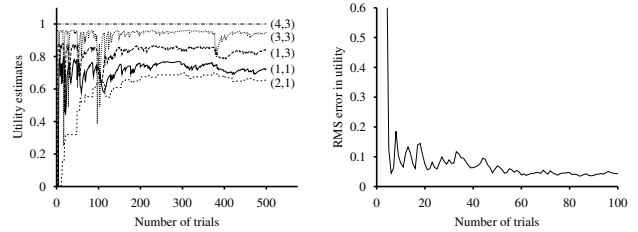
For this to occur the learning parameter α should decrease with the number of processed trials. More precisely, the subsequent values of this parameter should satisfy the requirements:

$$\sum_{n=1}^{\infty} \alpha(n) = \infty$$

and also:

$$\sum_{n=1}^{\infty} \alpha^2(n) < \infty$$

The convergence of another example learning experiment for the 4×3 environment:



Active reinforcement learning

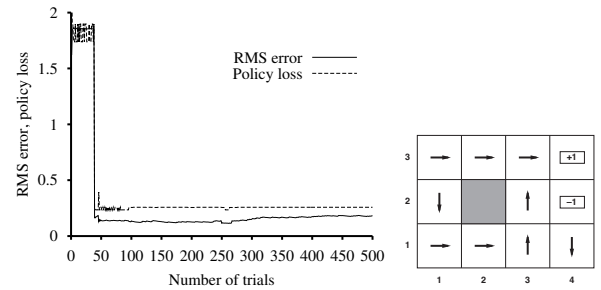
What should do an agent, which does not have a predefined policy, or who would like to find an optimal one?

First she should compute the complete transition model for all the actions. The ADP algorithm for a passive agent can be used for that. After that, the optimal policy, satisfying the Bellman equation, can be determined, like for a regular Markov decision process:

$$U(s) = R(s) + \gamma \max_a \sum_{s'} P(s'|s, a) U(s')$$

The agent can use the value iteration or the policy iteration algorithm. Then, having determined the optimal policy she could simply execute this policy.

But is this really what she should do?



The chart on the left shows the result of learning for a sample experiment. The agent found a direct path to the $[+1]$ solution in the 39th trial, but this was the worse path, along the states: $(2,1)$, $(3,1)$, $(3,2)$, $(3,3)$. This has, however, determined the agent's computed optimal policy, on the right. It turns out to be typical in this environment, that the agent only rarely arrives at the optimal policy preferring the "upper" path: $(1,2)$, $(1,3)$, $(2,3)$, $(3,3)$.

Why is this happening?

Exploration

If the agent, while learning the rules of the world, performed an arbitrarily imposed policy, then she did not build a correct model of the world ($P(s'|s, a)$), because she did not learn all her possibilities and the consequences of her possible actions. She did not learn them because she did not perform them, because they were not part of the policy.

If, instead, the agent had not only followed the imposed policy, or after initially following it, had not gone straight to calculating the optimal policy, but had performed other actions for a while, she would have had a chance to learn more about her world.

This is the trade-off between **exploitation** of existing knowledge and **exploration** of the environment. An agent should not too quickly accept the learned world model, and the calculated associated policy. If she wants to learn an accurate model of the world, and then determine the truly optimal policy, she should try different possibilities.

Moreover, she must repeatedly try all the actions in all the states if she wants to avoid having a random unlucky series prevent her from discovering some particularly good move. However, eventually she will want to start executing the optimal policy in order to reap the full benefits of the knowledge she has developed.

The exploration policy

In order to combine an effective world exploration with the exploitation of the acquired knowledge, an agent must have some **exploration policy**. This is necessary to ensure that the agent is able to learn all available actions to the extent permitting her to calculate the globally optimal policy.

A simple exploration policy would be to execute some random actions in all states some fraction of the time, and follow the optimal policy the rest of the time.

This approach works, but is slow to converge. It would be better to prefer the exploration of relatively unexplored state-action pairs, while at the same time avoiding the pairs already better known and believed to be of low utility.

The exploration function

A reasonable exploration policy can be achieved by introducing the optimistic approximation of utilities $U^+(s)$:

$$U^+(s) \leftarrow R(s) + \gamma \max_a f \left(\sum_{s'} P(s'|s, a) U^+(s'), N(a, s) \right)$$

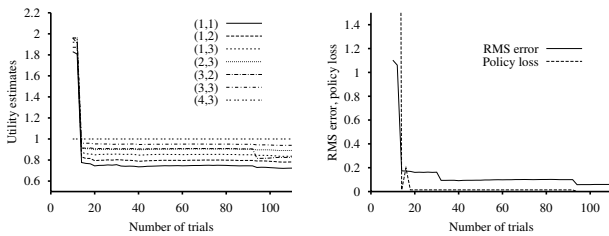
where $N(a, s)$ is the number of previous choices of action a in state s , and $f(u, n)$ is the **exploration function**, trading off the greed (large values of u) against curiosity (small values of n).

Obviously the f function should be increasing with respect to u and decreasing with respect to n . A simple definition of f can be:

$$f(u, n) = \begin{cases} R^+ & \text{if } n < N_e \\ u & \text{otherwise} \end{cases}$$

where the R^+ denotes the optimistic estimate of the best reward possible to obtain in any state, and the N_e is the minimal number of times that the agent will want to try each state-action pair.

The fact, that the update formula for U^+ in the right hand side also uses U^+ is important. Since the states and actions surrounding the starting state will be executed multiple times, if the agent used non-optimistic utilities for updating, she might get discouraged from such states, and start to avoid “setting out”. Using U^+ instead means, that the optimistic values generated around the newly explored regions will be propagated back, and the actions leading to unknown regions will be favored.



The chart on the left shows the results of learning with exploration. The policy close to optimal was reached after 18 trials. Note that the utility values converge more slowly (RMS error) than the optimal policy is determined (policy loss denotes the maximum loss from executing a suboptimal policy).

Active temporal difference (TD-)learning

The method of temporal differences can also be applied to active learning. The agent might not have a fixed policy, and still calculate the state utilities using the same update formula as in the passive case:

$$U^\pi(s) \leftarrow U^\pi(s) + \alpha(R(s) + \gamma U^\pi(s') - U^\pi(s))$$

Given the computed utilities the agent can determine the optimal actions for each state using the utilities of successor states. It can be proved, that the active TD agent will eventually arrive at the same resulting utility values as the active ADP agent.

The Q-function

An alternative to temporal difference learning of utilities is the Q -learning method, which learns an **action-value function** $Q(s, a)$, which represents the expected utility of taking a given action in a given state. The Q -function is related to the utility function U with the formula:

$$U(s) = \max_a Q(s, a)$$

Furthermore, the optimal policy can be extracted from the Q -function as follows:

$$\pi^*(s) = \operatorname{argmax}_a Q(s, a)$$

This represents a significant advantage of using — and learning — the Q -function instead of the utilities U . Even having the full distribution of utilities, the agent must also know the probability distribution $P(s'|s, a)$ to compute the optimal action in each state.

The Q -function gives this directly (almost, just by looking up the max). The $P(s'|s, a)$ distribution need not be learned or known.

A version of a Bellman equation can also be written for Q -values, deriving from the fact that the expected total reward for taking an action is its immediated reward plus the discounted utility of the outcome state:

$$Q(s, a) = R(s) + \gamma \sum_{s'} P(s'|s, a) U(s')$$

and the equation with respect to the Q -values only is:

$$Q(s, a) = R(s) + \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q(s', a')$$

The above equation can further be extended to the more general case of the reward function depending on the previous state and action taken therein, in addition to the goal state, so the expected reward has to be computed, instead of taking just the fixed value:

$$Q(s, a) = \sum_{s'} P(s'|s, a) [R(s, a, s') + \gamma \max_{a'} Q(s', a')]$$

Q -learning — updating by temporal differences

The above form of the Bellman equation can be used to derive the TD update formula for the Q values, by analogy to the TD update for utilities:

$$Q(s, a) \leftarrow Q(s, a) + \alpha (R(s) + \gamma \max_{a'} Q(s', a') - Q(s, a))$$

This update is calculated whenever action a is executed in state s leading to the result state s' . Here the term $(R(s) + \gamma \max_{a'} Q(s', a') - Q(s, a))$ represents an error in the known Q values and the update tries to minimize it.

We should once again note that the TD Q -learning agent does not need or learn the transition model $P(s'|s, a)$, so this represents a **model-free** approach to learning. This method can be applied even in very complex domains, but on the other hand the Q -learning agent has no way of looking into the future, so it may have difficulty when rewards are sparse and long action sequences must be taken to reach them.

Q -learning with temporal difference updating converges to the result much slower than the ADP algorithm since it does not enforce computing the full consistency of the model, which it does not have. The consequence it a much faster computation.

The full Q -learning algorithm with exploration

function Q-LEARNING-AGENT(*percept*) **returns** an action
inputs: *percept*, a percept indicating the current state s' and reward signal r'
persistent: Q , a table of action values indexed by state and action, initially zero
 N_{sa} , a table of frequencies for state–action pairs, initially zero
 s, a, r , the previous state, action, and reward, initially null
if TERMINAL?(s) **then** $Q[s, None] \leftarrow r'$
if s is not null **then**
 increment $N_{sa}[s, a]$
 $Q[s, a] \leftarrow Q[s, a] + \alpha (N_{sa}[s, a]) (r + \gamma \max_{a'} Q[s', a'] - Q[s, a])$
 $s, a, r \leftarrow s', \text{argmax}_{a'} f(Q[s', a'], N_{sa}[s', a']), r'$
return a

In general an active Q -learning agent requires exploration, just as it is in the case with the ADP method. That is why the algorithm uses the exploration function f and the action frequency table N . With a simpler exploration function (eg. executing some fraction of random moves) the N table might not be necessary.

SARSA — State-Action-Reward-State-Action

There exists a variant of the Q -learning algorithm with a temporal difference update method called SARSA (*State-Action-Reward-State-Action*):

$$Q(s, a) \leftarrow Q(s, a) + \alpha (R(s) + \gamma Q(s', a') - Q(s, a))$$

SARSA updates take into account five elements: s, a, r, s', a' . While the Q -learning algorithm adjustments are based on the best action selected for the state reached by the action a , SARSA considers which action was in fact selected. Therefore, eg. for a greedy agent executing only exploitation, these two approaches would give the same result.

However, in case of exploration the difference is significant. Q -learning is an **off-policy** learning algorithm, computing the best possible Q values, regardless of where the current policy makes the agent go. On the other hand, SARSA is an **on-policy** algorithm, which is more appropriate for the agent acting according to the policy.

Q -learning is a more flexible algorithm, as it permits the agent to learn the optimal behavior even if she currently executes a policy different from the best learned patterns. On the other hand, SARSA is more realistic, since, for example, if the agent was unable to fully control her policy, then it would be better for her to learn the patterns corresponding to what will in fact happen to her, than to learn the best possible patterns.

Both Q -learning and SARSA are capable of learning the optimal policy for the 4x3 example environment, albeit more slowly than the ADP (in terms of the number of iterations). This is because the local updates they make do not enforce the full consistency of the Q function.

Generalization in reinforcement learning

The above reinforcement learning algorithms assume an explicit representation of the $U(s)$ or $Q(s)$ function, such as, eg. table representation. This may be practical only up to some size of the problem.

For example, for problems with a very large number of states (eg. $\gg 10^{20}$ for games such as chess or backgammon), it is hard to envision executing a sufficient number of trials to visit each state frequently enough. It is necessary to use some generalization method, which would permit to determine an effective policy based on a small part of the state space explored.

Function approximation

One of such methods is the **function approximation**, based on the representation of the function under examination (like U or Q) in some nontabular form, like a finite formula. Similarly as was the case with the heuristic functions, some linear combination of some features (also called the state attributes) can be used:

$$\hat{U}_\theta(s) = \theta_1 f_1(s) + \theta_2 f_2(s) + \dots + \theta_n f_n(s)$$

The reinforcement learning algorithm would learn the vector of coefficients $\theta = \langle \theta_1, \theta_2, \dots, \theta_n \rangle$ so that the evaluation function $\hat{U}_\theta$ would closely enough approximate the state utility function. And this provides a very compact representation of the utility U (or Q) function for even very large state spaces.

This approach is called a function approximation because there is no proof that the real evaluation function can be expressed by this kind of a formula. However, while it seems doubtful that, for example, the optimal policy for chess can be expressed by a function with just a few coefficients, it is entirely possible that a good level of playing can be achieved this way.

Approximating direct utility estimation

We can try to apply the idea of function generalization to the method of direct utility estimation. Previously, we used this method directly on state utilities. Now we want to convert to the parametric representation of the utilities:

$$\hat{U}_\theta(s) = \theta_1 f_1(s) + \theta_2 f_2(s) + \dots + \theta_n f_n(s)$$

Given a collection of trials, we obtain a set of sample values approximating $\hat{U}_\theta(s)$ and a set of equations from the above formula, which can be solved for the optimal values of $\langle \theta_1, \theta_2, \dots, \theta_n \rangle$ in the sense of minimizing the squared error, for example, using the linear regression.

Alternatively, we may try to use **on-line learning**, with some way of updating the parameters based on the reinforcements obtained after each trial (or each step). We want to express an update formula for each of the approximation parameters θ_i based on the observed rewards. For example, if $u_j(s)$ is the reward-to-go for state s in j -th trial, then the utility function approximation error can be computed as:

$$E_j = \frac{(\hat{U}_\theta(s) - u_j(s))^2}{2}$$

The main idea of this approach is not an approximation, using fewer coefficients, of a function, which in fact requires many more of them, but the generalization. This is, we want to generate a policy valid for all the states based on the analysis of a small fraction of them.

For examples, in the experiments with the game backgammon, it was possible to train a player to a level of play comparable to human players based on examining one in 10^{12} states.

Obviously, the success in reinforcement learning in such cases depends on the correct selection of the approximation function. If no combination of the selected features can give a good strategy for a game, then no method of learning the coefficients will lead to one. On the other hand, selecting a very elaborate function, with a large number of features and coefficients, increases the chance for a success, but at the expense of a slower convergence and, consequently, the learning process.

The rate of change of this error with respect to the parameter θ_i can be written as $\partial E_j / \partial \theta_i$, so in order to adjust this parameter toward decreasing the error, the proper adjustment formula may be, where α is the usual learning coefficient:

$$\theta_i \leftarrow \theta_i - \alpha \frac{\partial E_j(s)}{\partial \theta_i} = \theta_i + \alpha (u_j(s) - \hat{U}_\theta(s)) \frac{\partial \hat{U}_\theta(s)}{\partial \theta_i}$$

The above formula is known as the **Widrow-Hoff** or the **delta** rule.

Approximating direct utility estimation — an example

As an example, for the 4x3 environment the state utility function could be approximated using a linear combination of the coordinates:

$$\hat{U}_\theta(x, y) = \theta_0 + \theta_1 x + \theta_2 y$$

According to the delta rule the corrections will be given by:

$$\begin{aligned}\theta_0 &\leftarrow \theta_0 + \alpha(u_j(s) - \hat{U}_\theta(s)) \\ \theta_1 &\leftarrow \theta_1 + \alpha(u_j(s) - \hat{U}_\theta(s))x \\ \theta_2 &\leftarrow \theta_2 + \alpha(u_j(s) - \hat{U}_\theta(s))y\end{aligned}$$

Assuming for an example $\theta = \langle \theta_0, \theta_1, \theta_2 \rangle = \langle 0.5, 0.2, 0.1 \rangle$ we get the initial approximation $\hat{U}_\theta(1, 1) = 0.8$. If, after executing a trial, we computed eg. $u_j(1, 1) = 0.72$, then all the coefficients $\theta_0, \theta_1, \theta_2$ would be reduced by 0.08α , which in turn would decrease the error for the state (1,1). Obviously, all the values of $\hat{U}_\theta(s)$ would then change, which is the idea of generalization.

Approximating temporal-difference learning

It is also possible to apply the idea of parametric representation to the temporal-difference learning. As in the previous approach, we need a formula to adjust the representation parameters $\theta = \langle \theta_1, \theta_2, \dots, \theta_n \rangle$ to reduce the approximation error. This time, instead of computing this error based on the reward-to-go values, we try to reduce the temporal difference between successive states after each step.

The formula for TD-learning:

$$\theta_i \leftarrow \theta_i + \alpha[R(s) + \gamma \hat{U}_\theta(s') - \hat{U}_\theta(s)] \frac{\partial \hat{U}_\theta(s)}{\partial \theta_i}$$

and for Q -learning:

$$\theta_i \leftarrow \theta_i + \alpha[R(s) + \gamma \max_{a'} \hat{Q}_\theta(s', a') - \hat{Q}_\theta(s, a)] \frac{\partial \hat{Q}_\theta(s, a)}{\partial \theta_i}$$

For passive TD learning, these update rules can be shown to converge to the closest possible approximation to the true function when the function approximator is linear in its features. But with active learning and non-linear approximation functions in many cases the parameters do not converge, even when good solutions exist.

Deep reinforcement learning

In many cases it is difficult to come up sufficiently good features for an acceptable linear approximation of U or Q . Sometimes it is easier to develop a complex, non-linear function, which works better.

In this role, neural networks can be used, and particularly deep neural networks. They can be used on data such as raw images, with no feature extraction at all. The deep neural network is able to discover the useful features. A reinforcement learning system that uses a deep network as a function approximator is called a deep reinforcement learning system.

The deep neural network is a function parametrized by θ , elements of which are all the parameters of the network (all the weights and biases in all the layers of the network). The feed-forward neural network updates these parameters in the process called **backpropagation**, which is their update rule.

Useful resources

Stuart J. Russell, Peter Norvig: Artificial Intelligence: A Modern Approach (Fourth Edition), Prentice-Hall, 2020
<https://aima.cs.berkeley.edu/global-index.html>

David Silver: Reinforcement Learning course (2015, University College London):
<http://www0.cs.ucl.ac.uk/staff/d.silver/web/Teaching.html>