

Reinforcement Learning Algorithms and Equations

Robert J. Moss

MOSSR@CS.STANFORD.EDU

Stanford University, Stanford, CA 94305

This write-up is intended as a collection of common algorithms and equations in reinforcement learning, deep reinforcement learning, decision making under uncertainty, approximate dynamic programming, and stochastic optimization methods.

1 Reinforcement Learning

Reinforcement learning is a form of machine learning that maximizes the cumulative discounted reward an agent receives from taking actions in an environment.¹

1.1 Bellman Equation

The *Bellman equation*² recursively computes the value of a decision problem.³ The next state s' comes from taking action a in state s using the transition model T , namely $s' = T(s, a)$.

$$\underbrace{U(s)}_{\text{state value}} \leftarrow \max_a \left(\underbrace{R(s, a)}_{\text{reward}} + \underbrace{\gamma U(s')}_{\substack{\text{discounted value} \\ \text{of next state}}} \right) \quad (1)$$

estimate of optimal discounted value

¹ R. S. Sutton, "Dyna, an Integrated Architecture for Learning, Planning, and Reacting," *SIGART Bulletin*, vol. 2, no. 4, pp. 160–163, 1991.

² R. E. Bellman, *Dynamic Programming*. Princeton University Press, 1957.

³ https://en.wikipedia.org/wiki/Bellman_equation

1.2 Q-learning

The core of the *Q-learning* algorithm⁴ is the Bellman equation.⁵ *Q-learning* is model-free and does not rely on the transition model T and the reward model R , but use samples of reward r and the next state s' .

$$Q(s, a) \leftarrow \underbrace{Q(s, a)}_{\text{old value}} + \underbrace{\alpha}_{\text{learning rate}} \cdot \underbrace{\left(\underbrace{r}_{\text{reward}} + \underbrace{\gamma \cdot \max_{a'} Q(s', a')}_{\substack{\text{estimate of optimal} \\ \text{future values}}} - \underbrace{Q(s, a)}_{\text{old value}} \right)}_{\substack{\text{temporal difference} \\ \text{new value (temporal difference target)}}} \quad (2)$$

⁴ C. J. C. H. Watkins, "Learning from Delayed Rewards," PhD thesis, University of Cambridge, 1989.

⁵ <https://en.wikipedia.org/wiki/Q-learning#Algorithm>

1.3 Sarsa

The *Sarsa* algorithm⁶ is a modification of *Q-learning* that uses (s, a, r, s', a') . It uses the actual next action a' to update Q instead of maximizing over all actions.

$$Q(s, a) \leftarrow \underbrace{Q(s, a)}_{\text{old value}} + \underbrace{\alpha}_{\text{learning rate}} \cdot \underbrace{\left(\underbrace{r}_{\text{reward}} + \underbrace{\gamma Q(s', a')}_{\substack{\text{discounted} \\ \text{next value}}} - \underbrace{Q(s, a)}_{\text{old value}} \right)}_{\substack{\text{temporal difference} \\ \text{new value} \\ \text{(temporal difference target)}}} \quad (3)$$

⁶ G. A. Rummery and M. Niranjan, *On-Line Q-Learning Using Connectionist Systems*. University of Cambridge, 1994, vol. 37.

2 Dynamic Programming

Optimal strategies can be found using a computation technique called *dynamic programming*.⁷ Dynamic programming relies on optimal substructures of the problem to reuse already computed information.

⁷ M. J. Kochenderfer, *Decision Making Under Uncertainty: Theory and Application*. MIT Press, 2015.

2.1 Policy Evaluation

Policy evaluation computes the expected utility U obtained from executing a policy π for n steps (in the finite horizon case with discount γ).⁸

⁸ M. J. Kochenderfer, *Decision Making Under Uncertainty: Theory and Application*. MIT Press, 2015.

$$U_t^\pi(s) = R(s, \pi(s)) + \gamma \sum_{s'} T(s' | s, \pi(s)) U_{t-1}^\pi(s') \quad (4)$$

```

function POLICYEVALUATION( $\pi, n$ )
   $U_0^\pi(s) \leftarrow 0$  for all states  $s$ 
  for  $t \leftarrow 1$  to  $n$ 
     $U_t(s) \leftarrow R(s, \pi(s)) + \gamma \sum_{s'} T(s' | s, \pi(s)) U_{t-1}^\pi(s')$  for all  $s$ 
  return  $U_n$ 

```

Algorithm 2.1. Policy evaluation.

2.2 Policy Iteration

Policy iteration computes the optimal policy π^* using policy evaluation.⁹ Starting with any policy π_0 , policy iteration performs two steps:

⁹ M. J. Kochenderfer, *Decision Making Under Uncertainty: Theory and Application*. MIT Press, 2015.

1. *Policy evaluation*: given the current policy, compute the value.
2. *Policy improvement*: using the value, compute a new policy.

Policy improvement uses the value function returned from policy evaluation.

$$\pi_{k+1}(s) = \arg \max_a \left(R(s, a) + \gamma \sum_{s'} T(s' | s, a) U^{\pi_k}(s') \right) \quad (5)$$

```

function POLICYITERATION( $\pi_0, k \leftarrow 0$ )
  repeat
    Compute  $U^{\pi_k}$  through POLICYEVALUATION
     $\pi_{k+1}(s) = \arg \max_a (R(s, a) + \gamma \sum_{s'} T(s' | s, a) U^{\pi_k}(s'))$  for all states  $s$ 
     $k \leftarrow k + 1$ 
  until  $\pi_k = \pi_{k-1}$ 
  return  $\pi_k$ 

```

Algorithm 2.2. Policy iteration.

2.3 Value Iteration

Value iteration uses dynamic programming and the Bellman equation to compute the optimal value (or utility) U iteratively.¹⁰

$$U^*(s) = \max_a \left(R(s, a) + \gamma \sum_{s'} T(s' | s, a) U^*(s') \right) \quad (6)$$

The optimal policy π can be extracted using the value function.

$$\pi^*(s) \leftarrow \arg \max_a \left(R(s, a) + \gamma \sum_{s'} T(s' | s, a) U^*(s') \right) \quad (7)$$

¹⁰ M. J. Kochenderfer, *Decision Making Under Uncertainty: Theory and Application*. MIT Press, 2015.

function VALUEITERATION

$k \leftarrow 0$

$U_0(s) \leftarrow 0$ for all states s

repeat

$U_{k+1}(s) \leftarrow \max_a \left(R(s, a) + \gamma \sum_{s'} T(s' | s, a) U_k(s') \right)$ for all states s

$k \leftarrow k + 1$

until convergence

return U_k

Algorithm 2.3. Value iteration.

2.4 Asynchronous Value Iteration (Gauss-Seidel Value Iteration)

Asynchronous value iteration updates the value for only a subset of states per iteration.¹¹ An example would be Gauss-Seidel value iteration that iterates through some ordering of the states and updates the values in place.¹²

$$U(s) \leftarrow \max_a \left(R(s, a) + \gamma \sum_{s'} T(s' | s, a) U(s') \right) \quad (8)$$

¹¹ M. J. Kochenderfer, *Decision Making Under Uncertainty: Theory and Application*. MIT Press, 2015.

¹² D. P. Bertsekas, "A New Value Iteration Method for the Average Cost Dynamic Programming Problem," *SIAM Journal on Control and Optimization*, vol. 36, no. 2, pp. 742–759, 1998.

2.5 Local Approximation Value Iteration

For large or continuous state spaces, approximate dynamic programming is useful in finding approximately optimal policies for the substructure of the problem.¹³ For a finite set of states, a function $\beta(s)$ is used to weigh "nearby" states, where $\sum_{i=1}^n \beta_i(s) = 1$. The value λ_i of each state s_i is used to approximate the value of an arbitrary state as $U(s) = \sum_{i=1}^n \lambda_i \beta_i(s) = \lambda^\top \beta(s)$.

¹³ M. J. Kochenderfer, *Decision Making Under Uncertainty: Theory and Application*. MIT Press, 2015.

function LOCALAPPROXIMATIONVALUEITERATION

$\lambda \leftarrow 0$

loop

$u_i \leftarrow \max_a \left(R(s_i, a) + \gamma \sum_{s'} T(s' | s_i, a) \lambda^\top \beta(s') \right)$ for $i \leftarrow 1$ to n

$\lambda \leftarrow \mathbf{u}$

return λ

Algorithm 2.4. Local approximation value iteration.

3 Deep Reinforcement Learning

Deep reinforcement learning combines deep learning (i.e. neural networks) with reinforcement learning algorithms (e.g., Q-learning) and can scale to larger problems.¹⁴

3.1 Trust Region Policy Optimization

Trust region policy optimization (TRPO)¹⁵ is a policy gradient method that guarantees policy improvement.¹⁶

$$\mathcal{L}_\pi(\tilde{\pi}) = \underbrace{\eta(\tilde{\pi})}_{\text{loss}} + \underbrace{\sum_s \rho_\pi(s)}_{\text{visitation frequency}} \underbrace{\sum_a \tilde{\pi}(a|s)}_{\text{advantage improving policy}} \underbrace{A_\pi(s,a)}_{\text{advantage}} \quad (9)$$

¹⁴ K. Arulkumaran, M.P. Deisenroth, M. Brundage, and A.A. Bharath, "Deep Reinforcement Learning: A Brief Survey," *IEEE Signal Processing Magazine*, vol. 34, no. 6, pp. 26–38, 2017.

¹⁵ J. Schulman, S. Levine, P. Moritz, M. I. Jordan, and P. Abbeel, "Trust Region Policy Optimization," *CoRR*, vol. abs/1502.05477, 2015. arXiv: 1502.05477.

¹⁶ Advantage function:

$$A(s,a) = Q(s,a) - U(s)$$

3.2 Actor-Critic with Experience Replay

Actor-critic with experience replay (ACER)¹⁷ uses batch off-policy updates to improve stability with optional trust region updates.¹⁸

¹⁷ Z. Wang, V. Bapst, N. Heess, V. Mnih, R. Munos, K. Kavukcuoglu, et al., "Sample Efficient Actor-Critic with Experience Replay," *CoRR*, vol. abs/1611.01224, 2016. arXiv: 1611.01224.

¹⁸ https://nervanasystems.github.io/coach/components/agents/policy_optimization/acer

$$Q^{\text{ret}}(s,a) = \underbrace{r}_{\text{reward}} + \underbrace{\gamma \tilde{\rho}}_{\text{discounted next truncated importance weight}} \left[\underbrace{Q^{\text{ret}}(s',a') - Q(s',a')}_{\text{estimate of } Q^\pi} \right] + \underbrace{\gamma U(s')}_{\text{discounted state value}} \quad \text{where } \underbrace{\tilde{\rho} = \min(c, \rho)}_{\text{truncated importance weight}}, \quad \rho = \frac{\pi(a|s)}{\mu(a|s)} \quad (10)$$

$$\underbrace{\hat{g}^{\text{policy}}}_{\text{policy gradient (w/ bias correction)}} = \underbrace{\tilde{\rho} \nabla \log \pi(a|s)}_{\text{importance weighted gradient of log-policy}} \left[\underbrace{Q^{\text{ret}}(s,a)}_{\text{Q-value retrace}} - \underbrace{U(s)}_{\text{state value}} \right] + \underbrace{\mathbb{E}_{\tilde{a} \sim \pi} \left(\left[\frac{\rho(a) - c}{\rho(a)} \right] \nabla \log \pi(\tilde{a}|s) \left[Q(s,\tilde{a}) - U(s) \right] \right)}_{\text{bias correction}} \quad (11)$$

3.3 Proximal Policy Optimization

Proximal policy optimization (PPO)¹⁹ is a deep reinforcement learning algorithm that implements a trust region²⁰ update compatible with stochastic gradient descent.²¹

¹⁹ J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov, "Proximal Policy Optimization Algorithms," *CoRR*, vol. abs/1707.06347, 2017. arXiv: 1707.06347.

$$\mathcal{L}^{\text{clip}}(\theta) = \underbrace{\mathbb{E}_t}_{\text{empirical expectation over time}} \left[\underbrace{\min(r_t(\theta) \hat{A}_t, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon) \hat{A}_t)}_{\text{probability ratio of new and old policy}} \right] \quad (12)$$

²⁰ D.C. Sorensen, "Newton's Method with a Model Trust Region Modification," *SIAM Journal on Numerical Analysis*, vol. 19, no. 2, pp. 409–426, 1982.

²¹ <https://openai.com/blog/openai-baselines-ppo>

4 Online Search

4.1 Monte Carlo Tree Search

Monte Carlo tree search (MCTS)²² is an anytime algorithm that uses rollouts of a random policy to estimate the value of each state-action node in the tree.

²² R. Coulom, “Efficient Selectivity and Backup Operators in Monte-Carlo Tree Search,” in *International Conference on Computers and Games*, 2006.

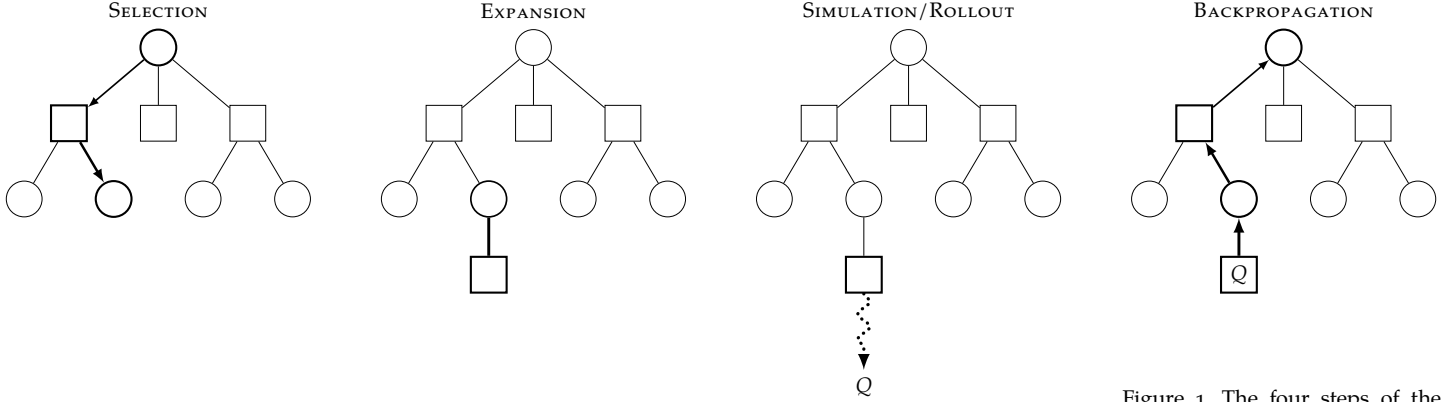


Figure 1. The four steps of the Monte Carlo tree search algorithm.

```

function MONTECARLOTREESearch( $s, d$ )
    loop SIMULATE( $s, d, \pi_0$ )
    return  $\arg \max_{a \in A(s)} Q(s, a)$                                 ▷ select action
    
```

Algorithm 4.1. Top-level Monte Carlo tree search algorithm.

```

function SIMULATE( $s, d, \pi_0$ )
    if  $d = 0$  return 0
    if  $s \notin \mathcal{T}$ 
         $\mathcal{T} \leftarrow \mathcal{T} \cup \{s\}$ 
         $(N(s, a), Q(s, a)) \leftarrow (N_0(s, a), Q_0(s, a))$  for all  $a \in A(s)$ 
    return ROLLOUT( $s, d, \pi_0$ )

     $a \leftarrow \arg \max_a Q(s, a) + c \sqrt{\frac{\log N(s)}{N(s, a)}}$                                 ▷ selection (UCT)
     $(s', r) \sim G(s, a)$                                                                 ▷ expansion
     $q \leftarrow r + \gamma \text{SIMULATE}(s', d - 1, \pi_0)$                                 ▷ simulation/rollout
     $N(s, a) \leftarrow N(s, a) + 1$ 
     $Q(s, a) \leftarrow Q(s, a) + \frac{q - Q(s, a)}{N(s, a)}$                                 ▷ backpropagation
    return  $q$ 
    
```

Algorithm 4.2. Monte Carlo tree search simulation.

```

function ROLLOUT( $s, d, \pi_0$ )
    if  $d = 0$  return 0
     $a \sim \pi_0(s)$ 
     $(s', r) \sim G(s, a)$ 
    return  $r + \gamma \text{ROLLOUT}(s', d - 1, \pi_0)$ 
    
```

Algorithm 4.3. Rollout evaluation.

5 Stochastic Optimization

Stochastic optimization methods use randomness for minimization or maximization.²³ They are often useful when dealing with black-box functions where no gradient information is available.

²³ M. J. Kochenderfer and T. A. Wheeler, *Algorithms for Optimization*. MIT Press, 2019.

5.1 Cross-Entropy Method

The *cross-entropy method* (CEM)²⁴ aims to minimize the cross-entropy between the unknown true distribution f and a proposal distribution g parameterized by θ . This technique reformulates the minimization problem as a probability estimation problem, and uses adaptive importance sampling to estimate the unknown expectation.

²⁴ R. Rubinstein, "The Cross-Entropy Method for Combinatorial and Continuous Optimization," *Methodology and Computing in Applied Probability*, vol. 1, no. 2, pp. 127–190, 1999.

```

function CROSSENTROPYMETHOD( $S, g, m, m_{\text{elite}}, k_{\text{max}}$ )
  for  $k \leftarrow 1$  to  $k_{\text{max}}$ 
     $\mathbf{X} \sim g(\cdot \mid \theta_k)$  where  $\mathbf{X} \in \mathbb{R}^{|g| \times m}$  ▷ draw  $m$  samples from  $g$ 
     $\mathbf{Y} \leftarrow S(\mathbf{x})$  for  $\mathbf{x} \in \mathbf{X}$  ▷ evaluate samples  $\mathbf{X}$  using objective  $S$ 
     $\mathbf{e} \leftarrow$  store top  $m_{\text{elite}}$  from  $\mathbf{Y}$  ▷ select elite samples output from objective
     $\theta_{k'} \leftarrow \text{FIT}(g(\cdot \mid \theta_k), \mathbf{e})$  ▷ re-fit distribution  $g$  using elite samples
  return  $g(\cdot \mid \theta_{k_{\text{max}}})$ 

```

Algorithm 5.1. The cross-entropy method.

References

1. K. Arulkumaran, M. P. Deisenroth, M. Brundage, and A. A. Bharath, "Deep Reinforcement Learning: A Brief Survey," *IEEE Signal Processing Magazine*, vol. 34, no. 6, pp. 26–38, 2017.
2. R. E. Bellman, *Dynamic Programming*. Princeton University Press, 1957.
3. D. P. Bertsekas, "A New Value Iteration Method for the Average Cost Dynamic Programming Problem," *SIAM Journal on Control and Optimization*, vol. 36, no. 2, pp. 742–759, 1998.
4. R. Coulom, "Efficient Selectivity and Backup Operators in Monte-Carlo Tree Search," in *International Conference on Computers and Games*, 2006.
5. M. J. Kochenderfer, *Decision Making Under Uncertainty: Theory and Application*. MIT Press, 2015.
6. M. J. Kochenderfer and T. A. Wheeler, *Algorithms for Optimization*. MIT Press, 2019.
7. R. Rubinstein, "The Cross-Entropy Method for Combinatorial and Continuous Optimization," *Methodology and Computing in Applied Probability*, vol. 1, no. 2, pp. 127–190, 1999.
8. G. A. Rummery and M. Niranjan, *On-Line Q-Learning Using Connectionist Systems*. University of Cambridge, 1994, vol. 37.
9. J. Schulman, S. Levine, P. Moritz, M. I. Jordan, and P. Abbeel, "Trust Region Policy Optimization," *CoRR*, vol. abs/1502.05477, 2015. arXiv: 1502.05477.
10. J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov, "Proximal Policy Optimization Algorithms," *CoRR*, vol. abs/1707.06347, 2017. arXiv: 1707.06347.
11. D. C. Sorensen, "Newton's Method with a Model Trust Region Modification," *SIAM Journal on Numerical Analysis*, vol. 19, no. 2, pp. 409–426, 1982.
12. R. S. Sutton, "Dyna, an Integrated Architecture for Learning, Planning, and Reacting," *SIGART Bulletin*, vol. 2, no. 4, pp. 160–163, 1991.
13. Z. Wang, V. Bapst, N. Heess, V. Mnih, R. Munos, K. Kavukcuoglu, and N. de Freitas, "Sample Efficient Actor-Critic with Experience Replay," *CoRR*, vol. abs/1611.01224, 2016. arXiv: 1611.01224.
14. C. J. C. H. Watkins, "Learning from Delayed Rewards," PhD thesis, University of Cambridge, 1989.