

* Policy Gradient Algorithms

→ Instead of maintaining estimates of Value Functions, Search in the space of policies.

+ Simpler Description, better convergence,

+ When you have continuous actions, hard to do a network, Value Functions are difficult to make it work so usually they are discretized. No need to do that for policy gradient methods.

+ Robust to partial observability: when you don't have access to state s but have access to $f(s)$. Policy Gradient methods are more robust because they don't use the notion of Value Functions and hence aren't tied to states.

→ Policy depends on some parameters θ . Idea is to modify policy parameters directly instead of estimating the action values.

$$\hookrightarrow \text{Max } J(\theta) = E[r_t] = \sum_a q_{\pi}(a) \pi_{\theta}(a)$$

$$\hookrightarrow \theta = \theta + \alpha \nabla_{\theta} J(\theta) \text{ (Gradient Ascent)}$$

For a particular policy π , $J(\theta)$ gives some kind of performance metric.

So, $J(\theta) = E[r_t] = \sum_a q_{\pi}(a) \pi_{\theta}(a)$ $\Rightarrow$ We can think of the Sum as the expected value of q_{π} over π_{θ} i.e.

$$J(\theta) = E[r_t] \quad \sum_a q_{\pi}(a) \pi(a) = E_{\pi_{\theta}}[q_{\pi}(a)]$$

Since q_{π} is unknown, fix π_{θ} and repeatedly sample

You have $(t_1, r_1, a_1), (t_2, r_2, a_2), \dots$ samples.

$$\text{Thus } E[r_t] \approx \frac{1}{N} \sum_i r_i = J(\theta) = E_{\pi_{\theta}}[q_{\pi}(a)]$$

$$\text{Now, } J(\theta) = \sum_a q_{\pi}(a) \pi_{\theta}(a) = \sum_a q_{\pi}(a) \frac{\pi_{\theta}(a)}{\pi_{\theta}(a)} = E_{\pi_{\theta}} \left[q_{\pi} \frac{\pi_{\theta}(a)}{\pi_{\theta}(a)} \right]$$

$$\approx \sum_i r_i \frac{\pi_{\theta}(a_i)}{\pi_{\theta}(a_i)} =$$

* REINFORCE

Learning Rules should not depend on the reward

$$\Delta \theta_t = \alpha r_t \frac{\nabla \pi_{\theta}(a_t)}{\pi_{\theta}(a_t)} = \alpha r_t \frac{\partial \ln \pi_{\theta}(a_t)}{\partial \theta}$$

(with baseline): $\Delta \theta_t = \alpha (r_t - b_t) \frac{\partial \ln \pi_{\theta}(a_t)}{\partial \theta}$

↳ baseline should not depend on action

* REINFORCE: MC POLICY GRADIENT

$$\nabla J(\theta) \propto \sum_s \mu(s) \sum_a q_{\pi}(s, a) \nabla \pi(a|s, \theta)$$

Fraction of time
Agent will spend
in state s while
executing policy π

REINFORCE is essentially an
instance of the Policy gradient
theorem applied to the full MDP

$$= E_{\pi} \left[G_t \frac{\nabla \pi(A_t | S_t, \theta)}{\pi(A_t | S_t, \theta)} \right]$$

↳ mc because full return is used

$$\theta_{t+1} = \theta_t + \alpha G_t \frac{\nabla \pi(A_t | S_t, \theta)}{\pi(A_t | S_t, \theta)} \quad \nabla \ln(\pi(A_t | S_t, \theta))$$

Algo:

↳ Input: Differentiable policy $\pi(a|s, \theta)$, $\alpha > 0$, $\theta \in \mathbb{R}^d$ (0 eg)

↳ For each episode, generate Trajectory,

↳ For each step of episode calculate disc return

$$G_t = \sum_{k=t}^T \gamma^{k-t} R_k$$

$$\theta = \theta + \alpha \gamma^t G_t \nabla \ln \pi(A_t | S_t, \theta)$$

(no bootstrapping)

- Unbiased estimate of ^{Return} ~~of~~ because you're not using any q -values.

- slow due to high variance in estimates

↳ Related to "recurrence time" of episode

- For large state spaces, variance Unacceptably high.

* REINFORCE: MC with Baseline

$$\nabla J(\theta) \propto \sum p(s) \sum (q_{\pi}(s,a) - b(s)) \nabla \pi(a|s, \theta)$$

cannot be a function of the action because

$$\sum b(s) \nabla \pi(a|s, \theta) = b(s) \sum \nabla \pi(a|s, \theta) = b(s) \nabla \sum \pi(a|s, \theta) = b(s) \nabla 1 = 0$$

That means gradients are not being affected.

$$\text{So } \theta_{t+1} = \theta_t + \alpha (G_t - b(s_t)) \frac{\nabla \pi(A_t|s_t, \theta_t)}{\pi(A_t|s_t, \theta_t)}$$

+ Doesn't change expected value, but improves Variance.

* Actor-Critic Methods:

→ Learn both a policy & state-value function

actor
Param: θ

→ Critic (evaluates state-actions)
Param: w

$$b(s) = E_{\pi_{\theta}} [G_t | s_t = s] = V_{\pi_{\theta}}(s)$$

$$\text{So, } \theta_{t+1} = \theta_t + \alpha \left(\underbrace{G_t - \hat{V}(s_t, w)}_{\text{TD error}} \right) \frac{\nabla \pi(A_t|s_t, \theta_t)}{\pi(A_t|s_t, \theta_t)}$$

$$= \theta_t + \alpha \left(\underbrace{R_{t+1} + \gamma \hat{V}(s_{t+1}, w) - \hat{V}(s_t, w)}_{\delta_t} \right) \frac{\nabla \pi(A_t|s_t, \theta_t)}{\pi(A_t|s_t, \theta_t)}$$

$$= \theta_t + \alpha \delta_t \frac{\nabla \pi(A_t|s_t, \theta_t)}{\pi(A_t|s_t, \theta_t)}$$

Input: differentiable $\hat{V}(s, w)$, $\pi(a|s, \theta)$, $\alpha^{\theta}, \alpha^w \in [0, 1]$

For each episode, initialize s , $i = 1$

For each step: $A \sim \pi(\cdot|s, \theta)$ - Take A , observe s', R

$$\delta = R + \gamma \hat{V}(s', w) - \hat{V}(s, w)$$

$$w = w + \alpha^w \delta \nabla \hat{V}(s, w), \quad \theta = \theta + \alpha^{\theta} \delta \nabla \pi(A|s, \theta)$$

$$i = i + 1, s = s'$$

→ Gre Although unbiased causes high variance. Now by using one-step bootstrapped estimate, $\uparrow$ bias $\downarrow$ variance, accelerate learning.

→ use $\hat{V}$ as baseline as estimate for q_{π} → stabilizes learning
→ $\alpha^o < \alpha^w$ so critic can converge to actor. otherwise target keeps moving

* Basic Actor Critic Algorithm

- | | |
|--------|--|
| | 1- Take $a \sim \pi_{\theta}(a s)$. Receive S, A, S', R |
| Critic | 2- Update w using $(G, R + \gamma \hat{V}_w(s'))$ (Batch mode, online, Replay Buffer) |
| | 3- $\hat{S}(s, a) = R + \gamma \hat{V}_w(s') - \hat{V}_w(s)$ |
| Actor | 4- $\theta \leftarrow \theta + \alpha \hat{S}(s, a) \cdot \nabla_{\theta} \ln \pi_{\theta}(a s)$ |

→ Step 2 & 4 are not always distinct. A shared network could update both.

Advantage
Actor

Critic.

Note that the Advantage Function $A_{\pi}(s, a) = q_{\pi}(s, a) - V_{\pi}(s) = \delta_{\pi}(s, a)$ for that part of TD-Error. Not expectation over all actions.

* A3C - Async Advantage Actor Critic

→ In order to learn π_{θ} , we want generate a lot of samples. Instead of doing it in sequence we can generate a lot of samples in parallel.

→ Varying episode lengths can cause the parameter updates to be inconsistent. Still works if the change is small enough.

→ Each agent is independent and hence sequence-actors are independent and thus the correlated sequence problem does not exist. as each agent will be using different sets of experiences.

→ A3C does not need to freeze targets either as it does not bootstrap. Uses n-step return until the end. If episodes are nonterminal then you are bootstrapping much later.

* A2C - Synchronous Advantage Actor Critic

→ Remove Async updates. Train agents independently. update only after all threads finish.

→ True gradients are always computed & applied.

* Compatible Parametrization

→ Substituting $\hat{q}(s,a,w)$ for $q(s,a)$ may introduce bias. It can be proved the approximator has no bias if the following two conditions are satisfied.

$\hat{q}$ has to be linear in characteristic eligibility.

$$1. \hat{q}(s,a,w) = \nabla_{\theta} \log \pi_{\theta}(a|s)^T w$$

↳ w s.t. when we project characteristic elig along w , we get smth along dir of $\hat{q}$

2. w minimizes MSE

$$w = \operatorname{argmin}_{w} E_{\pi_{\theta_0}, a \sim \pi_{\theta_0}} [\hat{q}(s,a,w) - q_{\pi_{\theta_0}}(s,a)]^2$$

* (Deep) Deterministic Policy Gradient

→ DPG is typically used with continuous actions. If policy is deterministic then gradients become 0 & cannot really use policy gradient.

→ Expectation over state & actions complicates gradient estimation

For deterministic policies, the performance objective is

$$J(\theta) = \int_{\mathcal{S}} \mu(s) q_{\pi_{\theta}}(s, \pi_{\theta}(s)) ds =$$

$$\rightarrow \text{In a DPG, } \nabla_{\theta} J(\theta) = \int \mu(s) \nabla_a q_{\pi}(s,a) \cdot \nabla_{\theta} \pi_{\theta}(s) |_{a=\pi_{\theta}(s)} ds$$

→ we also avoid taking a max action by moving in the direction of q_{π} . Value of the new policy action is guaranteed to be higher than previous action.

3 → 1 → 4 → 2

___/___/___

→ For a wide class of problems DPG is a limiting case of the SPG,
UPDATE EQ:

$$\delta_t = r_t + \gamma \hat{q}(s_{t+1}, a_{t+1}, \omega) - \hat{q}(s_t, a_t, \omega)$$

$$\omega = \omega + \alpha_\omega \delta_t \nabla_\omega \hat{q}(s_t, a_t, \omega)$$

$$\theta = \theta + \alpha_\theta \nabla_\theta \hat{q}(s_t, a_t, \omega) \nabla_{\theta_0} \pi_\theta(s_t) |_{a \sim \pi_{\theta_0}(s)}$$

→ To ensure exploration,

(i) environment should be stochastic (or)

ii) use off-policy actor-critic where behaviour policy differs from estimation policy

* Hierarchical RL

Instead of thinking about actions repeating over & over in the State Space, think of subproblems to solve.

* Temporal Abstraction

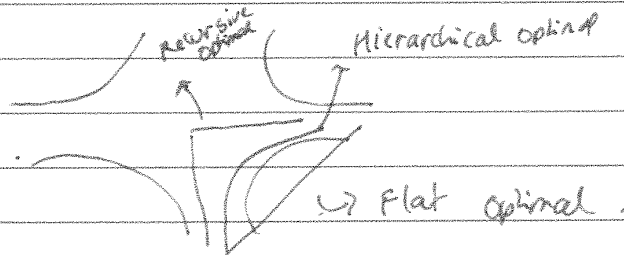
+ Ability transfer subproblems to other bigger problems.

+ More Powerful / Meaningful State Abstraction

↳ Reduced statespace, Greater usability

* Types of Optimality

→ Flat optimality < Hierarchical optimality < Recursive optimality
(Lowest cost)



* Semi-MDP

→ Previously $S_t \rightarrow A_t \rightarrow R_{t+1} \rightarrow S_{t+1} \dots$

For Semi MDP's,

$$S_t \rightarrow A_t \dots R_{t+T} \rightarrow S_{t+T} \rightarrow A_{t+T} \dots S_{t+T+T'}$$

→ Actions take different amount of time to complete. So optimality depends not only on where the action takes but also on How long it takes.

$$\langle S, A, P, R \rangle$$

$$P(S', T | S, A), \quad E[R | S, A, S', T]$$

* SMDP Q-Learning

$$Q(s,a) = Q(s,a) + \alpha [r_{\text{env}} + \gamma \max_{a'} Q(s',a') - Q(s,a)]$$

* Options $\Theta = \langle I_0, \pi_0, \beta_0 \rangle$

- Macro Actions

- Encapsulates Policy

I_0 : Initiation set

π_0 : Policy (Between start $\rightarrow$ end state Policy)

$\beta_0: S \rightarrow [0,1]$

$\rightarrow$ Markov options: π_0 depends on the current state

$\rightarrow$ Semi-Markov Options: π_0 depends on history since options started

MDP + OPTIONS = SMDP

* Intra Option Q-learning

π_0 :

s_2	s_2	s_n
a_2	$a_2 \dots$	a_n
r_2	r_2	r_n

$$Q(s_2, a_2) = Q(s_2, a_2) + \alpha [r_2 + \gamma \max_a Q(s_2, a) - Q(s_2, a_2)]$$

$\rightarrow$ Cannot use SARSA here because never had the choice to choose action in s_2 . Following π_0 . Hence you have to be off-policy and use Q-learning.