

Function Approximate:

$$Go \rightarrow |S| = 10^{170}$$

Helicopter

$\rightarrow$ state is cts.

Drone

Instead of represent $V(s)$, $Q(s, a)$ in a look up table /
Tabular case

we use a function approximation.

$\hat{V}(s; \vec{w})$ the value fn is determined by parameter w .

$$\hat{q}(s, a; \vec{w})$$

- Even if the data doesn't see the state, using fn approx will generalise to unseen state.

1).
$$\begin{array}{c} S \rightarrow \boxed{w} \rightarrow V(s; w) \\ \uparrow \text{input} \quad \downarrow \\ \text{is determined by } w \end{array}$$

$$\begin{array}{c} S \rightarrow \boxed{w} \rightarrow Q(s, a; w) \\ a \rightarrow \end{array}$$

$$S \rightarrow \boxed{\quad} \begin{array}{l} \rightarrow Q(s, a_1; w) \\ \rightarrow Q(s, a_2; w) \\ \vdots \\ \rightarrow Q(s, a_{|A|}; w) \end{array}$$

Type of approx:

- linear combination of feature. $\leftarrow$
 - NN (nonlinear fun approx) $\leftarrow$
 - Fourier / wavelet bases. $\leftarrow$
- $\leftarrow$ linear approx

Q: How to determine the parameter w ?

1. Recall SGD:

$$\min_w J(w)$$

$$w_{k+1} = w_k - \alpha \nabla_w J(w_k)$$

$$\min_w \mathbb{E}_{\pi \sim p(\pi)} J(w, \pi)$$

$$w_{k+1} = w_k - \alpha \nabla_w J(w_k, \pi_k) \leftarrow \text{stochastic GD}.$$

2. Imagine in policy prediction case, we know $V^\pi(s)$

$$\min_w J(w) = \mathbb{E}_s \left[\frac{1}{2} (V^\pi(s) - \hat{V}(s; w))^2 \right]$$

$$\text{GD: } w_{k+1} = w_k + \alpha_k \mathbb{E} \left[(V^\pi(s) - \hat{V}(s; w_k)) \nabla_w \hat{V}(s; w_k) \right]$$

$$\text{SGD: } w_{k+1} = w_k + \alpha_k (V^\pi(s_t) - \hat{V}(s_t; w_k)) \nabla_w \hat{V}(s_t; w_k)$$

e.g. $\hat{V}(s; \omega)$ is represented in $\{ \chi_j(s) \}_{j=1}^n$

$$\hat{V}(s; \omega) = \sum_{j=1}^n w_j \chi_j(s) = (\chi_1(s), \dots, \chi_n(s)) \begin{pmatrix} w_1 \\ \vdots \\ w_n \end{pmatrix} = (\pi(s))^T \omega \quad \downarrow \in \mathbb{R}^n$$

$$\text{GD: } w_{k+1} = w_k + \alpha_k \mathbb{E} \left[(V^\pi(s) - \chi(s)^T w_k) \chi(s) \right]$$

$$\text{SGD: } w_{k+1} = w_k + \alpha_k (V^\pi(s_t) - \chi(s_t)^T w_k) \chi(s_t)$$

e.g. Tabular is a special case of linear approx:

$$\text{with } \chi_i(s) = \begin{pmatrix} 0 \\ \vdots \\ 0 \\ 1 \\ \vdots \\ 0 \end{pmatrix} \leftarrow i\text{-component}$$

$$\hat{V}(s; \omega) = \sum_{i=1}^{|S|} \chi_i(s) w_i$$

3. However, in RL, we don't really know the true $V^\pi(s)$.

1). We can substitute $V^\pi(s)$ by an unbiased estimate $G_t^{(\infty)}$

$$G_t^{(\infty)} = r_t + \gamma r_{t+1} + \dots + \gamma^T r_{t+T}$$

$$\text{MC learning: } w_{k+1} = w_k + \alpha_k \left(\underbrace{G_t^{(\infty)} - \hat{V}(s_t, w_k)}_{\downarrow \mathbb{E}[G_t^{(\infty)}] = V^\pi(s_t)} \right) \pi_w \hat{V}(s_t, w_k)$$

$\mathbb{E}[G_t^{(\infty)}] = V^\pi(s_t)$
 an unbiased estimate for the true $V^\pi(s)$.

2). TD learning: substitute $G_t^{(1)} = r_t + \gamma \hat{V}(S_{t+1}; w_k)$.

$$w_{k+1} = w_k + \alpha_k \left(\underline{G_t^{(1)}} - \underline{\hat{V}(S_t; w_k)} \right) \underline{\nabla_w \hat{V}(S_t, w_k)}$$

* Note: here $G_t^{(1)}$ is not necessarily an unbiased estimate for $V(S)$.

3). TD(λ): $w_{k+1} = w_k + \alpha_k \left(G_t^\lambda - \hat{V}(S_t; w_k) \right) \nabla_w \hat{V}(S_t, w_k)$

$$G_t^\lambda = (1-\lambda) \sum_{i=0}^{\infty} \lambda^i G_t^{(i)}$$

e.g. linear approx:

$$w_{k+1} = w_k + \alpha_k \left(r_t + \gamma x(S_{t+1})^T w_k - x(S_t)^T w_k \right) x(S_t)$$

$$w_{k+1} = w_k + \alpha_k \underbrace{d_t}_{\text{TD: } r_t + \gamma \hat{V}(S_{t+1}; w_k) - \hat{V}(S_t; w_k)} x(S_t)$$

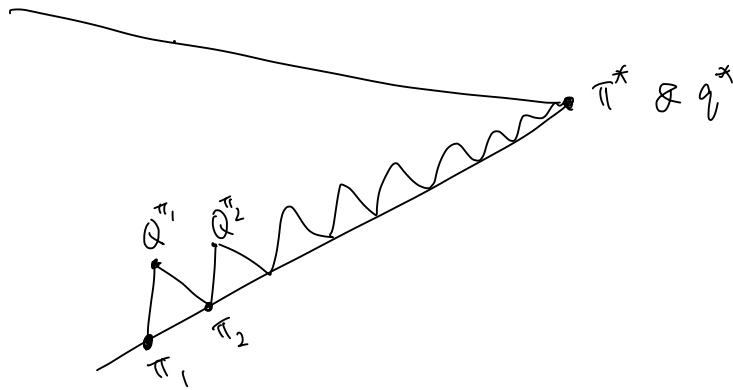
"Convergence Thm"

MC $\rightarrow$ converges to a local mm for both linear & nonlinear approx

TD $\rightarrow$ converge (close to) a local mm
 $\uparrow \gamma$
 for linear approx.

TD(λ) $\rightarrow$ similar as TD.

function approximation for control:



Algo: update $q(s, a; w)$ with one step update + ϵ -greedy policy update

(SARSA in fcn Approx).

Initialize s, a, w

$\rightarrow$ observe s', r
 $a' \sim \epsilon$ -greedy π from $\hat{q}(s, a; w)$.
 $w \leftarrow w + \alpha (r + \gamma \hat{q}(s', a'; w) - \hat{q}(s, a; w)) \nabla_w \hat{q}(s, a; w)$

⊗ Gradient TD.

- Prediction:

$$V(s) = r(s) + \gamma \mathbb{E}[V(s_{t+1}) | s_t = s]$$

$$\pi V(s) = r(s) + \gamma \mathbb{E}[V(s_{t+1}) | s_t = s]$$

Δ
Bellman

$$\pi V = r + \gamma P^\pi V$$

$\hat{\mathcal{T}}$ has contraction property.

$V^\pi(s) = \pi V^\pi(s) \leftarrow$ view $V^\pi(s)$ as the fixed point of the operator.

- Instead of view the solution $V^*(s)$ as the fixed pt of T , we can view it as the minimizer of $\int (V(s) - TV(s))^2 ds$.
- When $V_\theta(s)$ is parametrized by $\theta \in \mathbb{R}^d$

$$\min_{\theta} \int_S \left(r(s) - \gamma \mathbb{E}[V_\theta(s_{t+1}) | s_t = s] - V_\theta(s) \right)^2 ds$$

- One can solve it by GD.

$$\theta_{k+1} = \theta_k - \alpha_k \underbrace{\mathbb{E}_S \left[\left(r(s) - \gamma \mathbb{E}[V_\theta(s_{t+1}) | s_t = s] - V_\theta(s) \right) \left(-\gamma \mathbb{E}[\nabla V_\theta(s_{t+1}) | s_t = s] - \nabla V_\theta(s) \right) \right]}_{(x)}$$

- SGD:

unbiased estimate for (x)

$$\begin{aligned} & \left(r(s_t) - \gamma \mathbb{E}[V_\theta(s_{t+1}) | s_t] - V_\theta(s_t) \right) \left(-\gamma \mathbb{E}[\nabla V_\theta(s_{t+1}) | s_t] - \nabla V_\theta(s_t) \right) \\ &= \left[r(s_t) - V_\theta(s_t) \right] \left(-\gamma \mathbb{E}[\nabla V_\theta(s_{t+1}) | s_t] \right) - \gamma \mathbb{E}[V_\theta(s_{t+1}) | s_t] \left(-\nabla V_\theta(s_t) \right) \\ & \quad + \gamma \mathbb{E}[V_\theta(s_{t+1}) | s_t] \mathbb{E}[\nabla V_\theta(s_{t+1}) | s_t] \end{aligned}$$

$$\begin{aligned} & \overset{\text{unbiased}}{\mathbb{E}} \left(r(s_t) - V_\theta(s_t) \right) \left(-\gamma \nabla V_\theta(s_{t+1}) \right) - \gamma V_\theta(s_{t+1}) \left(-\nabla V_\theta(s_t) \right) \\ & \quad + \gamma^2 V_\theta(s_{t+1}) \nabla V_\theta(s'_{t+1}) \end{aligned}$$

where s_{t+1} & s'_{t+1} are two independent sample starting from s_t .

① when the underlying dynamics is deterministic.

$$S_{t+1} = S'_{t+1}$$

② when the underlying dynamics is stochastic, when we only given one trajectory data,

$\{S_t\}_{t=0}^T$, then it is hard to find two samples from S_t

"double sampling problem".

One way to avoid "DS"

$$\min_{\theta} \mathbb{E}_s \left(r(s) - \gamma \mathbb{E}[V_{\theta}(S_{t+1}) | S_t = s] - V_{\theta}(s) \right)^2$$

$$\min_{\theta} f^2(\theta)$$

$$\min_{\theta} \max_y f(\theta) y - \frac{1}{2} y^2$$

$$\min_{\theta} \max_y \mathbb{E}_s \left(r(s) - \gamma \mathbb{E}[V_{\theta}(S_{t+1}) | S_t = s] - V_{\theta}(s) \right) y - \frac{1}{2} y^2$$

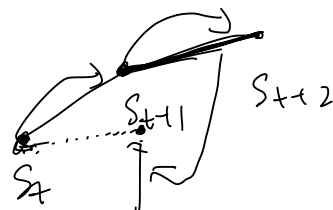
given fixed y :

$$\theta_{k+1} = \theta_k - \partial_k \left[\left(-\gamma \mathbb{E}[\nabla V_{\theta}(S_{t+1}) | S_t = s] - \nabla_{\theta} V_{\theta}(s) \right) y - \frac{1}{2} y^2 \right]$$

$$\begin{cases} \theta_{k+1} = \theta_k - \alpha_k \left[(-\gamma \nabla V_\theta(s_{t+1}) - \nabla_\theta V_\theta(s_t)) / y - \frac{1}{2} y^2 \right] \\ y_{k+1} = y_k + \beta_k (r(s) - \gamma V_\theta(s_{t+1}) - V_\theta(s_t)) y_k - \frac{1}{2} y_k^2 \end{cases}$$

Another way: Borrowing from the future.

$$s_t, s_{t+1}, s_{t+2}, s_{t+3}, \dots$$



$$\hat{s}_{t+1}^* = s_t + (s_{t+2} - s_{t+1})$$

↓
This could be a good approximation when the transition dynamics is smooth.

$$ds_t = \mu(s_t, a_t) dt + \sigma(s_t, a_t) dB_t$$

$$\|\nabla \mu\|, \|\nabla \sigma\| \leq L$$

$$\theta_{k+1} = \theta_k - \alpha_k \left(r(s_t) - \gamma V_\theta(s_{t+1}) - V_\theta(s_t) \right) \left(-\gamma \nabla V_\theta(s_t + s_{t+2} - s_{t+1}) - \nabla V_\theta(s_t) \right)$$

— control:

$$\min_{\theta} \mathbb{E}_{s,a} \left(r(s,a) - \gamma \max_a \mathbb{E} [Q_\theta(s_{t+1}, a) | s_t = s, a_t = a] - Q_\theta(s, a) \right)^2$$

(convergence for fun approx.

(prediction)

Tabular

(linear

non-linear

MC

✓

✓

✓

TD(0)

✓

✓

X

TD(λ)

✓

✓

X

∇ TD

✓

✓

✓

(control)

MC

✓

✓

X

SARSA

✓

✓

X

Q-learning

✓

X

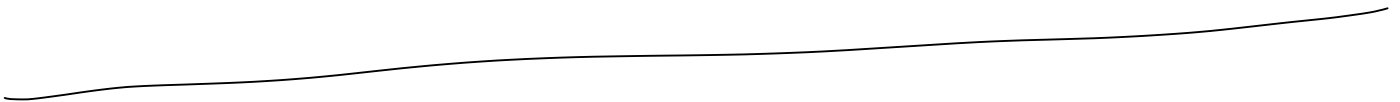
X

∇ TD

✓

✓

X



- Disadvantage of Value-based method.

$\max_a Q_{\theta}(s, a)$ $\rightarrow$ could be hard to solve
when $|A| \gg 1$ or its action
space

- Policy ∇ :

Value-based Algo

- learn Value fn

- Implicit policy based on value fn

Policy-based

- No value fn

- learn policy.