

Textbooks:- ① CS6700 in my website

② Tsitsiklis & Bertsekas "Neuro-dynamic Programming", 1996

③ Sutton & Barto, "Intro. to RL"

④ Bertsekas "DP & OC, Vol I & II".

Example: Machine replacement/repair

States = $\{1, 2, \dots, n\}$

↓

good

↓

worst

Actions = Do nothing, Repair

On repair: machine goes to "1". Cost = $\mathbb{R}$

Operating cost: $g(r)$

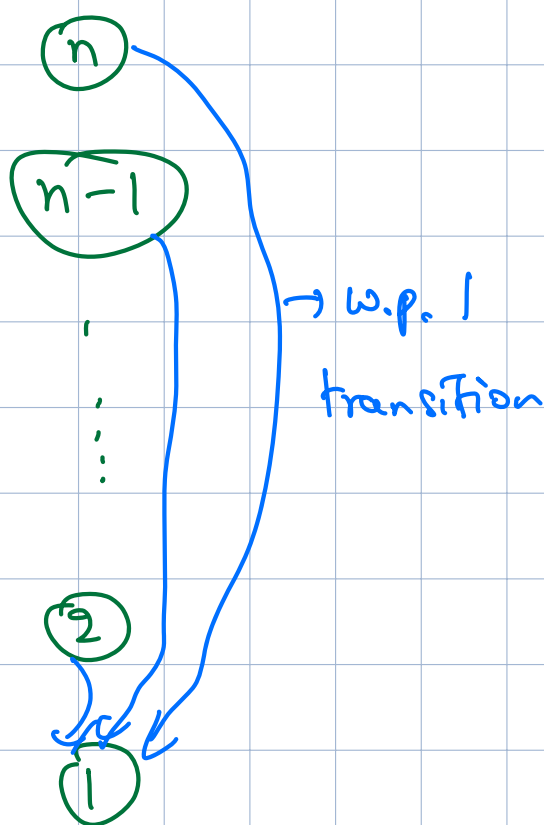
$$g(1) \subseteq g(2) \subseteq g(3) \subseteq \dots \subseteq g(n)$$

State transitions:

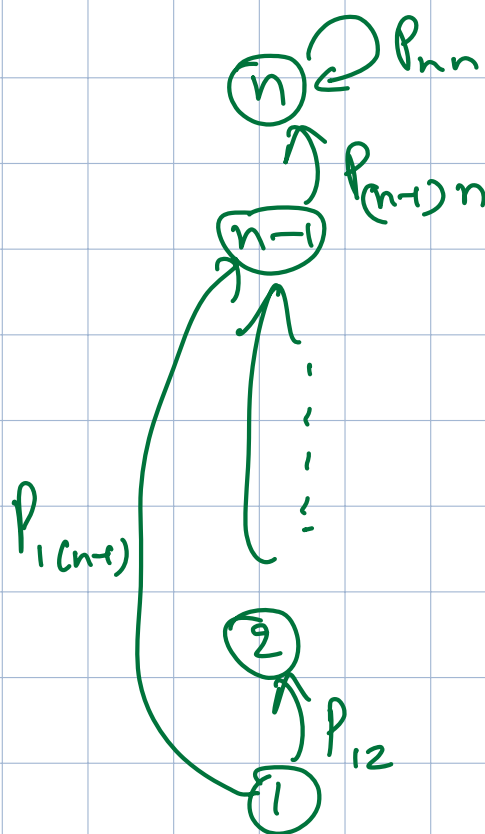
P_{ij} : probability of state transitioning from i to j

$$P_{ij} = 0 \quad \text{if } j < i$$

Action: repair σ



Action: Do nothing



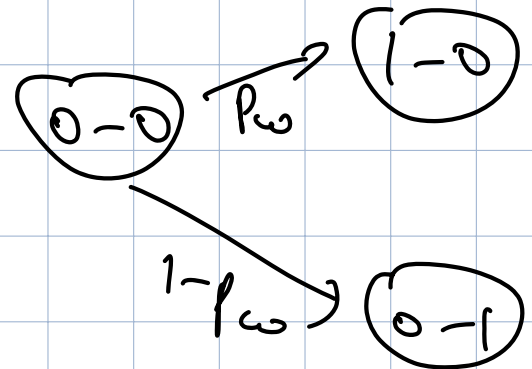
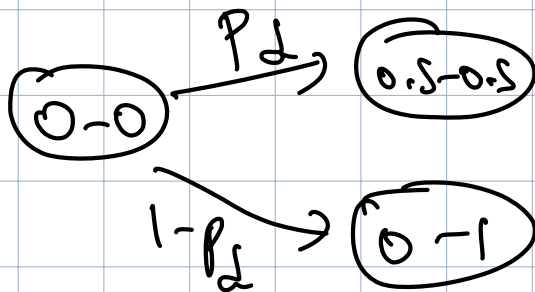
Example 2: Chess match

Playing styles (actions): Timid, Bold

Timid	draw P_d	lose $1 - P_d$
Bold	win P_w	lose $1 - P_w$

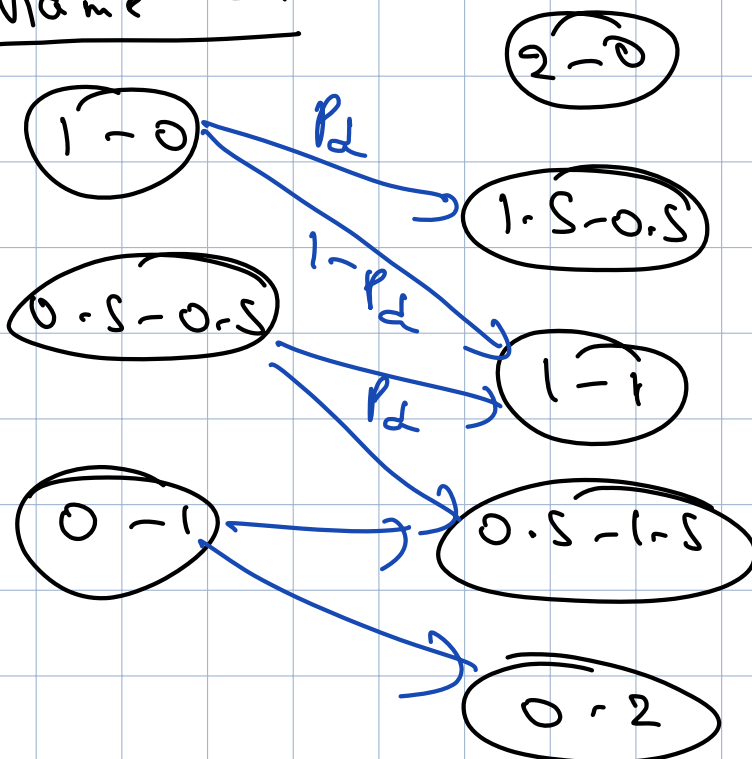
$$P_d > P_w$$

Game 1

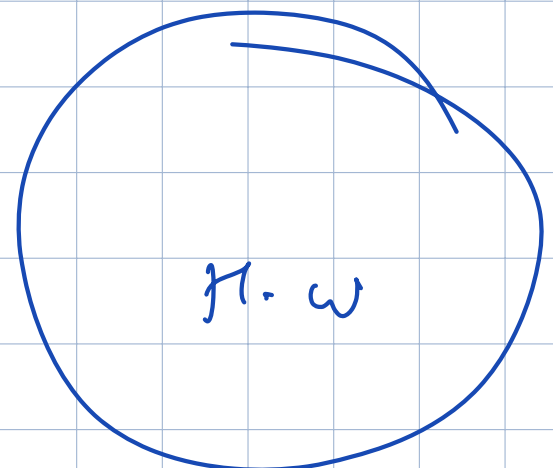


Timid

Game 2:



Timid play

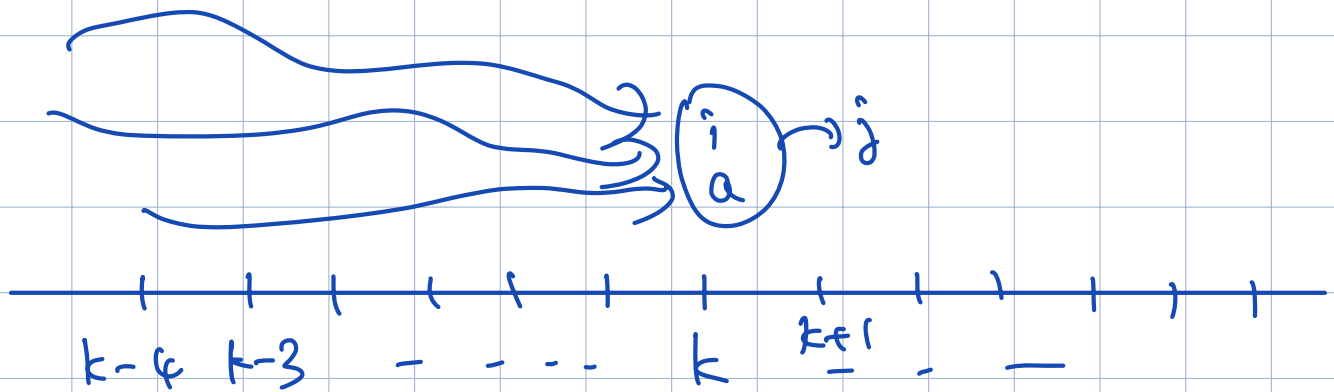


Bold play

Framework

S : State space A : action space

$$P_{ij}(a) = P(S_{k+1} = j \mid S_k = i, a_k = a)$$



Single-stage ~~stage~~ cost! $g_k(i, a, j)$
"slot"

Policy π : $\{\mu_0, \dots, \mu_{N-1}\}$

$N \leftarrow \# \text{ of slots}, \mu_i: S \rightarrow A$
 $\downarrow$
State space Action space

Performance metric

$$J_\pi(s_0) = E \left[\underbrace{g_N(s_N)}_{\text{Terminal cost}} + \sum_{k=0}^{N-1} g_k(s_k, \mu_k(s_k), S_{k+1}) \right]$$

Goal: Find π^* s.t.

$$J_{\pi^*}(s_0) = \min_{\pi} J_{\pi}(s_0)$$

Optimality principle:

$$\pi^* = (\mu_0^*, \mu_1^*, \dots, \mu_{N-1}^*)$$

A tail sub-problem

$$\min_{\pi^i = (\mu_i, \dots, \mu_{N-1})} E \left(g_N(s_N) + \sum_{k=i}^{N-1} g_k(s_k, \mu_k(s_k), s_{k+1}) \right)$$

Claim: For this tail sub-problem

$\{\mu_i^*, \dots, \mu_{N-1}^*\}$ is optimal

DP algorithm :

Idea: Solve tail sub-problems & progress backwards

Algorithm :

Set $J_N(s_N) = g_N(s_N) \quad \forall s_N \in S.$

For $k = N-1, \dots, 0$

{

$$J_k(s_k) = \min_{a_k \in A} E_{s_{k+1}} \left[g_k(s_k, a_k, s_{k+1}) + J_{k+1}(s_{k+1}) \right],$$

$\forall s_k \in S.$

}

Apply DP algo to "Machine Example".

States: $\{1 \dots n\}$

P.T.O.

Suppose terminal cost is zero

$$J_N(i) = 0$$

$$J_k(i) = \min \left\{ \underbrace{R + g(i) + J_{k+1}(i)}_{\text{Repair}}, \right.$$

$$\underbrace{g(i) + \sum_{j=i}^n p_{ij} J_{k+1}(j)}_{\text{Do nothing}} \}$$

Chess-example (re-visited)

State = net-score (#wins - #losses)

$$J_k(s_k) = \max \left(\underbrace{p_d J_{k+1}(s_k) + (1-p_d) J_{k+1}(s_k-1)}_{\text{Timid}}, \right.$$

$$\underbrace{p_w J_{k+1}(s_k+1) + (1-p_w) J_{k+1}(s_k-1)}_{\text{Bold}} \}$$

Play bold if $\frac{p_w}{p_d} \geq \frac{J_{k+1}(S_k) - J_{k+1}(S_k^{-1})}{J_{k+1}(S_{k+1}) - J_{k+1}(S_k^{-1})}$

Initial condition:

$$J_N(S_N) = \begin{cases} 1 & \text{if } S_N > 0 \\ p_w & \text{if } S_N = 0 \\ 0 & \text{if } S_N < 0 \end{cases}$$

S_{N-1}	$J_{N-1}(S_{N-1})$	Best play
> 1	1	does not matter
1	$\max(\underbrace{p_d + (1-p_d)p_w}_{\text{timid}}, \underbrace{p_w + (1-p_w)p_w}_{\text{bold}})$	
		Timid play is best
0	p_w	Bold play is best
-1	p_w^2	———,———

< -1 0 Leave it

For the 2-game match

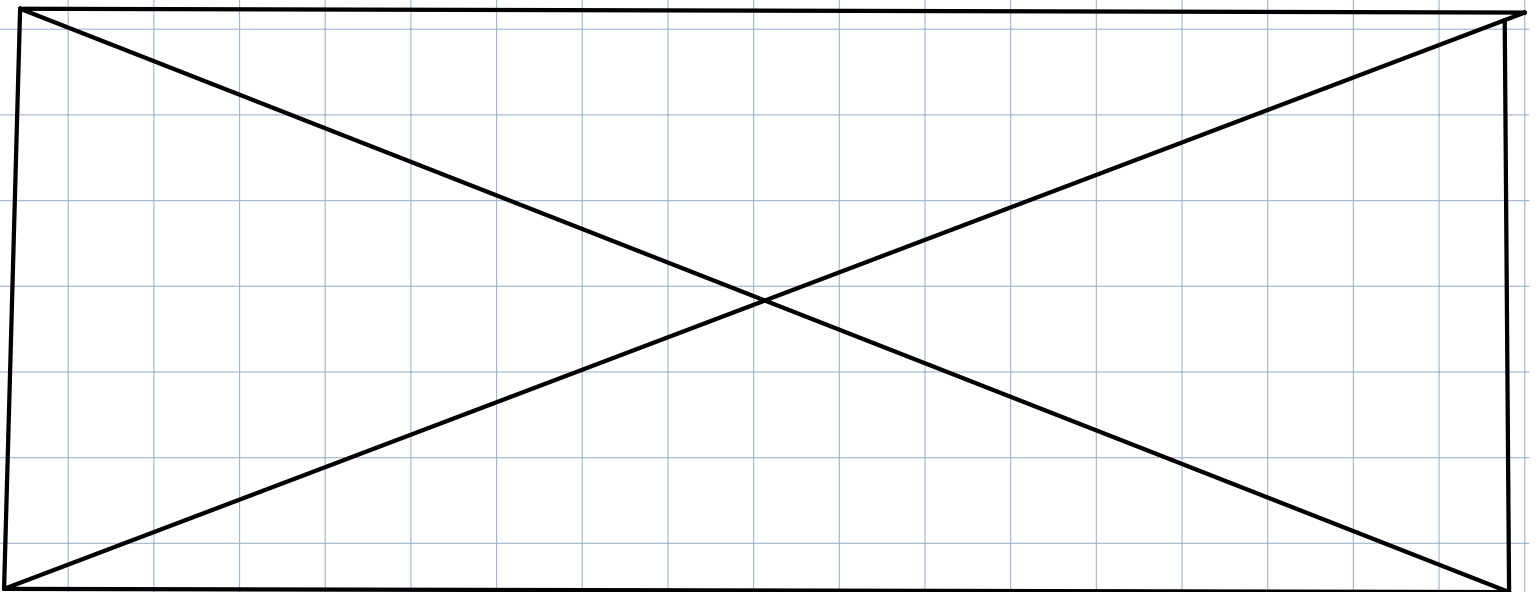
$$J_{N-2}(0) = \max \left(\underbrace{P_d P_w + (1 - P_d) P_w^2}_{\text{Timid}}, \right.$$

$$\left. \underbrace{P_w (P_d + (1 - P_d) P_w) + (1 - P_w) P_w^2}_{\text{Bold}} \right)$$

Check

$\Rightarrow$

$$P_w (P_d + (1 - P_d) P_w) + (1 - P_w) P_w^2$$



Lecture - 21

Infinite horizon discounted cost MDPs:

States: $\{1, \dots, n\}$

Stationary deterministic policy (SDP)

$\pi = \{\mu_0, \mu_1, \dots\}$ $\mu_i: S \rightarrow A$

Stationary policy $\pi = \{\mu, \mu, \dots\}$

$$J_\pi(i) = E \left(\sum_{k=0}^{\infty} \alpha^k g(S_k, \pi(S_k), S_{k+1}) \mid S=i \right)$$

Annotations:

- α : discount factor $0 < \alpha < 1$
- $g(S_k, \pi(S_k), S_{k+1})$: single-slot cost
- $\pi(S_k)$: policy
- $S=i$: start state

Goal: $J^*(i) = \min_{\pi} J_\pi(i), \forall i \in S.$

Let π^* be an optimal policy
corresponding to J^*

states

$$J^* = [J^*(1), \dots, J^*(n)]$$

Bellman operator

$$J = (J(1), \dots, J(n))$$

$$TJ = (TJ(1), \dots, TJ(n)), \text{ where}$$

$$(TJ)(i) = \min_{a \in A} \sum_{j=1}^n P_{ij}(a) [g(i, a, j) + \alpha J(j)], \quad \forall i \in S$$

$$(T_{\pi} J)(i) = \sum_{j=1}^n P_{ij}(\pi(i)) (g(i, \pi(i), j) + \alpha J(j)), \quad \forall i \in S$$

$$\text{Let } P_{\pi} = \begin{bmatrix} P_{11}(\pi(1)) & \dots & \dots \\ \vdots & \ddots & \vdots \\ \dots & \dots & P_{nn}(\pi(n)) \end{bmatrix}$$

$$= \left[\left[P_{i\hat{j}}(\pi(i)) \right]_{i, \hat{j}=1 \dots n} \right]$$

$$g_{\pi} = \begin{bmatrix} \sum_{\hat{j}=1}^n P_{1\hat{j}}(\pi(1)) g(1, \pi(1), \hat{j}) \\ \vdots \\ \sum_{\hat{j}=1}^n P_{n\hat{j}}(\pi(n)) g(n, \pi(n), \hat{j}) \end{bmatrix}$$

$$\boxed{T_{\pi} J = g_{\pi} + \alpha P_{\pi} J}$$

Properties of T, T_{π} :

For any J, J' s.t. $J(i) \leq J'(i) \quad \forall i$

- ① $T^k J(i) \leq T^k J'(i), \forall i$
 - ② $T_{\pi}^k J(i) \leq T_{\pi}^k J'(i), \forall i$
- } $\forall k \geq 1$

Proposition:- (Basis for Value Iteration)

For any (bounded) $T: S \rightarrow \mathbb{R}$

$$J^*(i) = \lim_{N \rightarrow \infty} (T^N J)(i), \quad \forall i \in S$$

Remark:-

$$\begin{array}{c} J_0 \xrightarrow{T} J_1 = T J_0 \xrightarrow{T} J_2 = T^2 J_0 \\ \qquad \qquad \qquad \downarrow \\ \qquad \qquad \qquad \vdots \quad \text{as } N \rightarrow \infty \\ \qquad \qquad \qquad T^N J_0 \rightarrow J^* \end{array}$$

(Corollary:-

$$J_\pi(i) = \lim_{N \rightarrow \infty} (T_\pi^N J)(i), \quad \forall i \in S$$

for any bounded J .

Value iteration for solving a discounted MDP:

Fix J_0

Repeatedly apply T

$$J_0 \xrightarrow{T} J_1 \rightarrow \dots \rightarrow J^*$$

Bellman optimality equation

The optimal cost $J^*(i) = \min_{\pi} J_{\pi}(i), \forall i$

Satisfies

$$J^*(i) = T J^*(i)$$

Fixed-point equation

$$J^*(i) = \min_a \sum_{j=1}^n P_{ij}(a) (g(i, a, j) + \gamma J^*(j))$$

Claim:- J^* is the unique fixed point

Pf: Suppose not,

$$\exists J', \quad J' = T J'$$

$$\begin{aligned} J' &= T J' = T^2 J' = \dots = \lim_{N \rightarrow \infty} T^N J' \\ &= J^* \end{aligned}$$

Corollary: For any stationary π ,

$$J_\pi = T_\pi J_\pi \text{ \& } J_\pi \text{ is the unique solution.}$$

Illustration of Value Iteration:

$$S = \{1, 2\}$$

$$A = \{a, b\}$$

$$P(a) = \begin{bmatrix} P_{11}(a) & P_{12}(a) \\ P_{21}(a) & P_{22}(a) \end{bmatrix} = \begin{bmatrix} 3/4 & 1/4 \\ 3/4 & 1/4 \end{bmatrix}$$

$$P(b) = \begin{bmatrix} P_{11}(b) & P_{12}(b) \\ P_{21}(b) & P_{22}(b) \end{bmatrix} = \begin{bmatrix} 1/4 & 3/4 \\ 1/4 & 3/4 \end{bmatrix}$$

costs: $g(1,a) = 2, \quad g(1,b) = 0.5,$
 $g(2,a) = 1, \quad g(2,b) = 3$

$$\alpha = 0.9$$

$$J_0 = (0, 0)$$

D_0 VI for 2 steps

$$J_1 = (0.5, 1)$$

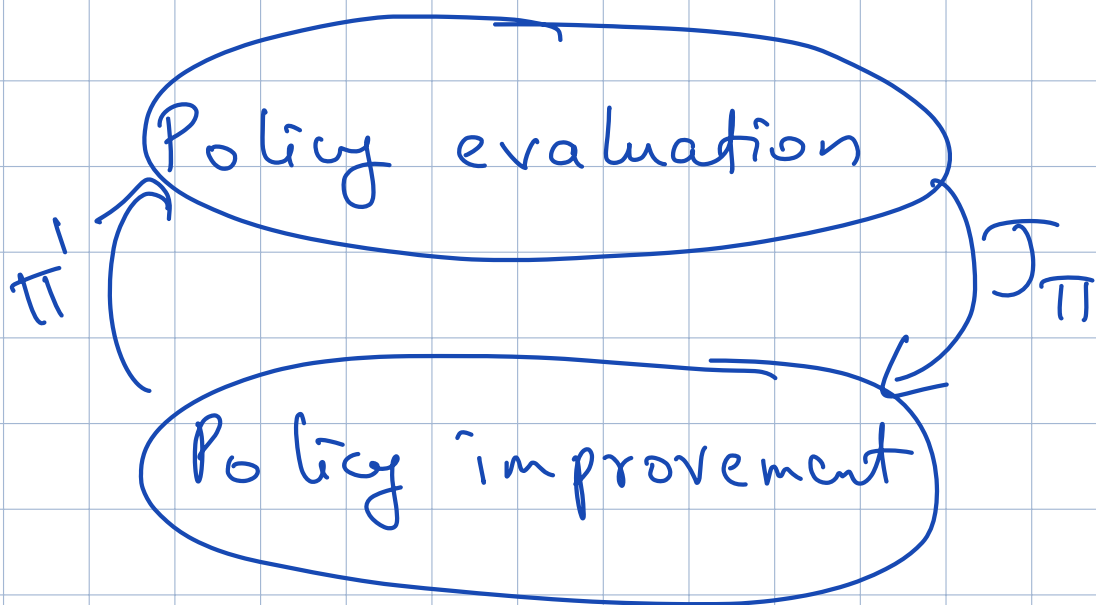
$$TJ(i) = \min \left(g(i,a) + \alpha \sum_{j=1}^2 P_{ij}(a) T(j) \right),$$

$$g(i, b) + \alpha \sum_{j=1}^2 p_{ij}^{(b)} J(j)$$

$$J_2 = (1.287, 1.562)$$

Lecture - 22

Policy iteration



$$\pi_0 \xrightarrow{\text{P.E.}} J_{\pi_0} \xrightarrow{\text{P.I.}} \pi_1 \quad (J_{\pi_1}(i) \leq J_{\pi_0}(i))$$

with strict ineq.
for at least one
state.)

P.E.

$$\pi^* \dots \leftarrow \pi_2 \xleftarrow{P.I} T \pi_1$$

Claim:- A policy π is optimal if and only if $\pi(i)$ attains the minimum in the Bellman equation

$$g(i, \pi(i)) + \alpha \sum_{\bar{j}=1}^n P_{i\bar{j}}(\pi(i)) J^*(\bar{j})$$

$$= \min_a \left[g(i, a) + \alpha \sum_{\bar{j}=1}^n P_{i\bar{j}}(a) J^*(\bar{j}) \right]$$

(or)

$$T_{\pi} J^* = T J^*$$

Proof: Assume $TJ^* = T_{\pi}J^*$ for some π

We know $J^* = TJ^*$

$$\text{So, } J^* = TJ^* = T_{\pi}J^*$$

$\Rightarrow J^* = J_{\pi}$ & hence, π is optimal.

Converse: π is optimal

$$\Rightarrow J^* = J_{\pi}$$

$$\Rightarrow J^* = T_{\pi}J^*$$

$$\Rightarrow TJ^* = T_{\pi}J^*$$



Policy iteration algorithm

Step 1: Fix π_0

Step 2: Find J_{π_k} corresponding to π_k

↓
Policy evaluation

Step 3: Obtain π_{k+1} as
↓
Policy improvement

$$T_{\pi_{k+1}} J_{\pi_k} = T J_{\pi_k}$$

If $\pi_{k+1} = \pi_k$, stop, else goto Step 2.

Claim: Let π, π' satisfy

$$T_{\pi'} J_{\pi} = T J_{\pi}$$

Then, $J_{\pi'}(i) \leq J_{\pi}(i), \forall i$

& if π is not optimal, then

$J_{\pi'}(k) < J_{\pi}(k)$ for at
least one state k .

Illustration of policy iteration:

$$S = \{1, 2\}, A = \{a, b\}$$

$$P_a = \begin{bmatrix} 3/4 & 1/4 \\ 3/4 & 1/4 \end{bmatrix}$$

$$P_b = \begin{bmatrix} 1/4 & 3/4 \\ 1/4 & 3/4 \end{bmatrix}$$

$$g(1, a) = 2, g(1, b) = 0.5, g(2, a) = 1, g(2, b) = 3$$

$$\alpha = 0.9$$

$$\pi_0(1) = a, \pi_0(2) = b$$

$$J_{\pi_0}(1) = g(1, a) + \alpha P_{11}(a) J_{\pi_0}(1) + \alpha P_{12}(a) J_{\pi_0}(2)$$

$$J_{\pi_0}(2) = g(2, b) + \alpha P_{21}(b) J_{\pi_0}(1) + \alpha P_{22}(b) J_{\pi_0}(2)$$

$$J_{\pi_0}(1) = 2 + 0.9 \times \frac{3}{4} J_{\pi_0}(1) + 0.9 \times \frac{1}{4} J_{\pi_0}(2)$$

$$J_{\pi_0}(2) = 3 + 0.9 \times \frac{1}{4} J_{\pi_0}(1) + 0.9 \times \frac{3}{4} J_{\pi_0}(2)$$

$$J_{\pi_0}(1) = 24.12, \quad J_{\pi_0}(2) = 25.96$$

Policy Improvement:

$$T_{\pi} J_{\pi_0} = T J_{\pi_0}$$

$$(T J_{\pi_0})(1) = \min \text{ of}$$

$$\left(2 + 0.9 \left(\frac{3}{4} \times 24.12 + \frac{1}{4} \times 25.96 \right) \right)$$

action a

$$0.5 + 0.9 \left(\frac{1}{4} \times 24.12 + \frac{3}{4} \times 25.96 \right)$$

$$= \min(24.12, 23.45) = 23.45 (= \text{action } b)$$

$$(T J_{\pi_0})(2) = \min \left(\frac{23.12}{25.95} \right),$$

$$\pi_1(1) = b, \quad \pi_1(2) = a$$

H.W:- Do one more step of policy iteration

i.e., find $J_{\pi_1}(1), J_{\pi_1}(2)$

& we that to perform policy improvement

You should get

$$\boxed{\pi_2 = \pi_1}$$

Sample path:

$$s_0 \xrightarrow{\pi(s_0)} s_1$$

$$s_1 \sim P_{s_0 s_1}(\pi(s_0))$$

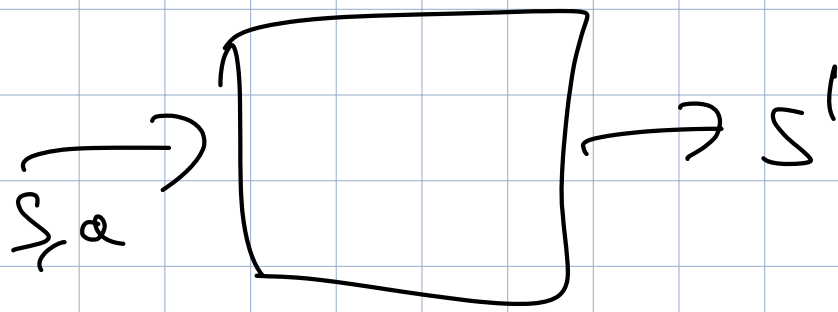
for e.g. $S = \{1, 2, 3\}$

$$s_0 = 1$$

$$P_{1,1}(\text{action}) = 0$$

$$P_{1,2}(\text{action}) = 0.4$$

$$P_{1,3}(\text{—u—}) = 0.6$$



$(s_0, a_0, s_1, a_1, s_2, \dots)$

Sample path

Policy evaluation (prediction)

$(s_0, \pi(s_0), s_1, \pi(s_1), \dots)$

Goal: Find J_π

Policy Control:-

$(s_0, a_0, s_1, \dots)$

Reinforcement Learning

Mean estimation:

R.v. $\mathcal{V}$ with mean μ and finite variance.

$\{\mathcal{V}_1, \dots, \mathcal{V}_n\} \leftarrow$ iid samples from the distribution of $\mathcal{V}$.

$$r_n = \frac{1}{n} \sum_{i=1}^n \mathcal{V}_i$$

$$r_{n+1} = \frac{1}{n+1} \sum_{i=1}^{n+1} \mathcal{V}_i$$

$$= \frac{n}{n+1} \left(\frac{1}{n} \sum_{i=1}^n \mathcal{V}_i \right) + \frac{1}{n+1} \mathcal{V}_{n+1}$$

$$= \frac{n}{n+1} r_n + \frac{1}{n+1} \mathcal{V}_{n+1}$$

$$r_{n+1} = r_n + \frac{1}{n+1} (\mathcal{V}_{n+1} - r_n)$$

More generally,

$$x_{n+1} = x_n + \eta_n (y_{n+1} - y_n)$$

↖ step-size

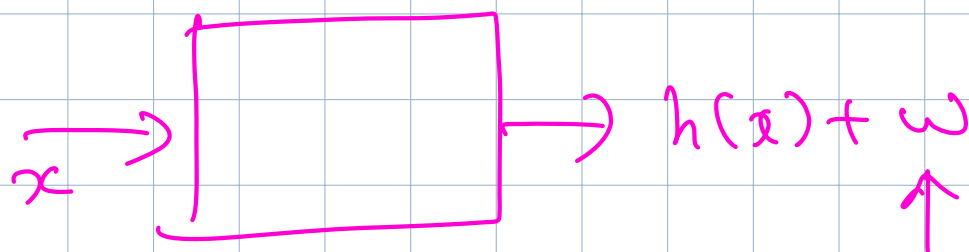
If $\sum_{n=1}^{\infty} \eta_n = \infty$, $\sum_{n=1}^{\infty} \eta_n^2 < \infty$, then

$x_n \rightarrow \mu$ w.p. 1 as $n \rightarrow \infty$.

Stochastic approximation

(Robbins Monro algorithm)

Suppose you want to solve $h(x) = 0$,
 h is Lipschitz



Zero-mean noise

e.g. $w \sim N(0,1)$

$$x_{n+1} = x_n + \eta_n (h(x_n) + w_n)$$

Under some conditions, it can be shown that $x_n \rightarrow x^*$ w.p.1 as $n \rightarrow \infty$, where $h(x^*) = 0$

Two special cases:-

① $\min_x f(x)$ Then $h = f'$
or in higher dimensions $h = \nabla f$
"Gradient descent"

② Want to solve $f(x) = x$
Set $h(x) = f(x) - x$ } Fixed point iteration.

In a MDP context, one usually wants to solve $Hx = r$

$$H: \mathbb{R}^n \rightarrow \mathbb{R}^n, \quad r \in \mathbb{R}^n$$

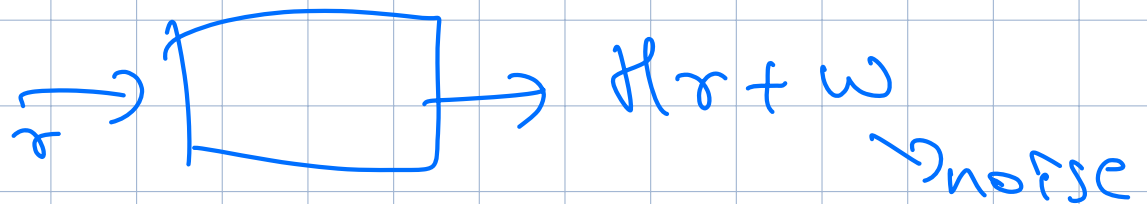
e.g. $H = T_\pi$

A direct method:

$$r_{n+1} = H r_n$$

(or) $r_{n+1} = (1 - \eta_n) r_n + \eta_n (H r_n)$

In a RL setting,



$$r_{n+1} = (1 - \eta_n) r_n + \eta_n (H r_n + w_n)$$

Assuming H is "well-behaved", &
noise w_n is zero-mean + bounded variance,

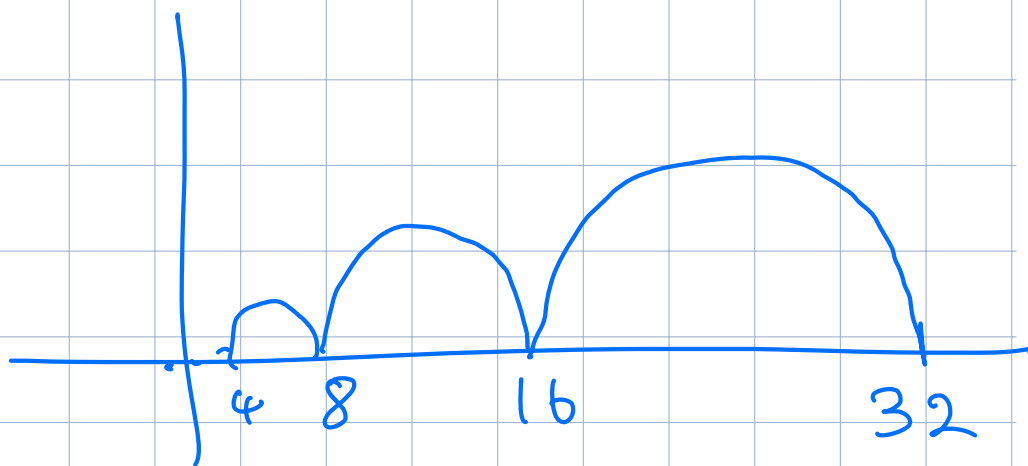
One can show that

$$r_n \rightarrow r^* \text{ a.s. as } n \rightarrow \infty$$

$$\text{where } \mathbb{1}_{r^*} = r^*$$

A brief tour of contraction mappings.

$$f(x) = \frac{x}{2}$$



$$\lim_{n \rightarrow \infty} f^n(x) = 0$$

$$f(0) = 0$$

Def: $f: \mathbb{R} \rightarrow \mathbb{R}$ is a contraction mapping with modulus $\beta \in (0,1)$ if

$$|f(x) - f(y)| \leq \beta |x - y|, \quad \forall x, y \in \mathbb{R}$$

Contraction properties of T, T_π

in a discounted MDP setting:

Max-norm:- $\|\cdot\|_\infty$ on $\mathbb{R}^n$ is

$$\|J\|_\infty = \max_{i=1 \dots n} |J(i)|$$

Proposition:- For any two bounded

J, J' , we have

$$\|TJ - TJ'\|_\infty \leq \alpha \|J - J'\|_\infty$$

$\uparrow$
discount factor

$$\|T_\pi J - T_\pi J'\|_\infty \leq \alpha \|J - J'\|_\infty,$$

for any stationary π .

Corollary:

$$\|T^k J - T^k J'\|_{\infty} \leq \alpha^k \|J - J'\|_{\infty}$$

Corollary:-

$$\|T^k J_0 - J^*\|_{\infty} \leq \alpha^k \|J_0 - J^*\|_{\infty}$$

Lecture-23

Stochastic iterative algorithm

for finding fixed point of a

contraction mapping

Goal: Solve $Hx^* = x^*$

$$H: \mathbb{R}^n \rightarrow \mathbb{R}^n, \quad x \in \mathbb{R}^n$$

Update rule:

$$r_{t+1}(i) = (1 - \eta_t) r_t(i) + \eta_t \left((H r_t)(i) + \underbrace{w_t(i)}_{\text{noise}} \right)$$

Assumption:

$$\begin{aligned} \mathcal{F}_t &= \{ r_0(i), \dots, r_t(i), w_0(i), \dots, w_{t-1}(i), \\ &\quad (i=1, \dots, n) \} \end{aligned}$$

(history)

(A1) $E(w_t(i) | \mathcal{F}_t) = 0$

$$E(w_t^2(i) | \mathcal{F}_t) \leq A + B \|r_t\|^2$$

(A2) $\sum_{t=1}^{\infty} \eta_t = \infty, \quad \sum_{t=1}^{\infty} \eta_t^2 < \infty.$

(A3)

H is a max-norm
pseudo-contraction

$$\|Hr - r^*\|_\infty \leq \alpha \|r - r^*\|_\infty, \forall r$$

Full contraction

$$\|Hr - Hr'\|_\infty \leq \alpha \|r - r'\|_\infty, \forall r, r'$$

Under A1-A3,

$$r_t \rightarrow r^* \text{ w.p.1 as } n \rightarrow \infty$$

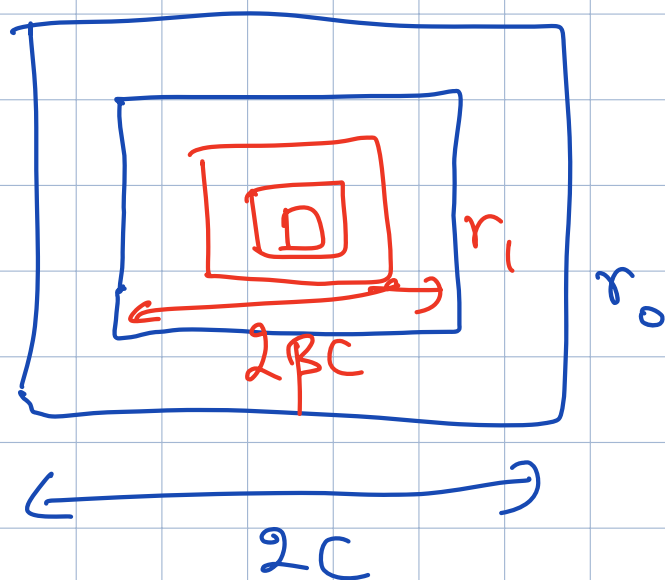
$$\text{where } Hr^* = r^*$$

An intuitive justification:-

$$H0 = 0, \quad r^* = 0$$

$$|r_0(i)| \leq C$$

$$r_{t+1}(i) = (Hr_t)(i)$$



$$r_{t+1}(i) = (1 - \eta_t) r_t(i)$$

$$+ \eta_t (Hr_t(i) + w_t(i))$$

Similar behavior, but shrinkage
is not immediate

Temporal-difference (TD) learning for policy evaluation:-

Fix a policy π .

Single stage wt: $g(\hat{i}, \hat{j})$

Want to estimate:-

$$J_{\pi}(\hat{i}) = \mathbb{E} \left[\sum_{m=0}^{\infty} \alpha^m g(\hat{i}_m, \hat{i}_{m+1}) \mid \hat{i}_0 = \hat{i} \right]$$

Bellman equation:-

$$J_{\pi} = T_{\pi} J_{\pi}$$

$$J_{\pi}(\hat{i}) = \mathbb{E} \left(g(\hat{i}, \bar{i}) + \alpha J_{\pi}(\bar{i}) \right)$$

TD(0) algorithm:-

$$J_{t+1}(\hat{i}) = (1 - \eta_t) J_t(\hat{i}) + \eta_t \left(g(\hat{i}, \bar{i}) + \alpha J_t(\bar{i}) \right)$$

$\bar{i} \sim P_{\hat{i}}(\pi(\hat{i}))$

$$J_{t+1}(i) = J_t(i) + \eta_t \left(\underbrace{g(i, \bar{i}) + \alpha J_t(\bar{i})}_{\text{Proxy for } E(g(i, \bar{i}) + \alpha J_t(\bar{i}))} - J_t(i) \right)$$

$$= (T_\pi J_t)(i)$$

$$J_{t+1}(i) = J_t(i) + \eta_t \left((T_\pi J_t)(i) - J_t(i) + \underbrace{g(i, \bar{i}) + \alpha J_t(\bar{i}) - (T_\pi J_t)(i)}_{w_t(i)} \right)$$

Using result for r_t above, we can infer

$$J_t \rightarrow J_\pi \text{ w.p.1 as } t \rightarrow \infty$$

$$\begin{aligned}
J_{\pi}(i) &= E(g(i, \bar{i}) + \alpha J_{\pi}(\bar{i})) \\
&= E\left(g(i, \underset{\substack{\downarrow \\ \text{next state}}}{\bar{i}}) + \alpha g(\underset{\substack{\downarrow \\ \text{next-to-next state}}}{\bar{i}}, \bar{\bar{i}}) + \alpha^2 J_{\pi}(\bar{\bar{i}})\right)
\end{aligned}$$

$T_D(\lambda), \lambda \in [0, 1]$:

$$J_{\pi}(i) = E\left(\sum_{m=0}^{\infty} \alpha^m g(\hat{i}_m, \hat{i}_{m+1}) + \alpha^{\infty} J_{\pi}(\hat{i}_{\infty})\right)$$

Fix $\lambda < 1$, and multiply $(1-\lambda)\lambda^l$ & sum over l :

$$J_{\pi}(i) = (1-\lambda) E \left(\sum_{l=0}^{\infty} \lambda^l \left(\sum_{m=0}^l \alpha^m g(i_m, i_{m+1}) + \alpha^{l+1} J_{\pi}(i_{l+1}) \right) \right)$$

Interchange the two summations & use $(1-\lambda) \sum_{l=m}^{\infty} \lambda^l = \lambda^m$, to obtain

$$\begin{aligned} J_{\pi}(i) &= E \left[(1-\lambda) \sum_{m=0}^{\infty} \alpha^m g(i_m, i_{m+1}) \sum_{l=m}^{\infty} \lambda^l + \sum_{l=0}^{\infty} \alpha^{l+1} J_{\pi}(i_{l+1}) (\lambda^l - \lambda^{l+1}) \right] \\ &= E \left(\sum_{m=0}^{\infty} \lambda^m \left(\alpha^m g(i_m, i_{m+1}) + \alpha^{m+1} J_{\pi}(i_{m+1}) - \alpha^m J_{\pi}(i_m) \right) \right) \\ &\quad + J_{\pi}(i) \end{aligned}$$

Letting $d_m = g(i_m, i_{m+1}) + \alpha J_{\pi}(i_{m+1}) - J_{\pi}(i_m)$

$$J_{\pi}(i) = E \left(\sum_{m=0}^{\infty} (\alpha \lambda)^m d_m \right) + J_{\pi}(i)$$

TD(λ) update rule:-

$$J_{t+1}(i) = J_t(i) + \eta_t \left(\sum_{m=0}^{\infty} (\alpha \lambda)^m d_m \right)$$

for $\lambda=0$, $d_0 = g(i, \bar{i}) + \alpha J_{\pi}(\bar{i}) - J_{\sigma}(i)$

$\Leftarrow$ we recover TD(0)

Stochastic shortest path

$$S = \{1, \dots, n, T\}$$

$$P_{TT}(a) = 1$$

$$g(\bar{1}, a, T) = 0$$

$$J_{\pi}(i) = E \left(\sum_{t=0}^{\infty} g(i_t, \pi(i_t), i_{t+1}) \mid \bar{i}_0 = i \right)$$

$$J_{\pi}(i) = E \left(g(i, \pi(i), \bar{i}) + J_{\pi}(\bar{i}) \right)$$

π : proper if it ensures that terminal state T is reached with positive probability.

TD(0): $J_{t+1}(\bar{i}) = J_t(i) + \eta_t \left(\underbrace{g(i, \pi(i), \bar{i}) + J_t(\bar{i}) - J_t(i)}_{\text{temporal-difference}} \right)$

for $\bar{i} = 1, \dots, n$

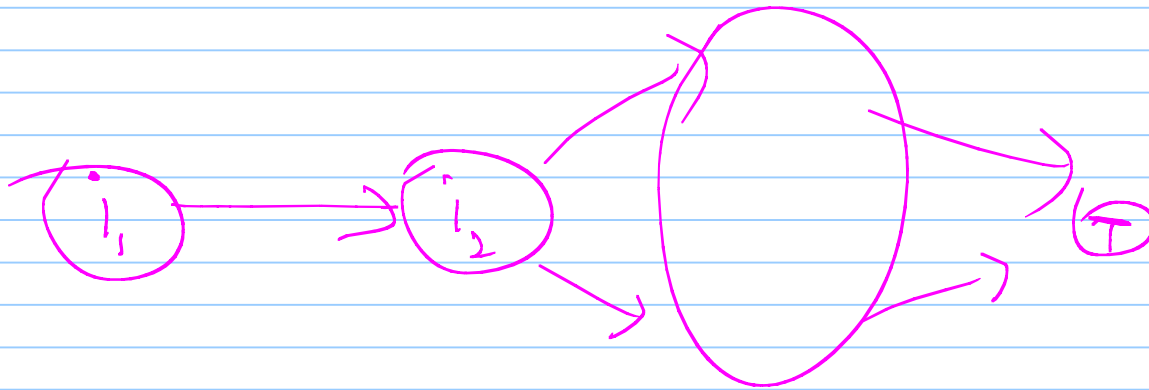
TD(i):- Simulate episodes of the underlying SSP w.r.t π .

$$(i_0^m, i_1^m, \dots, i_N^m), \quad i_N = T, \quad i_0 = i$$

$$\hat{J}^m = g(i_0^m, \pi(i_0^m), i_1^m) + \dots + g(i_{N-1}^m, \pi(i_{N-1}^m), i_N^m)$$

$$\hat{J}(i) = \frac{1}{M} \sum_{m=1}^M \hat{J}^m, \quad M = \# \text{ of episodes simulated}$$

TD(0) vs TD(1) :- A toy example



TD(1) : Start in i_1 , simulate episode, we total get, log $\hat{J}(i_1)$,
to estimate $J_\pi(i_1)$

TD(0) : Start in i_1 , simulate one transition and
we " $g(i_1, i_2) + J(i_2) - J(i_1)$ " to estimate $J_\pi(i_1)$

Lecture-24

Q-learning: SSP setting

Let J^* be the optimal cost-to-go vector

$$Q^*(i, a) = \sum_{\hat{j}=1-n, T} P_{ij}(a) (g(i, a, \hat{j}) + J^*(\hat{j})), \quad i=1 \dots n$$

Bellman's equation:

$$J^*(i) = \min_a \sum_{\hat{j}=1-n, T} P_{ij}(a) (g(i, a, \hat{j}) + J^*(\hat{j}))$$

$$J^*(i) = \min_a Q^*(i, a)$$

Q-Bellman equation:-

$$Q^*(i, a) = \sum_{\hat{j}=1-n, T} P_{ij}(a) (g(i, a, \hat{j}) + \min_b Q^*(\hat{j}, b))$$
$$Q^*(i, \mu) = E (g(i, \mu, \hat{j}) + \min_b Q^*(\hat{j}, b))$$

$$Q_{t+h}(i,a) = Q_t(i,a) + \eta_t \left(g(i,a,\bar{i}) + \min_b Q_t(\bar{i},b) - Q_t(i,a) \right)$$

Define the operator $\mathcal{H}$

$$(\mathcal{H}Q)(i,a) = \sum_{\bar{j}=1, \dots, N, T} P_{i,\bar{j}}(a) \left(g(i,a,\bar{j}) + \min_b Q(\bar{j},b) \right)$$

Under some conditions, $\mathcal{H}Q$ is a contraction mapping.

$$Q_{t+h}(i,a) = Q_t(i,a) + \eta_t \left((\mathcal{H}Q_t)(i,a) + w_t(i,a) \right)$$

$$w_t(i,a) = g(i,a,\bar{i}) + \min_b Q_t(\bar{i},b) - \sum_{\bar{j}=1, \dots, N, T} P_{i,\bar{j}}(a) \left(g(i,a,\bar{j}) + \min_b Q_t(\bar{j},b) \right)$$

$Q_t \rightarrow Q^*$ w.p.1 as $t \rightarrow \infty$ under conditions like (A1) - (A3)

Issue of exploration:-

For Q-learning to converge, all state-action pairs (i, a) have to be visited infinitely often.

Greedy policy: $\pi_{t+1}(i) = \arg \min_a Q_t(i, a)$

ϵ -Greedy policy:- $\pi_{t+1}(i) = \begin{cases} \text{greedy action w.p. } (1-\epsilon) \\ \text{random action w.p. } \epsilon \end{cases}$

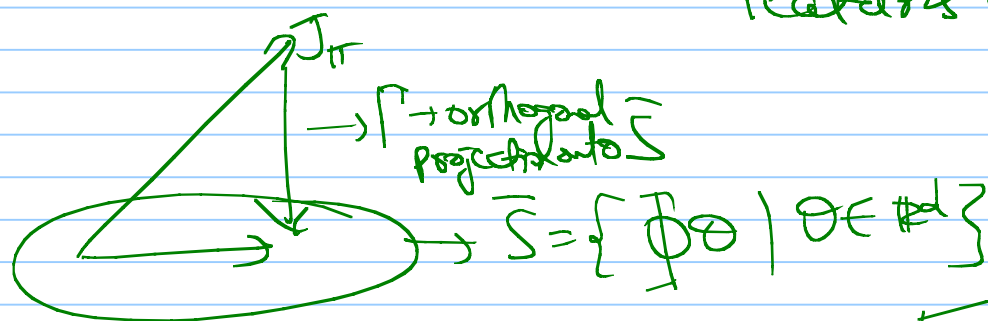
$\epsilon \rightarrow$ small number e.g. 0.05

Linear function approximation Fix a policy π .

$$J_{\pi}(i) \approx \phi(i)^T \theta$$

$\downarrow$ features $\in \mathbb{R}^d$ $\rightarrow \in \mathbb{R}^d$

$$d \ll |\mathcal{S}|$$



$$J_{\pi} \approx \Phi \theta$$

Cannot solve

$$J_{\pi} = T_{\pi} J_{\pi}$$

Approximate solution

$$\bar{\Phi} \theta = \Gamma (T_{\pi} \bar{\Phi} \theta) \rightarrow \text{has a unique solution since } \Gamma T_{\pi} \text{ is a contraction.}$$

$$\Theta_{t+1} = \Theta_t + \eta_t \left(g(i, \pi(i), \bar{i}) + \phi(\bar{i})^T \Theta_t - \phi(i)^T \Theta_t \right) \phi(i)$$

TD(0) with linear function approximation

Q-learning with linear function approximation

$$Q(i, a) \approx \phi(i, a)^T \Theta$$

$$\Theta_{t+1} = \Theta_t + \eta_t \left(g(\bar{i}, a, \bar{i}) + \min_b \phi(\bar{i}, b)^T \Theta_t - \phi(i, a)^T \Theta_t \right) \times \phi(i, a)$$

Greedy action $\min_a \Theta_t^T \phi(i, a)$

Policy gradient methods (PG methods)

The main idea behind PG methods is the "likelihood ratio" trick (aka score function).
An illustration of this trick in a very simple setting.

Let X be a ^{discrete} r.v. with mass function $p(\theta, \cdot)$

$\rightarrow$ parameter

i.e., $P(X=x)$ is parameterized by θ .

$$J(\theta) = E f(X)$$

e.g. $\theta \in \mathbb{R}^d$

Goal:

$$\min_{\theta} J(\theta)$$

$\rightarrow$ min among a class of parameterized r.v.s

Want to find best θ using a gradient method

$$\theta_{t+1} = \theta_t - \beta_t \nabla J(\theta_t) \quad \leftarrow \text{Gradient descent}$$

Need: " $\nabla J(\theta)$ "

Use "likelihood ratio" trick, which is shown below.

$$J(\theta) = \sum_x f(x) p(\theta, x) \quad \leftarrow \text{LOTUS}$$

$$\nabla J(\theta) = \nabla_{\theta} \left(\sum_x f(x) p(\theta, x) \right)$$

need a few conditions to justify interchange of $\sum$ & ∇

$$= \sum_x f(x) \nabla p(\theta, x)$$

need regularity conditions $\rightarrow$ usually let one invoke "dominated"

$$\nabla J(\theta) = \sum_x f(x) \nabla p(\theta, x)$$

$$\nabla J(\theta) = \sum_x \left(f(x) \frac{\nabla p(\theta, x)}{p(\theta, x)} \right) p(\theta, x) \rightarrow \text{PMS is an expectation}$$

$$\nabla J(\theta) = E \left(f \frac{\nabla p}{p} \right) = E(f \nabla \log p)$$

To get an estimate of $\nabla J(\theta)$, I can sample from $\underline{\hspace{2cm}}$

$$\frac{\nabla p(\theta, x)}{p(\theta, x)} \rightarrow \text{likelihood ratio.}$$

$$\nabla \log p(\theta) = \frac{\nabla p(\theta)}{p(\theta)}$$

Connecting likelihood ratio to RL:

Fix an SSP problem with state space X , action space A .

Π_{det} : class of admissible stationary deterministic policies

$\{ \pi : \pi : X \rightarrow A \text{ \& it is time invariant} \}$

$\pi \rightarrow$ unparameterized policy $\rightarrow \pi(x) \rightarrow$ an action $\pi: \begin{matrix} \text{states} \\ \text{actions} \end{matrix}$

for Ph method, we consider stationary randomized policies, i.e.,

$\pi : X \rightarrow \Delta(A)$ $\rightarrow$ set of all distributions over the actions.
For simplicity, assume all actions available in all states.

e.g. $\pi_{\theta}(x, a) = \frac{\exp(h(\theta, x, a))}{\sum_b \exp(h(\theta, x, b))}$ → randomized policy parameterized by " θ "

is the probability of choosing action a in state x

$\pi_{\theta}(x) = [\pi_{\theta}(x, a), \forall a \in \mathcal{A}]$

π_{θ} → distribution over $\mathcal{A}$.

$\pi_{\theta}(x, a) = \frac{\exp(\theta_{x,a})}{\sum_b \exp(\theta_{x,b})}$ → vector of size $|\mathcal{A}|$ (states x (action))

Simple example for h : $h(\theta, x, a) = \theta^T \phi(x, a)$ → $\theta \in \mathbb{R}^d$, $\phi(x, a) \in \mathbb{R}^d$

state-action features

$\pi_{\theta}(x, a) = \frac{\exp(\theta^T \phi(x, a))}{\sum_b \exp(\theta^T \phi(x, b))}$ Boltzmann distribution aka Soft-max

$\theta \in \mathbb{R}^d, \phi(x, \cdot) \in \mathbb{R}^d$

Assumption A1: Policy π_{θ} is a continuously differentiable function of θ and $\nabla \log \pi_{\theta}$ exists.

Note: Every π_{θ} is identified by its parameter $\theta \in \mathbb{R}^d$

Goal: $\min_{\theta \in \Theta} J_{\pi_{\theta}}(x^0)$

→ Find an approximately optimal policy in the class of parameterized policies

& we want to find the best parameter in a class

$\{\pi_{\theta} \mid \theta \in \Theta\}$

e.g. $\Theta \subset \mathbb{R}^d$

Want to find a $\theta^* \in \arg \min_{\theta \in \Theta} J_{\pi_{\theta}}(x^0)$

"—" is or necessarily a convex function over Θ

A formula for policy gradient in an SSP

SSP: $\emptyset$ is a special cost-free absorbing state.

Total cost Objective:

$$J_{\pi_{\theta}}(x^0) = E \left(\sum_{k=0}^{\infty} g(x_k, a_k, x_{k+1}) \mid x_0 = x^0 \right)$$

$a_k \sim \pi_{\theta}(x_k, \cdot)$

A trajectory (aka episode) is

$$\tau = x_0 \overset{\text{fixed}}{a_0} x_1 a_1 \dots x_T, \quad x_T = \emptyset, \text{ the terminal state.}$$

Probability of seeing a trajectory τ is

$$\pi(x_0, a_0) P(x_1 | x_0, a_0) \dots P(x_T | x_{T-1}, a_{T-1})$$

$J_{\pi_{\theta}} \rightarrow$ can be seen as an average of the total cost over episodes.

Let $D(\tau)$ denote the total cost r.v.

$$D(\tau) = \sum_{k=0}^{T-1} g(x_k, a_k)$$

Here $\tau = (x_0, a_0, \dots, x_{T-1}, a_{T-1}, x_T)$
in the episode

$$J_{\pi_\theta}(x^0) = E(D(\tau))$$

↓
expectation over episodes.

$$\nabla J_{\pi_\theta}(x^0) = \nabla_\theta \sum_{\tau} D(\tau) \text{Prob}(\tau \text{ under } \pi_\theta)$$

$D(\tau)$ is not a function of θ
interchange of ∇ & $\sum$
justified for finite state action spaces.

$$= \sum_{\tau} D(\tau) \nabla_\theta \text{Prob}(\tau \text{ under } \pi_\theta)$$

$$= \sum_{\tau} D(\tau) \text{Prob}(\tau \text{ under } \pi_\theta) \nabla \log(\text{Prob}(\tau \text{ under } \pi_\theta))$$

$$= \sum_{\tau} D(\tau) \text{Prob}(\tau \text{ under } \pi_\theta)$$

$$\propto \nabla \log \left\{ \pi_\theta(x_0, a_0) P(x_1 | x_0, a_0) \dots \pi_\theta(x_{T-1}, a_{T-1}) P(x_T | x_{T-1}, a_{T-1}) \right\}$$

why!?

$$= \sum_{\tau} D(\tau) \text{Prob}(\tau \text{ under } \pi_\theta) \sum_{k=0}^{T-1} \nabla \log \pi_\theta(x_k, a_k)$$

Main claim! PTO

$$\nabla J_{\pi_{\theta}}(x^0) = \mathbb{E} \left[D(\tau) \sum_{k=0}^{T-1} \nabla \log \pi_{\theta}(x_k, a_k) \right]$$

↓
expectation over trajectories.

An unbiased estimate of $\nabla J_{\pi_{\theta}}(x^0)$

given an episode $(x^0, a_0, \dots, x_T)$

Calculate total cost of this episode, say $\hat{D}$

Calculate likelihood ratio $\hat{\psi} = \sum_{k=0}^{T-1} \nabla \log \pi_{\theta}(x_k, a_k)$

$$\hat{\nabla} J = \hat{D} \times \hat{\psi}.$$

$$\theta_{t+1} = \theta_t - \alpha(t) \hat{\nabla} J$$

policy gradient update

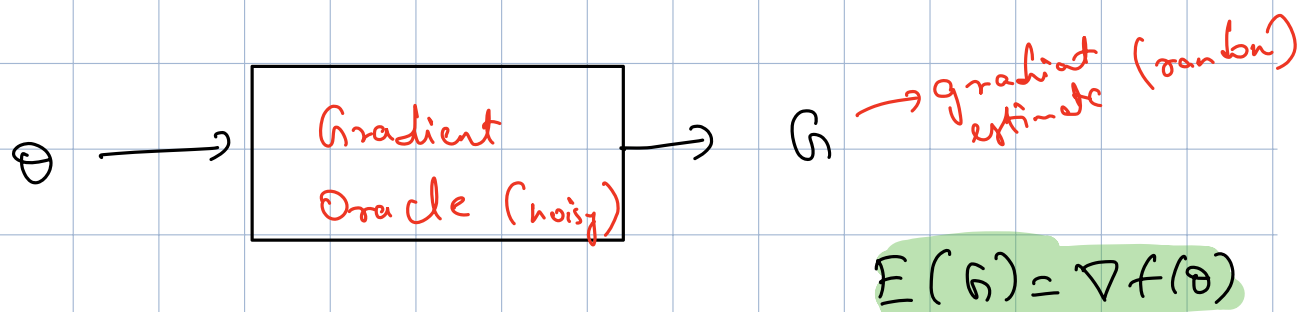
"REINFORCE". (Williams, 1992)

Lecture - 18

A detour $\rightarrow$ Non-asymptotic analysis of
Stochastic gradient algorithms

Aim: Solve $\min_{\theta} f(\theta)$
 $\nearrow$ smooth

Setting: Unbiased gradient information



SG algorithm: $\theta_{k+1} = \theta_k - \alpha(k) G(\theta_k, \xi_k)$ ①

$\nearrow$ noise

Assumption:

(A1) $E_{\xi_k} (G(\theta_k, \xi_k)) = \nabla f(\theta_k), \forall k \geq 1$

$E_{\xi_k} \| G(\theta_k, \xi_k) - \nabla f(\theta_k) \|^2 \leq \sigma^2$
for some $\sigma > 0$.

(A0) f is L -smooth

$$\|\nabla f(\theta_1) - \nabla f(\theta_2)\| \leq L \|\theta_1 - \theta_2\|, \quad \forall \theta_1, \theta_2 \in \mathbb{R}^d$$

Asymptotic convergence:

$$(A2) \quad \sum_k a(k) = \infty, \quad \sum_k a(k)^2 < \infty$$

e.g. $a(k) = \frac{1}{k} \rightarrow$ standard stochastic approximation conditions.

Under (A0) - (A2), the parameter θ_k governed by ① converges almost surely to the set $\{\theta^* \mid \nabla f(\theta^*) = 0\}$

Tool used for proving this claim: Kushner-Clark Lemma

Eq ① ($\Rightarrow$)

$$\theta_{k+1} = \theta_k - a(k) (\nabla f(\theta_k) + \xi_k)$$

$$\xi_k = G(\theta_k, \xi_k) - \nabla f(\theta_k)$$

We want to understand the non-asymptotic performance of the SG algorithm above.

Stationary point, say Θ^* , has $\nabla f(\Theta^*) = 0$

ϵ -stationary point: $\bar{\Theta}$ is an ϵ -stationary point if $E \|\nabla f(\bar{\Theta})\| \leq \epsilon$, where the expectation is over the randomness in the SG algorithm.

[Ghadimi-Lan, 2014, SIAM J. Opt, "Stochastic first/zeroth order"]

"Randomized Stochastic gradient" (RSG)

Suppose we run (1) for N iterations.

$\{\Theta_1, \Theta_2, \Theta_3, \dots, \Theta_N\} \xrightarrow{\text{(random)}} \text{set of iterates visited by SG algorithm.}$

RSG will pick a random iterate as follows!

Θ_R is picked uniformly at random from $\{\Theta_1, \dots, \Theta_N\}$

$$\text{i.e., } P(\theta_k = \theta_i) = \frac{1}{N} \quad \forall i$$

$$E \theta_k = \frac{1}{N} \sum_{i=1}^N \theta_i \quad \left. \vphantom{\sum_{i=1}^N} \right\} \text{average iterate.}$$

The non-asymptotic bound for RSG is of the form

$$E \|\nabla f(\theta_k)\|^2 \leq \frac{\text{const}}{\sqrt{N}}$$

under suitable choice of step-size.

RL connection: $\theta \rightarrow$ policy parameter

Objective $f \rightarrow J_{\pi_\theta}(x^0)$ or $J(\theta)$
 $\downarrow$ $\downarrow$
 Value function with state x^0 in a discounted/SSP $\downarrow$ Average cost

Proof of the non-asymptotic bound for RSG:

First $\rightarrow$ derive a bound for a general step-size

Second $\rightarrow$ specialize the bound above with a particular S1 S1 c.

$$a_i \sim \pi_\theta(\cdot | s_i) \xrightarrow{\text{term}} \\ \tau = (x_0, a_0, x_1, a_1, \dots, x_T)$$

$$\nabla J_{\pi_\theta}(x^0) = \mathbb{E} \left[D(\tau) \sum_{k=0}^{T-1} \nabla \log \pi_\theta(x_k, a_k) \right]$$

expectation over trajectories τ

"Each τ ends in a terminal state.

Random length trajectories.

but length bounded w.p.1
if policy θ is proper".

$\theta \rightarrow$ parameter

$$J(\theta) = J_{\pi_\theta}(x^0)$$

fixed start state

For RSG analysis/bound to be applicable,
we need to show

$$(1) \quad \|\nabla J(\theta_1) - \nabla J(\theta_2)\| \leq 2 \|\theta_1 - \theta_2\|$$

" J is 2-smooth".

(2) Gradient estimate from a sample trajectory τ .

$$\hat{G} = D(\tau) \sum_{k=0}^{T-1} \nabla \log \pi_\theta(x_k, a_k)$$

$$\mathbb{E}(\hat{G}) = \nabla J(\theta) \leftarrow \text{unbiased gradient estimate.}$$

To verify ① & ② in an RL context,
we make the following assumptions:

Ⓑ1 State & action spaces are finite

Ⓑ2 $\forall \theta$, π_θ is proper.

($\Rightarrow$) "in a finite, say M steps, there is a positive prob. of reaching the terminal state".

Ⓑ3 $\|\nabla \log \pi_\theta(x, a)\| \leq G$

$\|\nabla^2 \log \pi_\theta(x, a)\| \leq H$

$$\left(\begin{array}{l} K_i \log \pi_\theta(x, a) \leq G \\ \downarrow \\ \left| \frac{\partial^2}{\partial \theta_i \partial \theta_j} \right| \leq H \end{array} \right)$$

"Showing J is L -smooth as a function of θ "

Use this identity: for any twice diff'ble fn f :

$$\nabla^2 f(\theta) = f(\theta) \left(\nabla^2 \log f(\theta) + \nabla \log f(\theta) \nabla \log f(\theta)^\top \right)$$

(Scalar case: $f''(\theta) = f(\theta) \left(\frac{d^2}{d\theta^2} (\log f(\theta)) + \left(\frac{d}{d\theta} \log f(\theta) \right)^2 \right)$)

$$J(\theta) = \sum_{\tau} \text{Prob}(\tau \text{ under } \pi_{\theta}) \underbrace{D(\tau)}_{\text{Total cost for } \tau.}$$

$$= \sum_{\tau} p_{\theta}(\tau) D(\tau)$$

$$\begin{aligned} \nabla^2 J(\theta) &= \sum_{\tau} p_{\theta}(\tau) \nabla \log p_{\theta}(\tau) \nabla \log p_{\theta}(\tau)^{\top} D(\tau) \\ &\quad + \sum_{\tau} p_{\theta}(\tau) \nabla^2 \log p_{\theta}(\tau) D(\tau) \quad (*) \end{aligned}$$

Proper policies (see (B2))

Assume: Trajectory length is finite w.p. 1

Let T be a r.v. denoting the length of a random trajectory τ .

Then $|T| < C_1 < \infty$ w.p. 1.

Fact 2: Finite state-action spaces (B1)

$$\Rightarrow \sup_{x,a} |g(x,a)| < C_2 < \infty.$$

Fact 1 & 2 $\Rightarrow$

$$D(\tau) = \sum_{k=0}^{T-1} g(x_k, a_k)$$

$$|D(\tau)| \leq C_2 C_1$$

Using (B3) $(\|\nabla \log \pi\| \leq G, \|\nabla^2 \log \pi(\cdot, \cdot)\| \leq H)$

in (*),

$$\|\nabla^2 J(\theta)\| \leq C_1 C_2 \left(\sum_{\tau} p_{\theta}(\tau) \|\nabla \log p_{\theta}(\tau) \nabla \log p_{\theta}(\tau)^T\| \right.$$

$$\left. + \sum_{\tau} p_{\theta}(\tau) \|\nabla^2 \log p_{\theta}(\tau)\| \right)$$

$$\leq C_1 C_2 (G^2 C_1 + H C_1),$$

Since

$$\tau = (x_0, a_0, \dots, x_T)$$

$$p_{\theta}(\tau) = \pi_{\theta}(x_0, a_0) p(x_1 | x_0, a_0) \dots p(x_T | x_{T-1}, a_{T-1})$$

$$\nabla \log p_{\theta}(\tau) = \sum_{k=0}^{T-1} \nabla \log \pi_{\theta}(x_k, a_k) \leq G C_1$$

So, $J(\theta)$ is L -smooth, where

$$L = C_1 C_2 (G^2 C_1 + H C_1)$$

Homework:

$$E \|\hat{g} - E\hat{g}\|^2$$

$$= E \|\hat{g} - \nabla J(\theta)\|^2$$

$$\leq E \|\hat{g}\|^2$$

$$\hat{g} = D(\tau) \sum_{k=0}^{T-1} \nabla \log \pi_\theta(x_k)$$

$$\leq C_1^2 C_2^2 C_1 E \sum_{k=0}^{T-1} \|\nabla \log \pi(x_k, a_k)\|^2$$

$$\leq C_1^3 C_2^2 G^2 C_1$$

So, variance of the gradient estimate is bounded.

$\Rightarrow$ RSG analysis is applicable, leading to

$$E \|\nabla J(\theta_R)\|^2 \leq \frac{\text{const}}{\sqrt{N}},$$

where θ_k chosen unif. at. random from
 $\{\theta_1, \dots, \theta_N\}$

$$\& \theta_{k+1} = \theta_k - a(k) (\hat{G}_k)$$

↓
single trajectory estimate
of $\nabla J(\theta_k)$