

Theoretical foundations of RL

Not: Deep RL

Today: Theory of Markov decision processes (MDPs)

Tommorow: Reinforcement learning (RL)

Finite horizon MDPs:-

Example 1: Machine replacement/repair

States = $\{1, 2, \dots, n\}$
 $\downarrow$ $\downarrow$
 good worst

Actions = Do nothing, Repair

On repair, M/c goes to "1". Cost = R.

Operating cost: $g(i)$, $i=1 \dots n$

$$g(1) \leq g(2) \dots \leq g(n)$$

References:-

(1)

My webpage
↳ CS6700

(2)

(i) "Neuro-dynamic Programming",
Tsitsiklis & Bertsekas

(ii) "Introduction of RL"

Barto & Sutton

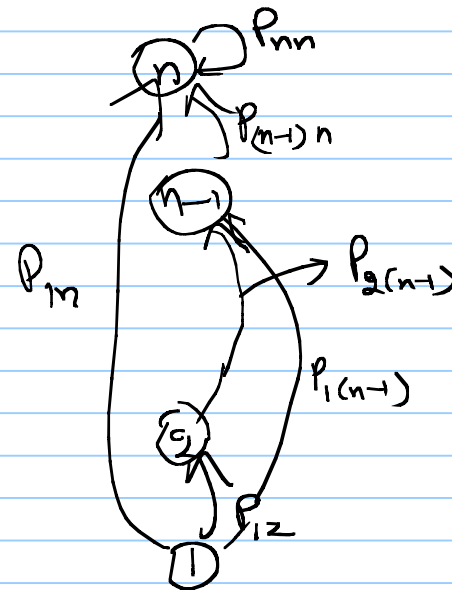
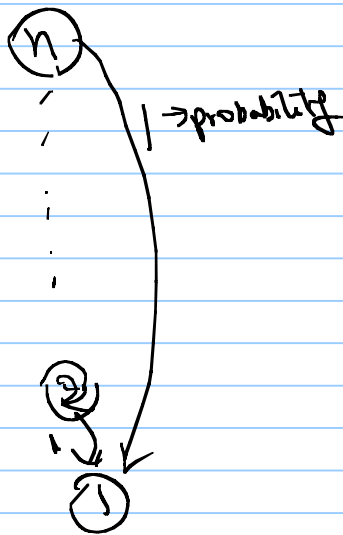
(iii) Bertsekas

"Dynamic programming
& optimal control"
Vol I & II.

State transitions:- $P_{ij} \rightarrow$ Probability of state transition from i to j .
 $P_{ij} = 0$ if $j < i$

Action: Do nothing

Action: repair



Example 2: Chess match

Playing styles (actions): Timid, Bold

draw: P_d

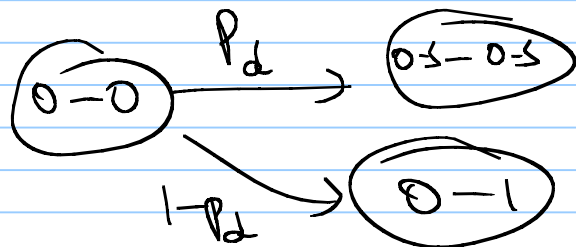
lose: $1 - P_d$

win: P_w

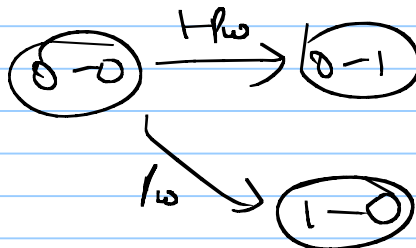
lose: $1 - P_w$

$P_d > P_w$

Game 1

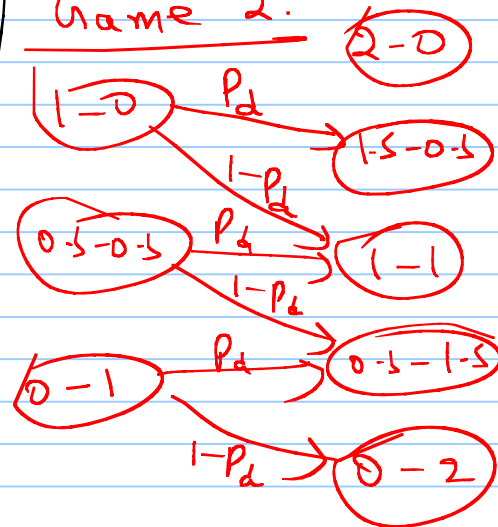


Timid



Bold

Game 2:



Timid

H.W:

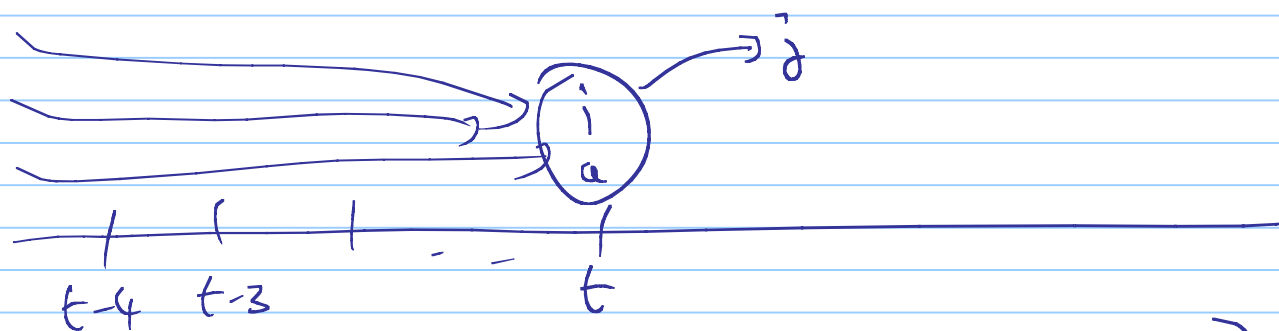
Transition
diagram
for

bold
play

Framework

S : State space, A : Action space

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



$$P(S_{k+1} = \hat{j} \mid S_k = \hat{i}, a_k = a, S_{k-1}, a_{k-1}, \dots, s_0) = P_{ij}(a)$$

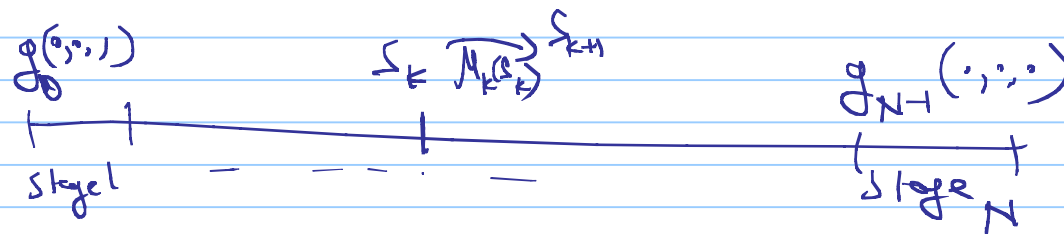
Single-stage cost: $\phi_k(\hat{i}, a, \hat{j})$

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

$\mu_i: \underset{\text{state space}}{S} \rightarrow \underset{\text{Action space}}{A}$

Performance metric

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



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

Optimality principle & dynamic programming (DP) algorithm

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

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

(N-i) stage tail subproblem

For this tail problem, $\{\mu_i^*, \dots, \mu_{N-1}^*\}$ is optimal.

DP algorithm:- Solve tail sub-problems, starting at the end, & progressing backwards.

DP algorithm

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

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

{

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

$\forall s_k \in S$

}

Example:- Applying DP algo to "Machine replacement" example

Suppose terminal cost is zero. $S = \{1 \dots n\}$

$$J_N(i) = 0$$

$$J_k(i) = \min \left[\underbrace{R + g(i) + J_{k+1}(i)}_{\text{Repair}}, \underbrace{g(i) + \sum_{j=1}^n \frac{P_{ij}}{j-i} J_{k+1}(j)}_{\text{Do nothing}} \right]$$

Chess-example (revisited)

State = net-score = (#wins - #losses) N stages.

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

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	$\max(P_d + (1 - P_d)P_w, P_w + (1 - P_w)P_w)$	So, timid play is best
0	P_w	bold
-1	P_w^2	bold

S_{N-1}	$J_{N-1}(S_{N-1})$
< -1	0
	Best play: does not matter

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

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

Example! Job-scheduling

N jobs. Job i takes T_i time to complete

If job i is completed at time t , then the reward is $\boxed{\alpha^t R_i}$

$$0 < \alpha < 1$$

Schedule $L = \{\bar{i}_0, \hat{i}_1, \dots, \hat{i}_{k-1}, \bar{i}, \hat{j}, \hat{i}_{k+2}, \dots, \bar{i}_{N-1}\}$

$L' = \{\text{---} \text{---} \text{---} \hat{j}, \bar{i}, \text{---} \text{---} \text{---}\}$

$$J_L = E \left[\alpha^{t_0} R_{\bar{i}_0} + \dots + \alpha^{t_{k-1}} R_{\hat{i}_{k-1}} + \alpha^{t_{k-1} + \bar{T}_i} R_{\bar{i}} + \alpha^{t_{k-1} + \bar{T}_i + \bar{T}_j} R_{\hat{j}} + \dots + \alpha^{t_{N-1}} R_{\bar{i}_{N-1}} \right]$$

$$J_{L'} = E \left[\alpha^{t_0} R_{\bar{i}_0} + \dots + \alpha^{t_{k-1}} R_{\hat{i}_{k-1}} + \alpha^{t_{k-1} + \bar{T}_j} R_{\hat{j}} + \alpha^{t_{k-1} + \bar{T}_j + \bar{T}_i} R_{\bar{i}} + \dots + \alpha^{t_{N-1}} R_{\bar{i}_{N-1}} \right]$$

So, L better than L' if

$$E \left[\alpha^{t_{k-1} + \bar{T}_i} R_{\bar{i}} + \alpha^{t_{k-1} + \bar{T}_i + \bar{T}_j} R_{\hat{j}} \right] \geq E \left[\alpha^{t_{k-1} + \bar{T}_j} R_{\hat{j}} + \alpha^{t_{k-1} + \bar{T}_j + \bar{T}_i} R_{\bar{i}} \right]$$

$$(or) \quad \frac{E(\alpha^{\bar{T}_i}) R_{\bar{i}}}{1 - E(\alpha^{\bar{T}_i})} \geq \frac{E(\alpha^{\bar{T}_j}) R_{\hat{j}}}{1 - E(\alpha^{\bar{T}_j})}$$

Optimal policy
is an
index-based
policy

Infinite horizon discounted cost MDPs

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

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

Stationary policy $\pi = \{\mu, \dots, \mu\}$ $\mu: S \rightarrow A$

$$J_{\pi}(i) = \mathbb{E} \left[\sum_{k=0}^{\infty} \alpha^k g(s_k, \pi(s_k), s_{k+1}) \mid s_0 = i \right]$$

discount factor

$$0 < \alpha < 1$$

Single-stage cost

(Same for all stages)

Initial state.

Policy

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

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

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

Bellman-operator $TJ = ((TJ)(1), \dots, (TJ)(n))$

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

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

Let $P_{\pi} = \left[\left[P_{i\bar{j}}(\pi(i)) \right]_{i, \bar{j}=1}^n \right]_{i=1}^n = \begin{bmatrix} P_{11}(\pi(1)) & & \\ & \ddots & \\ & & P_{nn}(\pi(n)) \end{bmatrix}$

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

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

Properties of T, T_π : ① For any n -vectors J, J' s.t. $J(i) \leq J'(i), i=1-n$

and stationary policy π , we have

$$\textcircled{1} (T^k J)(i) \leq (T^k J')(i), \forall i$$

$$\textcircled{2} (T_\pi^k J)(i) \leq (T_\pi^k J')(i), \forall i$$

② Let $e = n$ -vector of ones. For any J and scalar α ,

$$\textcircled{1} (T^k(J + \alpha e))(i) = (T^k J)(i) + \alpha^k, \forall i, \forall k \geq 1$$

$$\textcircled{2} (T_\pi^k(J + \alpha e))(i) = (T_\pi^k J)(i) + \alpha^k, \quad \text{---} \quad \text{---}$$

Proposition (Basis for Value-iteration)

For any bounded $J: S \rightarrow \mathbb{R}$,

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

Proof: For any policy $\pi = \{\mu_0, \mu_1, \dots\}$ and state $i \in S$,

$$\begin{aligned} J_\pi(i) &= \lim_{N \rightarrow \infty} \mathbb{E} \left(\sum_{k=0}^{N-1} \alpha^k g(s_k, \mu_k(s_k), s_{k+1}) \mid s_0 = i \right) \\ &= \mathbb{E} \left(\sum_{k=0}^{K-1} \alpha^k g(s_k, \mu_k(s_k), s_{k+1}) \right) + \underbrace{\lim_{N \rightarrow \infty} \mathbb{E} \left(\sum_{k=K}^{N-1} \alpha^k g(s_k, \mu_k(s_k), s_{k+1}) \mid s_0 = i \right)}_{\text{term-(B)}} \end{aligned}$$

Since $|g(s_k, \mu_k(s_k), s_{k+1})| \leq M \quad \forall s_k, s_{k+1} \in S, \mu_k(s_k) \in A(s_k)$

$$|\text{term-LB})| \leq M \sum_{k=K}^{\infty} \alpha^k \leq \frac{\alpha^K M}{1-\alpha}$$

$$E \left(\sum_{k=0}^{K-1} \alpha^k g(s_k, \mu_k(s_k), s_{k+1}) \mid s_0 = \hat{i} \right)$$

$$= J_{\pi}(\hat{i}) - \text{term}(B)$$

$$E \left(\alpha^K J(s_K) + \sum_{k=0}^{K-1} \alpha^k g(s_k, \mu_k(s_k), s_{k+1}) \mid s_0 = \hat{i} \right)$$

$$= J_{\pi}(\hat{i}) - \text{term}(B) + E(\alpha^K J(s_K))$$

$$\begin{aligned}
& J_{\pi}(i) - \frac{\alpha^k M}{1-\alpha} - \alpha^k \max_j |J(j)| \\
& \leq \mathbb{E} \left[\alpha^k J(s_k) + \sum_{k=0}^{k-1} \alpha^k g(s_k, \mu_k(s_k), s_{k+1}) \mid s_0 = i \right] \\
& \leq J_{\pi}(i) + \frac{\alpha^k M}{1-\alpha} + \alpha^k \max_j |J(j)|
\end{aligned}$$

Take min over π on both sides,

$$\begin{aligned}
J^*(i) - \frac{\alpha^k M}{1-\alpha} - \alpha^k \max_j |J(j)| & \leq (T^k J)(i) \\
& \leq J^*(i) + \frac{\alpha^k M}{1-\alpha} + \alpha^k \max_j |J(j)|
\end{aligned}$$

Claim follows by taking $k \rightarrow \infty$ on both sides, i.e., $\lim_{k \rightarrow \infty} (T^k J)(i) = J^*(i)$ $\square$

Value iteration

Fix J_0 .

Repeatedly apply T i.e.,

$$J_0 \xrightarrow{T} J_1 = TJ_0 \xrightarrow{T} J_2 = TJ_1 + \dots \rightarrow J^*$$

Corollary: For any stationary policy π ,

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

where J is
any bounded
 n -vector.

Bellman optimality equation

from the proof above,

$$J^*(i) - \frac{\alpha^k M}{1-\alpha} - \alpha^k \max_j |J(j)| \leq (T^k J)(i) \leq J^*(i) + \frac{\alpha^k M}{1-\alpha} + \alpha^k \max_j |J(j)|$$

Apply T to both sides

$$\begin{aligned} (TJ^*)(i) - \frac{\alpha^{k+1} M}{1-\alpha} - \alpha^{k+1} \max_j |J(j)| &\leq (T^{k+1} J)(i) \\ &\leq (TJ^*)(i) + \frac{\alpha^{k+1} M}{1-\alpha} + \alpha^{k+1} \max_j |J(j)| \end{aligned}$$

Taking $k \rightarrow \infty$ on both sides, we obtain

$$\boxed{J^*(i) = TJ^*(i), \forall i}$$

Uniqueness: Suppose J' satisfies $J' = TJ'$

$$J' = TJ' = T^2 J' = \dots = \lim_{k \rightarrow \infty} T^k J' = J^*$$

Corollary: For any Stationary π ,

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

Illustration of Value Iteration

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

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

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

$$J_0 = (0, 0), \quad \alpha = 0.1$$

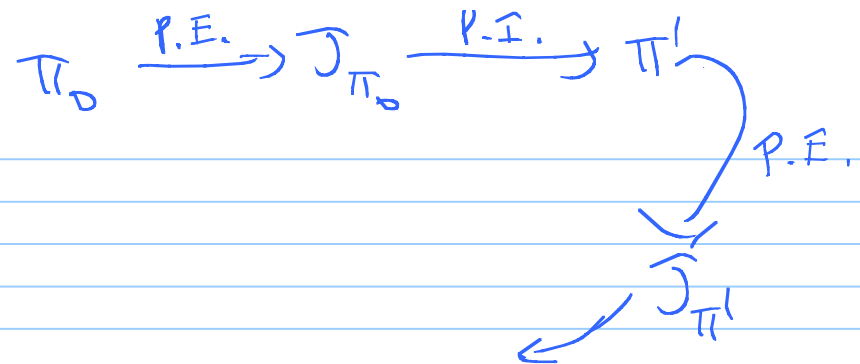
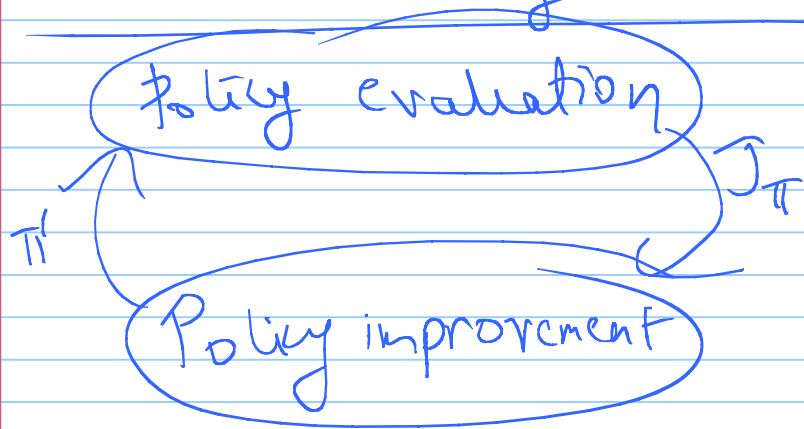
Do Value Iteration for 2 steps.

$$J_1 = (0.5, 1)$$

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

$$J_2 = (1.287, 1.562)$$

Prelude to policy iteration



$\pi^* \leftarrow \text{finite steps}$

Claim: A policy π is optimal if and only if $T_\pi(i)$ attains the minimum in the Bellman equation, $\forall i \in S$, i.e.,

$$\boxed{T_\pi J^* = T J^*} \quad \text{i.e.,} \quad g(i, \pi(i)) + \alpha \sum_{j=1}^n P_{ij}(\pi(i)) J^*(j) = \min_a g(i, a) + \alpha \sum_{j=1}^n P_{ij}(a) J^*(j), \quad \forall i \in S$$

Pf. Assume $TJ^* = T_\pi J^*$

we know $J^* = TJ^*$

So, $J^* = TJ^* = T_\pi J^* \Rightarrow J^* = J_\pi \Rightarrow \pi$ is optimal.

Converse: π is optimal $\Rightarrow J^* = J_\pi \Rightarrow J^* = T_\pi J^*$

From BE, $J^* = TJ^*$, So, $\boxed{T_\pi J^* = TJ^*}$

Example:- Machine replacement

$S = \{1, \dots, n\}$ $g(1) \leq \dots \leq g(n)$ P_{ij} : transition probabilities

Actions: Do nothing, repair

Bellman equation $J^*(i) = \min \left(\underbrace{R + g(i)}_{\text{Repair}}, \underbrace{g(i) + \alpha \sum_{j=1}^n P_{ij} J^*(j)}_{\text{Do nothing}} \right)$

Assume $p_{ij} \leq p_{(i+1)j}$ if $i < j$, Also, $p_{ij} = 0$ if $j < i$

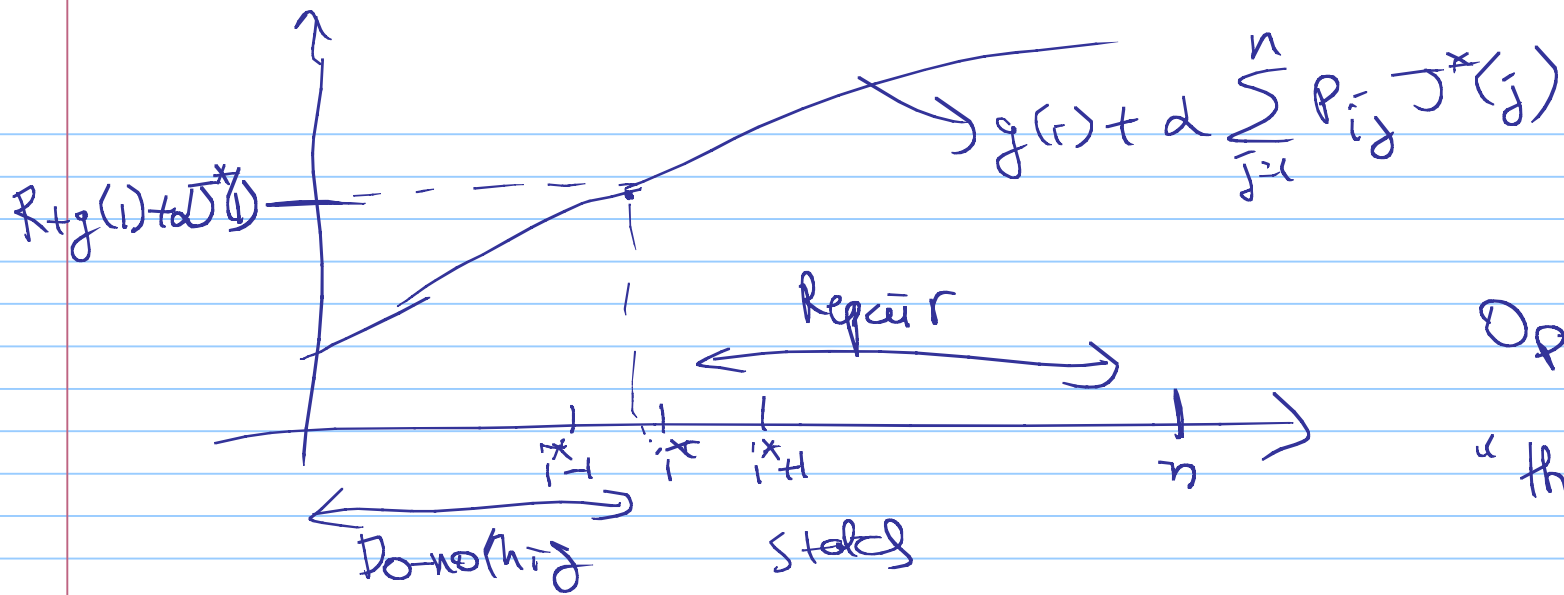
So,
$$\sum_{\bar{j}=1}^n p_{i\bar{j}} J(\bar{j}) \leq \sum_{\bar{j}=1}^n p_{(i+1)\bar{j}} J(\bar{j}), \quad i=1, \dots, n-1$$

for J that is monotonically non-decreasing

$$\sum_{\bar{j}=1}^n p_{i\bar{j}} (T^k J)(\bar{j}) \leq \sum_{\bar{j}=1}^n p_{(i+1)\bar{j}} (T^k J)(\bar{j})$$

So, $J^*(i) = \lim_{k \rightarrow \infty} (T^k J)(i)$ is monotonically non-decreasing in " i "

So, $g(i) + \alpha \sum_{\bar{j}=1}^n p_{i\bar{j}} J^*(\bar{j})$ is non-decreasing in " i ".



Optimal policy is
"threshold" based

Policy iteration algorithm

Step 1: Fix an initial policy π_0

Step 2: Find the cost-to-go J_{π_k} corresponding to π_k } Policy Evaluation

Step 3: Obtain a new policy π_{k+1} by solving

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

} Policy Improvement

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

Proposition: Let π and π' be two stationary policies that satisfy

$$T_{\pi'} J_{\pi} = T J_{\pi} \text{ or, } \\ g(i, \pi'(i)) + \alpha \sum_{j=1}^n P_{ij}(\pi'(i)) J_{\pi}(j) = \min_a \left(g(i, a) + \alpha \sum_{j=1}^n P_{ij}(a) J_{\pi}(j) \right),$$

Then, $J_{\pi'}(i) \leq J_{\pi}(i), i=1 \dots n$

Further, if π is not optimal, then $J_{\pi'}(i) < J_{\pi}(i)$ for at least one i .

Proof: $J_{\pi} = T_{\pi} J_{\pi}$. Also, $T_{\pi'} J_{\pi} = T J_{\pi}$

$$J_{\pi}(i) = g(i, \pi(i)) + \alpha \sum_{j=1}^n P_{ij}(\pi(i)) J_{\pi}(j)$$

$$\geq g(i, \pi'(i)) + \alpha \sum_{j=1}^n P_{ij}(\pi'(i)) J_{\pi}(j) = (T_{\pi'} J_{\pi})(i)$$

$$J_{\pi} \geq T_{\pi'} J_{\pi} \geq \dots \geq T_{\pi'}^k J_{\pi} \geq \dots \geq \lim_{k \rightarrow \infty} T_{\pi'}^k J_{\pi} = J_{\pi'}$$

Uniqueness: Skipped.

Policy iteration: Illustration

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

$$P(a) = \begin{bmatrix} \frac{3}{4} & \frac{1}{4} \\ \frac{3}{4} & \frac{1}{4} \end{bmatrix} \quad P(b) = \begin{bmatrix} \frac{1}{4} & \frac{3}{4} \\ \frac{1}{4} & \frac{3}{4} \end{bmatrix}$$

$$g(1, a) = 2, \quad g(1, b) = 0.5, \quad g(2, a) = 1, \quad g(2, b) = 3, \quad \gamma = 0.9$$

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

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

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

Using MDP data, we have

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

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

So, solving, we obtain $J_{\pi_0}(1) = 24.12, J_{\pi_0}(2) = 25.96$

Policy improvement

$$T_{\pi_1} J_{\pi_0} = T J_{\pi_0}$$

$$\begin{aligned} (T J_{\pi_0})(1) &= \min \left\{ \underbrace{2 + 0.9 \left(\frac{3}{4} \times 24.12 + \frac{1}{4} \times 25.96 \right)}_{\text{action a}}, \underbrace{0.5 + 0.9 \left(\frac{1}{4} \times 24.12 + \frac{3}{4} \times 25.96 \right)}_{\text{action b}} \right\} \\ &= \min(24.12, 23.45) = 23.45 \end{aligned}$$

$$\begin{aligned} (T J_{\pi_0})(2) &= \min \left(1 + 0.9 \left(\frac{3}{4} \times 24.12 + \frac{1}{4} \times 25.96 \right), 3 + 0.9 \left(\frac{1}{4} \times 24.12 + \frac{3}{4} \times 25.96 \right) \right) \\ &= \min(23.12, 25.95) = 23.12 \end{aligned}$$

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

Policy evaluation:-

$$J_{\pi_1}(1) = 7.33, \quad J_{\pi_1}(2) = 7.67$$

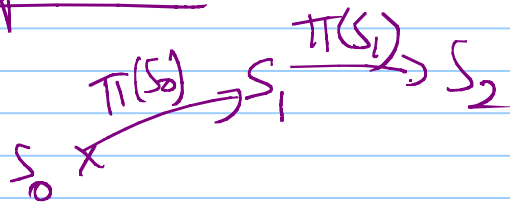
Policy improvement:-

$$\pi_2 = \pi_1$$

Tomorrow:- Reinforcement learning (RL)

Task: Policy evaluation π or π

Sample path:



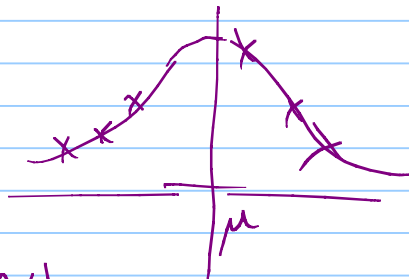
s_1 sampled from $P_{s_0 s_1}(\pi(s_0))$

TD-learning $\rightarrow$ Policy evaluation & Q-learning $\rightarrow$ Control
(to find the optimal policy)

Today: Reinforcement Learning

Mean estimation:- Random variable $\mathcal{V}$ with mean μ and finite variance.

$\{\mathcal{V}_1, \mathcal{V}_2, \dots, \mathcal{V}_n\}$ i.i.d. Samples from the distribution of $\mathcal{V}$.



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

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

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

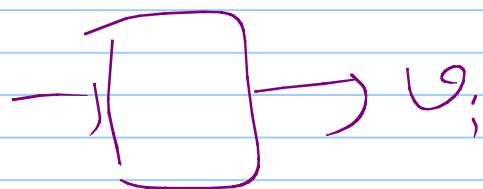
More generally, $r_{n+1} = r_n + \eta_n (\mathcal{V}_{n+1} - r_n),$

where $\sum_n \eta_n = \infty$, $\sum_n \eta_n^2 < \infty$

η_n : step-size

Then, $r_n \rightarrow \mu$ w.p.1 as $n \rightarrow \infty$

Want to estimate $E[v]$

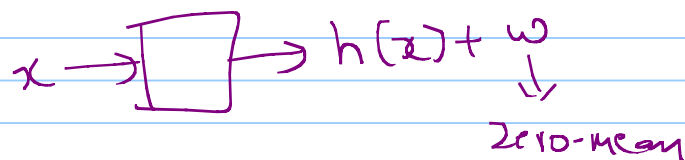


$$r_{n+1} = r_n + \eta_n (v_{n+1} - r_n)$$

$$\begin{aligned} h(x) &= f(x) - x \\ x_{n+1} &= x_n + \eta_n (f(x_n) + \omega_n - x_n) \\ &= (1 - \eta_n)x_n + \eta_n (f(x_n) + \omega_n) \end{aligned}$$

Robbins-Monro
Stochastic
approximation
Scheme (1951)

Generally, if you want to solve $h(x) = 0$



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

$x_n \rightarrow x^*$ w.p.1, where $h(x^*) = 0$
assuming $E\omega_n = 0$, $E\omega_n^2 < \infty$

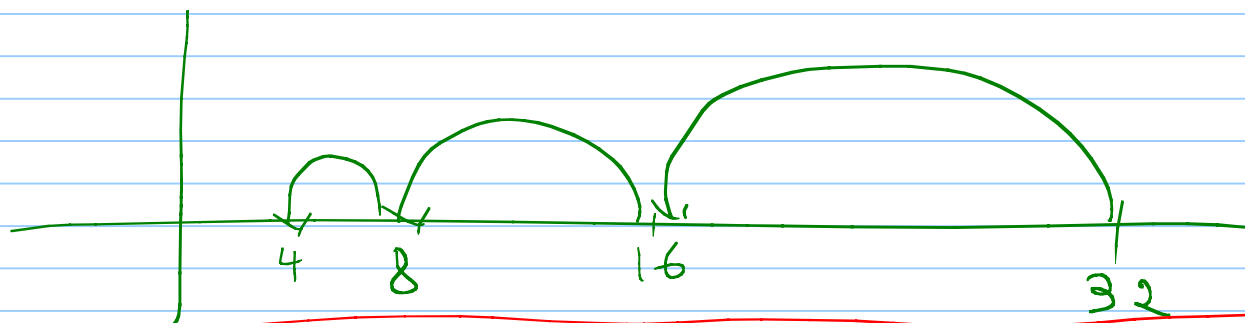
Assuming H is "well-behaved", & noise w_n is zero-mean $\oplus$ bounded variance,
 One can show $r_n \xrightarrow[n \rightarrow \infty]{} r^*$ w.p.1, where $\boxed{Hr^* = r^*}$

Brief introduction to contraction mappings:-

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

$$\boxed{f(0) = 0}$$

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



Def: $f: \mathbb{R} \rightarrow \mathbb{R}$ is a contraction mapping with modulus β if
 $|f(x) - f(y)| \leq \beta |x - y|$, for any x, y
 with $0 < \beta < 1$

Contraction properties of T, T_π in a discounted MDP setting:

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

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

Proposition:- For any two bounded J, J' , we have

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

↓
discount factor

$$\|T_\pi J - T_\pi J'\|_\infty \leq \alpha \|J - J'\|_\infty, \text{ for any stationary policy } \pi.$$

Further, $\forall k$, $\|T^k J - T^k J'\|_\infty \leq \alpha^k \|J - J'\|_\infty$ &

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

Corollary:

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

Stochastic iterative algorithm for finding fixed point of a contraction map:

Goal: - Solve $Hx^* = x^*$, $H: \mathbb{R}^n \rightarrow \mathbb{R}^n$, $x \in \mathbb{R}^n$

Fixed point iteration: - $x_{t+1}(i) = (1 - \eta_t) x_t(i) + \eta_t \left((Hx_t)(i) + w_t(i) \right)$

Assumptions: $\mathcal{F}_t = \{x_0(i), \dots, x_t(i), w_0(i), \dots, w_{t-1}(i), i=1, \dots, n\}$ $i=1, \dots, n$

history

(A1)

$$E[w_t(i) | \mathcal{F}_t] = 0, \quad E[w_t^2(i) | \mathcal{F}_t] \leq A + B \|x_t\|^2$$

noise,

(A2) $\sum_t \eta_t = \infty$, $\sum_t \eta_t^2 < \infty$. e.g. $\eta_t = \frac{c}{t}$

(A3) H is a max-norm pseudo-contraction, i.e.,

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

where $0 < \beta < 1$ & $Hr^* = r^*$.

Under (A1) - (A3)

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

Full-contraction

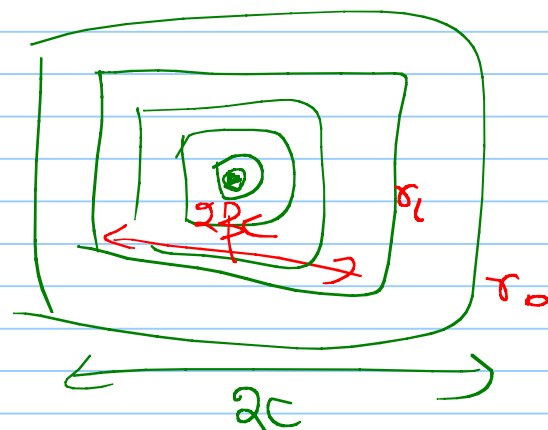
$$\begin{aligned} \|Hr - Hr'\|_\infty \\ \leq \beta \|r - r'\|_\infty \\ \forall r, r' \end{aligned}$$

$$r_{t+1}(i) = (1 - \eta_n) r_t(i) + \eta_n (H r_t)(i) + w_t(i)$$

$$\boxed{H0=0} \quad r^* \geq 0$$

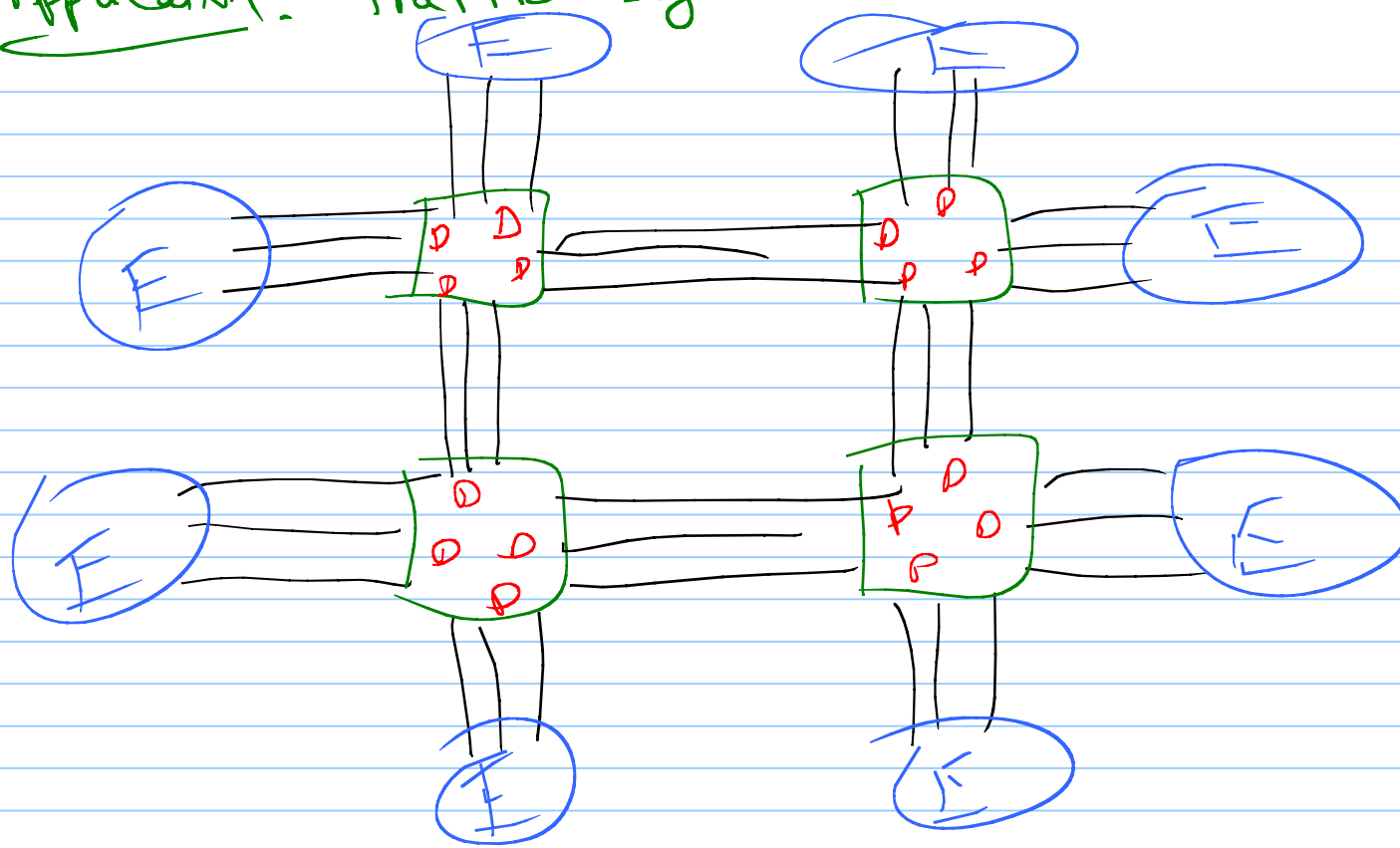
$$|r_0(i)| \leq C$$

$$r_{t+1}(i) = (H r_t)(i)$$

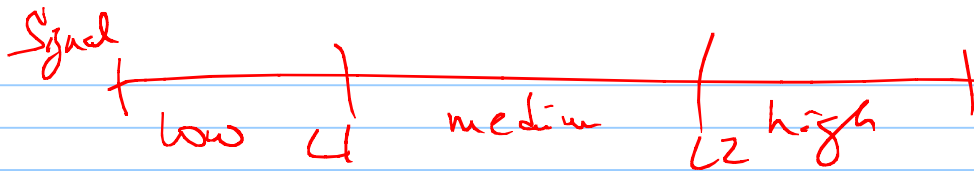


Now, with $r_{t+1}(i) = (1 - \eta_t) r_t(i) + \eta_t (H r_t)(i) + w_t(i)$
 Similar behaviour, but shrinkage is not immediate

Application: Traffic signal control



Discounted MDP



N lanes with signal (cross junction)

State = $s_t = (q_1(t), \dots, q_N(t), \underbrace{t_1(t), \dots, t_N(t)}_{\text{elapsed times}})$

Actions = feasible sign configurations

$$\text{Single stage cost } z = g(s_t, a_t) = \alpha_1 \left(\sum_{i=1}^N q_i(t) \right) + \alpha_2 \left(\sum_{i=1}^N t_i(t) \right)$$

normalized to be in $[0,1]$

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:-

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

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

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

$$\begin{aligned} & \text{Proxy for} \\ & E(g(i, \hat{i}) + \alpha J_t(\hat{i})) \\ & = (T_\pi J_t)(i) \end{aligned}$$

$$\begin{aligned} J_{t+1}(i) = J_t(i) + \eta_t \big(& (T_\pi J_t)(i) - J_t(i) \\ & + \underbrace{g(i, \hat{i}) + \alpha J_t(\hat{i}) - (T_\pi J_t)(i)}_{w_t(i)} \big) \end{aligned}$$

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

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

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

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

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

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

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

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

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

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

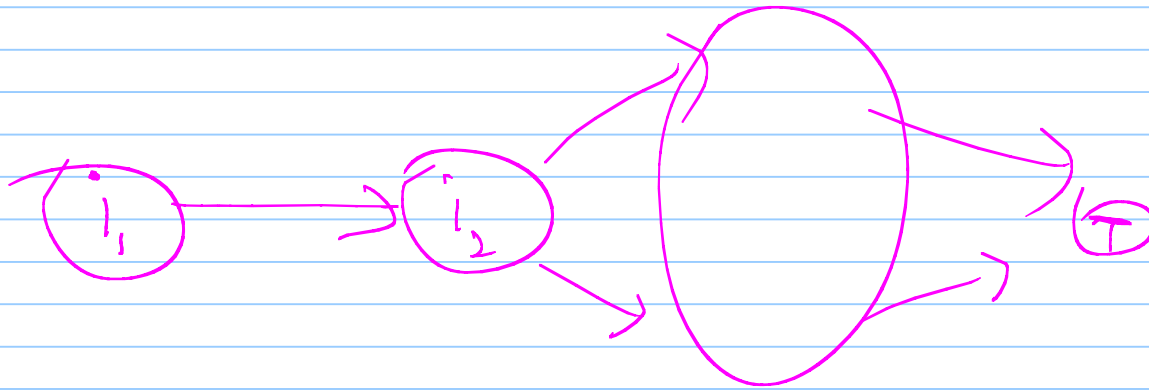
TD(i):- Simulate episodes of the underlying SSP w.r.t π .

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

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

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

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 $J_\pi(i_1)$

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 $\mathcal{H}$

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

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

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

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

$Q_t \rightarrow Q^*$ 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

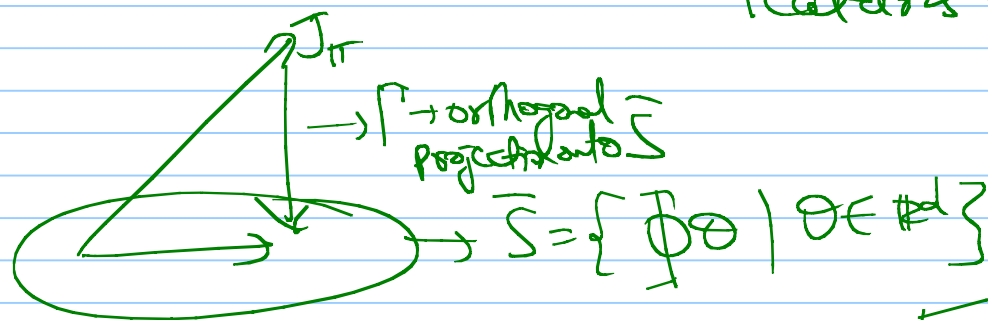
Linear function approximation

Fix a policy π .

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

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

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



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

Cannot solve

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

Approximate solution

$$\bar{J} \theta = \Gamma(T_{\pi} \bar{J} \theta)$$

$\rightarrow$ has a unique solution since ΓT_{π} 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 $\min_a \Theta_t^T \phi(i, a)$