

5. ALGORITHMS FOR MDPs

THINK OF AN ALGORITHM HAVING TWO STEPS

1. A POLICY IMPROVEMENT STEP: eg.

$$\hat{\pi}(x) \in \underset{a}{\operatorname{ARGMAX}} \quad r(x, a) + \beta \mathbb{E}_{x, a} [R(\hat{\pi}, \pi)]$$

2. A POLICY EVALUATION STEP: eg.

FIND

$$R(x, \pi) \doteq \mathbb{E}_x \left[\sum_{t=0}^{\infty} \beta^t r(X_t, \pi(X_t)) \right]$$

WE COVER TWO ALGORITHMS

1. VALUE ITERATION

2. POLICY ITERATION

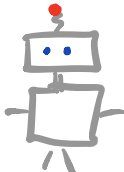
VALUE ITERATION

TAKE $V_0(x) = 0 \quad \forall x$

UNTIL CONVERGENCE DO:

$$V_{t+1}(x) = \max_a r(x, a) + \beta \mathbb{E}_{x, a} [V_t(s')]]$$

VALUE ITERATION EXAMPLE

END 1						END 2
----------	--	--	---	--	--	----------

$t=0$	0	0	0	0	0
$t=1$	1	0	0	0	2
$t=2$	1	1	0	2	2
$t=3$	1	1	2	2	2
$t=4$	1	2	2	2	2
$t=5$	2	2	2	2	2
$t=6$	2	2	2	2	2

SOME CODE :

```
def Value_Iteration(V,P,r,discount):  
    ''' Value Iteration - a numerical solution to a MDP  
  
    # Arguments:  
        P - P[a][x][y] gives probability of x -> y for action a  
        r - r[a][x][y] gives reward for x -> y for action a  
        V - V[x] gives value for state x  
        discount - a float. discount factor  
  
    # Returns:  
        Value function and policy from **one** value iteration  
    ...  
  
    number_of_actions = len(P)  
    number_of_states = len(P[0])  
  
    Q = np.zeros((number_of_actions, number_of_states))  
  
    for _ in range(time):  
        for a in range(number_of_actions):  
            for x in range(number_of_states):  
                Q[a][x] = np.dot(P[a][x], r[a][x]+discount*V)  
            V_new = np.amax(Q, axis=0)  
  
    pi = np.argmax(Q, axis=0)  
  
    return V_new, pi
```



A RESULT:

Thrm 59. For positive programming, i.e. where all rewards are positive and the discount factor β belongs to the interval $(0, 1]$, then

$$0 \leq V_s(x) \leq V_{s+1}(x) \nearrow V(x), \quad \text{as } s \rightarrow \infty.$$

Here $V(x)$ is the optimal value function.

PROOF IS SAME AS ~~THE~~ PROGRAMMING.

POLICY ITERATION:

Def 60 (Policy Iteration). Given the stationary policy Π , we may define a new (improved) stationary policy, $I\Pi$, by choosing for each x the action $I\Pi(x)$ that solves the following maximization

$$I\Pi(x) \in \operatorname{argmax}_{a \in \mathcal{A}} r(x, a) + \beta \mathbb{E}_{x, a} [R(\hat{X}, \Pi)]$$

UNTIL CONVERGENCE DO

$$\Pi_{n+1} = I \Pi_n$$

FINDING $R(x, \pi)$:

WE KNOW FROM MARKOV CHAINS (LECTURE 2)
THAT :

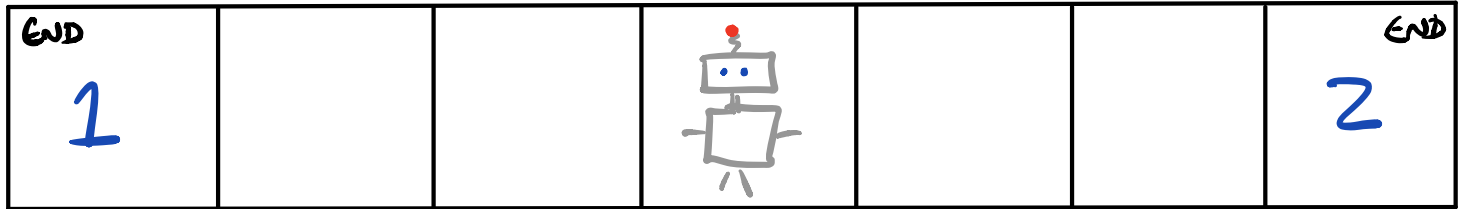
$$R(x) = r(x) + \beta \mathbb{E}_{\pi} R(\hat{x})$$

HERE $R(x) = R(x, \pi)$, $r(x) = r(x, \pi)$, $P_{xy} = P_{xy}^{\pi(x)}$.

INTERPRET AS VECTORS & MATRICES.

$$\underline{R} = \underline{r} + \beta \underline{P} \underline{R} \quad (\Leftrightarrow) \quad \underline{R} = (\underline{I} - \beta \underline{P})^{-1} \underline{r}$$

POLICY ITERATION EXAMPLE



$t=0$ $\leftarrow$ $\rightarrow$ $\leftarrow$ $\rightarrow$ $\leftarrow$

$t=1$ $\leftarrow$ $\leftarrow$ $\leftarrow$ $\rightarrow$ $\rightarrow$

$t=2$ $\leftarrow$ $\leftarrow$ $\rightarrow$ $\rightarrow$ $\rightarrow$

$t=3$ $\leftarrow$ $\rightarrow$ $\rightarrow$ $\rightarrow$ $\rightarrow$

$t=4$ $\rightarrow$ $\rightarrow$ $\rightarrow$ $\rightarrow$ $\rightarrow$

$t=5$ $\rightarrow$ $\rightarrow$ $\rightarrow$ $\rightarrow$ $\rightarrow$

SOME CODE :

```
def Policy_Iteration(pi,P,r,discount):  
    ''' Policy Iteration - a numerical solution to a MDP  
    IS JUST MATRIX ALGEBRA  
    # Arguments:  
        P - P[a][x][y] gives probability of x -> y for action a  
        r - r[a][x][y] gives reward for x -> y for action a  
        pi - pi[x] gives action for state x  
        discount - discount factor  
  
    # Returns:  
        policy from **one** policy iteration  
        value function of input policy  
        ...  
  
    # Collate array of states and actions  
    number_of_actions, number_of_states = len(P), len(P[0])  
    Actions, States = np.arange(number_of_actions), np.arange(  
        number_of_states)  
  
    # Get transitions and rewards of policy pi  
    P_pi = np.array([P[pi[x]][x] for x in States])  
    r_pi = np.array([r[pi[x]][x] for x in States])  
    Er_pi = [ np.dot(P_pi[x], r_pi[x]) for x in States]  
  
    # Calculate Value of pi  
    I = np.identity(number_of_states)  
    A = I - discount * P_pi  
    R_pi = np.linalg.solve(A, Er_pi)  
  
    # Calculate Q_factors of pi  
    Q = np.zeros((number_of_actions, number_of_states))  
    for a in range(number_of_actions):  
        for x in range(number_of_states):  
            Q[a][x] = np.dot(P[a][x], r[a][x]+discount*R_pi)  
  
    # policy iteration update  
    pi_new = np.argmax(Q, axis=0)  
  
    return pi_new, R_pi
```

A RESULT:

Thrm 61. Under Policy Iteration

$$R(x, \Pi_{n+1}) \geq R(x, \Pi_n)$$

and, for bounded programming,

$$R(x, \Pi_n) \nearrow V(x) \quad \text{as } n \rightarrow \infty$$

PROOF:

FROM MARKOV CHAINS



$$R(x) = \beta(PR)(x) + r(x), \quad x \in \mathcal{X}.$$

RECALL:

Moreover, if function $\tilde{R} : \mathcal{X} \rightarrow \mathbb{R}_+$ satisfies

$$\tilde{R}(x) \stackrel{\leq}{\geq} \beta(P\tilde{R})(x) + r(x), \quad x \in \mathcal{X}.$$

then $\tilde{R}(x) \stackrel{\leq}{\geq} R(x), x \in \mathcal{X}.$

NOTICE π FOR $\hat{\pi}$

CURRENT

NEXT

BY DEF

$$R(x, \pi) = r(x, \pi(x)) + \beta \mathbb{E}_{x, \pi} [R(x, \pi)] \leq r(x, \hat{\pi}) + \beta \mathbb{E}_{x, \hat{\pi}} [R(x, \hat{\pi})]$$

$\therefore R(x, \pi) \leq R(x, \hat{\pi}) \quad (+) \checkmark$ IS INCREASING.

PROOF (CONT)

FURTHER NOTICE $\forall a$

$$r(x, a) + \beta \mathbb{E}_{x, a} [R(\hat{x}, \pi)]$$

← BY DEF OF
 $\mathcal{L}_\pi$
 ↙

$$\leq r(x, \mathcal{L}_\pi(x)) + \beta \mathbb{E}_{x, \mathcal{L}_\pi(x)} [R(\hat{x}, \pi)]$$

$$\leq r(x, \mathcal{L}_\pi(x)) + \beta \mathbb{E}_{x, \mathcal{L}_\pi(x)} [R(\hat{x}, \mathcal{L}_\pi)]$$

↖ BY (7)
 ↙

$$= R(x, \mathcal{L}_\pi) \quad (\#)$$

TO PROVE CONVERGENCE TO OPTIMAL POLICY

LET

$$M_t = \sum_{s=0}^{t-1} \beta^s r(x, \pi^*(x)) + \beta^t R(x, \pi_{T-t})$$

EXPECTATION
UNDER
OPTIMAL
POLICY

DO OPTIMAL POLICY

THEY DO π_{T-t}

BY (H)

$$\mathbb{E}^*[M_{t+1} - M_t | \mathcal{F}_t] = \beta^t \mathbb{E}^* \left[\underbrace{\beta R(x_{t+1}, \pi_{T-t-1}) + r(x_t, \pi^*) - R(x_t, \pi_{T-t})}_{\leq 0} | \mathcal{F}_t \right]$$

BECAUSE POLICY π_{T-t} IS BETTER $\forall$ ACTIONS
UNTIL $R(\cdot, \pi_{T-t-1})$

$\therefore M_t$ IS SUPER M_0 .

$$\therefore R(x, \pi_t) = \mathbb{E}[M_0] \geq \mathbb{E}[M_T]$$

$$\geq R_T(x, \pi^*) \rightarrow R(x, \pi^*) \quad \square$$

LINEAR PROGRAMMING APPROACH:

WE CAN RECAST AS A LINEAR OPTIMIZATION
HERE THERE ARE LOTS OF ALGORITHMS:

- SIMPLEX ALGORITHM
- PROJECTED GRADIENT DESCENT
- INTERIOR POINT METHODS
- ...

- WE WANT TO FIND

$$V(x) = \max_{a \in A} r(x, a) + \beta \mathbb{E}_{\hat{x}|x,a} [V(\hat{x})]$$

- i.e. $(V(x) : x \in X)$ IS THE SMALLEST FUNCTION/VECTOR SUCH THAT

$$V(x) \geq r(x, a) + \beta \sum_{\hat{x}} P(\hat{x}|x, a) V(\hat{x})$$

- SO SOLVE

$$\min_{V \in \mathbb{R}^X} \sum_x \gamma(x) V(x)$$

$$\text{s.t.} \quad V(x) \geq r(x, a) + \beta \sum_{\hat{x}} P(\hat{x}|x, a) V(\hat{x}) \quad \forall x, a \quad (\text{PDP})$$

LINEAR PROGRAMMING (DUAL)

THE DUAL OF THE LINEAR PROGRAM (PDP) IS

$$\text{MAX}_{\gamma \geq 0} \sum_x \gamma(x, a) r(x)$$

$$\text{s.t.} \quad \sum_{\hat{a}} \gamma(\hat{x}, \hat{a}) = \{(\hat{x}) + \beta \sum_{x, a} \gamma(x, a) P(\hat{x} | x, a) \quad \forall \hat{x}$$

HERE $\gamma(\hat{x}, \hat{a})$ CAN BE INTERPRETED

AS THE OCCUPANCY MEASURE OF THE DYNAMIC PROGRAM

PROOF: THE LAGRANGIAN IS ≥ 0 FOR COMPLEMENT SACKING

$$L(x; \gamma) = \sum_x \xi(x) v(x) - \sum_{x,a} \gamma(x,a) \left(v(x) - r(x,a) - \beta \sum_{\hat{x}} P(\hat{x}|x,a) v(\hat{x}) \right)$$

$$= \sum_x v(x) \left[\xi(x) - \underbrace{\sum_a \gamma(x,a) + \beta \sum_{\hat{x},a} \gamma(\hat{x},a) P(x|\hat{x},a)}_{\text{THIS } \geq 0 \text{ FOR A WELL DEFINED LAGRANGIAN}} \right] + \sum_{x,a} \gamma(x,a) r(x,a)$$

THIS ≥ 0 FOR A WELL DEFINED
LAGRANGIAN

$\therefore$ THE DUAL IS

$$\max \sum_{x,a} \gamma(x,a) r(x,a)$$

$$\text{s.t. } \xi(x) + \beta \sum_{\hat{x},a} \gamma(\hat{x},a) P(x|\hat{x},a) = \sum_a \gamma(x,a)$$

□

