

FINITE STATE MARKOV CHAINS

$$W_t = W_t(\omega) \quad \omega \in (\Omega, \mathcal{F}, \mathbb{P})$$

MBP. 1

$$X_{t+1} = f(X_t, W_t)$$

discrete time stochastic systems

$$X_0 = X$$

X_t = state

exogenous input

W_t = stochastic process $\Rightarrow X_t$ = stochastic process

Hp: W_t is an i.i.d sequence

X_t = random variable $\forall t$

usually, $X_t \in \mathbb{R}^M$, suppose instead that $\forall t$

$$X_t \in \mathcal{X} \text{ a finite set w.l.o.g. } X_t \in \{1, 2, \dots, M\}$$

$\Rightarrow X_t$ = finite state Markov chain

$X_{t+1} = f(X_t, W_t)$ induces a distribution over $\{X_t\}$ and in particular the so called transition probabilities

$$P(i|j) = \mathbb{P}(X_{t+1}=i | X_t=j) \quad \& \text{ invariant}$$

transition probabilities enough to characterize the evolution of the Markov chain
 $\downarrow$ thanks to

"the future is independent of the past given the present"

$$X_{t+1} = f(X_t, W_t) \text{ i.i.d.} \Rightarrow$$

MARKOV property:

the past given the present

$\hookrightarrow$ depends on X_t and W_t only

$$P(X_{t+1}=i | X_t=j \wedge X_{t-1}=k \wedge \dots) = P(X_{t+1}=i | X_t=j)$$

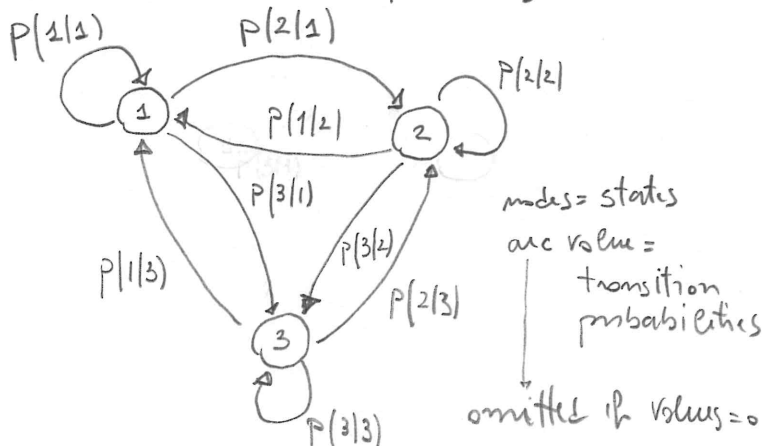
Clearly, $\sum_{i=1}^M P(i|j) = 1$ and note that $P(i|i) \neq 0$ in general

$[P(i|j)]_{M \times M}$ = transition matrix = $\mathbb{P}$ $\Leftarrow$ conveys the whole info

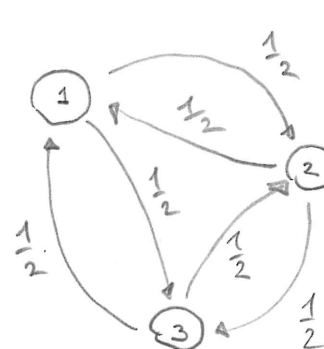
$$[P(X_t=i)]_n = \mathbb{P}^t \cdot [P(X_0=i)] \quad \text{for example}$$

example

$$\mathcal{X} = \{1, 2, 3\}$$



e.g.



$\mathbb{P} =$

$$\begin{bmatrix} 0 & \frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & 0 & \frac{1}{2} \\ \frac{1}{2} & \frac{1}{2} & 0 \end{bmatrix}$$

$$x_{t+1} = f(x_t, u_t, w_t)$$

↳ control action

$w_t = \text{i.i.d. sequence}$

(finite state/input)
Markov Decision
Process - MDP
(controlled
Markov chain)

$$\forall t, x_t \in X = \{1, \dots, n\} = \{x_1, x_2, \dots, x_n\}$$

$$u_t \in U = \{1, \dots, m\} = \{u_1, u_2, \dots, u_m\}$$

transition probabilities

$$P(i|j, u) = P\{x_{t+1}=i | x_t=j \wedge u_t=u\}$$

$$i, j \in X \quad u \in U$$

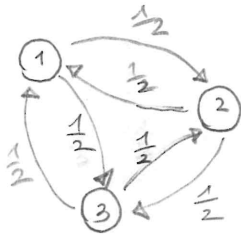
MARKOV property

$$\downarrow \\ P\{x_{t+1}=i | x_t=j \wedge u_t=u \wedge x_{t-1}=k \wedge u_{t-1}=l \wedge \dots\}$$

→ m transition matrices $P_u = [P(i|j, u)]_{n \times n}$

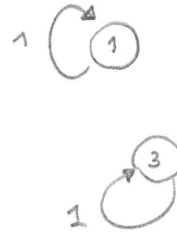
e.g. $X = \{1, 2, 3\} \quad U = \{1, 2\}$

$u=1$



$$P_1 = \begin{bmatrix} 0 & \frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & 0 & \frac{1}{2} \\ \frac{1}{2} & \frac{1}{2} & 0 \end{bmatrix}$$

$u=2$



$$P_2 = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

u_t to be chosen so as to impose a desired behavior → optimal control problem

⇒ EXAMPLE

every transition from x_t to x_{t+1} generates a (stage) cost

$g(x_t, u_t, w_t) \rightarrow$ e.g. $g(x_{t+1}) = g(f(x_t, u_t, w_{t+1}))$ but here more general to represent various situations

use the degree of freedom so as to minimize a "total" (average) cost
↓
to be precisely defined later

Example Inventory control

x_t = size of inventory at the end of day t (no. of items)

u_t = no. of items ordered at the end of day t

w_t = demand on day $t+1$

$x_t \in \{0, 1, \dots, M\}$ $M = \text{max no. of items that can be stored}$ MDP.3

$u_t \in \{0, 1, \dots, M\}$ ^{when $x_t = 0$} to pick up inventory when $x_t = 0$

$w_t \in \mathbb{Z}$

$$x_{t+1} = \min(\max(x_t + u_t, M) - w_t, 0)$$

$$g(x_t, u_t, w_t) = \underbrace{K \cdot \mathbb{1}_{u_t > 0}}_{\substack{\text{fixed cost} \\ \text{in placing an} \\ \text{order}}} + \underbrace{C \cdot u_t}_{\substack{\text{transportation} \\ \text{costs}}} + \underbrace{h \cdot x_t}_{\substack{\text{inventory} \\ \text{maintenance} \\ \text{costs}}} - \underbrace{p \cdot \min(\max(x_t + u_t, M), w_t)}_{\substack{\text{price} \\ \text{actually sold items,} \\ \text{income}}}$$

risk measure = expected value

expected stage cost : $E[g(x_t, u_t, w_t)]$

MDP. 4
we consider the evolution of the MDP up to time T

finite horizon T ← terminal cost

$$J(x_0, \{u_t\}_{t=0}^{T-1}) = \sum_{t=0}^{T-1} E[g(x_t, u_t, w_t)] + E[g_T(x_T)]$$

t-varying
↓
 $g = g_t$

OBS : for the finite horizon problem we may admit $f = f_t$ $g = g_t$
results are unchanged

→ to obtain good solutions, u_t must be stochastic and adapted to x_t

Hp : x_t is observable and u_t is chosen as a function of x_t
by means of a policy $\pi = \{\mu_0, \mu_1, \dots, \mu_{T-1}\}$ where
each $\mu_t: X \rightarrow U$. ← motivated by the Markov property

In other words, $u_t = \mu_t(x_t)$ (state feedback) so that

$$\begin{cases} x_{t+1} = f(x_t, \mu_t(x_t), w_t) \\ x_0 = x \end{cases}$$

$x \in X$ initial condition

$$\text{and } J(x, \{\mu_t(x_t)\}_{t=0}^{T-1}) = \sum_{t=0}^{T-1} E[g(x_t, \mu_t(x_t), w_t)] + E[g_T(x_T)]$$

Goal : $\min_{\pi} V^{\pi}(x)$ $= V^{\pi}(x)$ ← value function associated to policy π

well posed $\forall x \in X$ (finite no. of cases)

optimal value function $V^*(x) = \min_{\pi} V^{\pi}(x)$

as $x_0 = x$ deterministic, only $\mu_0(x)$ counts, while, as for $\mu_1, \dots, \mu_{T-1}$ potentially different x could lead to different π

⇒ there exists a unique policy $\pi^* = \{\mu_0^*, \mu_1^*, \dots, \mu_{T-1}^*\}$ such that

$$V^{\pi^*}(x) = V^*(x) \quad \forall x \in X \quad \hookrightarrow \text{simultaneously optimal for all initial conditions}$$

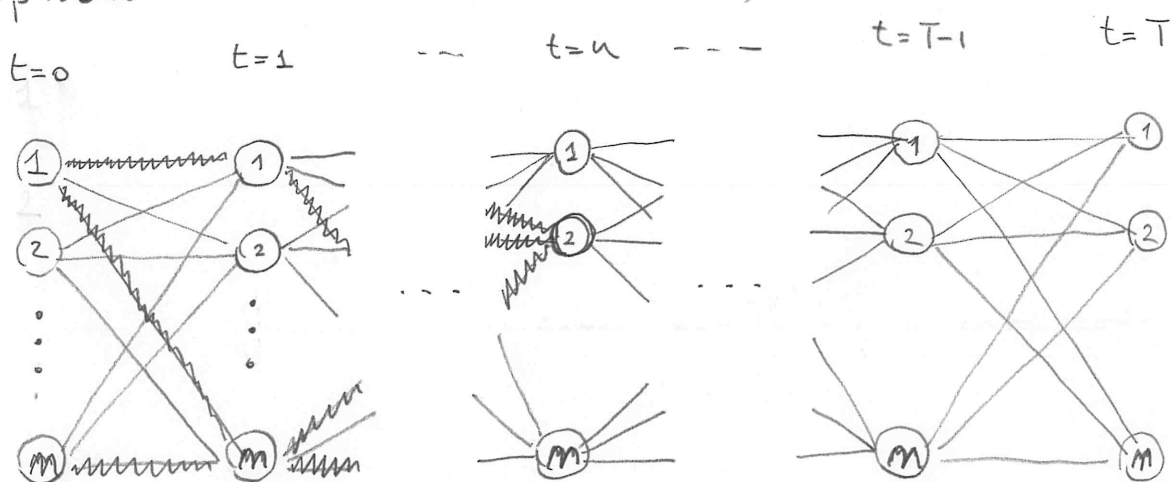
⇒ consequence of the Bellman optimality principle : given any $x \in X$, let

$\pi^* = \{\mu_0^*, \mu_1^*, \dots, \mu_{T-1}^*\}$ optimal for problem $\min_{\pi} V^{\pi}(x)$ and assume

that using π^* , $x_k = x'$ arises. Consider the sub-problem where

$x_k = x'$ and we wish to minimize the "cost-to-go" from $t=k$ to MDP.5
 $t=T$: $V_k^{\pi_k}(x') = \sum_{t=k}^{T-1} \mathbb{E}[g(x_t, \mu_t(x_t), w_t)] + \mathbb{E}[g_T(x_T)]$ where $x_k = x'$

and $\pi_k = \{\mu_k, \dots, \mu_{T-1}\}$. Then $\pi_k^* = \{\mu_k^*, \dots, \mu_{T-1}^*\}$ is optimal for this sub-problem



once at $t=k$ I'm arrived at $x_k=2$, the future evolution doesn't depend on the past (Markov), if I'm not optimal we can switch to an optimal policy and simultaneously improve all the situations

⇒ Dynamic Programming algorithm: break down multistage optimisation into single stage optimisations; generate backwards μ_{T-1}^* , solving the cost-to-go from $T-1$ to T subproblem for all $x \in X$, $x_{T-1}=x$, then μ_{T-2}^* , solving the cost-to-go from $T-2$ to T subproblem where μ_{T-1}^* is already implemented for all $x \in X$ and $x_{T-2}=x$, and so forth and so on till μ_0^*

$$\text{init: } V_T^*(x) = g_T(x) \quad \forall x \in X$$

for $t=T-1, \dots, 0$ compute $\forall x \in X$

$$V_t^*(x) = \min_{u \in U} \mathbb{E}[g(x, u, w_t) + V_{t+1}^*(f(x, u, w_t))] \quad (\star)$$

$$\mu_t^*(x) = \operatorname{argmin}_{u \in U} \mathbb{E}[g(x, u, w_t) + V_{t+1}^*(f(x, u, w_t))]$$

end

then $\pi^* = \{\mu_0^*, \dots, \mu_{T-1}^*\}$ is optimal for $\min_{\pi} V^{\pi}(x) \quad \forall x \in X$ and

$$V_0^*(x) = V^*(x).$$

$= T-1, \dots, 0$ MDP 6

→ FORMAL PROOF: it is enough to show that $\forall k$ and $\forall x \in X$

$$V_k^*(x) = \min_{\pi_k} V_k^{\pi_k}(x) = \min_{\pi_k} \sum_{t=k}^{T-1} \mathbb{E}[g(x_t, \mu_t(x_t), w_t)] + \mathbb{E}[g_T(x_T)]$$

as defined in (*) $x_k = x$

and that $\pi_k^* = \{\mu_k^*, \dots, \mu_{T-1}^*\}$ as defined in (*) is optimal for $\min_{\pi_k} V_k^{\pi_k}(x)$. This can be done by induction

$$\begin{aligned} \min_{\pi_{T-1}} V_{T-1}^{\pi_{T-1}}(x) &= \min_{\mu_{T-1}} \mathbb{E}[g(x, \mu_{T-1}(x), w_{T-1})] + \mathbb{E}[g_T(f(x, \mu_{T-1}(x), w_{T-1})))] = \\ &= \min_{\mu_{T-1}} \mathbb{E}[g(x, \mu_{T-1}(x), w_{T-1}) + V_T^*(f(x, \mu_{T-1}(x), w_{T-1})))] = \\ &= \min_{u \in U} \mathbb{E}[g(x, u, w_{T-1}) + V_T^*(f(x, \mu_{T-1}(x), w_{T-1})))] \end{aligned}$$

for $x \in X$, we have n distinct opt. problems

Clearly, $\mu_{T-1}^*(x) = \arg\min_{u \in U} \dots$ defines an optimal policy π_{T-1}^*

Suppose, property is true for $t = k+1$

$$\begin{aligned} \min_{\pi_k} V_k^{\pi_k}(x) &= \min_{\mu_k, \pi_{k+1}} \sum_{t=k}^{T-1} \mathbb{E}[g(x_t, \mu_t(x_t), w_t)] + \mathbb{E}[g_T(x_T)] = \\ &= \min_{\mu_k, \pi_{k+1}} \mathbb{E}[g(x, \mu_k(x), w_k)] + \sum_{t=k+1}^{T-1} \mathbb{E}[g(x_t, \mu_t(x_t), w_t)] + \mathbb{E}[g_T(x_T)] = \\ &= \min_{\mu_k, \pi_{k+1}} \mathbb{E}[g(x, \mu_k(x), w_k)] + \mathbb{E}[V_{k+1}^{\pi_{k+1}}(f(x, \mu_k(x), w_k))] = \\ &= \min_{\mu_k, \pi_{k+1}} \mathbb{E}[g(x, \mu_k(x), w_k) + V_{k+1}^{\pi_{k+1}}(f(x, \mu_k(x), w_k))] = \\ &\geq \min_{\mu_k} \mathbb{E}[g(x, \mu_k(x), w_k) + \min_{\pi_{k+1}} V_{k+1}^{\pi_{k+1}}(f(x, \mu_k(x), w_k))] = \\ &= \min_{\mu_k} \mathbb{E}[g(x, \mu_k(x), w_k) + V_{k+1}^*(f(x, \mu_k(x), w_k))] \rightarrow n \text{ distinct problems} \\ &= \min_{u \in U} \mathbb{E}[g(x, u, w_k) + V_{k+1}^*(f(x, \mu_k(x), w_k))] \rightarrow \arg\min \text{ defines a feedback } \mu_k^* \end{aligned}$$

$$= \min_{u \in \mathcal{U}} \mathbb{E} \left[g(x, u, w_n) + V_{n+1}^{\pi_{n+1}^*} (f(x, u, w_n)) \right] =$$

$$= \mathbb{E} \left[g(x, \mu_n^*(x), w_n) + V_{n+1}^{\pi_{n+1}^*} (f(x, \mu_n^*(x), w_n)) \right] = V_n^{\pi_n^*}(x)$$

altogether gives a policy π_n^*

$$\min_{\pi_n} V_n^{\pi_n}(x) \geq \underbrace{V_n^*(x)}_{\text{must be an equality}} = V_n^{\pi_n^*}(x)$$