

Reinforcement Learning

Prof. Christian Bauckhage



lecture 08

solving Bellman equations

in this lecture, we have a closer look at some of the concepts which we introduced last time

to begin with, we examine properties of *deterministic policies* (will be useful later)

then, we get to know *Bellman's principle of optimality* which allows us to scrutinize the *Bellman equation* for value functions

crucially, we will see that we are dealing with whole systems of equations and will learn how to *solve systems of Bellman equations*

this will then allow us to perform *policy evaluation*

outline

recap

working with deterministic policies

Bellman equations revisited

Bellman optimality equations

summary

recap

reinforcement learning

- ⇔ the problem faced by an agent that uses trial-and-error to learn how to behave in a dynamic environment
- ⇔ the problem faced by an agent that uses trial-and-error to learn to decide *what* to do *when* in order to achieve a long-term goal

“the” modeling tool for RL problems are *Markov decision processes* over sets of random variables S_t , A_t , and R_{t+1}

$MDP \Leftrightarrow$ a tuple $(\mathcal{S}, \mathcal{A}, T, R)$ with

set of states

$$\mathcal{S} = \{s_1, s_2, \dots\}$$

set of actions

$$\mathcal{A} = \{a_1, a_2, \dots\}$$

transition function / state transition probabilities

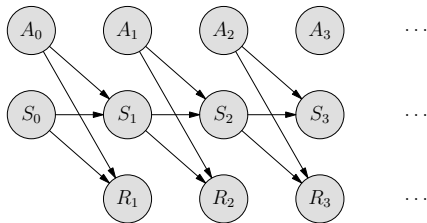
$$T(s', s, a) = p(S_{t+1} = s' \mid S_t = s, A_t = a)$$

reward function / expected immediate rewards

$$R(s, a) = \mathbb{E}[R_{t+1} \mid S_t = s, A_t = a]$$

$MDP \Leftrightarrow$ a stochastic process over S_t , A_t , and R_t where

$$p(S_{t+1} = s', R_{t+1} = r \mid S_t = s, A_t = a) = p(s', r \mid s, a)$$



a fully specified MDP $(\mathcal{S}, \mathcal{A}, T, R)$ *models* an environment

an MDP model $(\mathcal{S}, \mathcal{A}, T, R)$ **predicts an environment**, because
the transition function

$$T(s', s, a) = p(S_{t+1} = s' \mid S_t = s, A_t = a)$$

predicts next states

the reward function

$$R(s, a) = \mathbb{E}[R_{t+1} \mid S_t = s, A_t = a]$$

predicts next (immediate) rewards

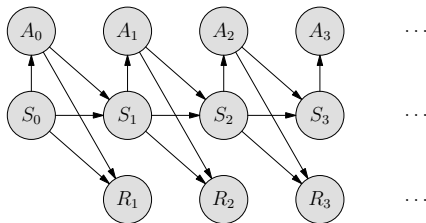
infinite-horizon discounted model of expected returns

$$G_0 = \sum_{t=0}^{\infty} \gamma^t R_{t+1} = \sum_{t=0}^{\infty} \gamma^t R(S_t, A_t) \quad \text{where} \quad 0 \leq \gamma < 1$$

stochastic policy

$$\pi : \mathcal{S} \times \mathcal{A} \rightarrow [0, 1] \quad \text{with} \quad \pi(s, a) \equiv \pi(a \mid s) \equiv p(a \mid s)$$

when introducing / working with a stochastic policy, the MDP model becomes



value function (for a given policy π)

$$V_{\pi}(s) = \mathbb{E}_{\pi} \left[\sum_{t=0}^{\infty} \gamma^t R_{t+1} \middle| S_0 = s \right] = \mathbb{E}_{\pi} \left[\sum_{t=0}^{\infty} \gamma^t R(S_t, A_t) \middle| S_0 = s \right]$$

quality function (for a given policy π)

$$Q_{\pi}(s, a) = \mathbb{E}_{\pi} \left[\sum_{t=0}^{\infty} \gamma^t R_{t+1} \middle| S_0 = s, A_0 = a \right] = \mathbb{E}_{\pi} \left[\sum_{t=0}^{\infty} \gamma^t R(S_t, A_t) \middle| S_0 = s, A_0 = a \right]$$

Bellman equations

$$V_{\pi}(s) = \sum_a R(s, a) \pi(a \mid s) + \gamma \sum_a \sum_{s'} V_{\pi}(s') p(s' \mid s, a) \pi(a \mid s)$$

$$Q_{\pi}(s, a) = R(s, a) + \gamma \sum_{s'} \sum_{a'} Q_{\pi}(s', a') \pi(a' \mid s') p(s' \mid s, a)$$

note the “next state” s' recursions



we also have

$$V_{\pi}(s) = \sum_a Q_{\pi}(s, a) \pi(a \mid s)$$

$$Q_{\pi}(s, a) = R(s, a) + \gamma \sum_{s'} V_{\pi}(s') p(s' \mid s, a)$$



$\Leftrightarrow$ we also have

$$V_{\pi}(s) = \mathbb{E}_{\pi} \left[Q_{\pi}(S_0, A_0) \mid S_0 = s \right]$$

$$Q_{\pi}(s, a) = \mathbb{E}_p \left[R(S_0, A_0) + \gamma V_{\pi}(S_1) \mid S_0 = s, A_0 = a \right]$$



these equations are **Bellman equations**

$$V_{\pi}(s) = \sum_a R(s, a) \pi(a \mid s) + \gamma \sum_a \sum_{s'} V_{\pi}(s') p(s' \mid s, a) \pi(a \mid s)$$

$$Q_{\pi}(s, a) = R(s, a) + \gamma \sum_{s'} \sum_{a'} Q_{\pi}(s', a') \pi(a' \mid s') p(s' \mid s, a)$$

because they relate $V_{\pi}(s)$ to $V_{\pi}(s')$ and $Q_{\pi}(s, a)$ to $Q_{\pi}(s', a')$, respectively

these equations are useful

$$V_{\pi}(s) = \sum_a Q_{\pi}(s, a) \pi(a \mid s)$$

$$Q_{\pi}(s, a) = R(s, a) + \gamma \sum_{s'} V_{\pi}(s') p(s' \mid s, a)$$

but they relate $V_{\pi}(s)$ to $Q_{\pi}(s, a)$ and $Q_{\pi}(s, a)$ to $V_{\pi}(s)$, respectively



all of this generalizes to

continuous time evolution

continuous state spaces

continuous action spaces

partially observable MDPs (POMDPs)

we will study generalizations later on ...

working with deterministic policies

disclaimer

what we will do next is nothing but a special case of what we did already !

our main reasons for considering this special case are to hit home a point about the reward function $R(s, a)$ and to be able to define the concept of a *Markov reward process*

we therefore recall from lecture 01 ...

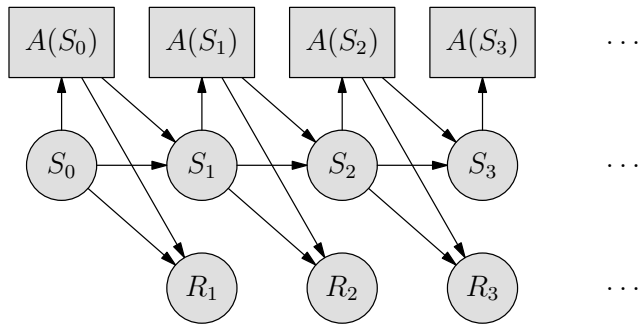
a **deterministic policy** is a *function*

$$\pi : \mathcal{S} \rightarrow \mathcal{A} \quad \text{with} \quad \pi(s) = a$$

which maps a state to an action

some people write $\mu(s)$ instead of $\pi(s)$

Bayesian network representation





under a deterministic policy π , an MDP transition function becomes

$$\begin{aligned} T(s', s, \pi(s)) &= p(S_{t+1} = s' \mid S_t = s, A_t = \pi(s)) \\ &= p_{\pi}(S_{t+1} = s' \mid S_t = s) \\ &= T_{\pi}(s', s) \end{aligned}$$

under a deterministic policy π , an MDP reward function becomes

$$\begin{aligned} R(s, \pi(s)) &= \mathbb{E}[R_{t+1} \mid S_t = s, A_t = \pi(s)] = \sum_{s'} \text{Util}(s') p_{\pi}(s' \mid s) \\ &= \mathbb{E}_{p_{\pi}}[R_{t+1} \mid S_t = s] \\ &= R_{\pi}(s) \end{aligned}$$

question

wait a minute . . .

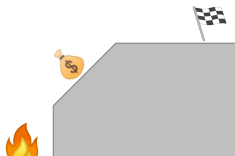
if action selection is deterministic, then why are state transitions still probabilistic ?

answer

the agent's environment may be antagonistic / non-deterministic / non-controllable

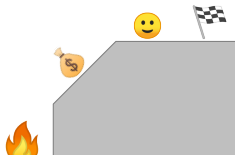
consider this . . .

recall the simple game environment

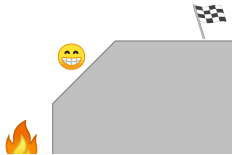


example inspired by [Elliot Waite](#)

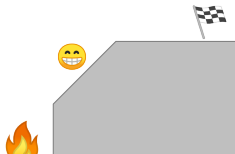
upon spawning, you go left ...



and collect the loot



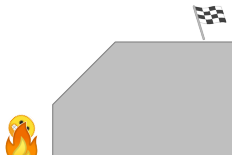
of course you next go right ...



with a 90% chance you succeed and then win



alas, with a 10% chance you slip into the fire





⇒ in general, selected actions can't be guaranteed to have the desired effect

this thus explains why it is common to work with *expected rewards*

$$R(s, a) = \mathbb{E}[R_{t+1} \mid S_t = s, A_t = a] = \sum_{s'} \text{Util}(s') p(s' \mid s, a)$$

since executing action a in state s may result in various successor states s'

one more thing ...

everybody writes the reward function as

$$R(s, a) = \mathbb{E}[R_{t+1} \mid S_t = s, A_t = a]$$

but it would of course be more precise to write

$$\begin{aligned} R(s, a) &= \sum_{s'} \text{Util}(s') p(s' \mid s, a) \\ &= \mathbb{E}_{\textcolor{red}{p}}[R_{t+1} \mid S_t = s, A_t = a] \end{aligned}$$

later in this lecture, we will be more precise and write $\mathbb{E}_{\textcolor{red}{p}}[\cdot]$



combining a Markov decision process with a deterministic policy leads to *deterministic action selection* in each state

$\Leftrightarrow$ a Markov decision process combined with a deterministic policy behaves like a Markov chain, because

$$p(s' \mid s, \pi(s)) = p_{\pi}(s' \mid s)$$

since there also are outputs $R_{\pi}(s)$, we are dealing with ...

Markov reward process (MRP)

$\Leftrightarrow$ a Markov chain with values

$\Leftrightarrow$ a tuple $(\mathcal{S}, T, R)$ such that

$\mathcal{S} \equiv$ set of states

$T(s', s) \equiv$ state transition probabilities

$$= p(S_{t+1} = s' \mid S_t = s)$$

$R(s) \equiv$ expected values / rewards

$$= \mathbb{E}[R_{t+1} \mid S_t = s]$$

using this, we may also say ...

Markov decision process (MDP)

$\Leftrightarrow$ a Markov reward process augmented by decisions

Bellman equations revisited

and back to stochastic policies $\pi(a \mid s) \dots$

an **optimal value function** / **optimal value of a state** is given by

$$\begin{aligned} V_*(s) &= \max_{\pi} \mathbb{E}_{\pi} \left[\sum_{t=0}^{\infty} \gamma^t R(S_t, \pi(A_t | S_t)) \middle| S_0 = s \right] \\ &= \mathbb{E} \left[\sum_{t=0}^{\infty} \gamma^t R(S_t, \pi_*(A_t | S_t)) \middle| S_0 = s \right] \end{aligned}$$

an **optimal policy** is a policy $\pi_*(a | s)$ such that for all states

$$V_*(s) = V_{\pi_*}(s) \geq V_{\pi}(s)$$



we seem to be going in circles, since

$V_*(s)$ is defined in terms of $\pi_*(s)$

$\pi_*(s)$ is defined in terms of $V_*(s)$



we seem to be going in circles, since

$V_*(s)$ is defined in terms of $\pi_*(s)$

$\pi_*(s)$ is defined in terms of $V_*(s)$

enter Richard Bellman ...

Bellman's principle of optimality

An optimal policy has the property that whatever the initial state and the initial decision are, the remaining decisions must constitute an optimal policy w.r.t. the state resulting from the first decision.



R.E. Bellman
(*1920, †1984)

Bellman equations

⇔ use the principle of optimality to relate a value function to itself

$$\begin{aligned} V_{\pi}(s) &= \sum_a R(s, a) \pi(a \mid s) + \gamma \sum_a \sum_{s'} V_{\pi}(s') p(s' \mid s, a) \pi(a \mid s) \\ &= \mathbb{E}_{\pi} \left[\mathbb{E}_p \left[R_1 \mid S_0 = s, A_0 \right] + \gamma \mathbb{E}_p \left[V_{\pi}(S_1) \mid S_0 = s, A_0 \right] \right] \end{aligned}$$

⇔ the (expected) value of a state is its immediate return plus the (discounted) average of the values of its possible successors

we have

$$\begin{aligned} V_{\pi}(s) &= \sum_a R(s, a) \pi(a | s) + \gamma \sum_a \sum_{s'} V_{\pi}(s') p(s' | s, a) \pi(a | s) \\ &= \sum_a \left[R(s, a) + \gamma \sum_{s'} V_{\pi}(s') p(s' | s, a) \right] \pi(a | s) \\ &= \mathbb{E}_{\pi} \left[R(s, a) + \gamma \sum_{s'} V_{\pi}(s') p(s' | s, a) \right] \\ &= \mathbb{E}_{\pi} \left[\mathbb{E}_p \left[R_1 \mid S_0 = s, A_0 \right] + \gamma \sum_{s'} V_{\pi}(s') p(s' | s, a) \right] \\ &= \mathbb{E}_{\pi} \left[\mathbb{E}_p \left[R_1 \mid S_0 = s, A_0 \right] + \gamma \mathbb{E}_p \left[V_{\pi}(S_1) \mid S_0 = s, A_0 \right] \right] \\ &= \mathbb{E}_{\pi} \left[\mathbb{E}_p \left[R_1 \mid S_0 = s, A_0 \right] \right] + \gamma \mathbb{E}_{\pi} \left[\mathbb{E}_p \left[V_{\pi}(S_1) \mid S_0 = s, A_0 \right] \right] \end{aligned}$$

from lecture 07, we recall

$$\mathbb{E}_{\pi} \left[\mathbb{E}_p \left[R_1 \mid S_0 = s, A_0 \right] \right] = \mathbb{E}_p \left[R_1 \mid S_0 = s \right] \equiv R_{\pi}(s)$$

recall

$$\mathbb{E}[\mathbb{E}[X]] = \mathbb{E}[X]$$

recall

$$\mathbb{E}_{\pi}[R(s, a) \mid S_0 = s] = \mathbb{E}_{\pi}[\mathbb{E}_p[R_1 \mid S_0, A_0] \mid S_0 = s] = \mathbb{E}_p[R_1 \mid S_0 = s] \equiv R(s)$$

observe

$$\mathbb{E}_{\pi}[R(s) \mid s] = R(s)$$

$$\mathbb{E}_p[R(s) \mid s] = R(s)$$

to reduce clutter, we next simply write

$$\mathbb{E}_{\pi}[\mathbb{E}_p[R_1 \mid S_0, A_0] + \gamma \mathbb{E}_p[V_{\pi}(s') \mid S_0, A_0] \mid S_0 = s] = R(s) + \gamma \mathbb{E}_{\pi}[\mathbb{E}_p[V_{\pi}(s')] \mid S_0 = s]$$

using *telescoping*, we find

$$\begin{aligned} V_{\pi}(s) &= R_{\pi}(s) + \gamma \mathbb{E}_{\pi} \left[\mathbb{E}_p \left[V_{\pi}(S_1) \mid S_0 = s, A_0 \right] \right] \\ &= R_{\pi}(s) + \gamma \mathbb{E}_{\pi} \left[\mathbb{E}_p \left[R_{\pi}(s') + \gamma \mathbb{E}_{\pi} \left[\mathbb{E}_p \left[V_{\pi}(S_2) \mid S_1 = s', A_1 \right] \right] \right] \right] \\ &= R_{\pi}(s) + \gamma R_{\pi}(s') + \gamma^2 \mathbb{E}_{\pi} \left[\mathbb{E}_p \left[V_{\pi}(S_2) \mid S_1 = s', A_1 \right] \right] && \text{(double expectations "cancel")} \\ &\vdots \\ &= R_{\pi}(s) + \gamma R_{\pi}(s') + \gamma^2 R_{\pi}(s'') + \gamma^3 R_{\pi}(s''') + \cdots \\ &= \mathbb{E}_{\pi} \left[\sum_{t=0}^{\infty} \gamma^t R(S_t, A_t) \mid S_0 = s \right] \end{aligned}$$

in short, we find

$$\begin{aligned} V_{\pi}(s) &= \sum_a R(s, a) \pi(a \mid s) + \gamma \sum_a \sum_{s'} V_{\pi}(s') p(s' \mid s, a) \pi(a \mid s) \\ &= \mathbb{E}_{\pi} \left[\sum_{t=0}^{\infty} \gamma^t R(S_t, A_t) \mid S_0 = s \right] \\ &= \mathbb{E}_{\pi} [G_0 \mid S_0 = s] \end{aligned}$$

$\Rightarrow$ Bellman's principle of optimality *implies* the definition of $V_{\pi}(s)$ we gave earlier

remark

above, we were excruciatingly pedantic and found

$$\begin{aligned} V_{\pi}(s) &= \mathbb{E}_{\pi} \left[G_0 \mid S_0 = s \right] \\ &= \mathbb{E}_{\pi} \left[R_1 + \gamma G_1 \mid S_0 = s \right] \end{aligned}$$

however, in the literature, it is very common to read

$$V_{\pi}(s) = \mathbb{E}_{\pi} \left[R_1 + \gamma V_{\pi}(S_1) \mid S_0 = s \right]$$

remark (continued)

the equality

$$\mathbb{E}_{\pi} \left[R_1 + \gamma G_1 \mid S_0 = s \right] = \mathbb{E}_{\pi} \left[R_1 + \gamma V_{\pi}(S_1) \mid S_0 = s \right]$$

can be confusing, because we actually have

$$G_1 \neq V_{\pi}(S_1)$$

we do, however, have *equality in expectation*

$$\mathbb{E}_{\pi} \left[G_1 \mid S_0 = s \right] = \mathbb{E}_{\pi} \left[V_{\pi}(S_1) \mid S_0 = s \right]$$

question

but this doesn't involve *optimal* $V_*(s)$ or $\pi_*(s)$, or does it ?

answer

no, it doesn't yet . . . we will get back to this point later . . .



there actually is one Bellman equation for each state $s \in \mathcal{S}$

$$V_{\pi}(s_1) = \sum_a R(s_1, a) \pi(a \mid s_1) + \gamma \sum_a \sum_{s'} V_{\pi}(s') p(s' \mid s_1, a) \pi(a \mid s_1)$$

$$V_{\pi}(s_2) = \sum_a R(s_2, a) \pi(a \mid s_2) + \gamma \sum_a \sum_{s'} V_{\pi}(s') p(s' \mid s_2, a) \pi(a \mid s_2)$$

$\vdots$

$\Rightarrow$ **Bellman equations always constitute whole systems of equations**

let us rearrange the Bellman equations as

$$\begin{aligned} V_{\pi}(s) &= \sum_a R(s, a) \pi(a \mid s) + \gamma \sum_a \sum_{s'} V_{\pi}(s') p(s' \mid s, a) \pi(a \mid s) \\ &= \underbrace{\sum_a R(s, a) \pi(a \mid s)}_{=T_1} + \gamma \underbrace{\sum_{s'} V_{\pi}(s') \sum_a p(s' \mid s, a) \pi(a \mid s)}_{=T_2} \end{aligned}$$

and examine the deeper meaning / consequences of the insight we just had

to begin with, we note ...

if the state set $\mathcal{S}$ is discrete and of finite size $|\mathcal{S}| = n$,
we can think of a value function $V_\pi(\cdot)$ as a **vector**

$$\mathbf{v}_\pi = \begin{bmatrix} V_\pi(s_1) \\ V_\pi(s_2) \\ \vdots \end{bmatrix}$$

next, we note ...

if the action set $\mathcal{A}$ is discrete and of finite size $|\mathcal{A}| = m$, we can think of reward functions $R(\cdot, \cdot)$ and stochastic policies $\pi(\cdot, \cdot)$ as **matrices**

$$\mathbf{R} = \begin{bmatrix} R(a_1, s_1) & R(a_1, s_2) & \cdots \\ R(a_2, s_1) & R(a_2, s_2) & \cdots \\ \vdots & \vdots & \ddots \end{bmatrix}$$

$$\mathbf{\Pi} = \begin{bmatrix} \pi(a_1 | s_1) & \pi(a_1 | s_2) & \cdots \\ \pi(a_2 | s_1) & \pi(a_2 | s_2) & \cdots \\ \vdots & \vdots & \ddots \end{bmatrix}$$

next, we note ...

$\Rightarrow$ w.r.t. T_1 , we realize that

$$\sum_a R(s, a) \pi(a \mid s)$$

computes a **vector**, namely

$$\mathbf{r}_\pi^\top = \mathbf{1}^\top [\mathbf{R} \odot \mathbf{\Pi}]$$

$$\Leftrightarrow \mathbf{r}_\pi = [\mathbf{R} \odot \mathbf{\Pi}]^\top \mathbf{1}$$

next, we note ...

for the inner sum of T_2 , we have the following equality

$$\sum_a p(s' | s, a) \cdot \pi(a | s) = \sum_a \frac{p(s', s, a)}{p(s, a)} \cdot \frac{p(s, a)}{p(s)} = \sum_a \frac{p(s', a | s) p(s)}{p(s)} = p(s' | s)$$

and can arrange all the resulting state transition probabilities $p(s' | s)$ in a **matrix**

$$\mathbf{P}_\pi = \begin{bmatrix} p(s_1 | s_1) & p(s_1 | s_2) & \cdots \\ p(s_2 | s_1) & p(s_2 | s_2) & \cdots \\ \vdots & \vdots & \ddots \end{bmatrix}$$

moreover, since this matrix $\mathbf{P}_\pi$ is column stochastic, its **spectral radius** is $\rho(\mathbf{P}_\pi) = 1$

$\Rightarrow$ w.r.t. T_2 , we realize that

$$\sum_{s'} V_{\pi}(s') \sum_a p(s' \mid s, a) \pi(a \mid s)$$

computes a **vector**, namely

$$\mathbf{u}^{\top} = \mathbf{v}_{\pi}^{\top} \mathbf{P}_{\pi}$$

$$\Leftrightarrow \mathbf{u} = \mathbf{P}_{\pi}^{\top} \mathbf{v}_{\pi}$$



⇒ we can rewrite a system of Bellman equations
in terms of just a single matrix-vector equation

$$\mathbf{v}_\pi = \mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top \mathbf{v}_\pi$$

we can solve this equation for $\mathbf{v}_\pi$ because

$$\mathbf{v}_\pi = \mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top \mathbf{v}_\pi$$

$$\Leftrightarrow \mathbf{v}_\pi - \gamma \mathbf{P}_\pi^\top \mathbf{v}_\pi = \mathbf{r}_\pi$$

$$\Leftrightarrow \left[\mathbf{I} - \gamma \mathbf{P}_\pi^\top \right] \mathbf{v}_\pi = \mathbf{r}_\pi$$

$$\Leftrightarrow \mathbf{v}_\pi = \left[\mathbf{I} - \gamma \mathbf{P}_\pi^\top \right]^{-1} \mathbf{r}_\pi$$

$$(\gamma < 1 \wedge \rho(\mathbf{P}_\pi^\top) = 1 \Rightarrow \rho(\gamma \mathbf{P}_\pi^\top) < 1)$$

now, let $\mathbf{v}_0$ be some initial value vector and consider this process

$$\mathbf{v}_1 = \mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top \mathbf{v}_0$$

$$\mathbf{v}_2 = \mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top \mathbf{v}_1$$

$$= \mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top [\mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top \mathbf{v}_0] = \mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top \mathbf{r}_\pi + \gamma^2 \mathbf{P}_\pi^{\top 2} \mathbf{v}_0$$

$$\mathbf{v}_3 = \mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top \mathbf{v}_2$$

$$= \mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top [\mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top \mathbf{r}_\pi + \gamma^2 \mathbf{P}_\pi^{\top 2} \mathbf{v}_0] = \mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top \mathbf{r}_\pi + \gamma^2 \mathbf{P}_\pi^{\top 2} \mathbf{r}_\pi + \gamma^3 \mathbf{P}_\pi^{\top 3} \mathbf{v}_0$$

$\vdots$

$$\mathbf{v}_t = \sum_{\tau=0}^{t-1} [\gamma \mathbf{P}_\pi^\top]^\tau \mathbf{r}_\pi + [\gamma \mathbf{P}_\pi^\top]^t \mathbf{v}_0$$

$$([\gamma \mathbf{P}_\pi^\top]^0 = \mathbf{I})$$

note

$$\lim_{t \rightarrow \infty} \mathbf{v}_t = \lim_{t \rightarrow \infty} \left[\sum_{\tau=0}^{t-1} [\gamma \mathbf{P}_{\pi}^{\top}]^{\tau} \mathbf{r}_{\pi} + [\gamma \mathbf{P}_{\pi}^{\top}]^t \mathbf{v}_0 \right]$$

$$= \lim_{t \rightarrow \infty} \sum_{\tau=0}^{t-1} [\gamma \mathbf{P}_{\pi}^{\top}]^{\tau} \mathbf{r}_{\pi} + \lim_{t \rightarrow \infty} [\gamma \mathbf{P}_{\pi}^{\top}]^t \mathbf{v}_0$$

$$= \left[\lim_{t \rightarrow \infty} \sum_{\tau=0}^{t-1} [\gamma \mathbf{P}_{\pi}^{\top}]^{\tau} \right] \mathbf{r}_{\pi} + \mathbf{0}$$

$$= [\mathbf{I} - \gamma \mathbf{P}_{\pi}^{\top}]^{-1} \mathbf{r}_{\pi}$$

$$= \mathbf{v}_{\pi}$$

$$(\gamma < 1 \wedge \rho(\mathbf{P}_{\pi}^{\top}) = 1 \Rightarrow \lim_{t \rightarrow \infty} [\gamma \mathbf{P}_{\pi}^{\top}]^t = \mathbf{0})$$

(Neumann series with $\rho(\gamma \mathbf{P}_{\pi}^{\top}) < 1$)



take home messages

$\Rightarrow$ given some policy π , we can determine V_π

we can always proceed iteratively

$$\mathbf{v}_{t+1} = \mathbf{r}_\pi + \gamma \mathbf{P}_\pi^\top \mathbf{v}_t$$

or, if it isn't too expensive, directly

$$\mathbf{v}_\pi = \left[\mathbf{I} - \gamma \mathbf{P}_\pi^\top \right]^{-1} \mathbf{r}_\pi$$

disclaimer

we just saw that the process

for $t \in \mathbb{N}$

for all $s \in \mathcal{S}$

$$V_{t+1}(s) = \sum_a R(s, a) \pi(a \mid s) + \gamma \sum_a \sum_{s'} V_t(s') p(s' \mid s, a) \pi(a \mid s)$$

is guaranteed to converge to a **fixed point** $V_\pi(s)$

next, we will reexamine this using different tools

to begin with, we note ...

value functions are points in an $|\mathcal{S}|$ -dimensional vector space over $\mathbb{R}$

we may measure the distance between any two value functions U, V in terms of the L_∞ norm

$$\|U - V\|_\infty = \max_{s \in \mathcal{S}} |U(s) - V(s)|$$

in terms of the above vector notation, we may write this distance as

$$\|\mathbf{u} - \mathbf{v}\|_\infty$$

next, we note ...

we may write the Bellman equation as

$$\boldsymbol{v}_\pi = \boldsymbol{B}_\pi(\boldsymbol{v}_\pi)$$

where the operator $\boldsymbol{B}_\pi : \mathbb{R}^{|\mathcal{S}|} \rightarrow \mathbb{R}^{|\mathcal{S}|}$ with

$$\boldsymbol{B}_\pi(\boldsymbol{x}) = \boldsymbol{r}_\pi + \gamma \boldsymbol{P}_\pi^\top \boldsymbol{x}$$

is called a **Bellman policy operator** or a **Bellman expectation operator**



the above equality $\boldsymbol{v}_\pi = \boldsymbol{B}_\pi(\boldsymbol{v}_\pi)$ implies that $\boldsymbol{v}_\pi$ is a *fixed point* of $\boldsymbol{B}_\pi$

next, we ask ...

question

what happens to the distance between any two $u, v \in \mathbb{R}^{|S|}$
if we subject both points to the same Bellman operator ?

answer

this ...

we have

$$\begin{aligned}\left\| \mathbf{B}_{\pi}(\mathbf{u}) - \mathbf{B}_{\pi}(\mathbf{v}) \right\|_{\infty} &= \left\| \mathbf{r}_{\pi} + \gamma \mathbf{P}_{\pi}^{\top} \mathbf{u} - \mathbf{r}_{\pi} - \gamma \mathbf{P}_{\pi}^{\top} \mathbf{v} \right\|_{\infty} = \left\| \gamma \mathbf{P}_{\pi}^{\top} [\mathbf{u} - \mathbf{v}] \right\|_{\infty} \\ &\leq \left\| \gamma \mathbf{P}_{\pi}^{\top} \mathbf{1} \left\| \mathbf{u} - \mathbf{v} \right\|_{\infty} \right\|_{\infty} && \left(\mathbf{P}_{\pi}^{\top} \text{ row stochastic} \Rightarrow \mathbf{P}_{\pi}^{\top} \mathbf{1} = \mathbf{1} \right) \\ &= \gamma \left\| \mathbf{u} - \mathbf{v} \right\|_{\infty} && (0 \leq \gamma < 1) \\ &\leq \left\| \mathbf{u} - \mathbf{v} \right\|_{\infty}\end{aligned}$$

$\Rightarrow$ applying $\mathbf{B}_{\pi}$ brings points $\mathbf{u}$ and $\mathbf{v}$ closer together $\Leftrightarrow \mathbf{B}_{\pi}$ is a contraction operator

given a complete metric space $(\mathcal{V}, d)$, an operator

$$f : \mathcal{V} \rightarrow \mathcal{V}$$

is a **contraction mapping** / **contraction** / **contractor**, if

$$\exists \gamma \in [0, 1) : \forall \mathbf{x}, \mathbf{y} \in \mathcal{V} : d(f(\mathbf{x}), f(\mathbf{y})) \leq \gamma d(\mathbf{x}, \mathbf{y})$$

the **Banach fixed point theorem** establishes that
every contraction has a *unique* fixed point

$$\mathbf{x}_* = f(\mathbf{x}_*)$$

the sequence $\mathbf{x}_t = f(\mathbf{x}_{t-1}) = f^t(\mathbf{x}_0)$ will
converge to this fixed point

$$\lim_{t \rightarrow \infty} \mathbf{x}_t = \mathbf{x}_*$$

for *any* $\mathbf{x}_0 \in \mathcal{V}$
the following useful upper bounds do hold

$$d(\mathbf{x}_t, \mathbf{x}_*) \leq \frac{\gamma^n}{1-\gamma} d(\mathbf{x}_0, f(\mathbf{x}_0))$$

($\Leftrightarrow$ “how many iterations from $\mathbf{x}_0$ to get close to $\mathbf{x}_*$?”)

$$d(\mathbf{x}_t, \mathbf{x}_*) \leq \frac{\gamma}{1-\gamma} d(\mathbf{x}_{t-1}, \mathbf{x}_t)$$

($\Leftrightarrow$ “how close are we currently to $\mathbf{x}_*$?”)

 note

we just saw twice in a row that we can perform **policy evaluation**

$\Leftrightarrow$ given any policy π and an arbitrary initial value function $V_0(s)$, we can determine the policy's value function $V_\pi(s)$

synchronous policy evaluation

function *SynchronousPolicyEvaluation*($MDP, \pi, \gamma, \varepsilon$, arbitrary V_{old})

while *TRUE*

for all $s \in \mathcal{S}$

$$V_{new}(s) \leftarrow \sum_a R(s, a) \pi(a \mid s) + \gamma \sum_a \sum_{s'} V_{old}(s') p(s' \mid s, a) \pi(a \mid s)$$

if $\max_{s \in \mathcal{S}} |V_{new}(s) - V_{old}(s)| \leq \varepsilon$

return V_{new}

$V_{old} \leftarrow V_{new}$

asynchronous policy evaluation

function *AsynchronousPolicyEvaluation*($MDP, \pi, \gamma, \varepsilon$, arbitrary V)

while $TRUE$

$\delta \leftarrow 0$

for all $s \in \mathcal{S}$

$h \leftarrow V(s)$

$V(s) \leftarrow \sum_a R(s, a) \pi(a \mid s) + \gamma \sum_a \sum_{s'} V(s') p(s' \mid s, a) \pi(a \mid s)$

$\delta \leftarrow \max(\delta, |V(s) - h|)$

if $\delta \leq \varepsilon$

return V

remarks

an iteration over the s is commonly called a **sweep** through the state space

asynchronous or “in place” policy evaluation often converges *faster* than its synchronous counterpart, i.e. it requires fewer sweeps

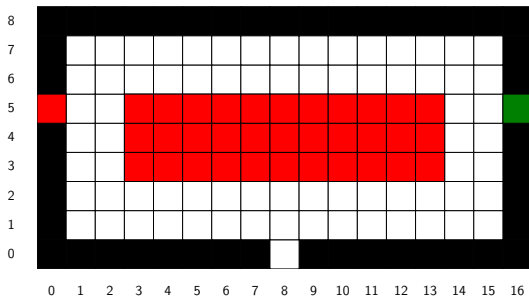
alas, the order in which states are swept over during asynchronous policy evaluation can significantly impact the rate of convergence

there are many clever ideas for how to find an optimal order but we will *not* discuss them . . . future lectures will reveal why this really would be overkill

instead, let's look an example with synchronous updates . . .

example

a simple grid world (RL instructors love grid worlds)



states

$$\begin{aligned}\mathcal{S} &= \{s_1, s_2, s_3, \dots, s_{108}\} \\ &= \left\{ \begin{bmatrix} 0 \\ 5 \end{bmatrix}, \begin{bmatrix} 1 \\ 1 \end{bmatrix}, \begin{bmatrix} 1 \\ 2 \end{bmatrix}, \dots, \begin{bmatrix} 16 \\ 5 \end{bmatrix} \right\}\end{aligned}$$

actions

$$\begin{aligned}\mathcal{A} &= \{a_1, a_2, a_3, a_4\} \\ &= \{\uparrow, \rightarrow, \downarrow, \leftarrow\} \\ &= \left\{ \begin{bmatrix} 0 \\ +1 \end{bmatrix}, \begin{bmatrix} +1 \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ -1 \end{bmatrix}, \begin{bmatrix} -1 \\ 0 \end{bmatrix} \right\}\end{aligned}$$

final states

$$\mathcal{F}_- = \{s_1, s_{18}, \dots, s_{91}\} \subset \mathcal{S}$$

$$\mathcal{F}_+ = \{s_{108}\} \subset \mathcal{S}$$

$$\mathcal{F} = \mathcal{F}_- \cup \mathcal{F}_+$$

“reward shaped” utilities $Util(s)$ of states s

$\Rightarrow$ utility vector

	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	
	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	
-10.0	-0.1	-0.1	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-0.1	+10.0
	-0.1	-0.1	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-0.1	-0.1
	-0.1	-0.1	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-10.0	-0.1	-0.1
	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1
	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1	-0.1
																-1.0

$$\mathbf{u} \in \mathbb{R}^{|S|}$$

with

$$u_s = Util(s)$$

reasoning

lava states $\in \mathcal{F}_-$, exit $\in \mathcal{F}_+$

entry should *not* be revisited

$\Leftrightarrow$ lava states have utility -10

the exit has a utility of $+10$

entry $[8, 0]^T$ has utility -1

setting up tensor T

$$T \leftarrow \mathbf{0} \in \mathbb{R}^{|\mathcal{S}| \times |\mathcal{S}| \times |\mathcal{A}|}$$

for all $a \in \mathcal{A}$
for all $s \in \mathcal{S}$

if $s \in \mathcal{F}$
 $T_{s,s,a} \leftarrow 1$
 continue

$s' \leftarrow \text{Succ}(s, a)$

if $s' \notin \mathcal{S}$
 $T_{s,s,a} \leftarrow 1$
else
 $T_{s',s,a} \leftarrow 1$

$$T_{:, :, a} \leftarrow \frac{T_{:, :, a}}{\sum_{s'} T_{s', :, a}}$$

preparing matrix R

$$R \leftarrow \mathbf{0} \in \mathbb{R}^{|\mathcal{A}| \times |\mathcal{S}|}$$

for all $a \in \mathcal{A}$

$$R_{a, :} = \sum_{s'} u_{s'} T_{s', :, a}$$

setting up a stochastic policy matrix Π

$$\Pi \leftarrow \mathbf{0} \in \mathbb{R}^{|\mathcal{A}| \times |\mathcal{S}|}$$

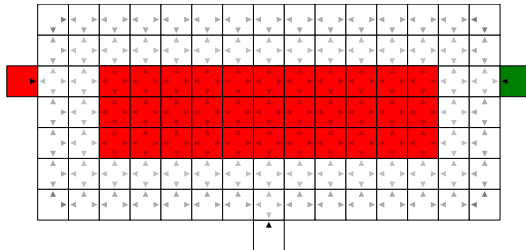
for all $a \in \mathcal{A}$
for all $s \in \mathcal{S}$

$$s' \leftarrow \text{Succ}(s, a)$$

if $s' \in \mathcal{S}$
 $\Pi_{a,s} \leftarrow 1$

$$\Pi_{:,i} \leftarrow \frac{\Pi_{:,i}}{\sum_a \Pi_{a,i}}$$

visualization

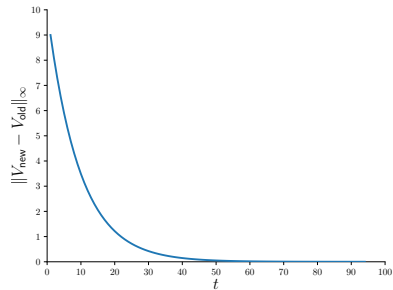


results

initial $V(s) = u(s)$, $\gamma = 0.9$, $\varepsilon = 0.0005$

	-36.4	-37.7	-43.2	-46.6	-48.3	-49.0	-49.3	-49.3	-49.1	-48.5	-47.0	-43.7	-37.0	-25.2	-16.7	
	-42.8	-45.8	-59.4	-63.6	-65.0	-65.5	-65.7	-65.7	-65.6	-65.2	-64.1	-61.5	-53.9	-29.9	-11.7	
	-100.0	-60.3	-63.3	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-41.6	+7.9	-100.0
	-50.4	-63.9	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-51.3	-22.5	
	-43.4	-58.7	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-100.0	-52.3	-31.4	
	-35.2	-42.1	-58.1	-63.0	-64.6	-65.1	-65.0	-64.6	-65.0	-65.0	-64.5	-62.5	-56.9	-38.5	-29.4	
	-31.6	-34.7	-41.7	-45.8	-47.7	-48.3	-47.8	-45.6	-47.7	-48.2	-47.4	-45.2	-40.4	-32.1	-27.8	
								-41.1								

convergence behavior



results

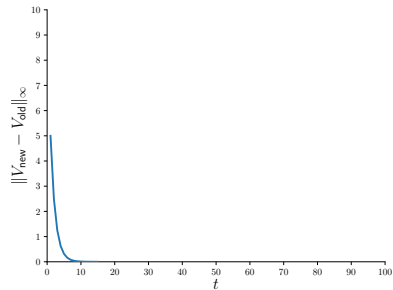
initial $V(s) = u(s)$, $\gamma = 0.5$, $\varepsilon = 0.0005$

	-0.7	-0.8	-1.6	-1.9	-1.9	-1.9	-1.9	-1.9	-1.9	-1.9	-1.9	-1.8	-1.6	-0.7	-0.2	
	-1.7	-2.1	-6.4	-7.0	-7.1	-7.1	-7.1	-7.1	-7.1	-7.1	-7.1	-7.0	-6.4	-1.6	+0.3	
-20.0	-6.5	-7.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-5.6	+4.2	-20.0
	-2.7	-7.1	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-6.7	-0.7	
	-1.7	-6.4	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-20.0	-6.3	-1.4	
	-0.8	-1.9	-6.4	-7.0	-7.1	-7.1	-7.1	-7.1	-7.1	-7.1	-7.1	-7.0	-6.4	-1.9	-0.7	
	-0.5	-0.8	-1.6	-1.8	-1.9	-1.9	-1.9	-1.8	-1.9	-1.9	-1.9	-1.8	-1.6	-0.8	-0.5	
								-1.0								

observe

many values seem identical, but this is due to rounding
all values do indeed differ after the 4th (!) decimal place

convergence behavior



question

why does the $\gamma = 0.5$ version converge faster than the version with $\gamma = 0.9$?

answer

given what we saw above, you might want to figure this out all by yourself ;-)

question

OK, but now for real: what about $V_*(s)$ and $\pi_*(s)$?

answer

...

Bellman optimality equations

Bellman optimality equation for $V_*(s)$

$$V_*(s) = \max_a Q_*(s, a)$$

$$= \max_a \mathbb{E} \left[R_1 + \gamma V_*(S_1) \mid S_0 = s, A_0 = a \right]$$

$$= \max_a R(s, a) + \gamma \sum_{s'} V_*(s') p(s' \mid s, a)$$

Bellman optimality equation for $Q_*(s, a)$

$$\begin{aligned} Q_*(s, a) &= \mathbb{E} \left[R_1 + \gamma \max_{a'} Q_*(S_1, a') \mid S_0 = s, A_0 = a \right] \\ &= R(s, a) + \gamma \sum_{s'} \max_{a'} Q_*(s', a') p(s' \mid s, a) \end{aligned}$$



because of the occurrence of the “max” operator neither system of equations is linear anymore

if we could determine $Q_*(s, a)$, we could easily determine $V_*(s)$

if we knew $Q_*(s, a)$ or $V_*(s)$, we could easily get $\pi_*(s)$ because ...

optimal policy $\pi_*(s)$

$$\pi_*(s) = \operatorname{argmax}_a Q_*(s, a)$$

$$= \operatorname{argmax}_a R(s, a) + \gamma \sum_{s'} V_*(s') p(s' \mid s, a)$$

question

so how can we proceed ?

answer

will be given next time ...

summary

we now know about

- aspects of working with deterministic policies

- Bellman's principle of optimality

- (iterative) policy evaluation

- Bellman optimality equations