

Building a game player $\rightarrow$ Evaluation function

$$f(b) = \sum w_i \phi_i(b)$$

$\uparrow$
 learned by playing games domain knowledge

Set of actions: for each a there is a reward $R(a)$

$R(a)$ might be random (in general, drawn from a distribution)

Ad placement: $\phi_i \rightarrow$ words/content in ad

Find ad a s.t. we obtain max reward $R(a)$

$a_i \rightarrow$ try it n_i times: $R_1 \dots R_{n_i}$

Find $\vec{w}_i$ (weight or parameter vector) $\rightarrow$ array

s.t. $\underbrace{w_i \cdot \phi_i}_{\text{correct-ish estimate of } R(a_i)}$

$$\sum_j w_{i,j} \phi_{i,j}(a)$$

Minimize the error between our estimates and what we get from the environment

Idea 1

www

Error function:

$$\min_{\mathbf{h}} \sum_{j=1}^n (\mathbf{h}(j) - R_j)^2$$

all data $\rightarrow$ $\mathbf{h}$ prediction on j th try $\leftarrow$ actual data R_j

E.g. $\mathbf{h} = \sum \mathbf{w} \cdot \Phi_i$ in our previous case.

$$\mathbf{h}(\mathbf{w}) \quad \min_{\mathbf{w}} \sum_{j=1}^n (\mathbf{h}(\mathbf{w}, j) - R_j)^2$$

take a derivative wrt $\mathbf{w}$
and set them to 0 to find $\mathbf{w}$

In linear case: we can solve this in closed form (nice!)

Look for action a_i that maximizes reward (according to our estimates)

$$\arg \max h(w_i)$$

$\rightarrow$ best action

Exploration - exploitation dilemma:

- try out something new? or

- use the best guess you have so far?

(137-3)

A1 $\rightarrow$ payoff of 100 | Big difference $\rightarrow$ A1
A2 $\rightarrow$ payoff of 10

Regret: how much loss do I get from not doing the optimal thing

A1 $\rightarrow$ payoff of 100 $\rightarrow$? } A2 not much
A2 $\rightarrow$ payoff of 99 } use

A1 $\rightarrow$ payoff of 100 } Both tried 10^6 times
A2 $\rightarrow$ payoff of 99 } Lots of data $\rightarrow$
more "greedy"

Randomised strategies for picking actions

$\rightarrow$ Best actions according to est est picked
more

$\rightarrow$ a_i and a_j have close reward est $\rightarrow p(a_i) \approx p(a_j)$

$\rightarrow$ Action that is best according to our est
should see into prob $\rightarrow$ 1 as data $\rightarrow \infty$

Boltzmann distribution

R^* - best est. so far

$(R^* - R_i) / \epsilon_{mi}$

$p(a_i)$ prop
to

$\uparrow$

"temperature"

Bandit problem $\rightarrow$ really used for ad placement

Example applications

L37-5

→ Game playing

→ Automatically flying helicopters

state/situation: set of pos/velocities +

sensors
action: how much torque on rotor?

reward: $\begin{cases} +100 & \text{for reaching destination} \\ -1000 & \text{for crashing} \\ 0 & \text{otherwise} \end{cases}$

⇒ Train! → simulations

choosing actions → max preferred (inmate bias of alg)

Robots → same

→ Stock market!