

09/10

for episode = 1, 2, ...

$S \sim d_0$

for $t = 0, 1, 2, \dots$

$a = \text{agent.getAction}(s)$

$s' = P(s, a, \cdot)$

$r = R(s, a, s')$ // Deterministic reward

$\# \text{agent.train}(s, a, r, s')$

if $(s' = s^*)$

break ;

$S = s'$

① Model Based Approach $\rightarrow$ Estimate P, R plan
Model Free Approach

② Model Free Approach $\rightarrow$ Value function based

③ Policy Gradient

④ Black Box Optimization (BBO)

Read CMA-ES

BBO for policy search:

- Simple
- Ignores MDP's structure.
- Phrases the problem as a Black Box Optimization problem.

$$\operatorname{argmax}_{\pi} J(\pi)$$

- Only have access to an estimate $\hat{J}(\pi)$
Run π for N episodes and average the returns.

$\gamma < 1 \Rightarrow$ Rewards are bounded.

$$\hat{J}(\pi) = \frac{1}{N} \sum_{i=1}^N G^i = \frac{1}{N} \sum_{i=1}^N \sum_{t=0}^{\infty} \gamma^t R_t^i$$

R_t^i = Reward at time t
at i^{th} episode

$$G_x^i = \sum_{t=0}^{\infty} \gamma^t R_t^i$$

How to represent π ?

Tabular representation.

	a_1	a_2	a_3
States	s_1		
	s_2		
	s_3		
	$\vdots$		
	$\vdots$		

"Parameters" are numbers we store to represent
Something $\rightarrow$ policy.

- Tabular policy has $|S| \times |A|$

All $|S| \times |A|$ matrices are valid policies if

- Rows all sum to 1

- All entries ≥ 0 .

Tabular Softmax Action Selection:-

Want to store π as an $|S| \times |A|$

- No constraints on p .

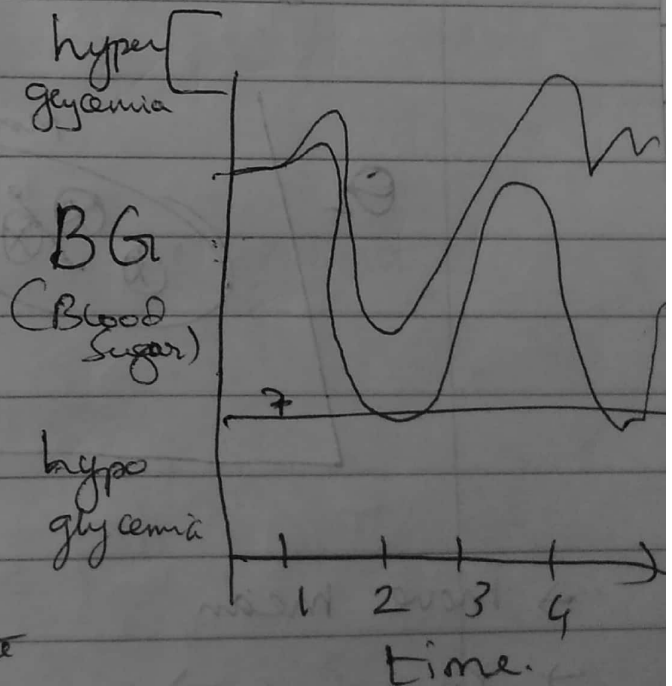
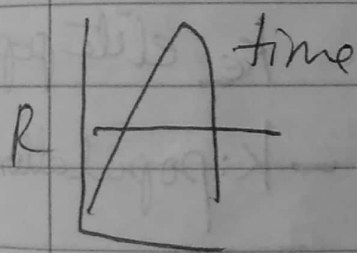
- Increasing $p(S, a)$ increases $\pi(S, a)$

$$\pi(s,a) = \frac{e^{p(s,a)}}{\sum_a e^{p(s,a)}} \rightarrow \text{Boltzmann distribution}$$

$$\sum_a \pi(s,a) = \sum_a \frac{e^{p(s,a)}}{\sum_{a'} e^{p(s,a')}} = \frac{\sum_a e^{p(s,a)}}{\sum_{a'} e^{p(s,a')}} = 1$$

$$e^{p(s,a)} > 0 \quad \forall p(s,a) \in \mathbb{R}$$

Policy Parameter Vector $\theta = \text{vector}(p)$.



$$\text{injection} = \frac{\text{State } \uparrow \text{ BG} - \text{State } \uparrow \text{ BG}_{\text{target}}}{\text{CF } \theta_1} + \frac{\text{Size-meal}}{\text{CF } \theta_2}$$

$$\theta = \begin{bmatrix} \theta_1 \\ \theta_2 \end{bmatrix}$$

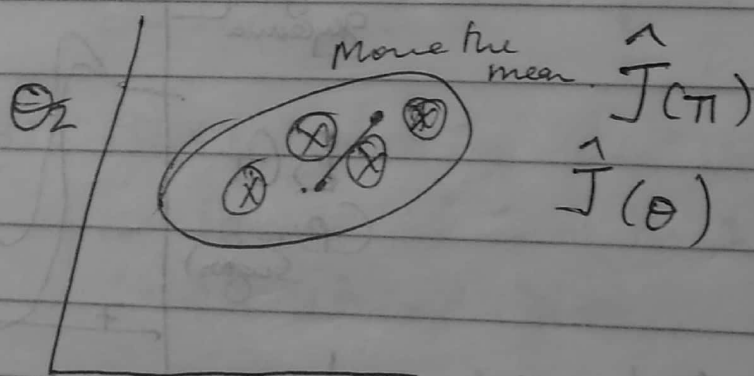
Deep learning: $ANN_{\theta}(S_t)$

Cross-Entropy Method:

- ① Generate a random data sample according to a parameterized mechanisms
- ② Update the parameters of the mechanisms based on the data to provide a better sample at the next iteration

— Mechanism = distribution over policies

— Random sample = histories generated by sample policies



k_e : elite population

k : population

$$k_e < k$$

θ_1

→ move mean

→ Covariance $\Rightarrow$ covariance of top k_e elite

N : Number of episodes on each sample policy

CEM

Input: θ : Mean parameter vector

Σ : Covariance matrix $\Sigma = \sigma^2 \mathbf{I}$

K : Population; K_e : Elite Population

Number of times to each policy,

1) for $k = 1: K$

2) $\theta_k \sim N(\theta, \Sigma)$

3) $\hat{J}_k = \text{evaluate}(\theta_k, N)$

4) Sort $(\theta_{1:k}, \hat{J}_{1:k}, \text{descending})$

5) $\theta \leftarrow \frac{1}{K_e} \sum_{k=1}^{K_e} \theta_k$

6) $\Sigma \leftarrow \frac{1}{K_e} \sum_{k=1}^{K_e} (\theta_k - \theta)(\theta_k - \theta)^T$

Numerical instability for small values.

$$\Sigma \leftarrow \frac{1}{K_e + \epsilon} \left(\epsilon \mathbf{I} + \sum_{k=1}^{K_e} (\theta_k - \theta)(\theta_k - \theta)^T \right)$$