

Optimization

Mahdi Roozbahani
Georgia Tech

Outline

Motivation

Entropy

Conditional Entropy and Mutual Information

Cross-Entropy and KL-Divergence



Let's work on this subject in our Optimization lecture

Cross Entropy

Cross Entropy: The expected number of bits when a wrong distribution Q is assumed while the data actually follows a distribution P

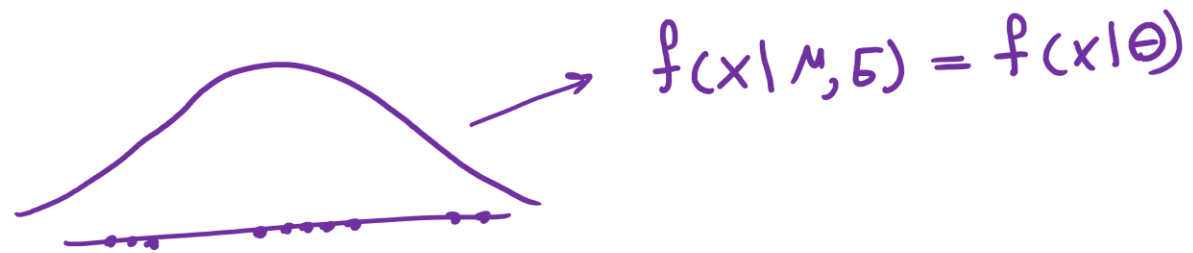
$$H(p, q) = - \sum_{x \in \mathcal{X}} p(x) \log q(x) = H(P) + KL[P][Q]$$

This is because:

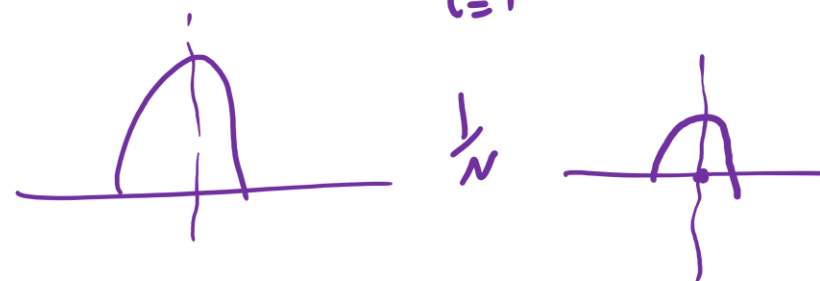
$$H(p, q) = \mathbb{E}_p[l_i] = \mathbb{E}_p \left[\log \frac{1}{q(x_i)} \right]$$

$$H(p, q) = \sum_{x_i} p(x_i) \log \frac{1}{q(x_i)}$$

$$H(p, q) = - \sum_x p(x) \log q(x).$$



$$\max l(\theta | x) = \sum_{i=1}^N \log f(x^{(i)} | \theta)$$



$$\min -\frac{1}{N} l(\theta | x) = -\frac{1}{N} \sum_{i=1}^N \log \underbrace{f(x^{(i)} | \theta)}_Q$$

$$\frac{1}{N} \sum (\cdot) = E[\cdot]$$

p — True actual pdf
 Q — predicted pdf model

$$= \left(\frac{1}{N} \sum -\log Q \right) = E[-\log Q] = \text{avg negative log likelihood}$$

$$E[-\log p + \log p - \log Q]$$

$$E[-\log P + \log P - \log Q]$$

$$\log a - \log b = \log \frac{a}{b} = -\log \frac{b}{a}$$

$$E[a+b+c] = E[a] + E[b+c]$$

$$E[g(x)] = \sum_{i=1}^N P(x_i) g(x_i)$$

$$E[-\log P] + E[-\log \frac{Q}{P}]$$

$$\sum_{i=1}^N P(x_i) \log \frac{1}{P(x_i)} + \sum P(x_i) \left(-\log \frac{Q(x_i)}{P(x_i)} \right)$$

$$H(x) + KL[P][Q] = CE = \overset{\text{avg}}{\text{Negative log likelihood}}$$

Labeling target values



Label encoding (ordinal) and One-hot encoding

$$X = \begin{bmatrix} h & w & a & \dots \\ \vdots & \vdots & \vdots & \ddots \end{bmatrix}_{n \times d}$$

$$Y_a = \begin{bmatrix} \text{cat} \\ \text{fish} \\ \text{dog} \\ \text{cat} \\ \vdots \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \\ 2 \\ 0 \\ \vdots \end{bmatrix} \leadsto \text{ML} \quad \hat{Y}_p = \begin{bmatrix} 0 \\ 0.5 \\ 1.2 \\ \vdots \end{bmatrix}$$

$$\text{cat} = [1 \ 0 \ 0] \quad \text{fish} = [0 \ 1 \ 0] \quad \text{dog} = [0 \ 0 \ 1]$$

$$Y_a = \begin{bmatrix} [1 \ 0 \ 0] \\ [0 \ 1 \ 0] \\ [0 \ 0 \ 1] \\ [1 \ 0 \ 0] \\ \vdots \end{bmatrix} \leadsto \text{ML} \Rightarrow \hat{Y}_p = \begin{bmatrix} [70 \ 20 \ 7] \\ [20 \ 137 \ 2] \\ [5 \ 6 \ 80] \\ \vdots \end{bmatrix} \leadsto \text{Softmax} \quad \hat{Y}_p = \begin{bmatrix} [0.7 \ 0.2 \ 0.1] \\ [0.3 \ 0.6 \ 0.1] \\ [0.1 \ 0.1 \ 0.8] \\ \vdots \end{bmatrix}$$

$$\text{Min } \|Y_a - \hat{Y}_p\|_2^2 = \sum_{i=1}^N \left[(1 - 0.7)^2 + (0 - 0.2)^2 + (0 - 0.1)^2 \right] + \left[(0 - 0.3)^2 + (1 - 0.6)^2 + (0 - 0.1)^2 \right] + \dots$$

$$\text{Max } Y_a \hat{Y}_p^T = [1 \ 0 \ 0] \begin{bmatrix} 0.7 \\ 0.2 \\ 0.1 \end{bmatrix} + [0 \ 1 \ 0] \begin{bmatrix} 0.3 \\ 0.6 \\ 0.1 \end{bmatrix} + \dots = \mathcal{N}$$

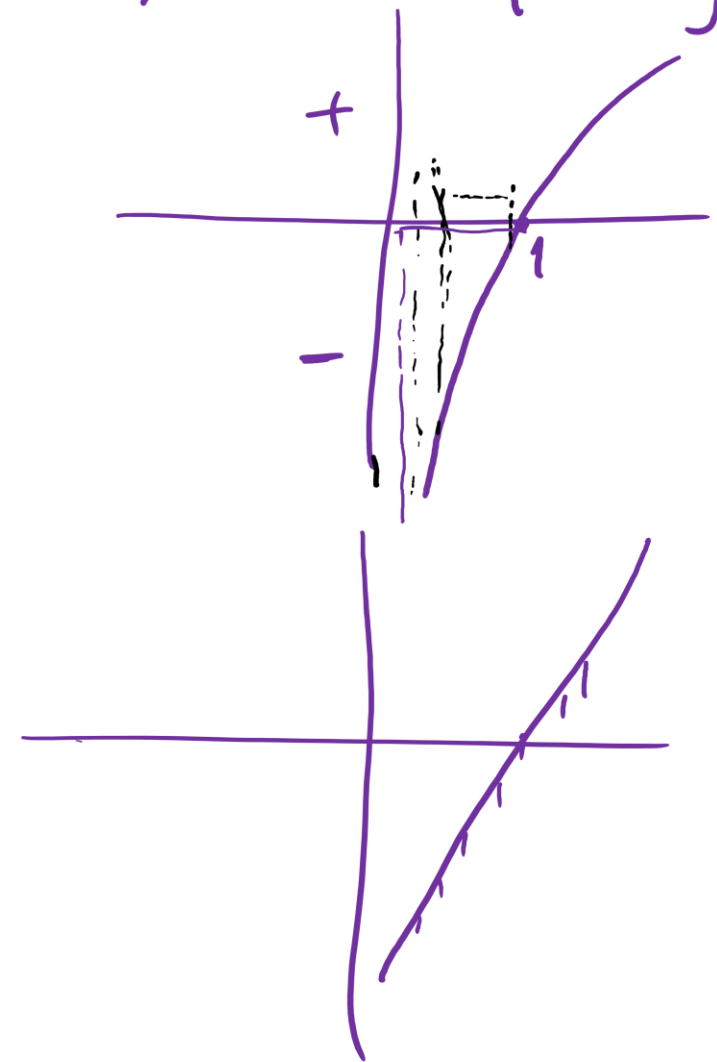
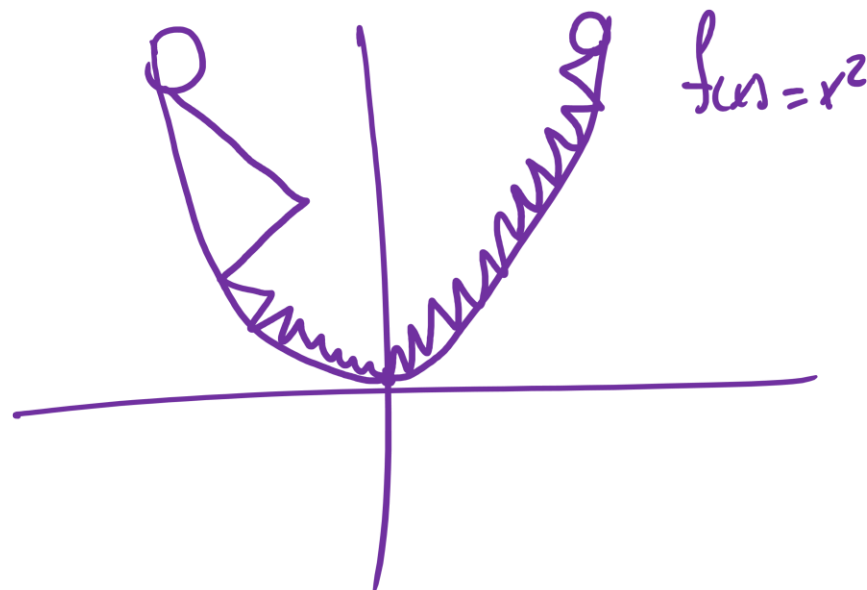
Why Cross entropy and not simply use dot product?

$$CE = - \sum P(x) \log Q(x) = - \sum y_a \log \hat{y}_p \quad y_a = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \end{bmatrix} \quad \hat{y}_p = \begin{bmatrix} 0.7 & 0.2 & 0.1 \\ 0.3 & 0.6 & 0.1 \end{bmatrix}$$

$$CE = - \left[\begin{bmatrix} 1 & 0 & 0 \end{bmatrix} \begin{bmatrix} \log 0.7 \\ \log 0.2 \\ \log 0.1 \end{bmatrix} + \begin{bmatrix} 0 & 1 & 0 \end{bmatrix} \begin{bmatrix} \log 0.3 \\ \log 0.6 \\ \log 0.1 \end{bmatrix} + \dots \right]$$

① $\text{Min } CE = y_a \log \hat{y}_p^T$

② $\text{Min } - y_a \hat{y}_p^T$



Kullback-Leibler Divergence

Another useful information theoretic quantity measures the difference between two distributions.

$$\begin{aligned}\mathbf{KL}[P(S)||Q(S)] &= \sum_s P(s) \log \frac{P(s)}{Q(s)} \\ &= \underbrace{\sum_s P(s) \log \frac{1}{Q(s)}}_{\text{cross entropy}} - \mathbf{H}[P] = H(P, Q) - H(P)\end{aligned}$$

Excess cost in bits paid by encoding according to Q instead of P .

KL Divergence is
a **KIND OF**
distance
measurement

$$-\mathbf{KL}[P||Q] = \sum_s P(s) \log \frac{Q(s)}{P(s)}$$

log function is
concave or
convex?

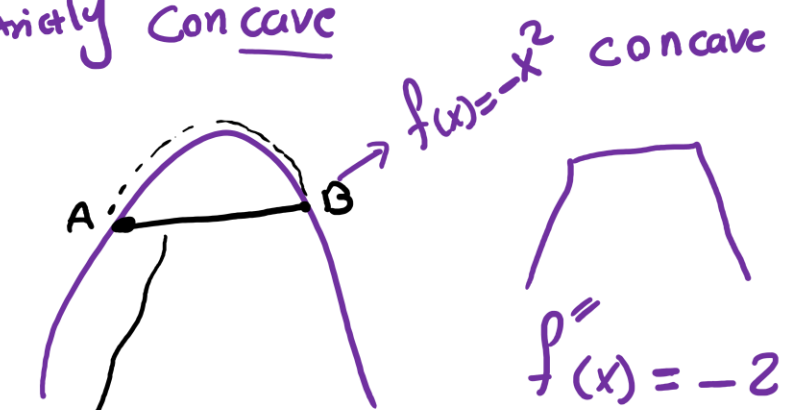
$$\begin{aligned}\sum_s P(s) \log \frac{Q(s)}{P(s)} &\leq \log \sum_s P(s) \frac{Q(s)}{P(s)} \\ &= \log \sum_s Q(s) = \log 1 = 0\end{aligned}$$

By Jensen Inequality

So $\mathbf{KL}[P||Q] \geq 0$. Equality iff $P = Q$

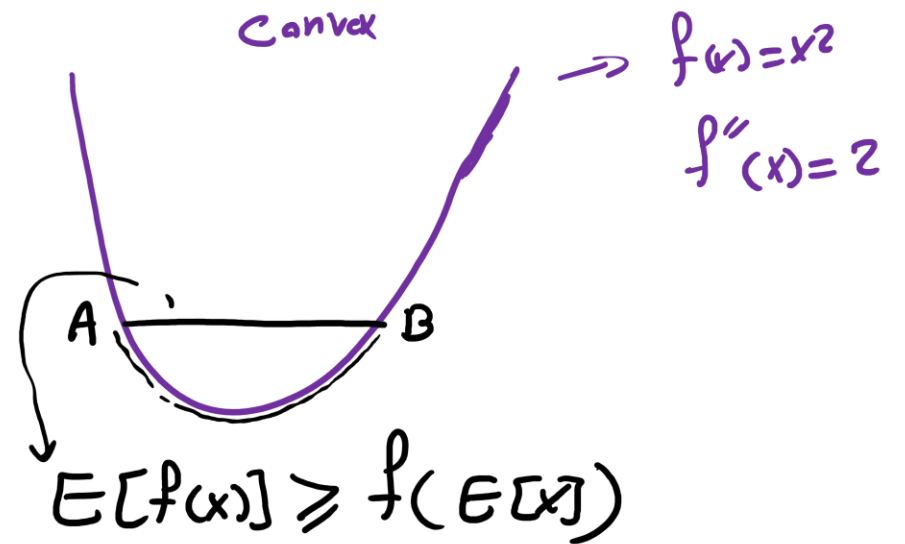
When $P = Q$, $KL[P||Q] = 0$

Strictly Concave



$$E[f(x)] \leq f(E[x])$$

Convex



$$E[f(x)] \geq f(E[x])$$

$$f(x) = x^2 \quad x = [1 \quad 2]$$

$$p(x) = [\frac{1}{2} \quad \frac{1}{2}]$$

$$E[f(x)] = \sum p(x) f(x)$$

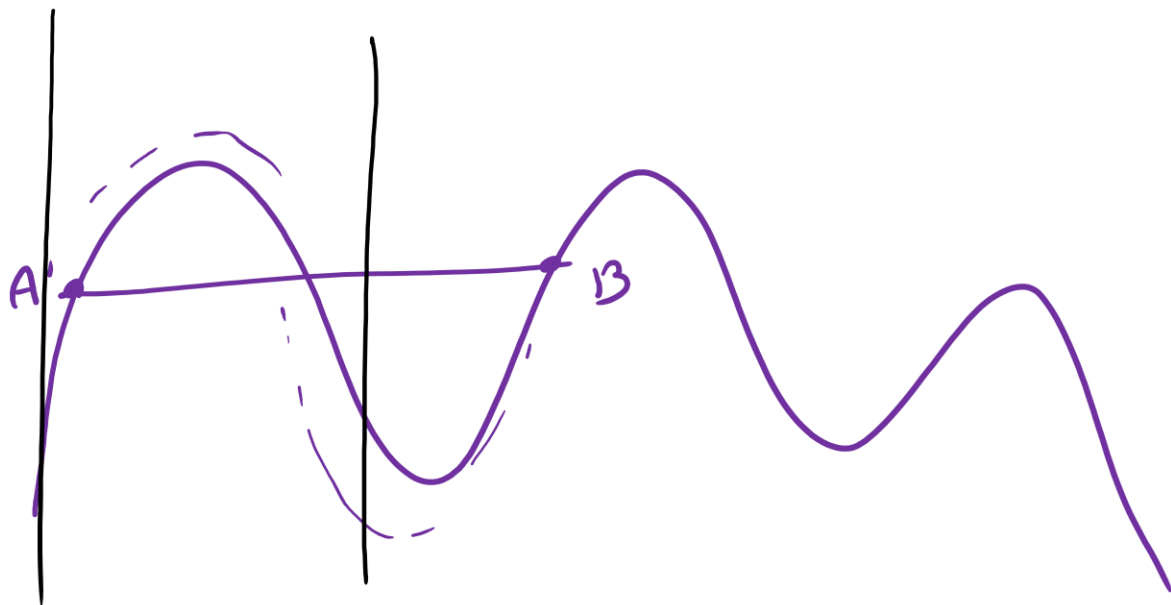
$$f(x=1) = 1 \quad f(x=2) = 4$$

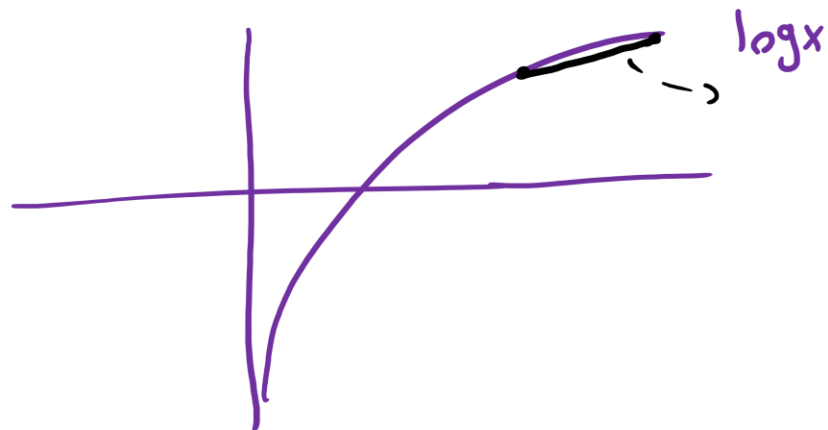
$$E[f(x)] = \frac{1}{2} * 1 + \frac{1}{2} * 4 = 2.5$$

$$f(E[x]) = ? \quad E[x] = \sum p(x) x$$

$$E[x] = \frac{1}{2} * 1 + \frac{1}{2} * 2 = 1.5$$

$$f(E[x] = 1.5) = 2.25$$





$$E[f(x)] \leq f(E[x]) \Rightarrow E[\log x] \leq \log(E[x])$$

$$E[g(x)] = \sum p(x) g(x)$$

$$-KL[P][Q] = \sum p(x) \log \left(\frac{Q(x)}{P(x)} \right) = \sum p(x) \log g(x)$$

$$\sum p(x) \log g(x) = E[\log g(x)]$$

$$E[\log g(x)] \leq \log(E[g(x)])$$

$$\leq \log \sum p(x) g(x)$$

$$\leq \log \sum \cancel{p(x)} \frac{Q(x)}{\cancel{p(x)}}$$

$$\leq \log \sum Q(x)$$

$$-KL[P][Q] \leq 0 \Rightarrow KL[P][Q] \geq 0$$

Optimization $\min f(x)$ $\rightarrow$ Objective function
 s.t $g(x) = c \rightarrow$ Equality constraint

$\left. \begin{array}{l} g(x) \leq c \\ g(x) \geq c \end{array} \right\} \Rightarrow$ Inequality constraint $\leadsto$ KKT conditions

Lagrange function

$\rightarrow f(x, y)$
 s.t $(x + y = 2)$

$g(x, y) = x + y - 2$

$$H = \begin{bmatrix} \frac{\partial^2 f(x, y)}{\partial x^2} & \frac{\partial^2 f(x, y)}{\partial x \partial y} \\ \frac{\partial^2 f(x, y)}{\partial y \partial x} & \frac{\partial^2 f(x, y)}{\partial y^2} \end{bmatrix}$$

- ① Linear programming $\rightarrow \left. \begin{array}{l} f(x, y) = x + y \\ \text{s.t. } x + y = 2 \end{array} \right\}$
- ② Quadratic programming $\leadsto \left. \begin{array}{l} f(x, y) = x^2 + y \\ \text{s.t. } x + y = 2 \end{array} \right\}$
- ③ Non linear programming $\leadsto \left. \begin{array}{l} f(x, y) = x^3 + y^2 \\ \text{s.t. } x^2 + y = 2 \end{array} \right\}$

$$\min f(M, S) = 6M^2 + 3S^2$$

$$\frac{\partial f(M, S)}{\partial M} = 0 \Rightarrow 12M + 0 = 0 \Rightarrow M = 0$$

$$\frac{\partial f(M, S)}{\partial S} = 0 \Rightarrow 0 + 6S = 0 \Rightarrow S = 0$$

$M = \#$ hours you study ML/day

$S = \#$ hours you sleep/day

$$f(M, S) = 6M^2 + 3S^2 \leadsto \text{objective function}$$

λ, α, β

$$\text{s.t. } M + S = 24 \leadsto g(M, S) = M + S - 24 \rightarrow \text{constraint function}$$

$$L(M, S, \lambda) \xrightarrow{\text{Lagrange multiplier}} = f(M, S) - \lambda g(M, S)$$

$$L(M, S, \lambda) = 6M^2 + 3S^2 - \lambda (M + S - 24)$$

$$\frac{\partial L}{\partial \lambda} = 0 \Rightarrow 0 + 0 - (M + S - 24) = 0 \Rightarrow M + S = 24 \Rightarrow \frac{\lambda}{12} + \frac{\lambda}{6} = 24 \Rightarrow \lambda = 96$$

$$\frac{\partial L}{\partial M} = 0 \Rightarrow 12M + 0 - \lambda(1) = 0 \Rightarrow M = \frac{\lambda}{12} = \frac{96}{12} = 8$$

$$\frac{\partial L}{\partial S} = 0 \Rightarrow 6S - \lambda = 0 \Rightarrow S = \frac{\lambda}{6} = \frac{96}{6} = 16$$

$$M + S = 24 \quad 8 + 16 = 24$$

$$f(M, S) = 6M^2 + 3S^2$$

$$\text{s.t. } M + S = 24 \Rightarrow g(M, S) = M + S - 24 \rightarrow \delta_1$$

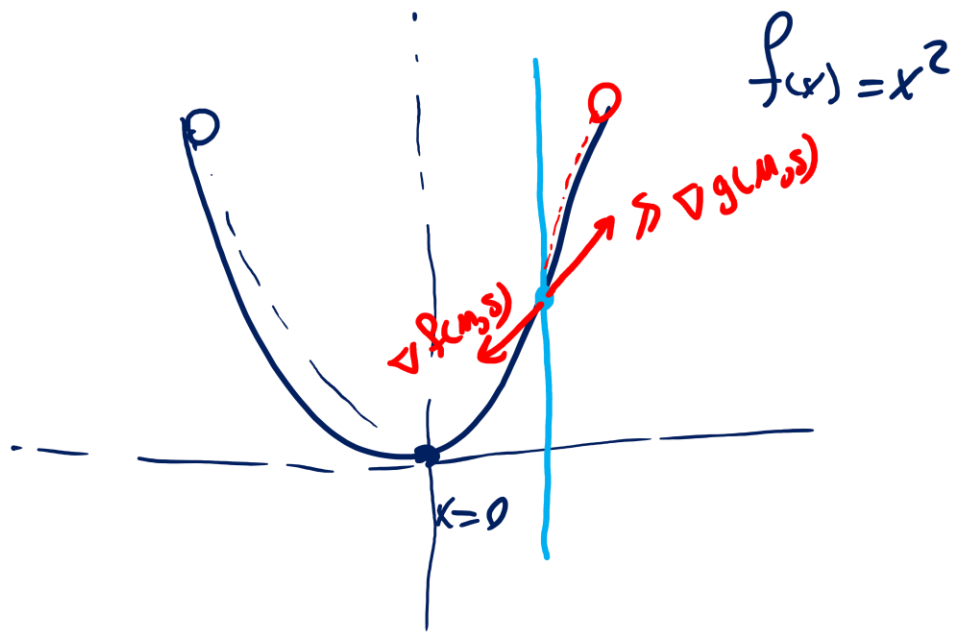
$$\text{s.t. } M - S = 12 \Rightarrow h(M, S) = M - S - 12 \sim \delta_2$$

$$L(M, S, \delta_1, \delta_2) = f(M, S) - \delta_1 g(M, S) - \delta_2 h(M, S)$$

$$L(u, s, \lambda) = f(u, s) - \lambda g(u, s)$$

$$\nabla L(u, s, \lambda) = 0 \Rightarrow \nabla (f(u, s) - \lambda g(u, s)) = 0$$

$$\nabla f(u, s) = \lambda \nabla g(u, s)$$



Min $f(m,s) = 6m^2 + 3s^2$
 s.t. $m+s \leq 24$ $g(m,s) = m+s-24 \leq 0$ — $\begin{cases} \rightarrow g(m,s)=0 \leadsto \text{Active solution} \\ \rightarrow g(m,s)<0 \leadsto \text{Inactive solution} \end{cases}$

$$L(m,s,\lambda) = 6m^2 + 3s^2 + \lambda(m+s-24)$$

① Stationary condition $\frac{\partial L}{\partial m} = 0 \dots$

② Primal feasibility $g(m,s) \leq 0$

③ Dual feasibility $\lambda \geq 0$

④ Complementary slackness $g(m,s)\lambda = 0$ — $\begin{cases} \rightarrow g(m,s) \neq 0 \leadsto \lambda = 0 \\ \rightarrow \lambda > 0 \leadsto g(m,s) = 0 \end{cases}$

Inactive

Active

convex $\leadsto$ min $\leadsto$ primal form

$$f(M, S) = \frac{1}{2}M^2 + \frac{1}{2}S^2$$

$$M + S = 24$$

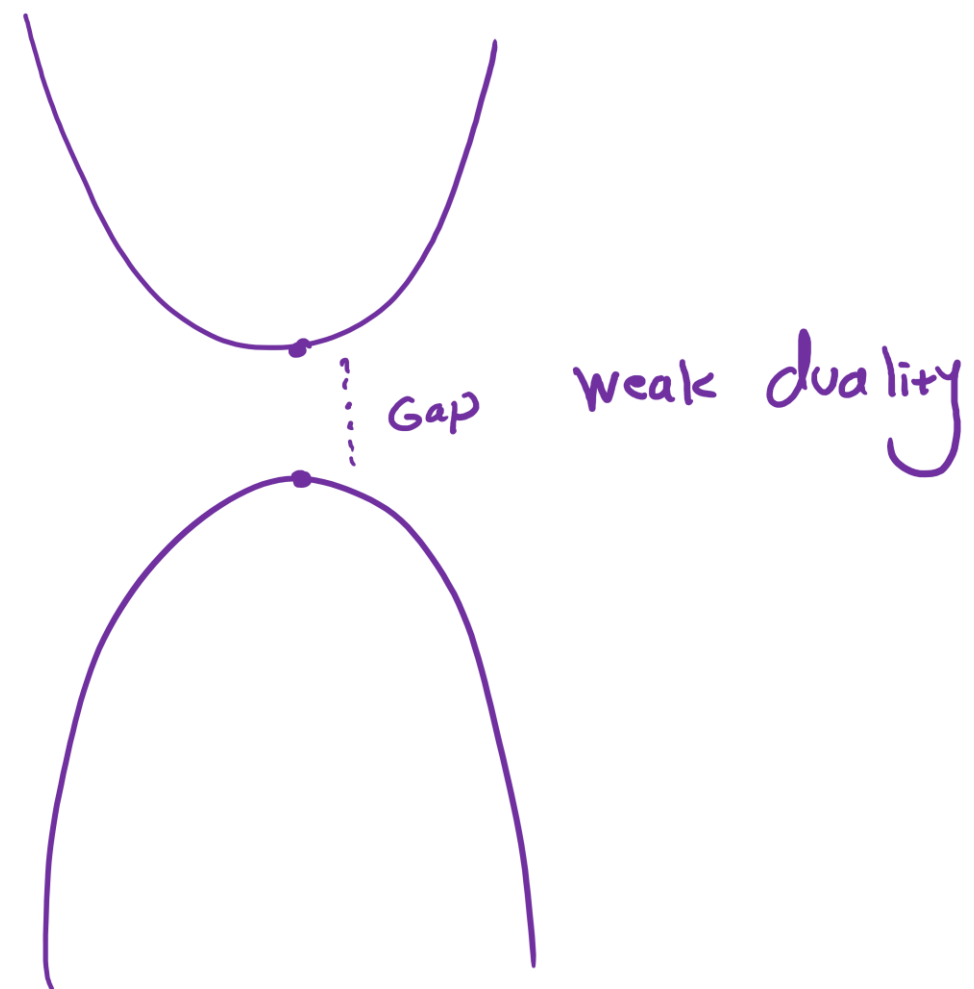
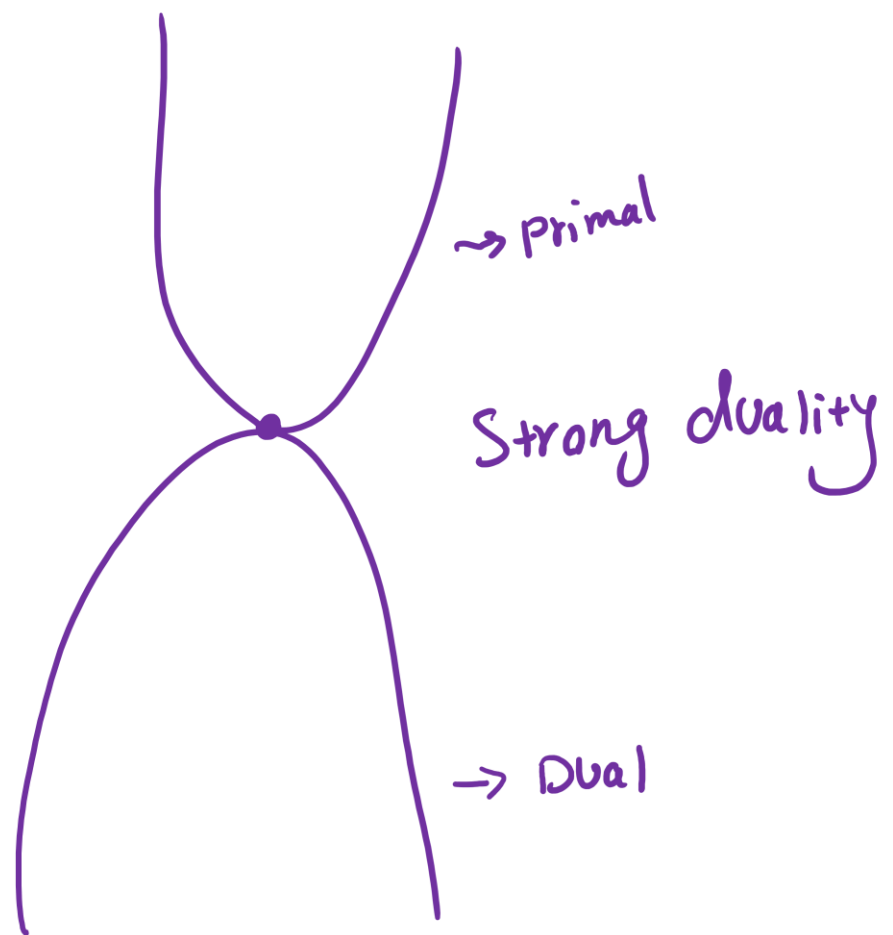
$$L(M, S, \lambda) = \frac{1}{2}M^2 + \frac{1}{2}S^2 - \lambda(M + S - 24)$$

$$\frac{\partial L}{\partial M} = 0 \Rightarrow M - \lambda = 0 \Rightarrow M = \lambda \quad \frac{\partial L}{\partial S} = 0 \Rightarrow S = \lambda$$

$$L(\lambda) = \frac{\lambda^2}{2} + \frac{\lambda^2}{2} - \lambda(\lambda + \lambda - 24) = \lambda^2 - 2\lambda^2 + 24\lambda = -\lambda^2 + 24\lambda$$

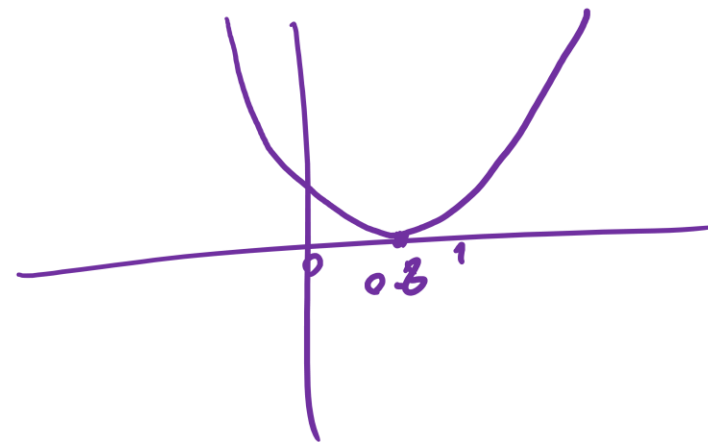
Dual form $\rightarrow$ concave $\rightarrow$ max

$$\frac{\partial L(\lambda)}{\partial \lambda} = 0 \Rightarrow -2\lambda + 24 = 0 \Rightarrow \lambda = 12$$



$$f(x) = e^x + x^2 \quad f''(x) = e^x + 2 > 0$$

$$f'(x) = e^x + 2x = 0$$



$$f(x) = x^2$$

$$f'(x) = 2x = 0$$

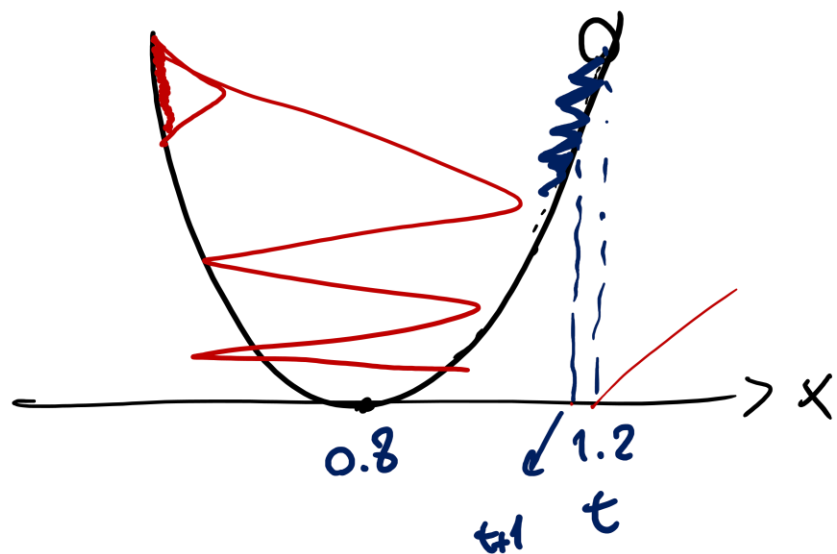
$$x = 0$$

Gradient descent ← $\{t+1\}$ $\{t\}$
 Gradient Ascent $X \leftarrow X$ $+$

$$\alpha (f'(x))$$

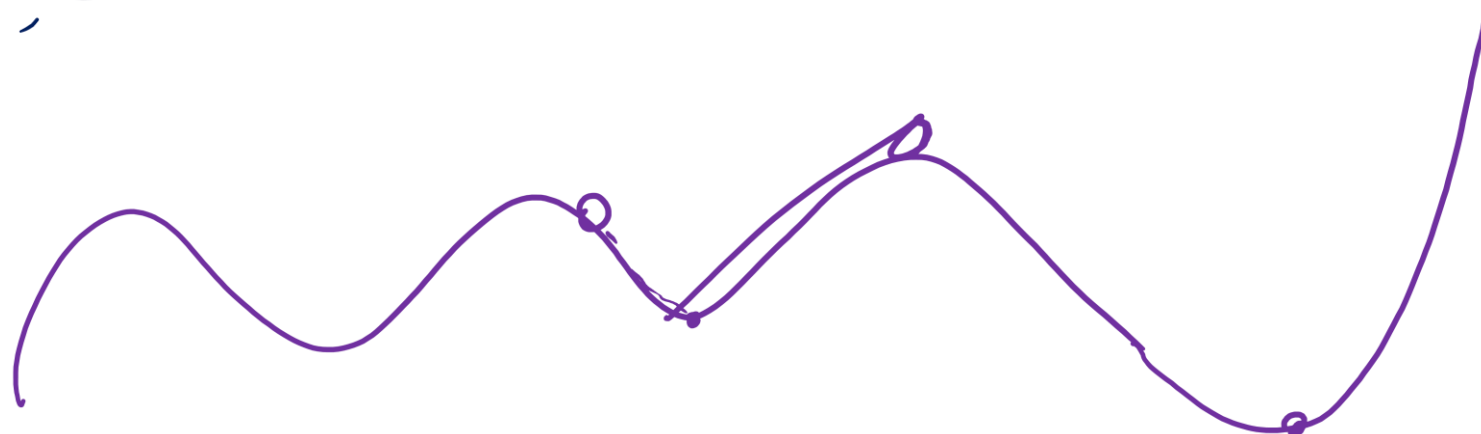
Dynamic Adam Optimizer

learning
step
0.01



$$\{t+1\}$$

$$X \leftarrow 1.2 - 0.01 (e^{1.2} + 2 \times 1.2)$$



Derivation rules:

Quotient rule: $\left(\frac{f}{g}\right)' = \frac{f'g - fg'}{g^2}$

Product rule: $(fg)' = f'g + fg'$

Chain rule: $(u^n)' = nu'u^{n-1}$