

Equivalent definitions of convexity
 Showing a function is convex
 Gradient Methods

"Strictly" convex
 $f(\theta x + (1-\theta)y) < \theta f(x) + (1-\theta)f(y)$
 —implies at most one g

"Concave" : $(-f)$ is co

"Log-convex" : $(\log f(x))$

— Pick up A2 to mark

— A3 due now.

— Pick A3 and course eval at end of class.

— No tutorials this week

(return now or Monday)

(A4 out tomorrow, due Oct 15)

— Why do support vectors make kernel methods more efficient

Usual

$$K(X, X) = \begin{bmatrix} & N \\ N & \end{bmatrix}$$

Support vectors

$$K(X_m, X_{sv}) = \begin{bmatrix} & \#sv \\ \#sv & \end{bmatrix}$$

$$K(\tilde{X}, X_{sv}) = \begin{bmatrix} & \#sv \\ T & \end{bmatrix}$$

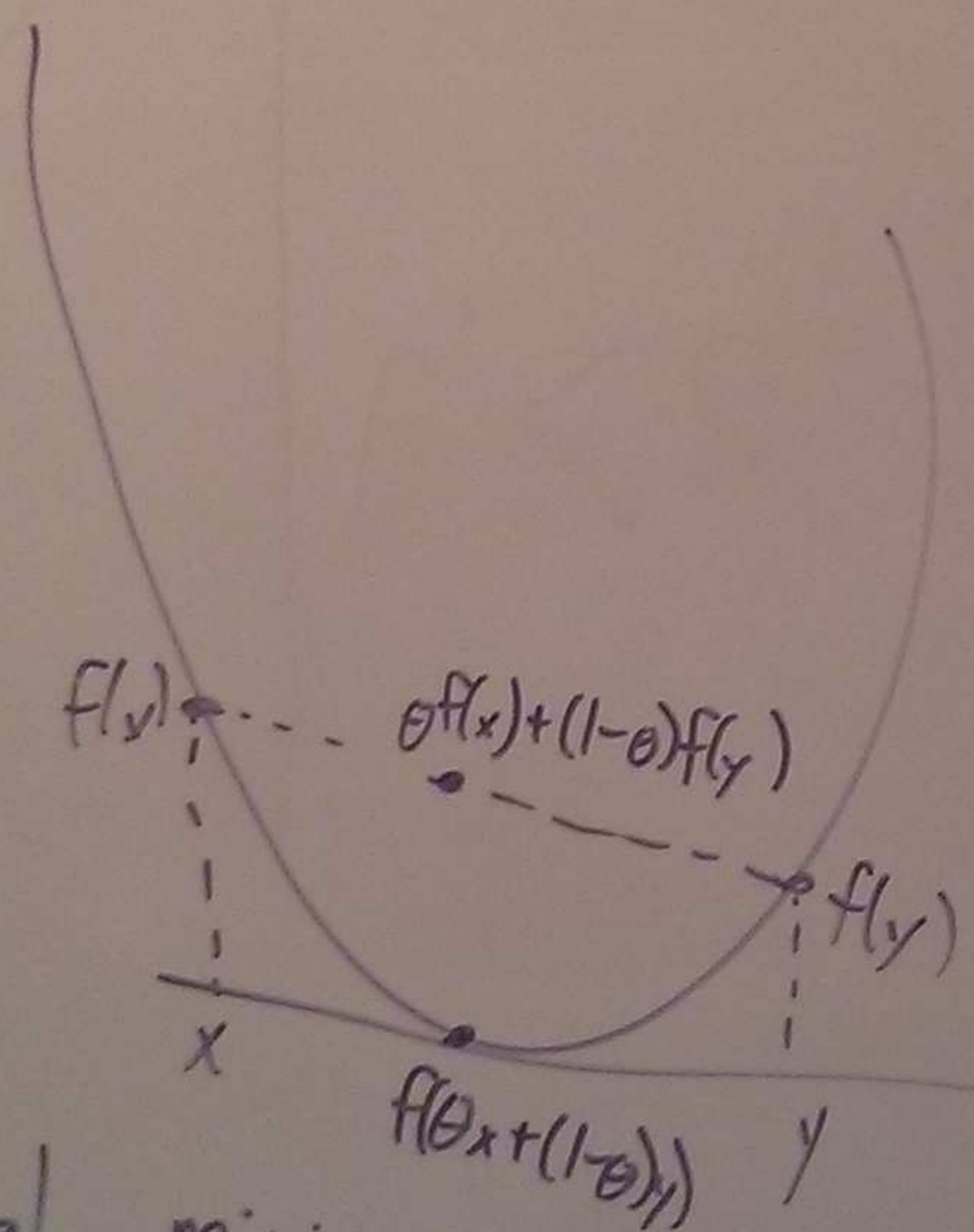


Convex Functions (zero-order definition)

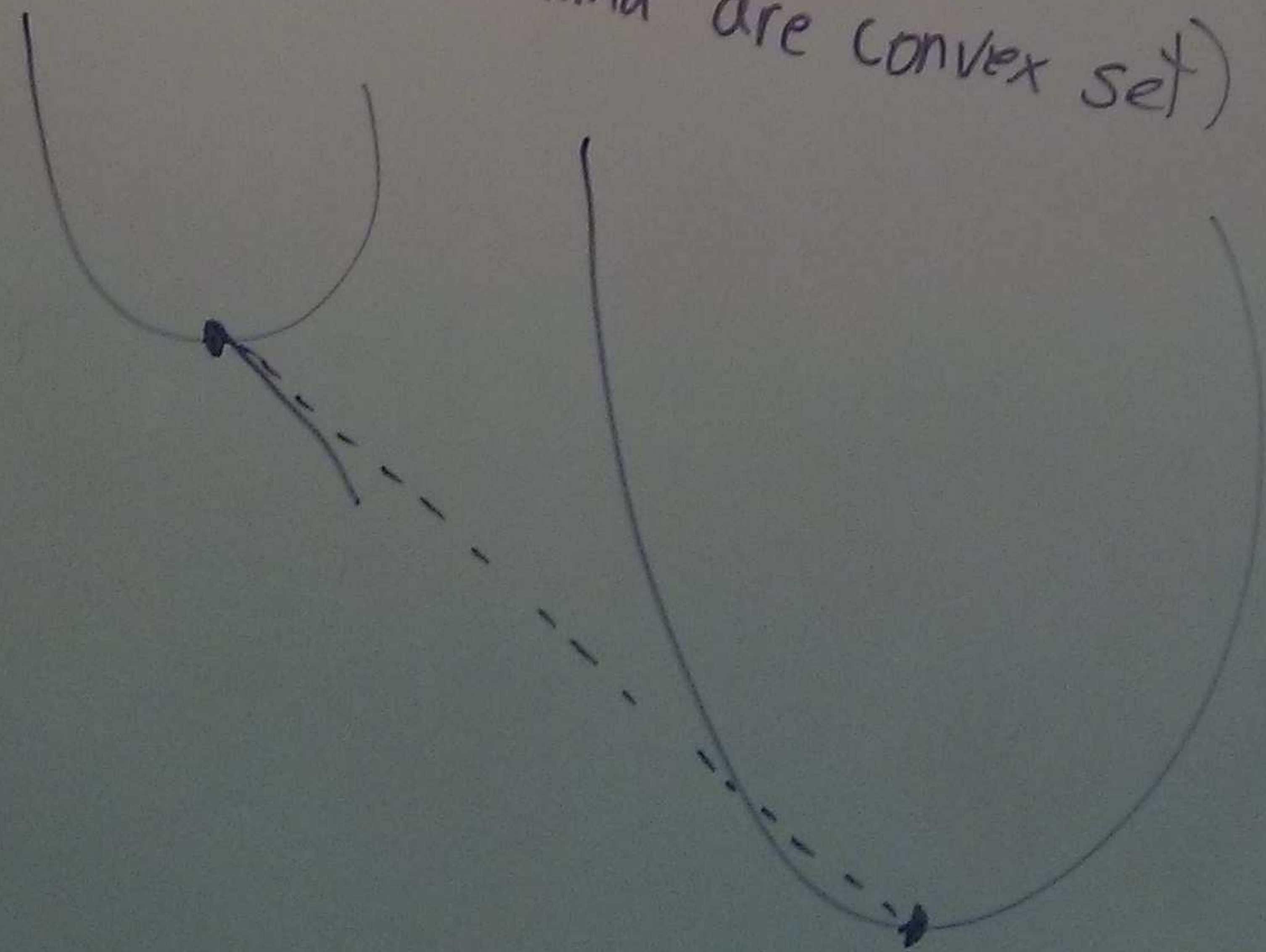
Function 'f' is convex if $\text{dom}(f)$ is convex, and

$$f(\theta x + (1-\theta)y) \leq \theta f(x) + (1-\theta)f(y)$$

"below the chord", "epigraph is convex set"



- Implies local minima are global minima
(and global minima are convex set)



Examples:

$$f(x) = \exp(ax)$$

$$f(x) = x \log x, x > 0$$

$$f(x) = ax^2 + bx, a \geq 0$$

$$f(\bar{x}) = \max_i \{x_i\}$$

$$f(\bar{x}) = \log\left(\sum_{i=1}^d \exp(x_i)\right)$$

$$f(x) = -\log(x)$$

Conve.

"Strictly" convex

$$f(\theta x + (1-\theta)y) < \theta f(x) + (1-\theta)f(y),$$

—implies at most one global minimum

"Concave" $\because (-f)$ is convex

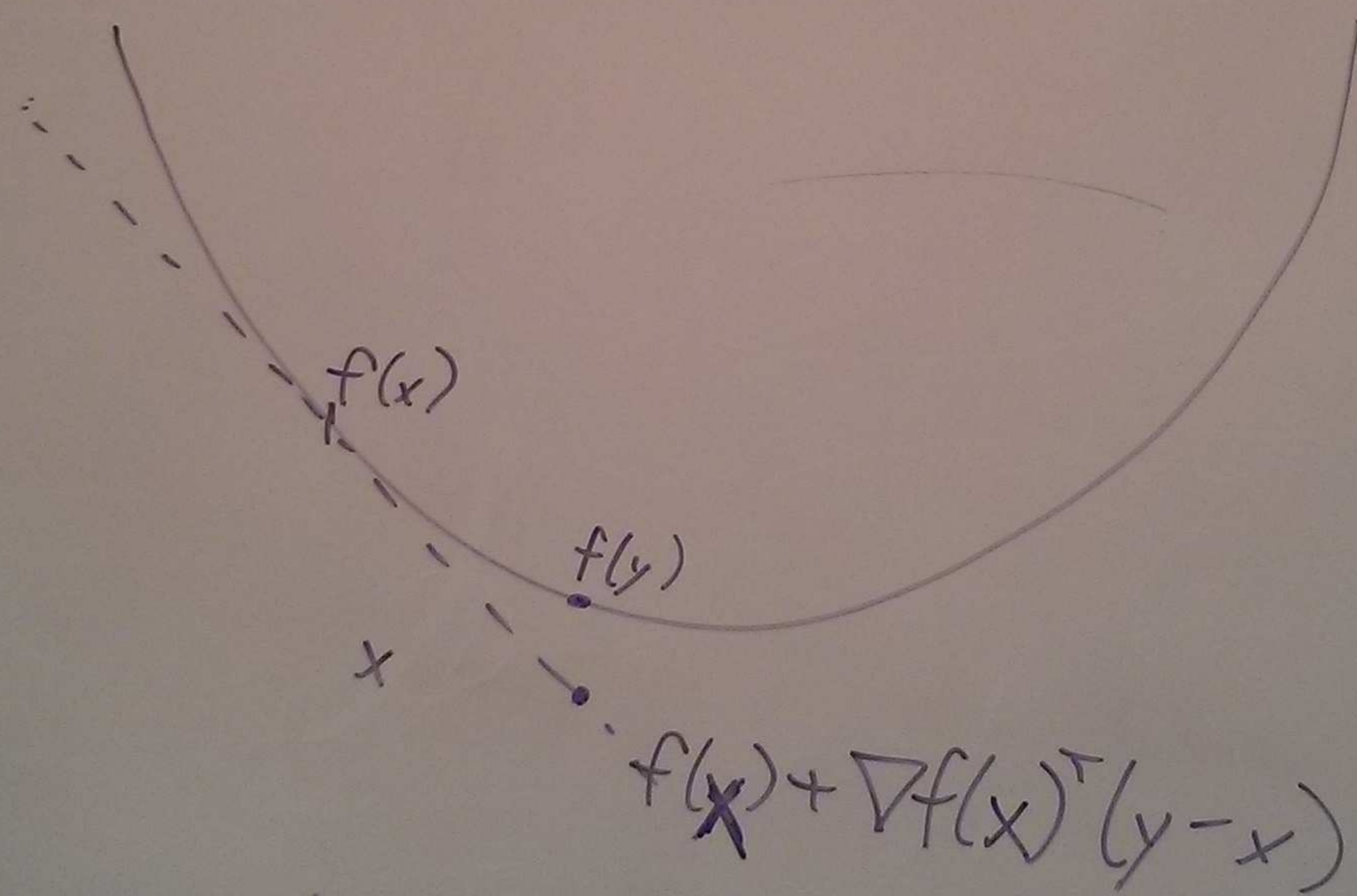
"Log-convex" $\because (\log f(x))$ is convex

Convex Functions (first-order definition)

If 'f' is differentiable, equivalent to

$$f(y) \geq f(x) + \nabla f(x)^T (y-x), \quad \forall y, x \in \text{dom}(f)$$

"f is above its tangent"



> 0

$a \geq 0$

}

$\exp(x_i)$

- Can see that any stationary point ($\nabla f(x^*) = 0$) is a global optimum,

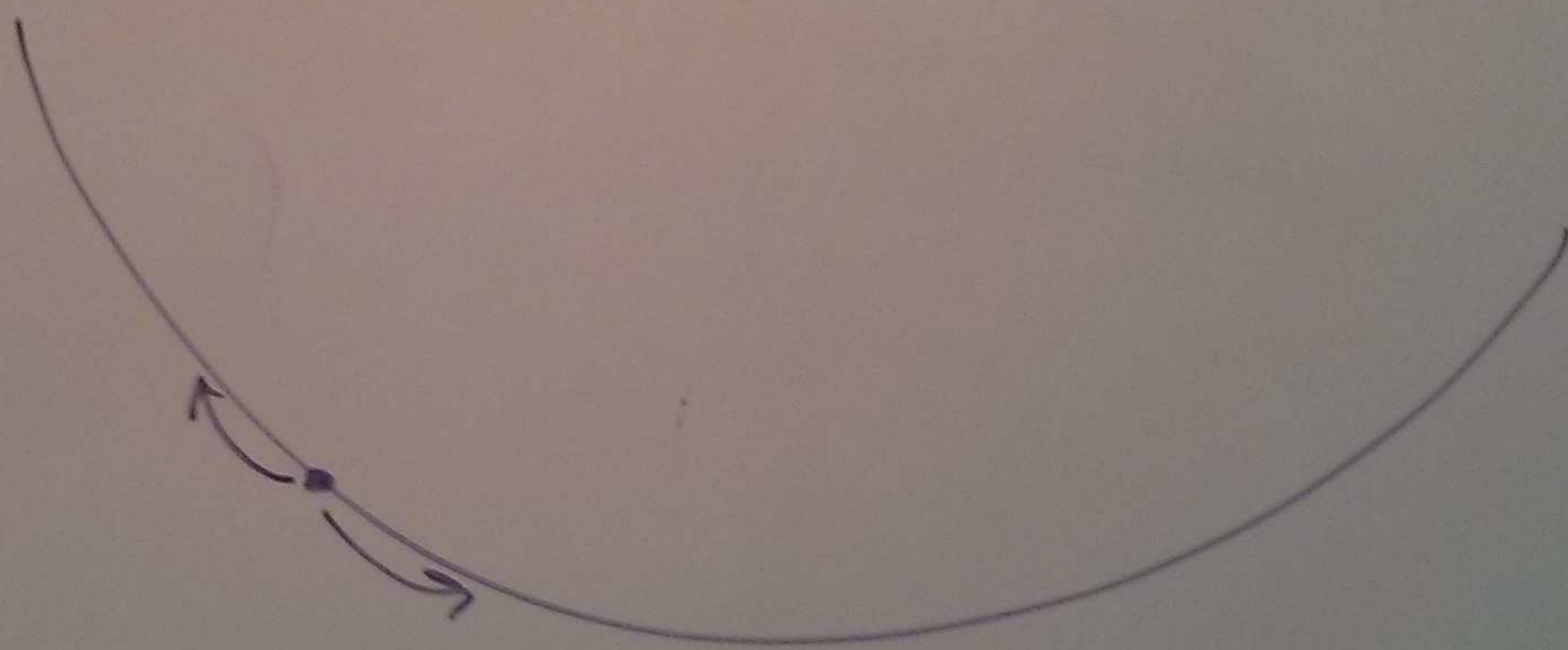
$$f(y) \geq f(x^*) + 0, \quad \forall y$$

Convex Functions (second order definition)

If 'f' is twice-differentiable, equivalent to

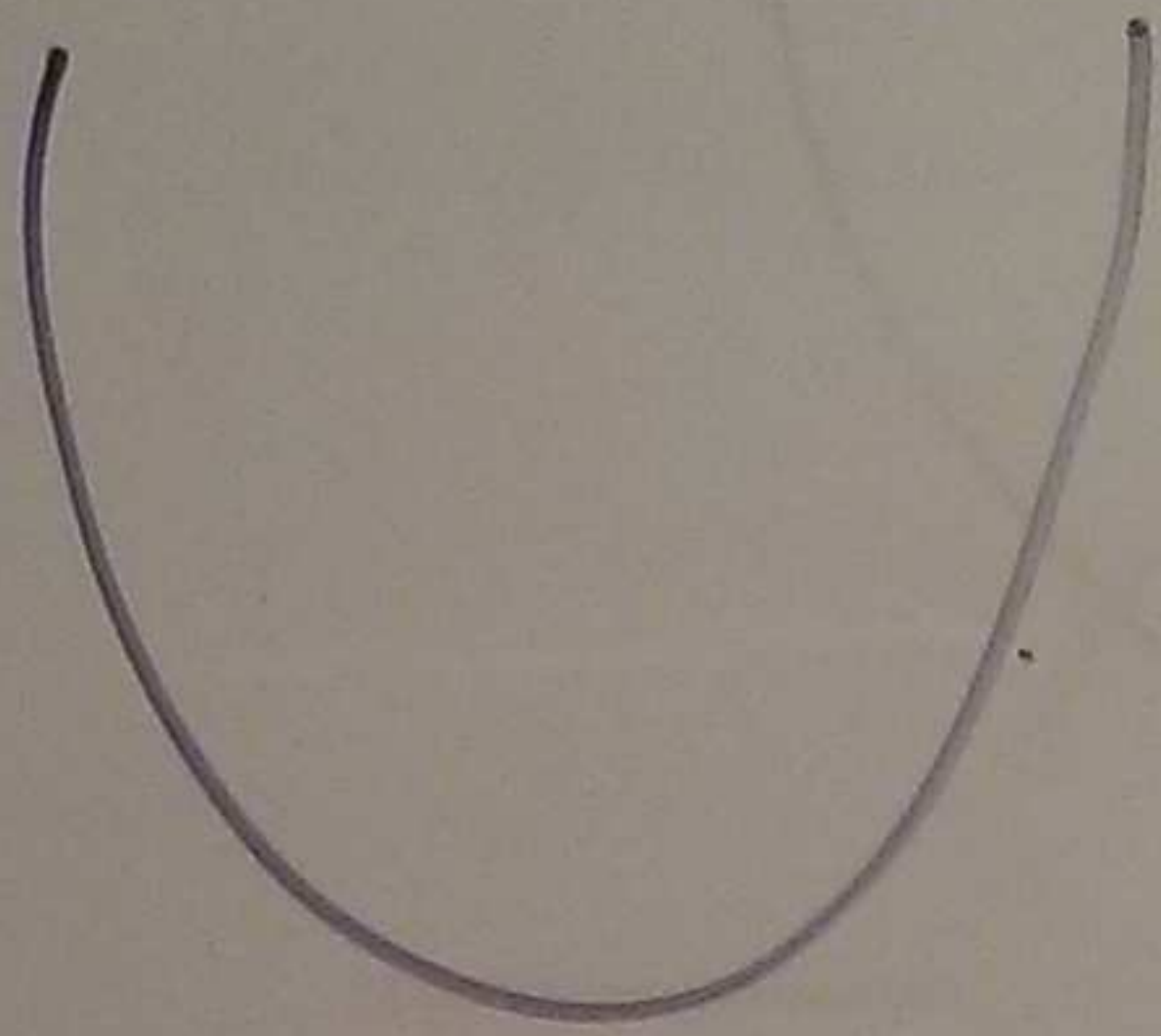
$$\nabla^2 f(x) \succeq 0, \underline{\forall x} \quad (f''(x) \geq 0, \text{ for } f \text{ scalar})$$

"f is curved upward everywhere"



0) - usually the easiest way to show a simple function is convex

a function is convex
ient Methods



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

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

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

$$f(x) = \frac{1}{2} x^T A x + b^T x + c, \quad A \succeq 0$$

$$\nabla f(x) = Ax + b$$

$$\nabla^2 f(x) = A \succeq 0$$

$$f(x) = \frac{1}{2} \|Ax - b\|^2$$

$$I_v = v$$

$$\nabla f(x) = A^T Ax - A^T b$$

$$\nabla^2 f(x) = \underline{A^T A}$$

$$\stackrel{(*)}{=} \sum_{i=1}^N a_i a_i^T$$

$$x^T A x \geq 0, \forall x$$

$$\begin{aligned} & x^T \left(\sum_{i=1}^N a_i a_i^T \right) x \\ &= \sum_{i=1}^N (x^T a_i) (a_i^T x) \geq 0 \end{aligned}$$

Operations that preserve convexity

Let f_1, f_2 be convex functions

① Non-negative weighted sum: $w_1 f_1(x) + w_2 f_2(x)$
 $w_1, w_2 \geq 0$

② Composition w/ Affine function:

$$f_1(Ax + b)$$

③ Pointwise maximum: $\max \{f_1(x), f_2(x)\}$

* ④ Other composition rules.

Show that SVMs are convex:

$$\frac{1}{2} \|w\|^2 + C \sum_{i=1}^N \max\{0, 1 - y_i \bar{w}^T \bar{x}_i\}$$

① ↓

$$w^T w = w^T I w$$

$$\nabla^2 [w^T w] = 2I$$

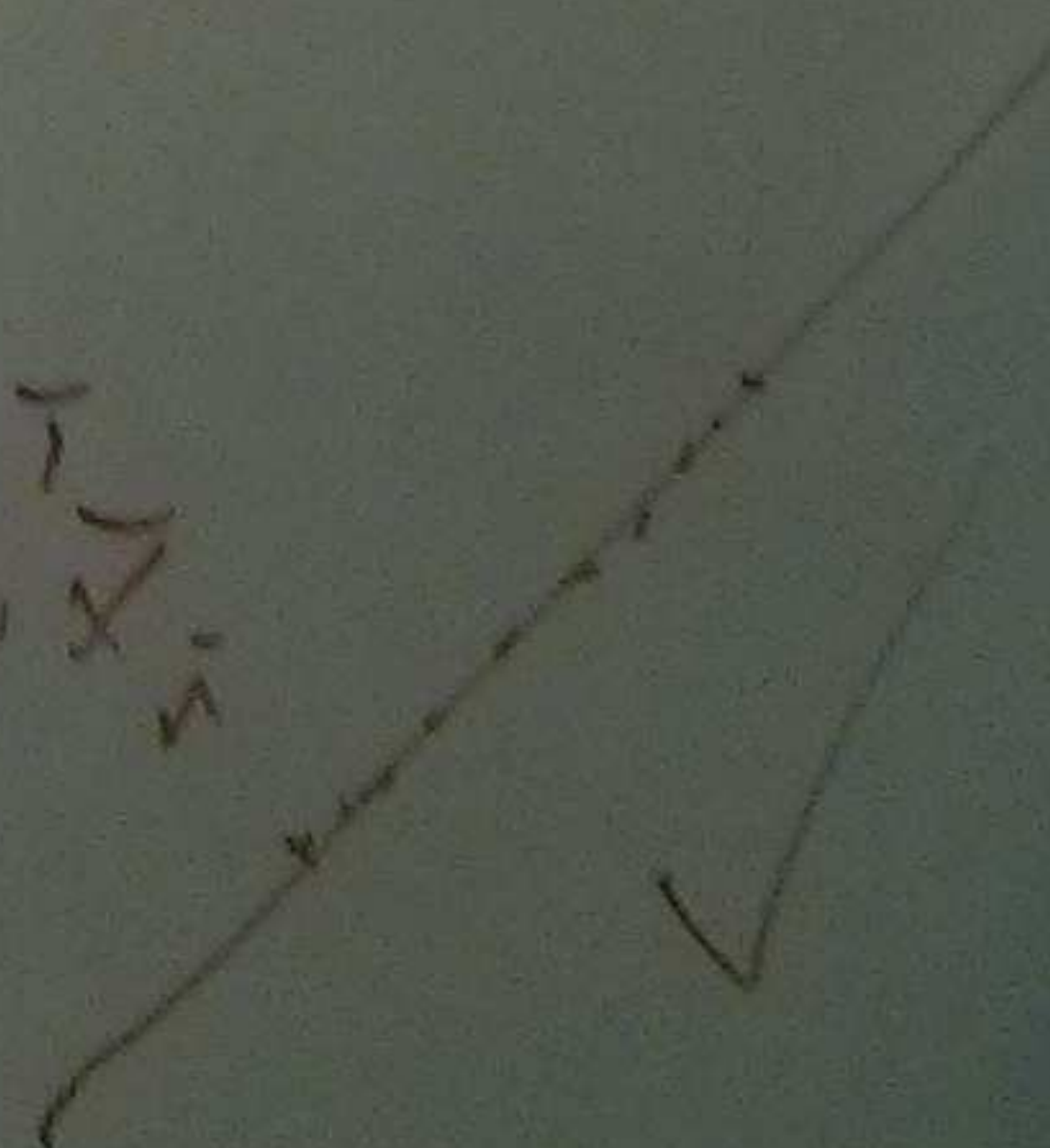
$$w^T A w$$

① $\max\{0, 1 - y_i \bar{w}^T \bar{x}_i\}$

③

0

$$1 - y_i \bar{w}^T \bar{x}_i$$



Gradient Method

- We want to minimize convex f !
(e.g., MLE, MAP, SVMs, etc.)

- Generate a sequence

$$x^0, x^1, x^2, x^3, \dots$$

(optimal value)

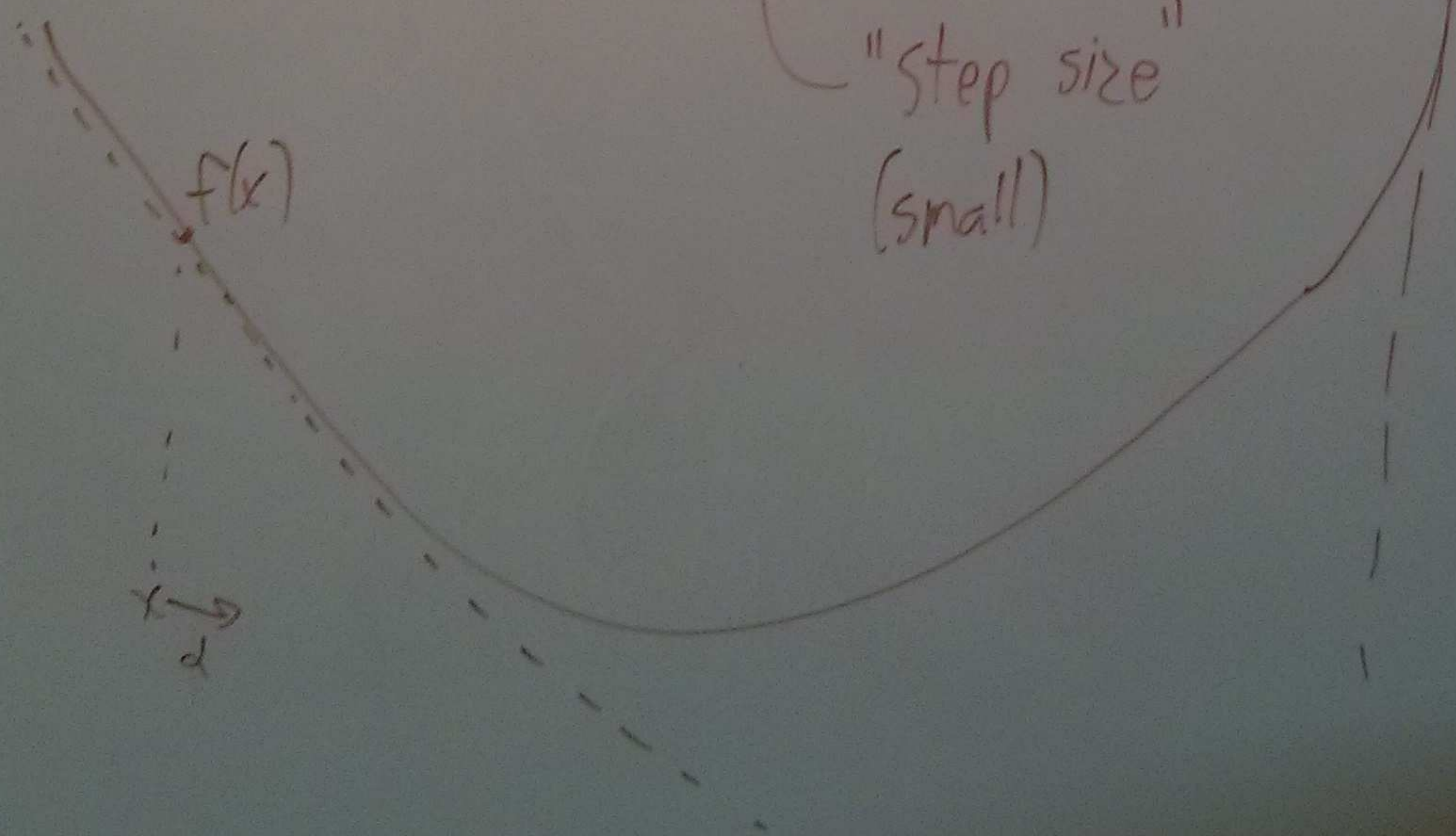
Such that $x^t \rightarrow x^*$ as $t \rightarrow \infty$

- Gradient method:

$$x^{t+1} = x^t - \alpha^t \nabla f(x^t)$$

"step size"
(small)

$$x^* = \underset{x}{\operatorname{argmin}} f(x)$$



Factors of convexity

Function is convex

Methods

Stochastic Gradient Method

$$\operatorname{argmin}_x \frac{1}{N} \sum_{i=1}^N f_i(x) = f(x)$$

$$x^{t+1} = x^t - \alpha^t \nabla f_i(x^t)$$

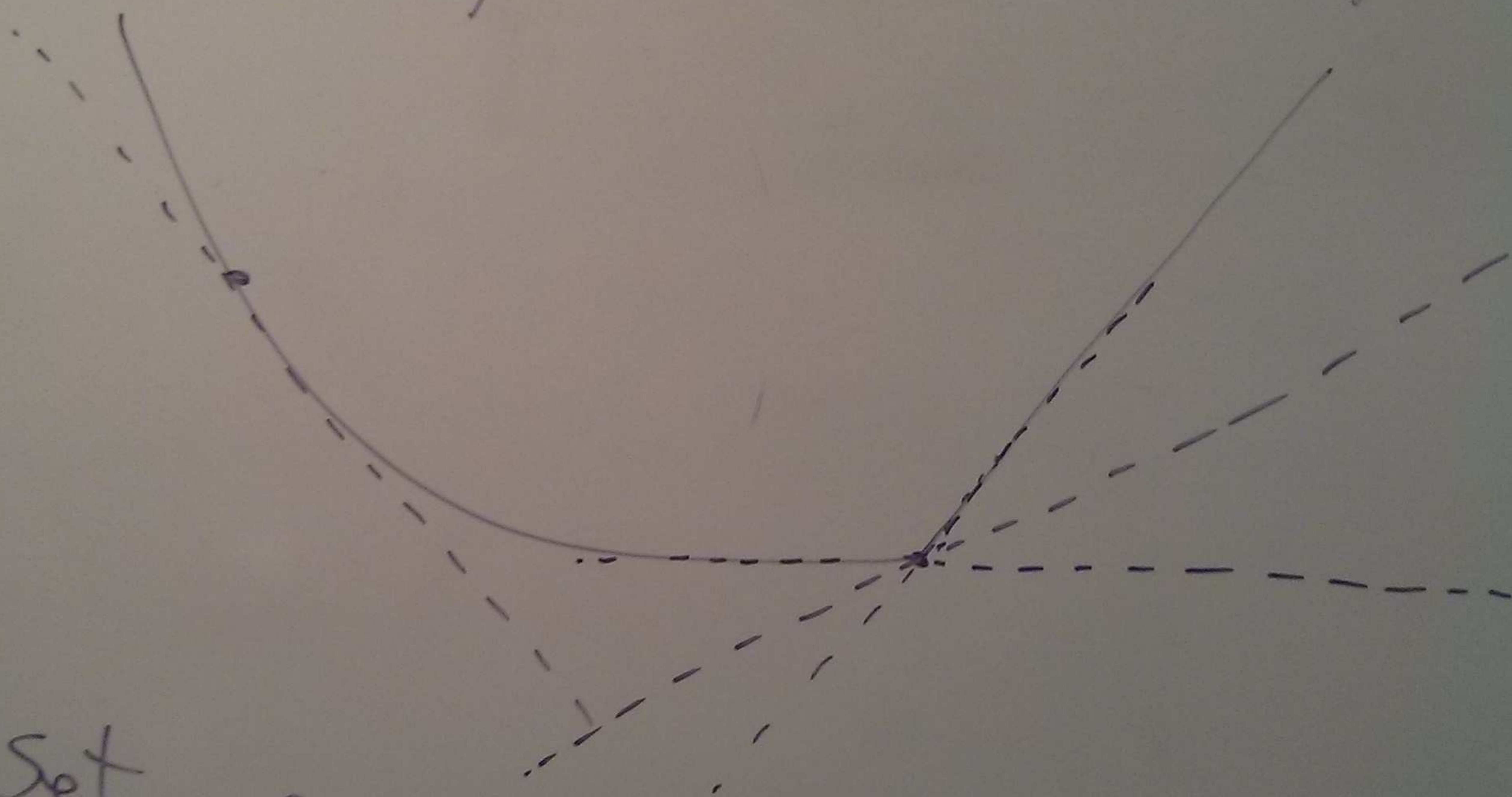
Is this reasonable? $\underbrace{\text{picked randomly}}_{i \in \{1, 2, 3, \dots\}}$

$$\begin{aligned} \mathbb{E}[\nabla f_i(x^t)] &= \sum_{i=1}^N p_i \nabla f_i(x^t) \\ &= \frac{1}{N} \sum_{i=1}^N \nabla f_i(x^t) \end{aligned}$$

Subgradient

- a 'subgradient' of 'f' at 'x',
is a vector 'd' such that

$$f(y) \geq f(x) + d^T (y - x)$$



- Set of subgradients is the "sub-differential"

$$\partial f(x) = \begin{cases} \{1\} & x > 0 \\ \{-1\} & x < 0 \\ [-1, 1] & x = 0 \end{cases}$$

x^* is Global optimum
iff $0 \in \partial f(x^*)$