

Gradient Descent

$$\mathbb{E}_x g(w) = w^2 + e^w$$

$$g'(w) = 2w + e^w$$

$$g''(w) = 2 + e^w \geq 0$$

Convex!

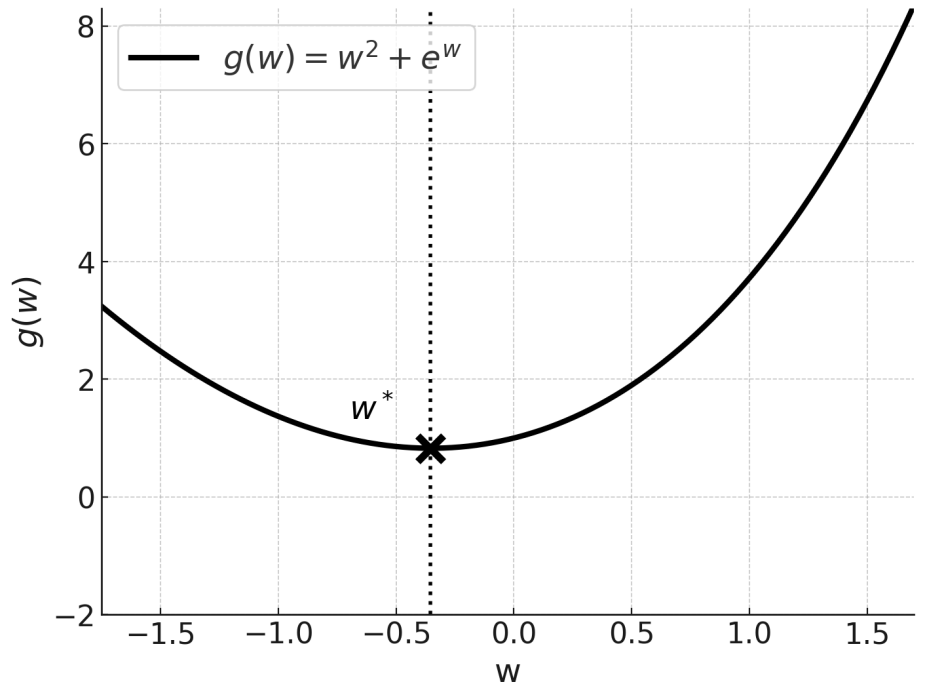
$$g'(w) = 2w + e^w = 0$$

$$e^w = -2w$$

$$w = \log(-2w) = \log(-2) + \log(w) \Rightarrow \log(-2) = w - \log(w)$$

No closed form solution for w^*

calculating $A^{-1} \approx d^3$ operations on computer
which is a lot



Second Order Gradient Descent (Newton's Method)

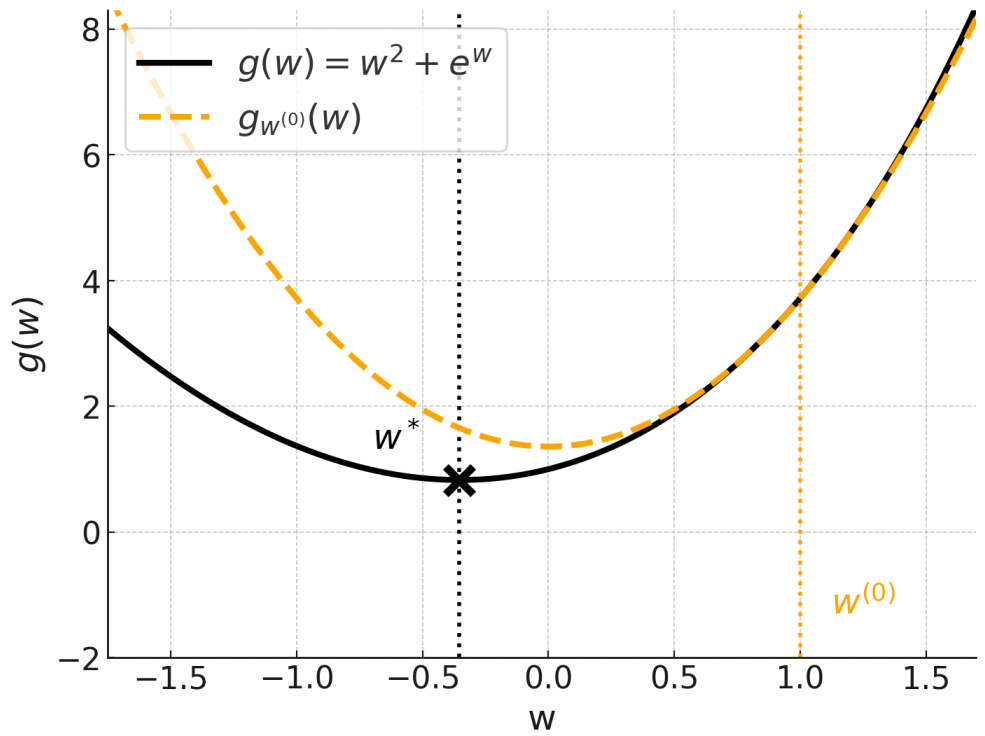
If $g(w)$ is at least a degree 5 polynomial

Then there does not exist a closed form solution for $g(w)=0$

Our approach: approximate $g(w)$ with a second order polynomial using Taylor's approximation

$$g(w) \approx g_{w^{(0)}}(w) = g(w^{(0)}) + g'(w^{(0)})(w - w^{(0)}) + \frac{g''(w^{(0)})}{2} (w - w^{(0)})^2$$

$g_{w^{(0)}}(w)$ where $w^{(0)} = 1$



Goal: $\arg\min_{w \in \mathcal{W}} g_{w^{(0)}}(w) = \arg\min_{w \in \mathcal{W}} g(w^{(0)}) + g'(w^{(0)})(w - w^{(0)}) + \frac{g''(w^{(0)})}{2} (w - w^{(0)})^2$

$$g'_{w^{(0)}}(w) = g'(w^{(0)}) + g''(w^{(0)})(w - w^{(0)}) = 0$$

$$\Rightarrow g''(w^{(0)}) w = g''(w^{(0)}) w^{(0)} - g'(w^{(0)})$$

$$\Rightarrow \underset{w^{(1)}}{w} = w^{(0)} - \frac{g'(w^{(0)})}{g''(w^{(0)})}$$

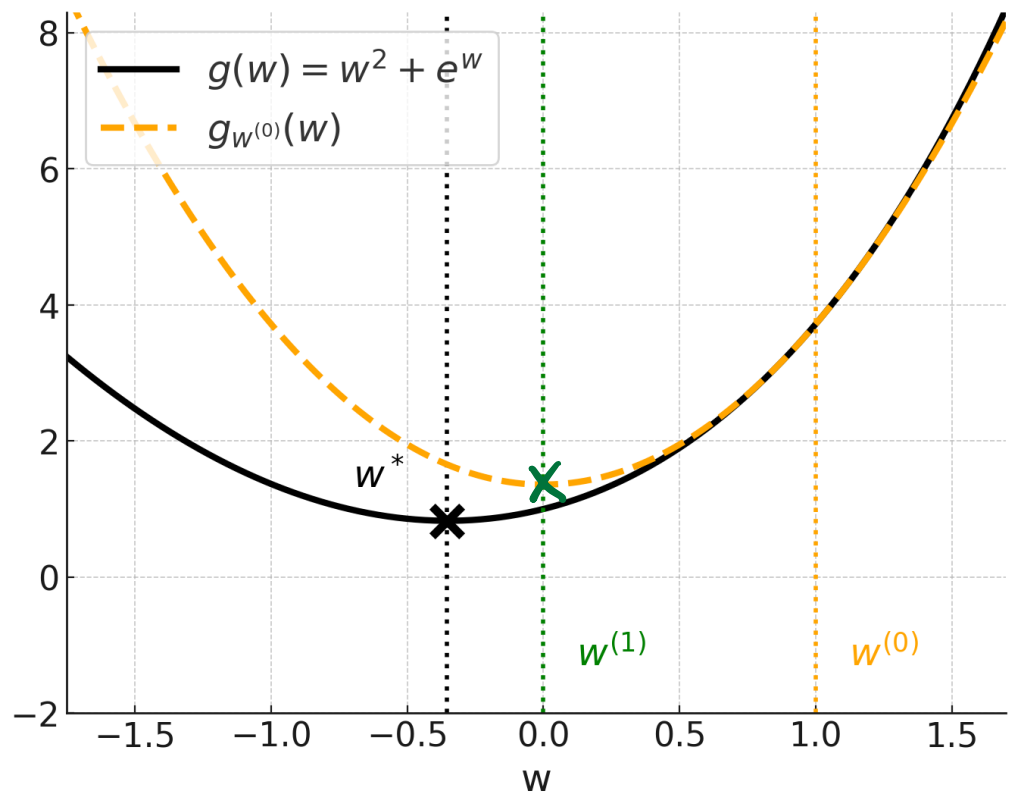
$$w^{(0)} = 1$$

$$g'(w) = 2w + e^w$$

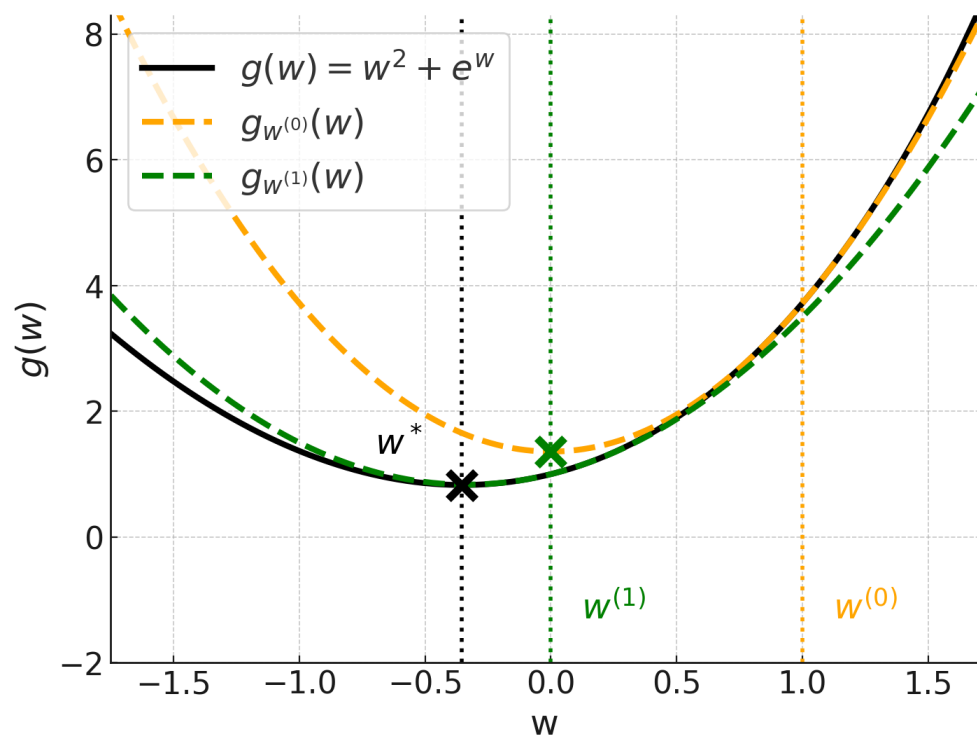
$$g''(w) = 2 + e^w \geq 0$$

$$g'(1) = 2 + e$$

$$g''(1) = 2 + e$$



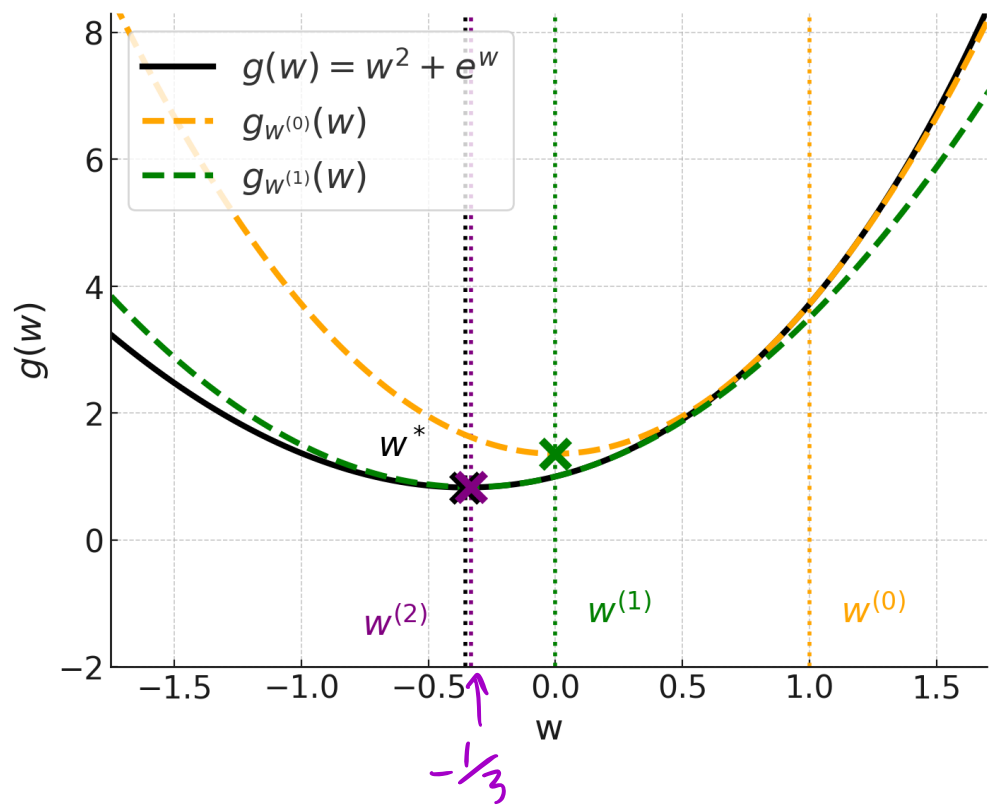
$$w^{(1)} = 1 - \frac{2+e}{2+e} = 0$$



$$w^{(2)} = w^{(1)} - \frac{g'(w^{(1)})}{g''(w^{(1)})}$$

$$= 0 - \frac{1}{3}$$

$$= -\frac{1}{3}$$



In general
$$w^{(t+1)} = w^{(t)} - \frac{g'(w^{(t)})}{g''(w^{(t)})}$$

for all $t \in \mathbb{N}$ and $w^{(0)}$ is your initial guess

$w^{(t+1)}$ approaches w^* as $t \rightarrow \infty$

(First-order) Gradient Descent

Calculating the second derivative (Hessian)
is computationally expensive in higher dimensions

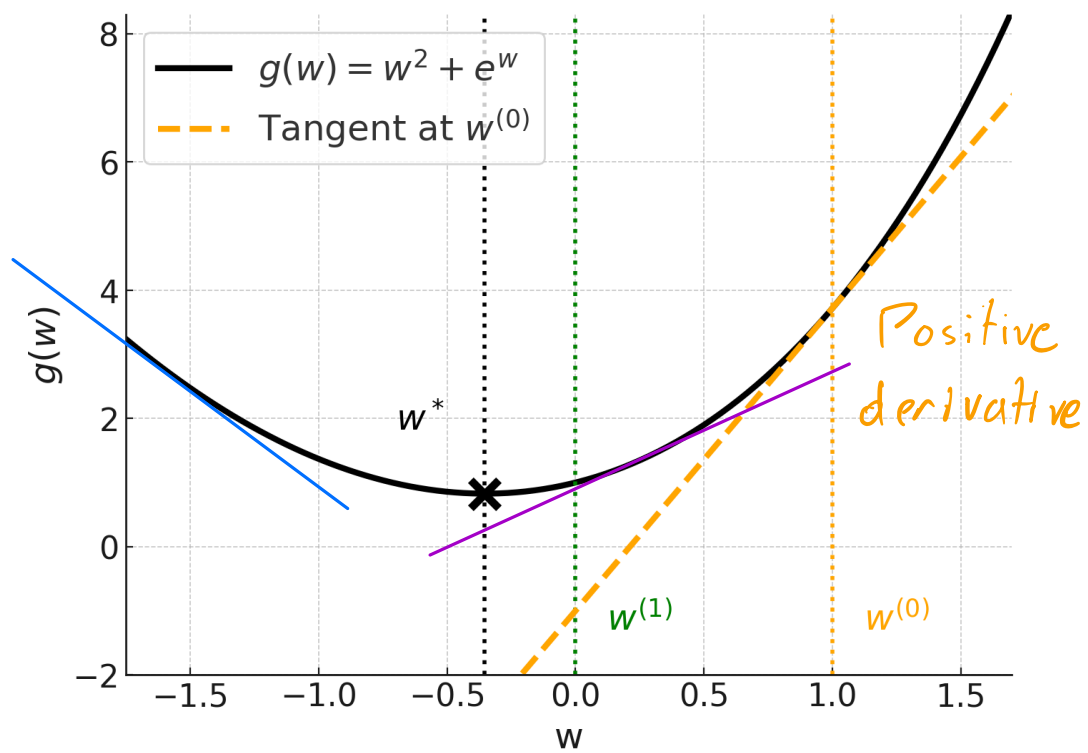
$$w^{(t+1)} = w^{(t)} - \eta^{(t)} g'(w^{(t)})$$

$\eta^{(t)}$ is the
"step-size"
"Learning rate"

if we knew $g''(w^{(t)})$ we should set

$$\eta^{(t)} = \frac{1}{g''(w^{(t)})}$$

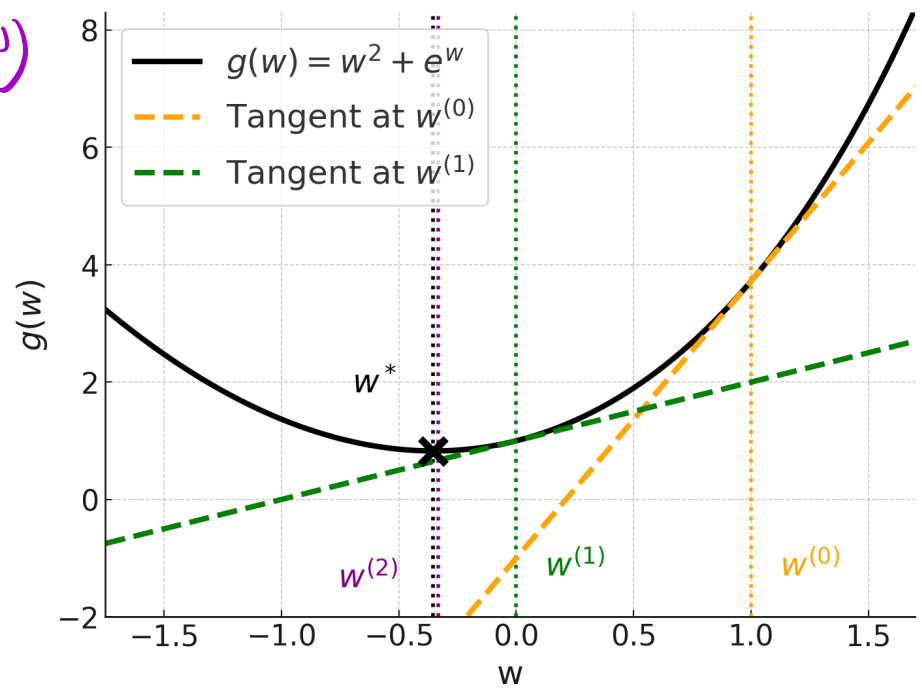
negative
derivative



$$w^{(1)} = w^{(0)} - \frac{1}{g'(w^{(0)})} g'(w^{(0)})$$

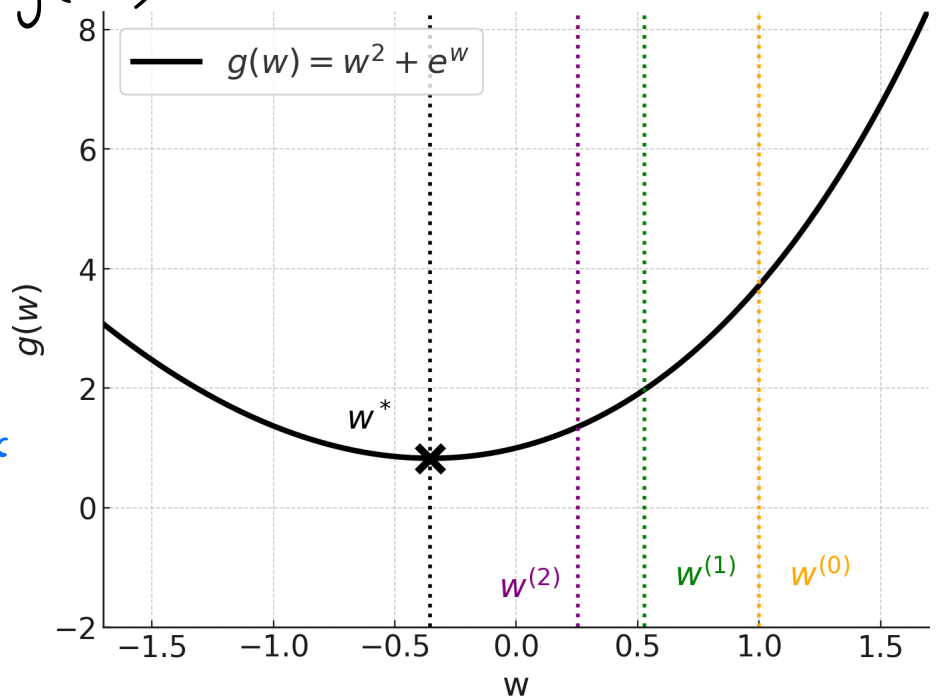
with $w^{(0)} = 1$

$$w^{(2)} = w^{(1)} - \frac{1}{g''(w^{(1)})} g'(w^{(1)})$$



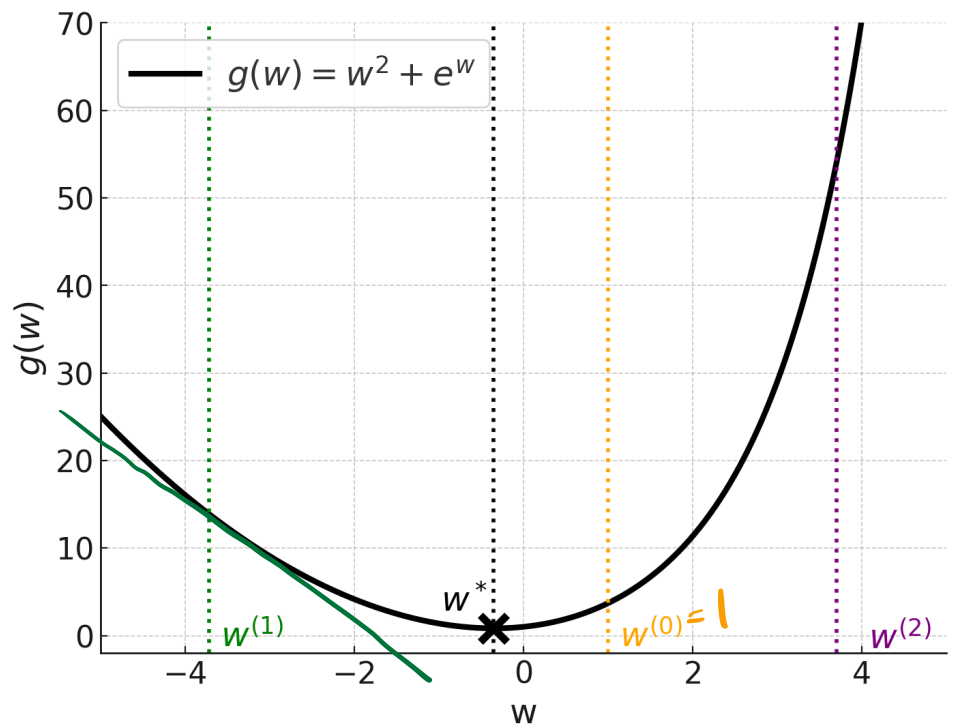
$$\eta^{(t)} = \frac{1}{10} \leq \left(\frac{1}{2+e} \right)^{\frac{1}{g''(w^{(0)})}}$$

Small $\eta^{(t)}$
slowly reaches w^*



$$\eta^{(t)} = 1 > \frac{1}{2+e}$$

if $\eta^{(t)}$ is large
then $w^{(t)}$ can
diverge from w^*
or oscillate



- Small step size is safe (you will get close to w^*); however, it can take long
 - Large step size can get you to w^* fast but sometimes it can diverge
-

Multivariate Gradient Descent

$$\vec{w} \in \mathcal{W} = \mathbb{R}^d, \quad d \geq 1 \qquad g: \mathbb{R}^d \rightarrow \mathbb{R}$$

$$\text{Obj: } \vec{w}^* = (w_1^*, \dots, w_d^*)^T = \underset{\vec{w} \in \mathbb{R}^d}{\operatorname{argmin}} g(\vec{w})$$

gradient descent update rule:

$$\vec{w}^{(t+1)} = \vec{w}^{(t)} - \eta^{(t)} \nabla g(\vec{w}^{(t)}) \quad \vec{\eta} \in \mathbb{R}^d$$

$$\text{where } \nabla g(\vec{w}^{(t)}) = \left(\frac{\partial g}{\partial w_1}, \dots, \frac{\partial g}{\partial w_d} \right)$$

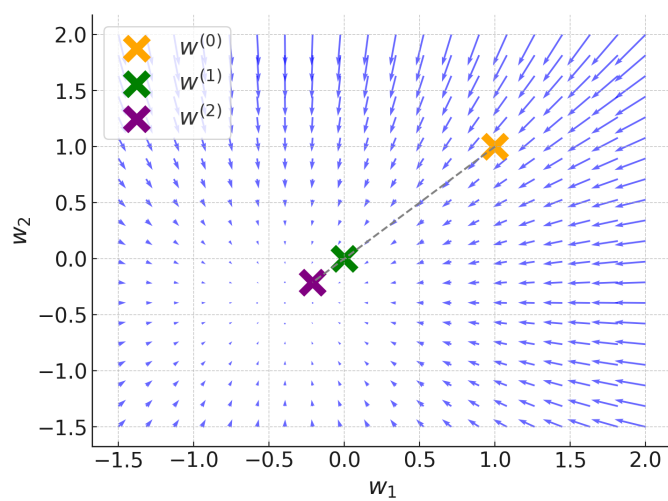
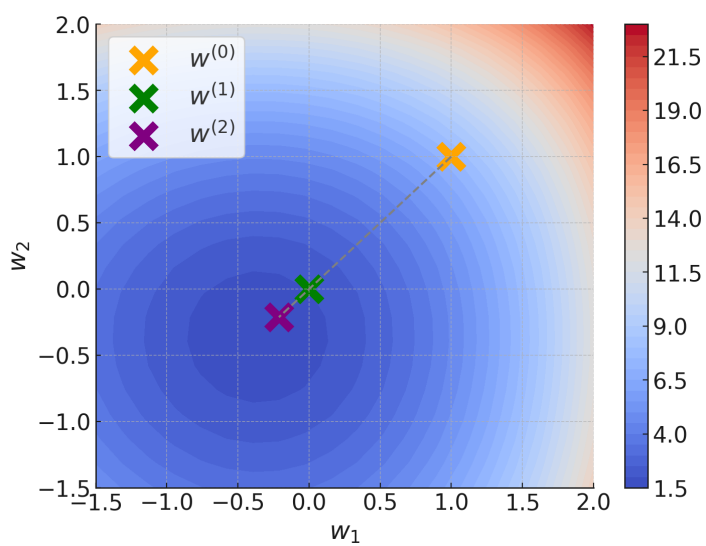
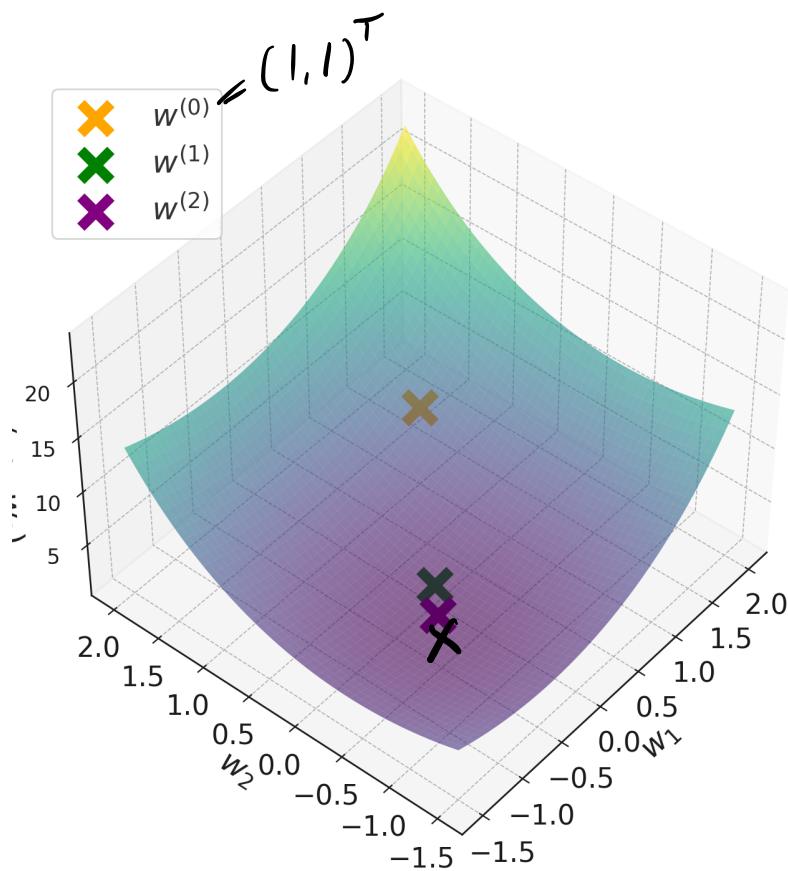
$$\underline{\underline{\text{Ex}}}: g(w_1, w_2) = w_1^2 + e^{w_1} + w_2^2 + e^{w_2}$$

$$\frac{\partial g}{\partial w_1} = 2w_1 + e^{w_1}, \quad \frac{\partial g}{\partial w_2} = 2w_2 + e^{w_2}$$

$$\begin{pmatrix} w_1^{(t+1)} \\ w_2^{(t+1)} \end{pmatrix} = \begin{pmatrix} w_1^{(t)} \\ w_2^{(t)} \end{pmatrix} - \eta^{(t)} \begin{pmatrix} 2w_1^{(t)} + e^{w_1^{(t)}} \\ 2w_2^{(t)} + e^{w_2^{(t)}} \end{pmatrix}$$

$$\vec{w}^{(0)} = (1, 1)^T, \quad \eta^{(t)} = \frac{1}{2+e}$$

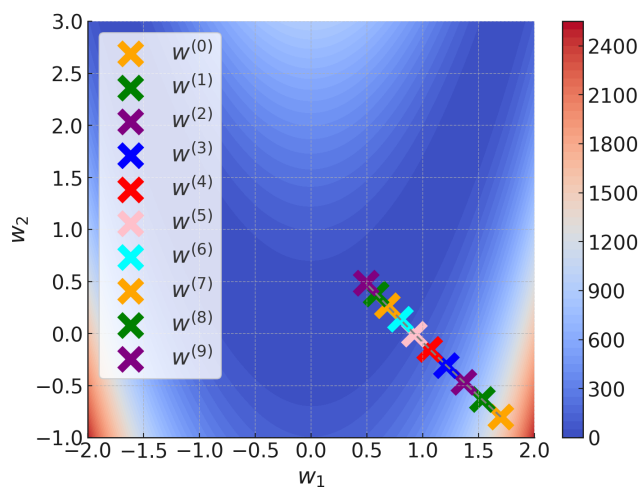
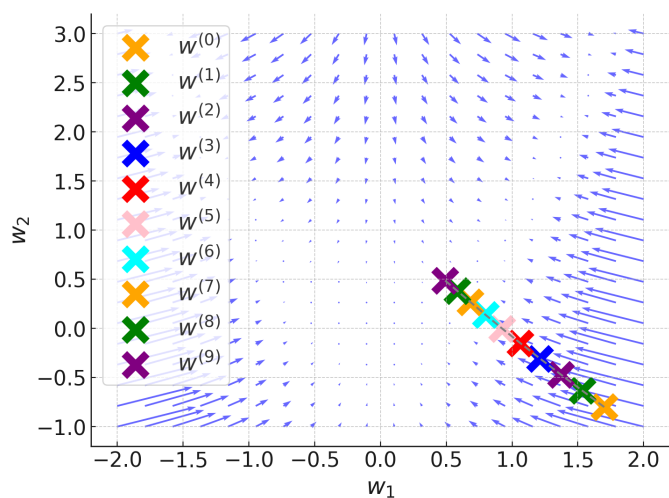
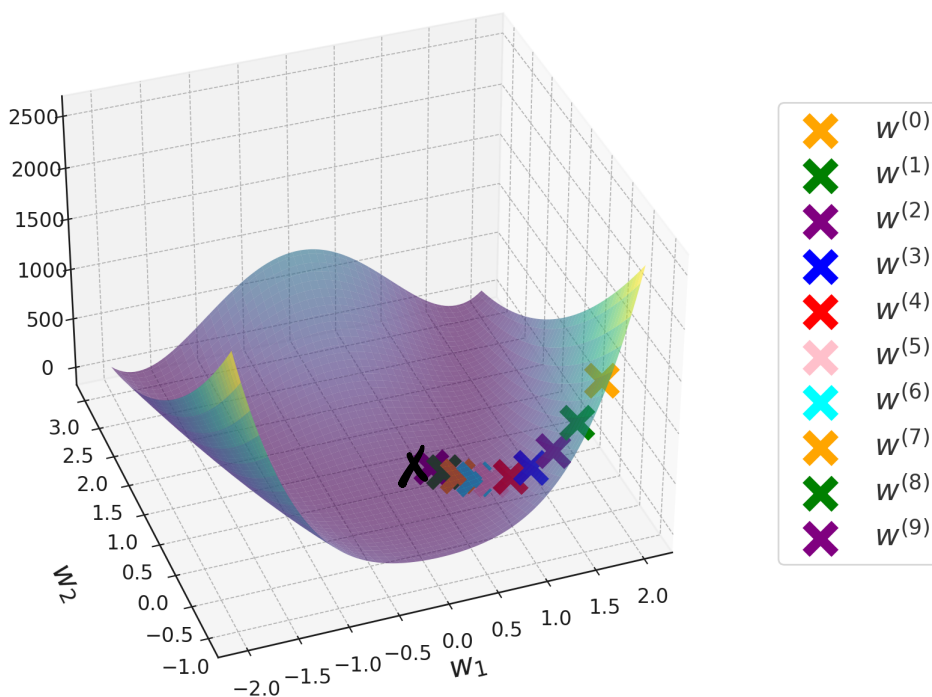
$$\vec{w}^{(1)} = \begin{pmatrix} 1 \\ 1 \end{pmatrix} - \frac{1}{2+e} \begin{pmatrix} 2+e \\ 2+e \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$



$$\underline{\underline{E_x:}} \quad g(w_1, w_2) = (1 - w_1)^2 + 100(w_2 - w_1^2)^2$$

" $g(\vec{w})$

$$\eta^{(t)} = \frac{1}{6}$$



Selecting Step Sizes

1. Constant: $\eta^{(t)} = \eta \in (0, \infty)$

2. Inverse decaying: $\eta^{(t)} = \frac{\eta}{1 + \lambda t} \propto \frac{1}{t}$ $\eta, \lambda \in (0, \infty)$

3. Exponential decaying: $\eta^{(t)} = \eta \frac{1}{e^{\lambda t}}$ $\eta, \lambda \in (0, \infty)$

4. Normalized gradient: $\eta^{(t)} = \frac{\eta}{\epsilon + \underbrace{\|\nabla g(\vec{w}^{(t)})\|_2}_{\text{size of your gradient}}}$

why not $\eta^{(t)} = \lambda t$

Because it increases with time
which you do not want

Gradient Descent with $\hat{L}(f)$

$$\mathcal{D} = ((\vec{x}_1, y_1), \dots, (\vec{x}_n, y_n)), \quad \mathcal{X} = \mathbb{R}^{d+1}, \quad \mathcal{Y} = \mathbb{R}$$

$$\mathcal{F} = \{f \mid f: \mathbb{R}^{d+1} \rightarrow \mathbb{R} \text{ s.t. } f(\vec{x}) = \vec{x}^T \vec{w}, \vec{w} \in \mathbb{R}^{d+1}\}$$

$\parallel$
 $(w_0, w_1, \dots, w_d)^T$

$$\ell(\hat{y}, y) = (\hat{y} - y)^2$$

$$\text{Obj: } \hat{\vec{w}} = \underset{\vec{w} \in \mathbb{R}^{d+1}}{\operatorname{argmin}} \quad \hat{L}(\vec{w})$$

$$= \underset{\vec{w} \in \mathbb{R}^{d+1}}{\operatorname{argmin}} \quad \underbrace{\frac{1}{n} \sum_{i=1}^n (\vec{x}_i^T \vec{w} - y_i)^2}_{g(\vec{w})}$$

(Batch) Gradient Descent (BGD)

$$\vec{w}^{(t+1)} = \vec{w}^{(t)} - \eta^{(t)} \nabla g(\vec{w}^{(t)})$$

$$\nabla \hat{L}(\vec{w}^{(t)}) = \left(\frac{2\hat{L}}{2w_0}, \dots, \frac{2\hat{L}}{2w_d} \right)^T \in \mathbb{R}^{d+1}$$

$$= \left(\frac{2}{n} \sum_{i=1}^n (\vec{x}_i^T \vec{w} - y_i) x_{i0}, \dots, \frac{2}{n} \sum_{i=1}^n (\vec{x}_i^T \vec{w} - y_i) x_{id} \right)^T$$

$$= \frac{2}{n} \sum_{i=1}^n (\vec{x}_i^T \vec{w} - y_i) (x_{i0}, \dots, x_{id})^T$$

$$= \frac{2}{n} \sum_{i=1}^n (\vec{x}_i^T \vec{w} - y_i) \vec{x}_i$$

Algorithm: BGD Linear Regression Learner
(with a constant size)

input: $D = ((\vec{x}_1, y_1), \dots, (\vec{x}_n, y_n))$, η , $T \leftarrow$ number of "epochs"

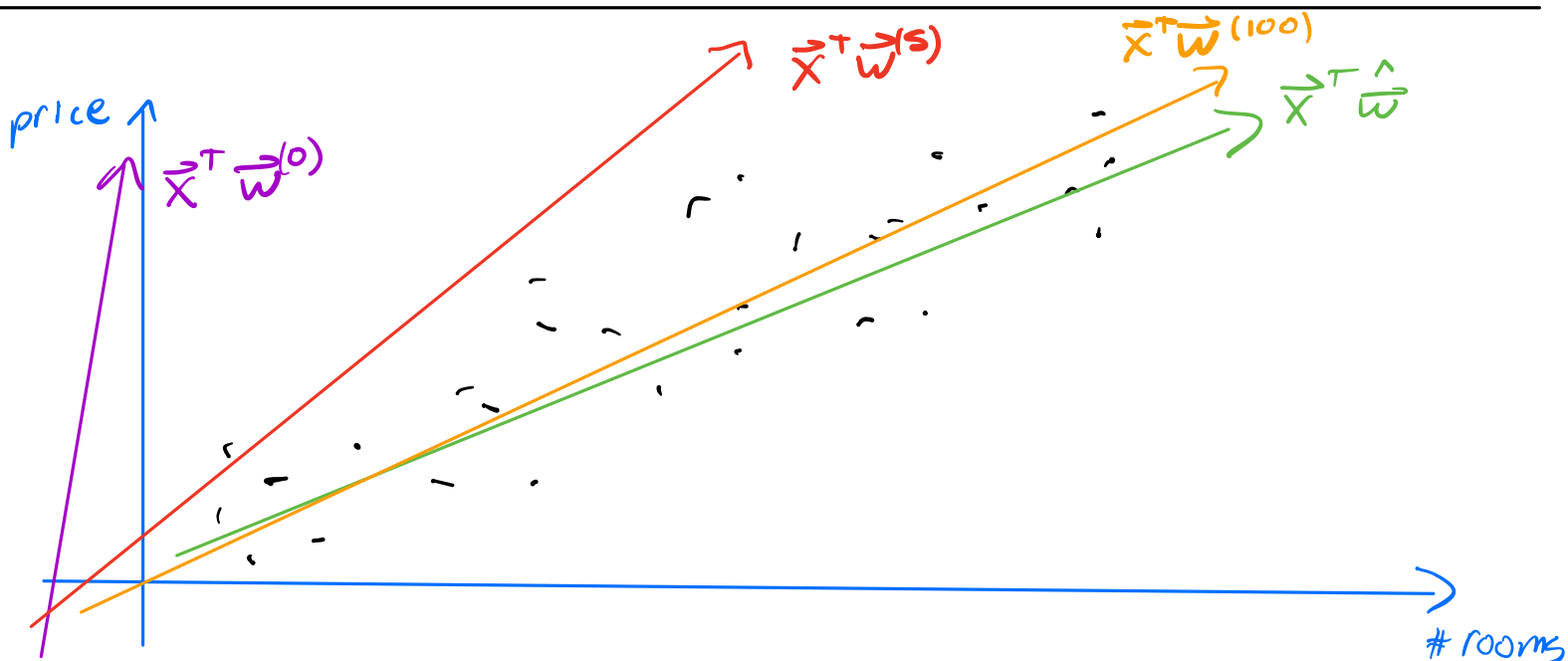
$\vec{w}$ = random vector in $\mathbb{R}^{d+1}$

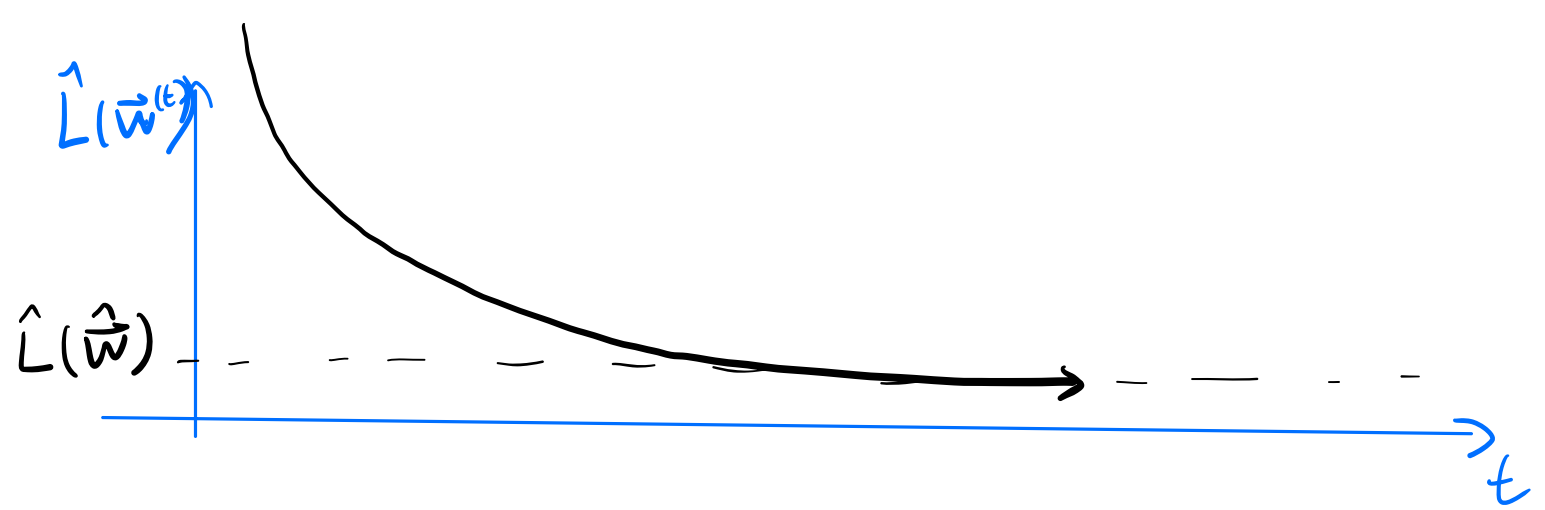
for $t = 1, \dots, T$

$$\nabla \hat{L}(\vec{w}) = \frac{2}{n} \sum_{i=1}^n (\vec{x}_i^T \vec{w} - y_i) \vec{x}_i$$

$$\vec{w} = \vec{w} - \eta \nabla \hat{L}(\vec{w})$$

return $\hat{f}: \mathbb{R}^{d+1} \rightarrow \mathbb{R}$ where $\hat{f}(\vec{x}) = \vec{x}^T \vec{w}$





Computation of Closed form soln. vs. BGD

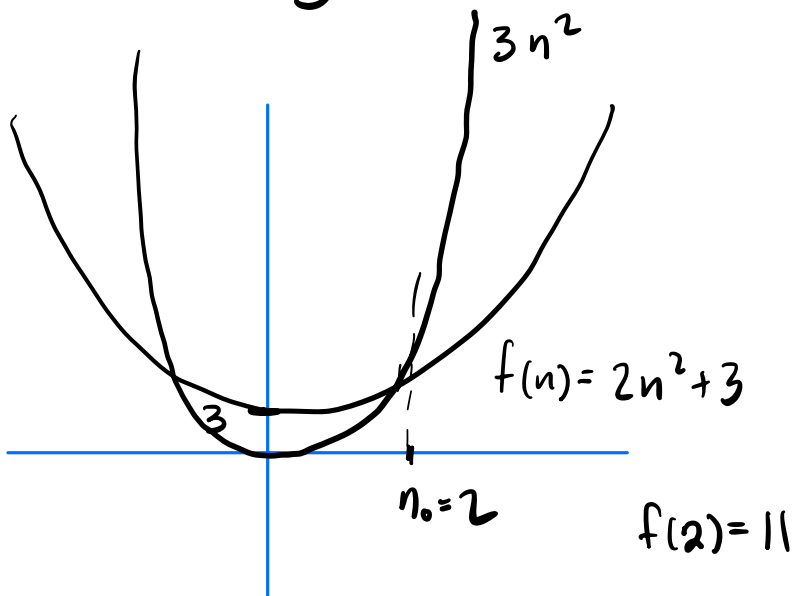
Number of computational steps $\approx$ number of elementary operations (add, subtracted, multiplied, etc.)

Big "O" notation

$$f(n) \in O(g(n))$$

$$f(n) \leq C g(n) \text{ for } \forall n \geq n_0$$

$$C \in (0, \infty)$$



Ex: $2n^2 + 3 = O(n^2)$

if $n_0 = 2$ and $C = 3$

then $2(2)^2 + 3 \leq 3(2)^2$

$11 \leq 12$

Closed form: $\hat{\vec{w}} = A^{-1} \vec{b}$ where $A = \sum_{i=1}^n \vec{x}_i \vec{x}_i^T$
 $\vec{b} = \sum_{i=1}^n \vec{x}_i y_i$

Computing $A: O(nd^2)$

$$\vec{b}: O(nd)$$

$$A^{-1}: O(d^3)$$

$$A^{-1} \vec{b}: O(d^2)$$

$$\begin{aligned} \text{Total compute: } O(nd^2 + nd + d^3 + d^2) \\ = O(nd^2 + d^3) \end{aligned}$$

BGD:

computing: $\nabla \hat{L}(\vec{w}): O(nd)$

$$\vec{w} \text{ update}: O(d)$$

$$\text{Total compute: } O(Tnd)$$

Comparing $O(Tnd)$ vs. $O(nd^2 + d^3)$

$$\text{if } T < d: O(Tnd) < O(nd^2) \leq O(nd^2 + d^3)$$

Mini-Batch Gradient Descent

$$D = \left((\vec{x}_1, y_1), \dots, (\vec{x}_b, y_b), \right. \\ (\vec{x}_{b+1}, y_{b+1}), \dots, (\vec{x}_{2b}, y_{2b}), \\ \vdots \\ \left. (\vec{x}_{(M-1)b+1}, y_{(M-1)b+1}), \dots, (\vec{x}_{Mb}, y_{Mb}) \right)$$

$b \in \mathbb{N}$: Mini-batch size

$$Mb \leq n$$

$$M = \left\lfloor \frac{n}{b} \right\rfloor : \# \text{ of mini batches}$$

$$\underline{\underline{Ex}}: n=6, b=2, M = \frac{6}{2} = 3$$

$$D = \left(\begin{array}{l} (\vec{x}_1, y_1), (\vec{x}_2, y_2), \\ (\vec{x}_3, y_3), (\vec{x}_4, y_4), \\ (\vec{x}_5, y_5), (\vec{x}_6, y_6) \end{array} \right) \begin{array}{l} \text{MB 1} \\ \text{MB 2} \\ \text{MB 3} \end{array}$$

$$\hat{L}_m(\vec{w}) = \frac{1}{b} \sum_{i=(m-1)b+1}^{mb} (\vec{x}_i^T \vec{w} - y_i)^2 \quad m \in \{1, \dots, M\}$$

if $b \leq n$ then $\hat{L}_m(\vec{w})$ is a worse estimate than $\hat{L}(\vec{w})$

Algorithm: MBGD Linear Regression Learner
(with a constant size)

input: $D = ((\vec{x}_1, y_1), \dots, (\vec{x}_n, y_n))$, η , T , b

$\vec{w}$ = random vector in $\mathbb{R}^{d+1}$

$M = \text{floor}(\frac{n}{b})$

for $t = 1, \dots, T$

Randomly shuffle D

for $m = 1, \dots, M$

$$\nabla \hat{L}(\vec{w}) = \frac{2}{b} \sum_{i=(m-1)b+1}^{mb} (\vec{x}_i^T \vec{w} - y_i) \vec{x}_i$$

$$\vec{w} = \vec{w} - \eta \nabla \hat{L}(\vec{w})$$

return $\hat{f}: \mathbb{R}^{d+1} \rightarrow \mathbb{R}$ where $\hat{f}(\vec{x}) = \vec{x}^T \vec{w}$

if $b = n \Rightarrow M = 1 \Rightarrow \text{BGD}$

if $b = 1 \Rightarrow M = n \Rightarrow \text{SGD}$ "stochastic gradient descent"

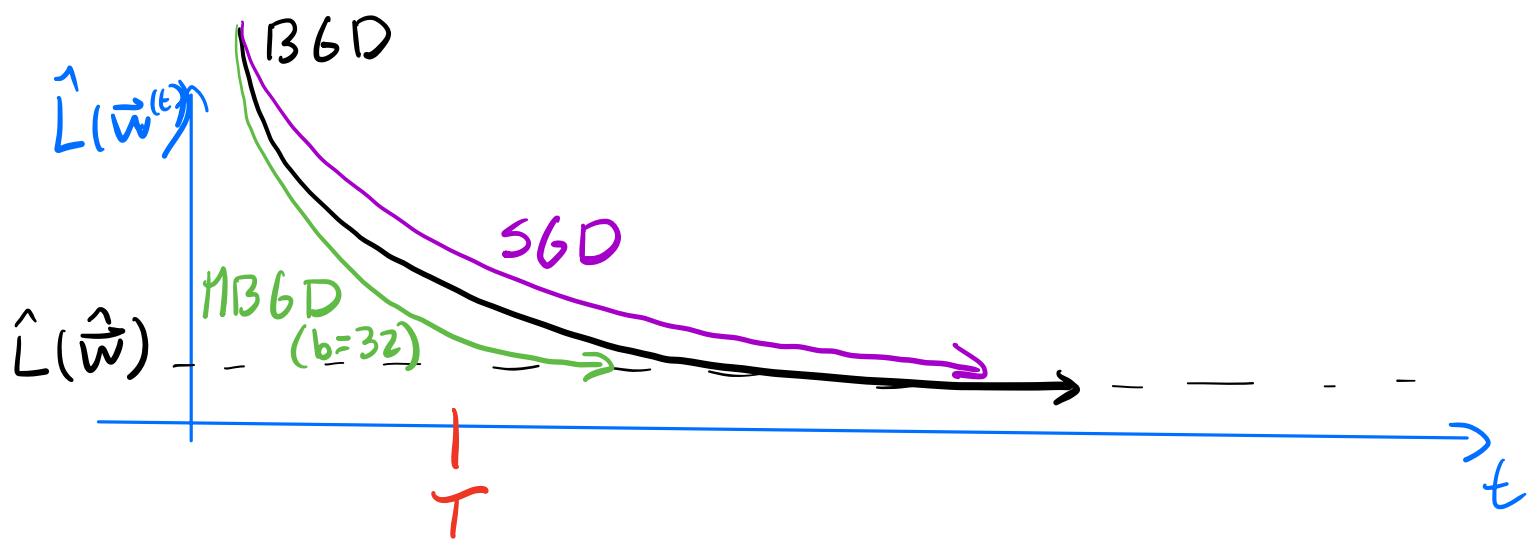
Computation

$$\text{MBGD} : O(T M b d)$$

$$\text{BGD} : O(T n d)$$

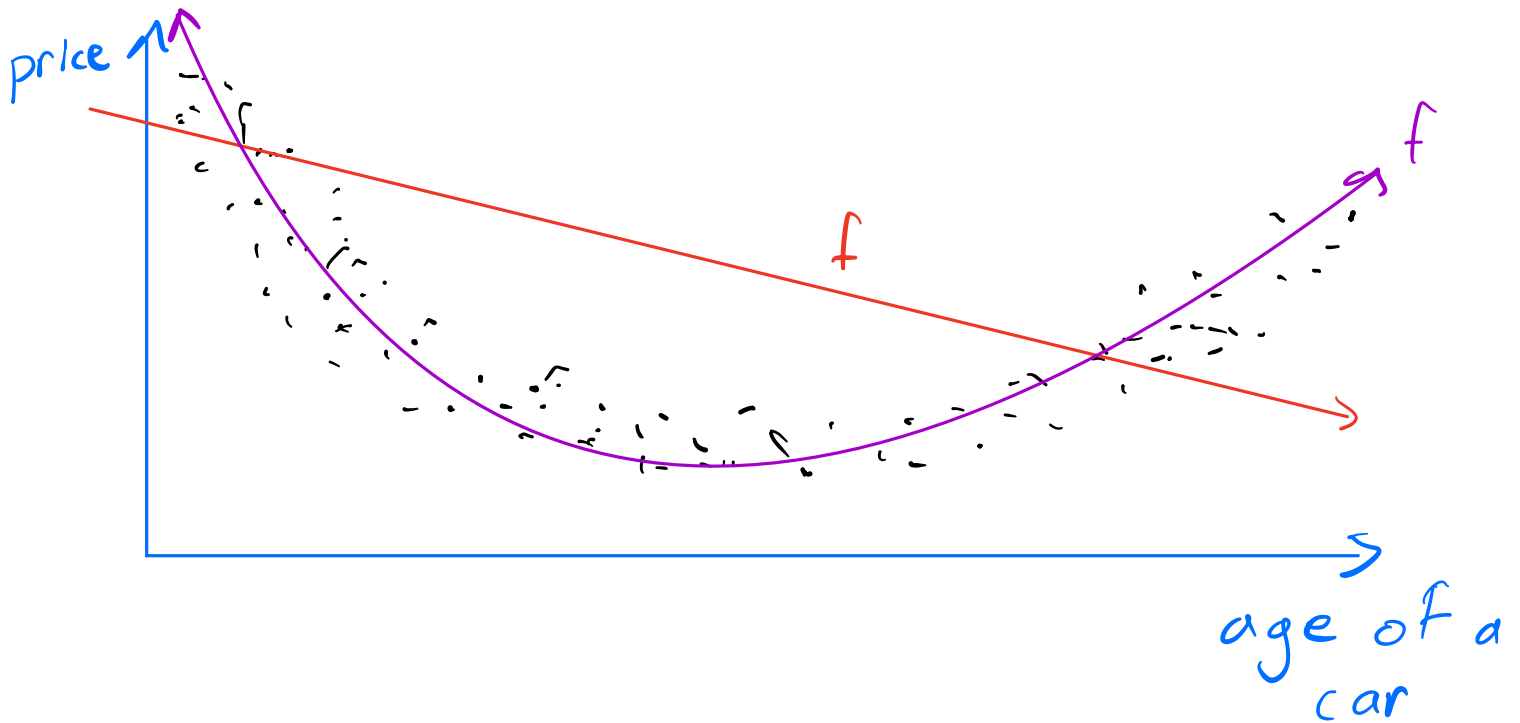
In practice for the same T

MBGD usually finds a better $\vec{w}$ (case $b < n$)



Polynomial Regression

$$\hat{L}(f) < \hat{L}(f)$$



$$\underline{\underline{E_x}}: d=1, \mathcal{X} = \mathbb{R}^{d+1} = \mathbb{R}^2, \mathcal{Y} = \mathbb{R}$$

linear function: $\vec{x}^T \vec{w} = (1, x_1) (w_0, w_1)^T = w_0 + w_1 x_1$

$$\widetilde{\mathcal{F}}_1 = \{f \mid f: \mathbb{R}^2 \rightarrow \mathbb{R} \text{ s.t. } f(\vec{x}) = \underbrace{\vec{x}^T \vec{w}}_{\phi_1(\vec{x})}, \vec{w} \in \mathbb{R}^2\}$$

degree 2 polynomial: $w_0 + w_1 x_1 + w_2 x_1^2$

$$\begin{aligned} &= (1, x_1, x_1^2) (w_0, w_1, w_2)^T \\ &\xrightarrow{\text{feature map}} \phi_2(\vec{x})^T (w_0, w_1, w_2)^T \end{aligned}$$

$$\widetilde{\mathcal{F}}_2 = \{f \mid f: \mathbb{R}^2 \rightarrow \mathbb{R} \text{ s.t. } f(\vec{x}) = \phi_2(\vec{x})^T \vec{w}, \vec{w} \in \mathbb{R}^3\}$$

degree 3 polynomial: $w_0 + w_1 x_1 + w_2 x_1^2 + w_3 x_1^3$

$$= (1, x_1, x_1^2, x_1^3) (w_0, w_1, w_2, w_3)^T$$

$$= \phi_3(\vec{x})^T (w_0, w_1, w_2, w_3)^T$$

$$\widetilde{F}_3 = \{f \mid f: \mathbb{R}^2 \rightarrow \mathbb{R} \text{ s.t. } f(\vec{x}) = \phi_3(\vec{x})^T \vec{w}, \vec{w} \in \mathbb{R}^4\}$$

Ex: $d=2$, $\mathcal{X} = \mathbb{R}^{d+1} = \mathbb{R}^3$, $\mathcal{Y} = \mathbb{R}$

linear function: $\vec{x}^T \vec{w} = (1, x_1, x_2) (w_0, w_1, w_2)^T$

$$= w_0 + w_1 x_1 + w_2 x_2$$

$$\widetilde{F}_1 = \{f \mid f: \mathbb{R}^3 \rightarrow \mathbb{R} \text{ s.t. } f(\vec{x}) = \vec{x}^T \vec{w}, \vec{w} \in \mathbb{R}^3\}$$

degree 2 polynomial:

$$w_0 + w_1 x_1 + w_2 x_2 + w_3 x_1^2 + w_4 x_2^2 + w_5 x_1 x_2$$

$$= (1, x_1, x_2, x_1^2, x_2^2, x_1 x_2) (w_0, w_1, w_2, w_3, w_4, w_5)^T$$

$$= \phi_2(\vec{x})^T (w_0, w_1, w_2, w_3, w_4, w_5)^T$$

$$\widetilde{F}_2 = \{f \mid f: \mathbb{R}^3 \rightarrow \mathbb{R} \text{ s.t. } f(\vec{x}) = \phi_2(\vec{x})^T \vec{w}, \vec{w} \in \mathbb{R}^6\}$$

$$\widetilde{F}_1 \subsetneq \widetilde{F}_2 \subsetneq \widetilde{F}_3 \dots$$

In general $\phi_p: \mathcal{X} \rightarrow \mathcal{Z}$ $p \in \mathbb{N}$

if $\mathcal{X} = \mathbb{R}^{d+1}$, $\mathcal{Z} = \mathbb{R}^{\bar{p}}$ where $\bar{p} = \binom{d+p}{p}$

Ex: $d=2$, $p=2 \Rightarrow \bar{p} = \binom{2+2}{2} = \frac{4 \cdot 3}{2} = 6$

$$\mathcal{F}_p = \{f \mid f: \mathbb{R}^{d+1} \rightarrow \mathbb{R} \text{ s.t. } f(\vec{x}) = \phi_p(\vec{x})^T \vec{w}, \vec{w} \in \mathbb{R}^{\bar{p}}\}$$

$$A_p(D) = \hat{f}_p = \arg \min_{f \in \mathcal{F}_p} \hat{L}(\hat{f}_p)$$

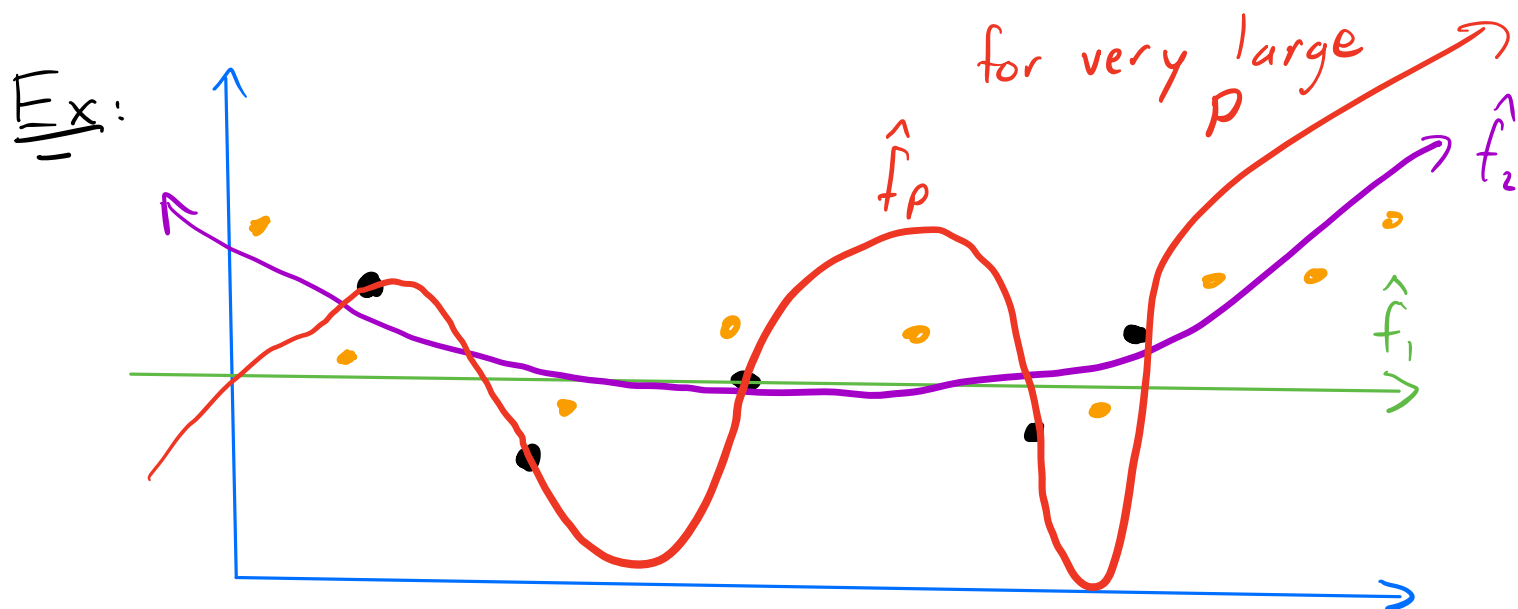
$$\min_{f \in \mathcal{F}_p} \hat{L}(\hat{f}_p) = \min_{\vec{w} \in \mathbb{R}^{\bar{p}}} \frac{1}{n} \sum_{i=1}^n \ell(\phi_p(\vec{x}_i)^T \vec{w}, y_i)$$

$$\Rightarrow \hat{f}_p(\vec{x}) = \phi_p(\vec{x})^T \hat{\vec{w}} \quad \text{where}$$

$$\hat{\vec{w}} = \arg \min_{\vec{w} \in \mathbb{R}^{\bar{p}}} \frac{1}{n} \sum_{i=1}^n \ell(\phi_p(\vec{x}_i)^T \vec{w}, y_i)$$

Solve for $\hat{\vec{w}}$ using closed form Soln
or BGD, MBGD, SGD

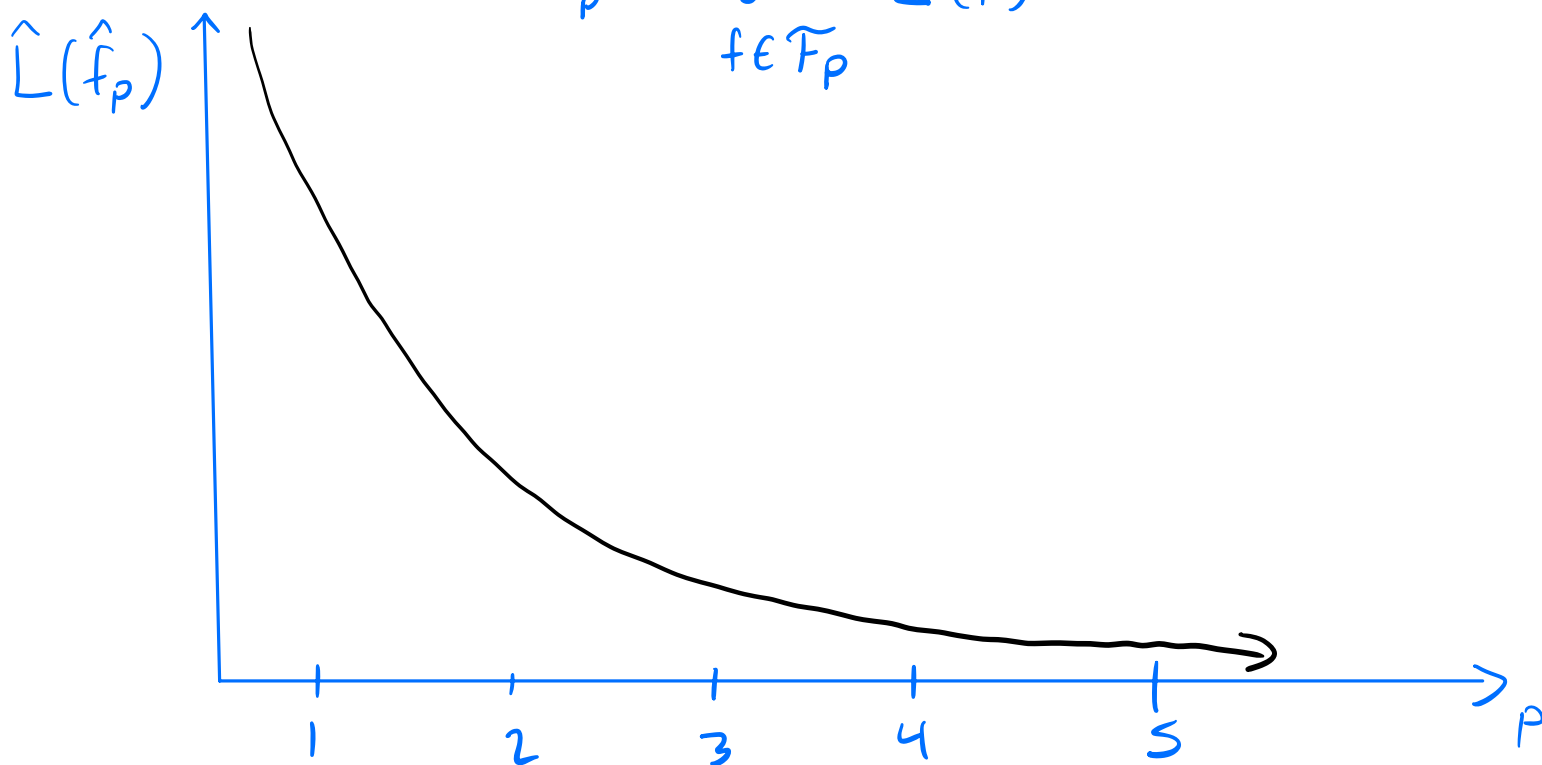
as p increase computation increases



$$\hat{L}(\hat{f}_1) \geq \hat{L}(\hat{f}_2) \geq \dots \geq \hat{L}(\hat{f}_p) \approx 0$$

$$\mathcal{F}_1 \subsetneq \mathcal{F}_2 \subsetneq \mathcal{F}_3 \subsetneq \dots$$

$$\hat{f}_p = \arg\min_{f \in \mathcal{F}_p} \hat{L}(f)$$



$X_1 \in \{1, \dots, 3\} \in \mathbb{R}$ # of rooms

$X_2 \in [0, 30] \in \mathbb{R}$ distance to downtown

$$200 X_1 + 300 X_1 \left(\frac{30 - X_2}{30} \right)$$

if $X_1 = 1, X_2 = 30$: 200

$X_2 = 15$: 350

$X_2 = 0$: 500

$X_1 = 5, X_2 = 30$: 1000

$X_2 = 15$: 1750

$X_2 = 0$: 2500