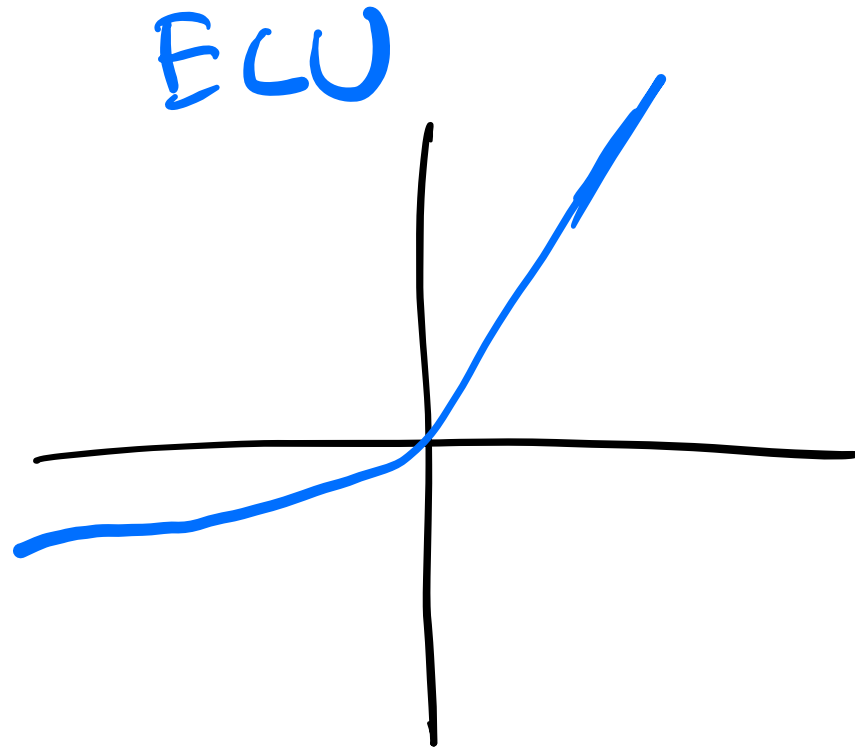
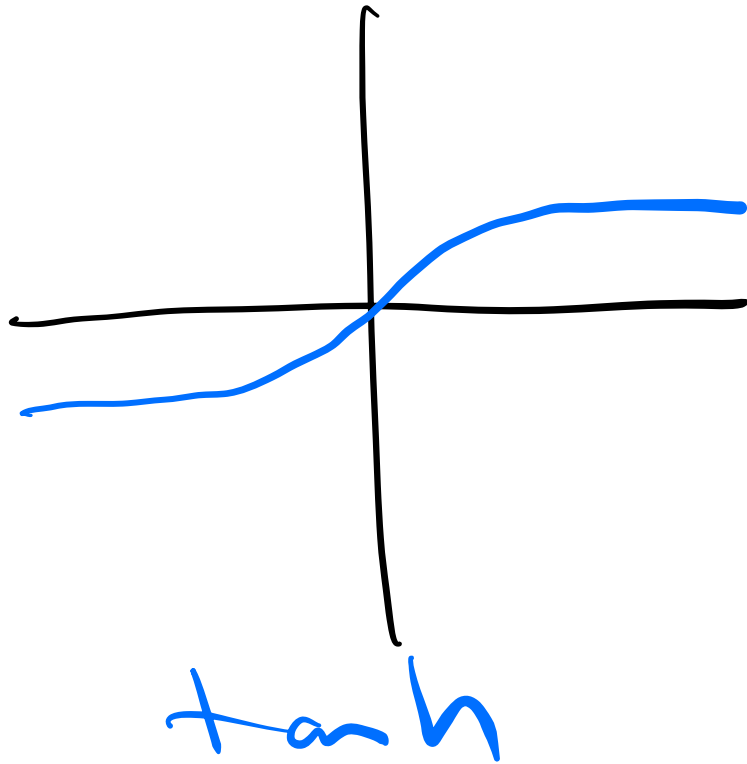


Activation functions



Neural Net Model

$$\text{linear layer} \rightarrow l = Wx + b$$

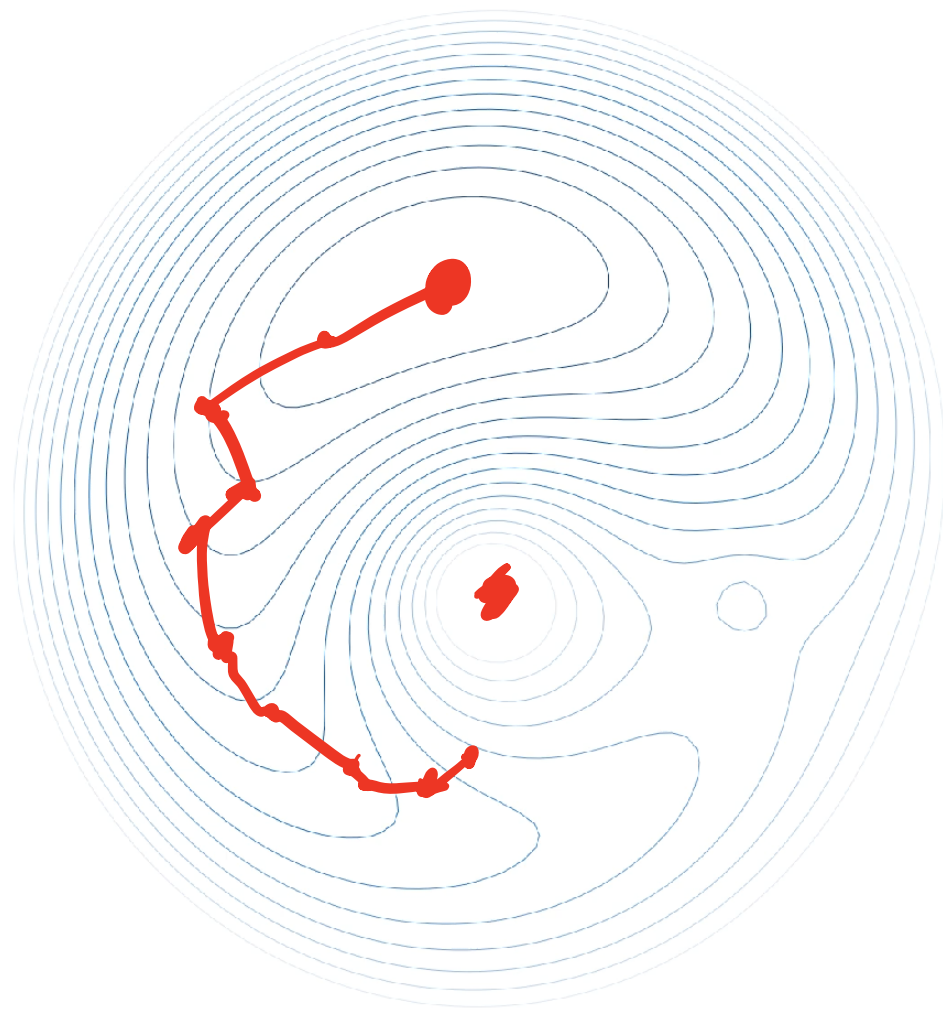
$a \rightarrow \text{activation}$.

$$\hat{y}(x; \phi) = (l^{(n)} \cdots a^{(2)} \circ l^{(2)} \circ a^{(1)} \circ l^{(1)})(x)$$

(unconstrained) Optimization

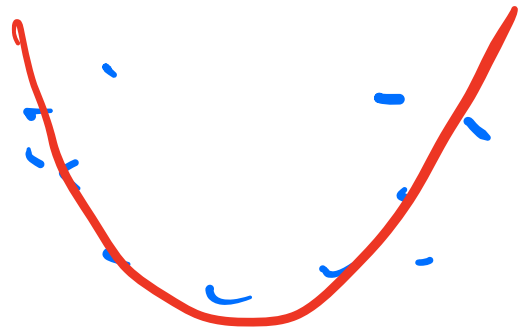
minimize $L(\theta)$

with respect to θ



Train Model: Objective or Loss Function

$$L(\theta) = \frac{1}{n} \sum_{i=1}^n \lambda(\hat{y}_i, y_i)$$



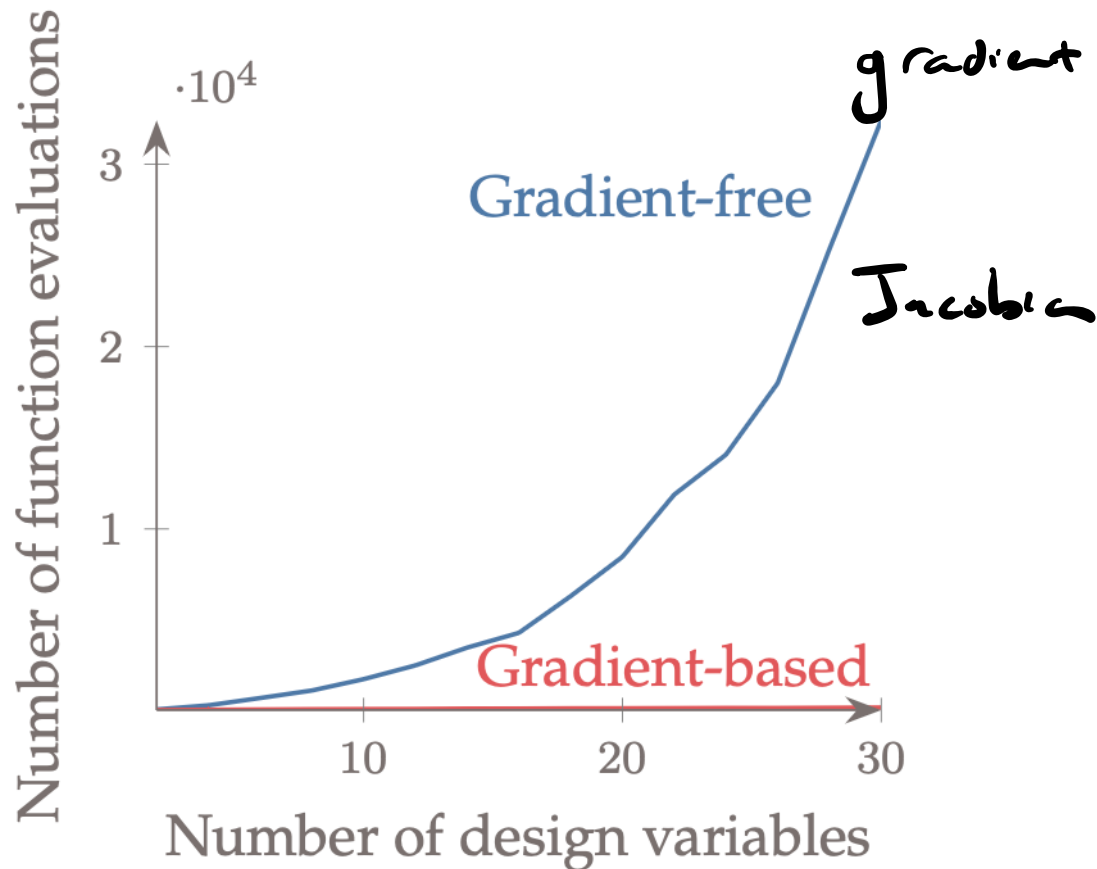
Mean square error:



$$\Rightarrow L = \frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i)^2$$

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

Gradient-based optimization



derivative.

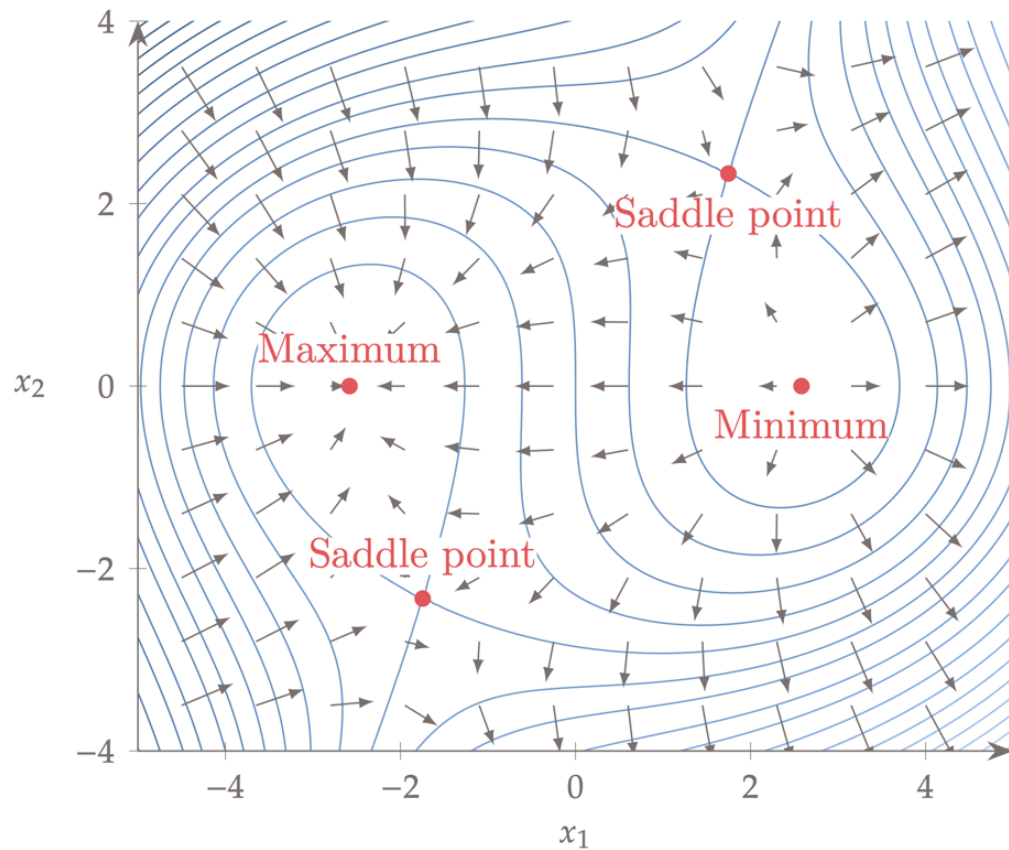
$$\frac{df}{dx}$$

$$\frac{df}{dx_i}$$

$$\frac{df_i}{dx_j}$$

Backpropagation

Gradient Descent



Gradient-Based Optimization

$$\mathcal{D}_{k+1} = \mathcal{D}_k + \alpha P_k$$

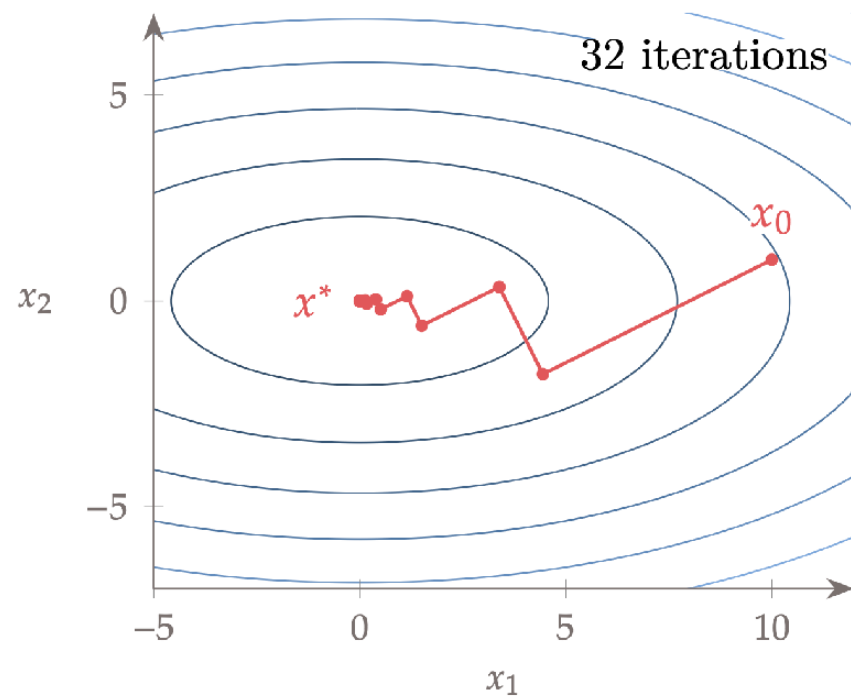
↙ direction.



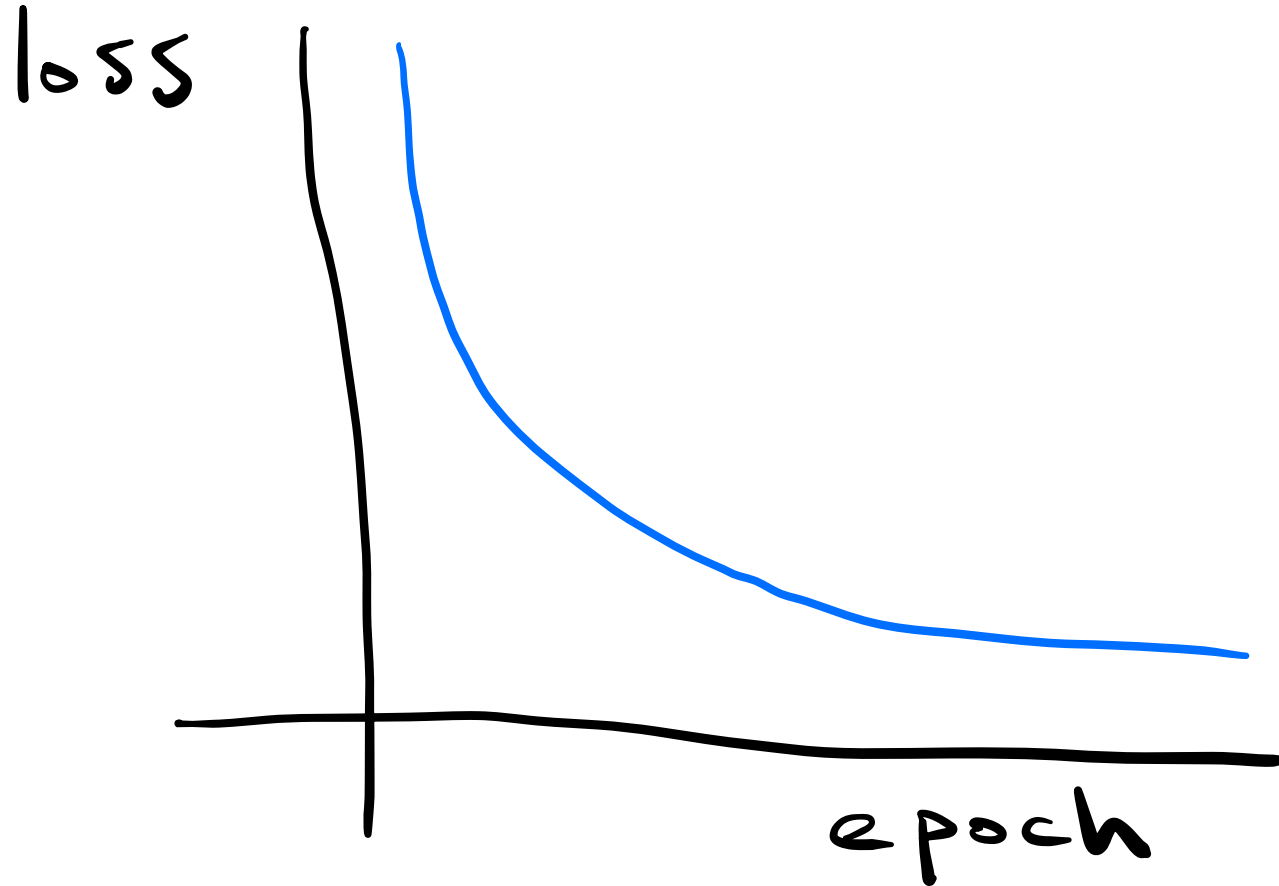
learning rate

$$\begin{bmatrix} \cdot \\ \cdot \\ \cdot \end{bmatrix} = \begin{bmatrix} \cdot \\ \cdot \\ \cdot \end{bmatrix} + \alpha \begin{bmatrix} \cdot \\ \cdot \\ \cdot \end{bmatrix}$$

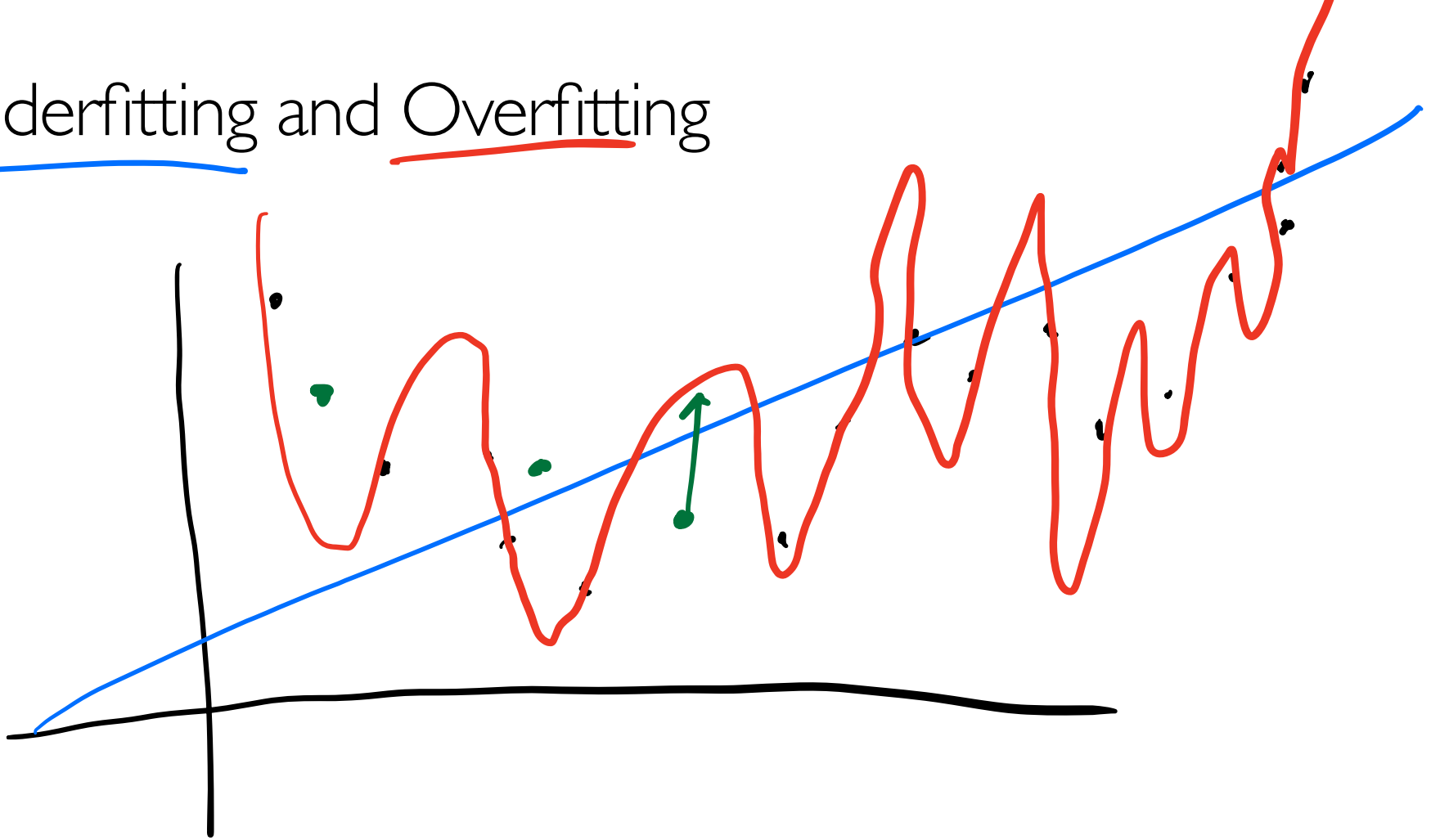
Gradient Descent (or Steepest Descent)



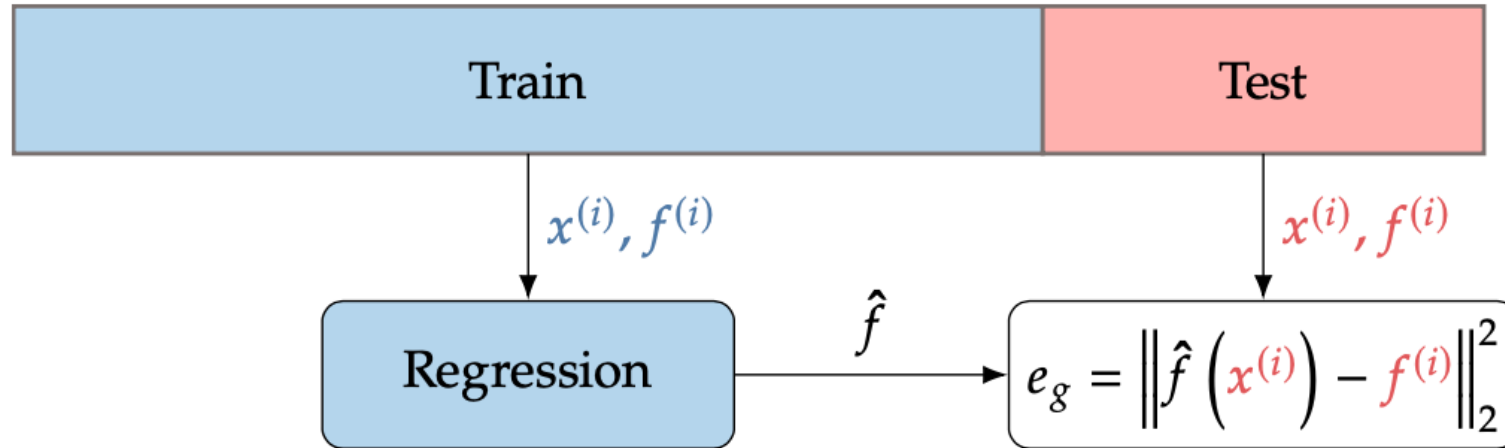
Convergence



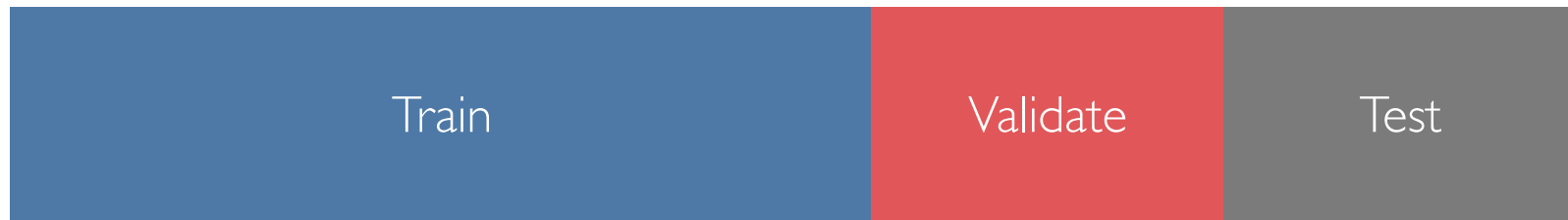
Underfitting and Overfitting



Cross Validation



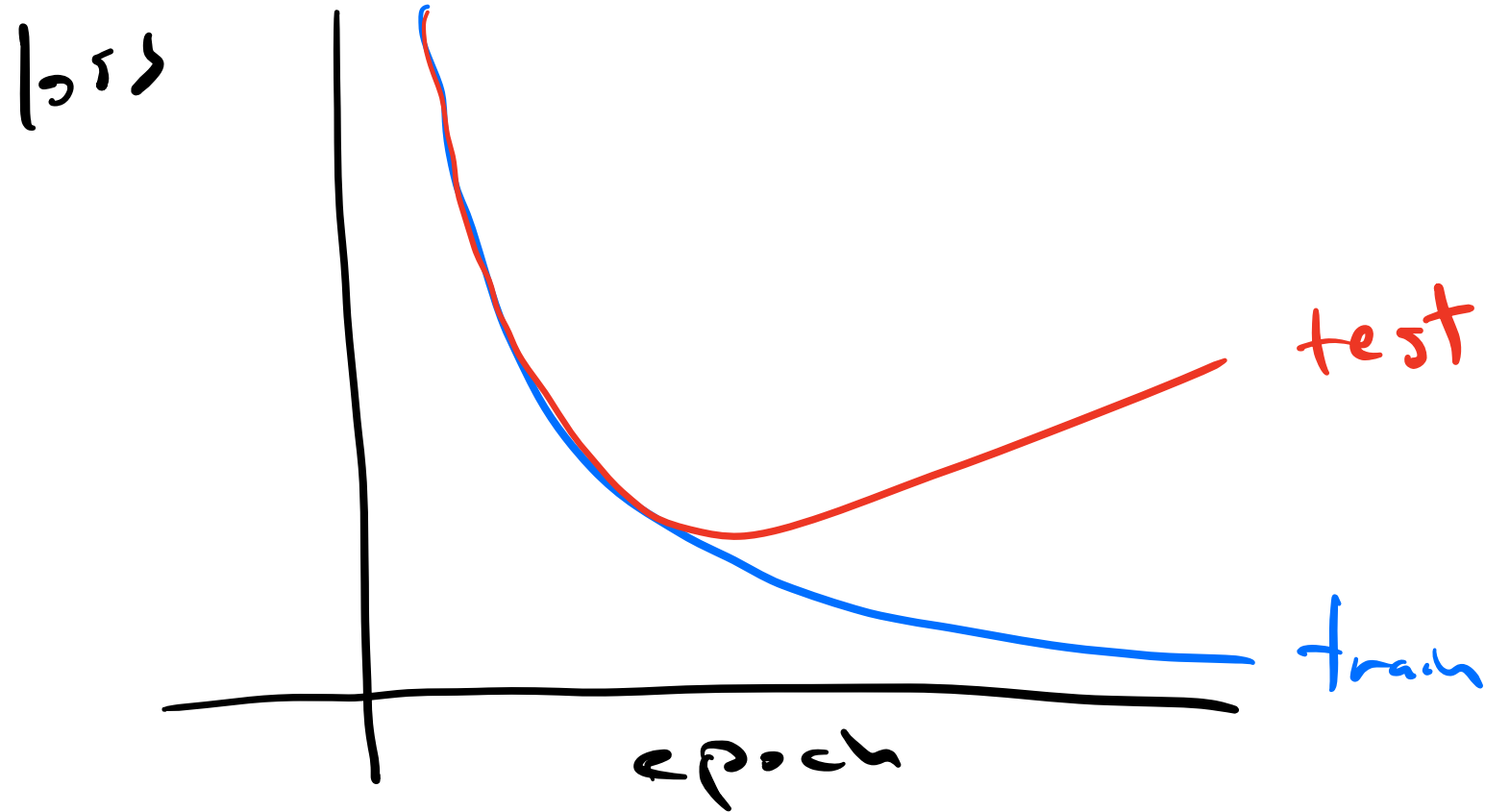
Train, Test, Validation



optimize

"optimize"
hyperparameters.

Underfitting and Overfitting



Batching

$$L = \frac{1}{m} \sum_{i=1}^m \lambda_i(\theta) = E(\lambda(\theta))$$

$$\frac{dL}{d\theta} = \frac{1}{m} \sum_{i=1}^m \frac{d\lambda_i}{d\theta} = E\left(\frac{d\lambda}{d\theta}\right)$$

$$\frac{dL}{d\theta} \approx \frac{1}{n} \sum_{i=1}^n \frac{d\lambda_i}{d\theta} \quad n \ll m$$

