

Foundations of Machine Learning

AI2000 and AI5000

FoML-21

Logistic Regression - Newton Raphson optimization

Dr. Konda Reddy Mopuri

Department of AI, IIT Hyderabad

July-Nov 2025



భారతీయ సాంకేతిక విజ్ఞాన సంస్థ హైదరాబాద్
भारतीय प्रौद्योगिकी संस्थान हैदराबाद
Indian Institute of Technology Hyderabad



So far in FoML

- Intro to ML and Probability refresher
- MLE, MAP, and fully Bayesian treatment
- Supervised learning
 - a. Linear Regression with basis functions (regularization, model selection)
 - b. Bias-Variance Decomposition (Bayesian Regression)
 - c. Decision Theory - three broad classification strategies
 - Probabilistic Generative Models - Continuous & discrete data
 - (Linear) Discriminant Functions - least squares solution, Perceptron
 - Probabilistic Discriminative Models - Logistic Regression

Logistic Regression - IRLS



Loss function approximation

- Taylor series expansion (univariate case)

$$E(w + \Delta w) = E(w) + \Delta w \cdot E'(w) + \frac{1}{2} \Delta w^2 E''(w) + \dots$$



Loss function approximation

- Taylor series expansion (multivariate case)

$$E(\mathbf{w} + \Delta \mathbf{w}) = E(\mathbf{w}) + \underline{\Delta \mathbf{w}}^T \cdot \underline{\nabla E} + \frac{1}{2} \underline{\Delta \mathbf{w}}^T \cdot \underline{H} \cdot \underline{\Delta \mathbf{w}} + \dots$$

$$\underline{H_{ij}} = \frac{\partial^2 E(\mathbf{w})}{\partial w_i \partial w_j}$$

Hessian of E at $\mathbf{w}$



Linear (first-order) approximation

$$E(\mathbf{w} + \Delta \mathbf{w}) \approx E(\mathbf{w}) + \Delta \mathbf{w}^T \nabla E(\mathbf{w})$$

for $\nabla E(\mathbf{w})$

G.S.

$\Rightarrow \Delta \mathbf{w} = -\eta \nabla E(\mathbf{w})$ { i.e., take a small step
in the opposite
direction of gradient



Quadratic (second-order) approximation

$$E(\mathbf{w} + \Delta \mathbf{w}) \approx E(\mathbf{w}) + \Delta \mathbf{w}^T \nabla E(\mathbf{w}) + \frac{1}{2} \Delta \mathbf{w}^T \nabla^2 E(\mathbf{w}) \Delta \mathbf{w}$$

what is the best $\Delta \mathbf{w}$ for minimizing $E(\mathbf{w})$

$$\Rightarrow \frac{\partial E(\mathbf{w} + \Delta \mathbf{w})}{\partial \Delta \mathbf{w}} = 0$$

$\Rightarrow$

$$\nabla E(\mathbf{w}) + \nabla^2 E(\mathbf{w}) \cdot \Delta \mathbf{w} = 0$$

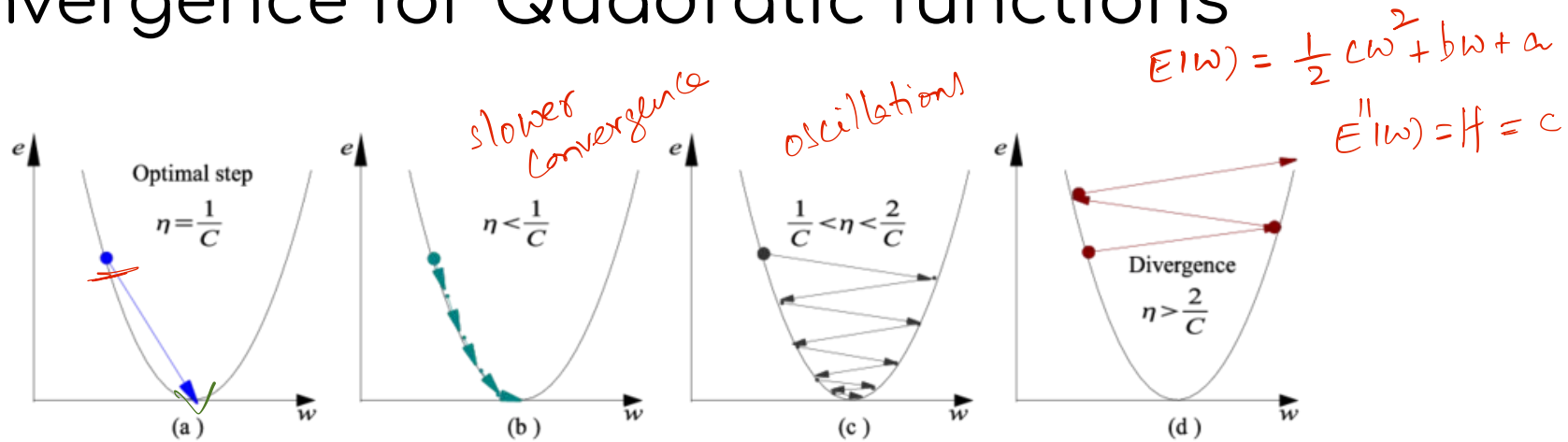
Hessian

$$\Delta \mathbf{w} = - \underbrace{\left[\nabla^2 E(\mathbf{w}) \right]^{-1}}_{M \times M} \underbrace{\nabla E(\mathbf{w})}_{M \times 1}$$

Best $\Delta \mathbf{w}$ to minimize $E(\mathbf{w})$

$$\mathbf{w}^{t+1} = \mathbf{w}^t - \mathbf{H}^{-1} \nabla E(\mathbf{w})$$

Convergence for Quadratic functions



quadratic loss function is convex
 $\Rightarrow$ 2nd order approximation will be exact for such functions
 $\Rightarrow$ single step convergence (irrespective of w^0)

$$w^* = w^0 - \frac{1}{H} \nabla E(w^0)$$



Generic Convex functions

second-order approx.

$$\begin{aligned} \rightarrow x_0 &= \text{initial guess} \\ \underline{x_1} &= x_0 - \frac{1}{f'(x_0)} f(x_0) \\ \underline{x_2} &= x_1 - \frac{1}{f'(x_1)} f(x_1) \end{aligned}$$

Newton-Raphson Iterative optimization

$$\mathbf{w}^{t+1} = \mathbf{w}^t - \mathbf{H}^{-1} \nabla E(\mathbf{w}^t)$$

$$\nabla E(\mathbf{w})^T = \sum_{n=1}^N (y_n - t_n) \phi_n = \Phi^T [\mathbf{y} - \mathbf{t}]$$

$$\phi_n = \underline{\phi}(x_n)$$

$$\Phi = \begin{bmatrix} \phi(x_1) \\ \phi(x_2) \\ \vdots \\ \phi(x_N) \end{bmatrix} \begin{matrix} \times (y_1 - t_1) \\ \vdots \\ \times (y_N - t_N) \end{matrix}$$

$N \times M$

$$\mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{bmatrix}_{N \times 1}$$

$$\mathbf{t} = \begin{bmatrix} t_1 \\ t_2 \\ \vdots \\ t_N \end{bmatrix}$$



Newton-Raphson Iterative optimization

$$\mathbf{w}^{t+1} = \mathbf{w}^t - \mathbf{H}^{-1} \nabla E(\mathbf{w}^t)$$

$$\Phi = \begin{bmatrix} \phi_1 \\ \phi_2 \\ \vdots \\ \phi_N \end{bmatrix} \quad N \times M$$

$$\mathbf{H}_{ij} = \frac{\partial^2 E(\mathbf{w}^t)}{\partial \mathbf{w}_i \partial \mathbf{w}_j} = \frac{\partial}{\partial \mathbf{w}_i} \sum_{n=1}^N (y_n - t_n) \phi_j(\mathbf{x}_n) = \sum_{n=1}^N \underbrace{y_n (1 - y_n)}_{\text{scalar}} \underbrace{\phi_i(\mathbf{x}_n) \phi_j(\mathbf{x}_n)}_{\text{scalar}}$$

$$\mathbf{H} = \begin{bmatrix} \cdot & \cdot & \cdot & \cdot & \cdot \\ \vdots & & & & \\ \cdot & \cdot & \cdot & \cdot & \cdot \end{bmatrix} \quad M \times M$$

$$\mathbf{H} = \mathbf{\Phi}^T \mathbf{R} \mathbf{\Phi} \quad \begin{matrix} M \times N & N \times N & N \times M \end{matrix}$$

$$R_{nn} = y_n (1 - y_n)$$



Newton-Raphson Iterative optimization

$$\mathbf{w}^{t+1} = \mathbf{w}^t - \mathbf{H}^{-1} \nabla E(\mathbf{w}^t)$$

$$\begin{aligned} \mathbf{w}^{t+1} &= \widetilde{\mathbf{w}}^t - (\Phi^T R \Phi)^{-1} \Phi^T (\mathbf{y} - \mathbf{t}) = \underbrace{(\Phi^T R \Phi)^{-1} \Phi^T}_{\text{Handwritten: } (\Phi^T \Phi)^{-1} \Phi^T} \underbrace{\left[\underbrace{\mathbf{I}}_{\text{Handwritten: } \mathbf{I}} \underbrace{(\Phi^T R \Phi)}_{\text{Handwritten: } (\Phi^T \Phi)} \right]}_{\text{Handwritten: } (\mathbf{y} - \mathbf{t})} \mathbf{w}^t - \underbrace{(\Phi^T R \Phi)^{-1} \Phi^T R \mathbf{z}}_{\text{Handwritten: } (\Phi^T \Phi)^{-1} \Phi^T \mathbf{t}} \end{aligned}$$

$$R_{nn} = y_n (1 - y_n) \quad \sigma(\mathbf{w}^t \Phi_n) = y_n = P(1/x)$$



$$E(w) = \sum_{i=1}^n \beta_i (t_i - y_i)^2$$

Next Neural Networks

