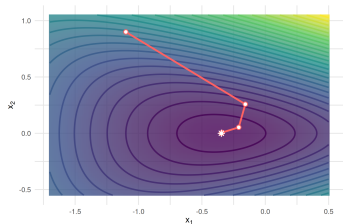
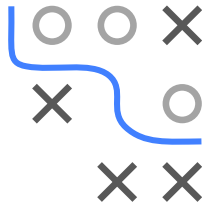


Optimization in Machine Learning

Second order methods Newton and IRLS

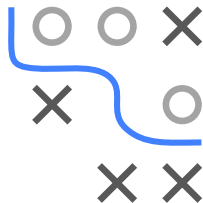


Learning goals

- Setup: Linear predictor + separable loss
- Structure of gradient and Hessian
- Newton as weighted least squares

Slides based on
► Aggarwal 2020 (Ch. 5.5)

SETUP: LINEAR PREDICTOR AND SEPARABLE LOSS



- We now look at L2-regularized ERM problems with a **linear predictor** and a loss that is **separable over the observations**:
 - **(A1) Linearity**: scalar score $f(\mathbf{x}^{(i)} \mid \boldsymbol{\theta}) = \boldsymbol{\theta}^\top \mathbf{x}^{(i)} = (\mathbf{X}\boldsymbol{\theta})_i$
 - **(A2) Separability**: obs. i only enters additively via its own loss term $L(y^{(i)}, f(\mathbf{x}^{(i)} \mid \boldsymbol{\theta}))$
- This yields objectives of the form

$$\mathcal{R}_{\text{emp}}(\boldsymbol{\theta}) = \sum_{i=1}^n L(y^{(i)}, f(\mathbf{x}^{(i)} \mid \boldsymbol{\theta})) + \frac{\lambda}{2} \|\boldsymbol{\theta}\|_2^2$$

THREE EXAMPLES FITTING THIS PATTERN

The following three (regularized) models are of exactly this form – they only differ in the point-wise loss $L(y^{(i)}, f(\mathbf{x}^{(i)} | \theta))$:

- **Least squares regression** ($y^{(i)} \in \mathbb{R}$): quadratic loss

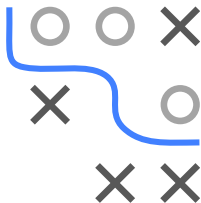
$$\mathcal{R}_{\text{emp}}(\theta) = \sum_{i=1}^n \frac{1}{2} \left(f(\mathbf{x}^{(i)} | \theta) - y^{(i)} \right)^2 + \frac{\lambda}{2} \|\theta\|_2^2$$

- **Logistic regression** ($y^{(i)} \in \{0, 1\}$): Bernoulli loss with $\pi^{(i)} = s(f(\mathbf{x}^{(i)} | \theta))$

$$\mathcal{R}_{\text{emp}}(\theta) = - \sum_{i=1}^n \left[y^{(i)} \log(\pi^{(i)}) + (1 - y^{(i)}) \log(1 - \pi^{(i)}) \right] + \frac{\lambda}{2} \|\theta\|_2^2$$

- **L2-SVM** ($y^{(i)} \in \{-1, +1\}$): squared hinge loss

$$\mathcal{R}_{\text{emp}}(\theta) = \sum_{i=1}^n \max \left(0, 1 - y^{(i)} f(\mathbf{x}^{(i)} | \theta) \right)^2 + \frac{\lambda}{2} \|\theta\|_2^2$$



NEWTON'S METHOD FOR THIS PROBLEM CLASS

- **Recall** the Newton-Raphson update applied to $\mathcal{R}_{\text{emp}}(\theta)$:

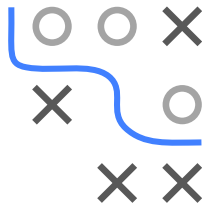
$$\theta^{[t+1]} = \theta^{[t]} - (\nabla^2 \mathcal{R}_{\text{emp}}(\theta^{[t]}))^{-1} \nabla \mathcal{R}_{\text{emp}}(\theta^{[t]})$$

- **Goal:** We will now show that for problems of this kind, each step of Newton's method can be seen as solving a **weighted least squares** problem – with weights and “working responses” that change over the iterations and can be interpreted:

Iteratively Reweighted Least Squares (IRLS)

- **Plan:**

- **Step 1:** derive the structure of the gradient $\nabla \mathcal{R}_{\text{emp}}(\theta)$
- **Step 2:** derive the structure of the Hessian $\nabla^2 \mathcal{R}_{\text{emp}}(\theta)$
- **Step 3:** rewrite the Newton step as a weighted LS problem



STEP 1: STRUCTURE OF THE GRADIENT

- Quantity of interest: $\nabla \mathcal{R}_{\text{emp}}(\boldsymbol{\theta}) \in \mathbb{R}^{1 \times d}$, with

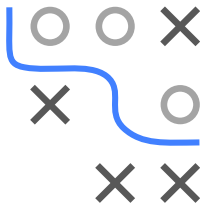
$$\nabla \mathcal{R}_{\text{emp}}(\boldsymbol{\theta}) = \sum_{i=1}^n \nabla_{\boldsymbol{\theta}} L\left(y^{(i)}, f(\mathbf{x}^{(i)} | \boldsymbol{\theta})\right) + \lambda \boldsymbol{\theta}^{\top}$$

- Chain rule on each loss term:

$$\nabla_{\boldsymbol{\theta}} L\left(y^{(i)}, f(\mathbf{x}^{(i)} | \boldsymbol{\theta})\right) = \underbrace{\frac{\partial L\left(y^{(i)}, f(\mathbf{x}^{(i)} | \boldsymbol{\theta})\right)}{\partial f(\mathbf{x}^{(i)} | \boldsymbol{\theta})}}_{=: r_i \in \mathbb{R} \text{ (residual)}} \underbrace{\nabla_{\boldsymbol{\theta}} f(\mathbf{x}^{(i)} | \boldsymbol{\theta})}_{= (\mathbf{x}^{(i)})^{\top} \text{ by (A1)}} \in \mathbb{R}^{1 \times d}$$

- Summing over the observations, with $\mathbf{r} := (r_1, \dots, r_n)^{\top} \in \mathbb{R}^n$:

$$\boxed{\nabla \mathcal{R}_{\text{emp}}(\boldsymbol{\theta}) = \mathbf{r}^{\top} \mathbf{X} + \lambda \boldsymbol{\theta}^{\top}} \in \mathbb{R}^{1 \times d}$$



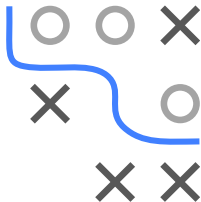
STEP 2: STRUCTURE OF THE HESSIAN

- Differentiate the gradient from step 1, using that the only θ -dependence in $\mathbf{r}^\top \mathbf{X}$ sits inside $\mathbf{r}$:

$$\nabla^2 \mathcal{R}_{\text{emp}}(\boldsymbol{\theta}) = \frac{\partial}{\partial \boldsymbol{\theta}^\top} \left(\underbrace{\mathbf{r}^\top \mathbf{X}}_{=\sum_{i=1}^n r_i (\mathbf{x}^{(i)})^\top} + \lambda \boldsymbol{\theta}^\top \right) = \sum_{i=1}^n \frac{\partial r_i}{\partial \boldsymbol{\theta}^\top} (\mathbf{x}^{(i)})^\top + \lambda \mathbf{I}$$

- Differentiating one residual entry:

$$\frac{\partial r_i}{\partial \boldsymbol{\theta}^\top} = \underbrace{\frac{\partial^2 L(y^{(i)}, f(\mathbf{x}^{(i)} | \boldsymbol{\theta}))}{\partial f(\mathbf{x}^{(i)} | \boldsymbol{\theta})^2}}_{=: d_i \in \mathbb{R}} \underbrace{\frac{\partial f(\mathbf{x}^{(i)} | \boldsymbol{\theta})}{\partial \boldsymbol{\theta}^\top}}_{=\mathbf{x}^{(i)} \text{ by (A1)}} = d_i \mathbf{x}^{(i)} \in \mathbb{R}^d$$



STEP 2: STRUCTURE OF THE HESSIAN

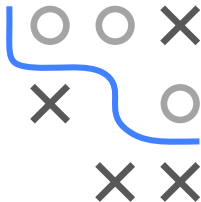
- Plugging in yields the Hessian in sum form:

$$\nabla^2 \mathcal{R}_{\text{emp}}(\boldsymbol{\theta}) = \sum_{i=1}^n d_i \mathbf{x}^{(i)} \left(\mathbf{x}^{(i)} \right)^\top + \lambda \mathbf{I} \in \mathbb{R}^{d \times d}$$

- In matrix notation, stacking the d_i into $\mathbf{D} := \text{diag}(d_1, \dots, d_n) \in \mathbb{R}^{n \times n}$:

$$\boxed{\nabla^2 \mathcal{R}_{\text{emp}}(\boldsymbol{\theta}) = \mathbf{X}^\top \mathbf{D} \mathbf{X} + \lambda \mathbf{I}}$$

- **Note:** the loss only impacts the Hessian through the diagonal entries d_i of $\mathbf{D}$, while the structure $\mathbf{X}^\top \mathbf{D} \mathbf{X}$ itself is fixed by the setup (A1) + (A2); $\mathbf{D}$ is diagonal here because the loss is separable over the observations (A2)



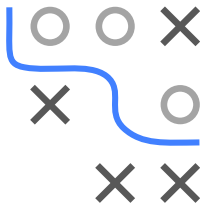
STEP 3: NEWTON STEP AS WEIGHTED LEAST SQUARES

- Plug gradient and Hessian into the Newton-Raphson update, where $\mathbf{r}$ and $\mathbf{D}$ are evaluated at the current iterate $\boldsymbol{\theta}^{[t]}$:

$$\begin{aligned}\boldsymbol{\theta}^{[t+1]} &= \boldsymbol{\theta}^{[t]} - (\mathbf{X}^\top \mathbf{D} \mathbf{X} + \lambda \mathbf{I})^{-1} (\mathbf{X}^\top \mathbf{r} + \lambda \boldsymbol{\theta}^{[t]}) \\ &= (\mathbf{X}^\top \mathbf{D} \mathbf{X} + \lambda \mathbf{I})^{-1} ((\mathbf{X}^\top \mathbf{D} \mathbf{X} + \lambda \mathbf{I}) \boldsymbol{\theta}^{[t]} - \mathbf{X}^\top \mathbf{r} - \lambda \boldsymbol{\theta}^{[t]}) \\ &= (\mathbf{X}^\top \mathbf{D} \mathbf{X} + \lambda \mathbf{I})^{-1} (\mathbf{X}^\top \mathbf{D} \mathbf{X} \boldsymbol{\theta}^{[t]} - \mathbf{X}^\top \mathbf{r}) \\ &= (\mathbf{X}^\top \mathbf{D} \mathbf{X} + \lambda \mathbf{I})^{-1} \mathbf{X}^\top \mathbf{D} \underbrace{(\mathbf{X} \boldsymbol{\theta}^{[t]} - \mathbf{D}^{-1} \mathbf{r})}_{=: \tilde{\mathbf{f}} \in \mathbb{R}^n}\end{aligned}$$

- Here, $\tilde{\mathbf{f}}$ is called the **working response** (assuming $d_i > 0$), which will take the role of a pseudo-target; entry-wise

$$\tilde{f}^{(i)} = f(\mathbf{x}^{(i)} | \boldsymbol{\theta}^{[t]}) - r_i / d_i$$

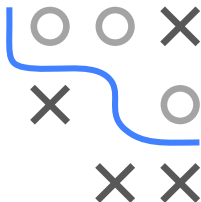


STEP 3: NEWTON STEP AS WEIGHTED LEAST SQUARES

- This Newton iterate solves a **regularized weighted least squares** problem with working responses $\tilde{f}^{(i)}$ as targets and weights d_i :

$$\boldsymbol{\theta}^{[t+1]} = \arg \min_{\boldsymbol{\theta}} \frac{1}{2} \sum_{i=1}^n d_i \left(\tilde{f}^{(i)} - \boldsymbol{\theta}^\top \mathbf{x}^{(i)} \right)^2 + \frac{\lambda}{2} \|\boldsymbol{\theta}\|_2^2$$

- Since $\mathbf{r}$ and $\mathbf{D}$ depend on the current iterate $\boldsymbol{\theta}^{[t]}$, the weights d_i and working responses $\tilde{f}^{(i)}$ change in every iteration
- The iterative Newton method thus repeatedly solves *reweighted* LS problems – hence the name **Iteratively Reweighted Least Squares (IRLS)**



EXAMPLE: LEAST SQUARES REGRESSION

- Quadratic loss $L(y^{(i)}, f(\mathbf{x}^{(i)} | \theta)) = \frac{1}{2} (f(\mathbf{x}^{(i)} | \theta) - y^{(i)})^2$, so

$$r_i = f(\mathbf{x}^{(i)} | \theta) - y^{(i)}, \quad d_i = 1$$

- Weights and working responses **collapse** to $\mathbf{D} = \mathbf{I}$ and $\tilde{\mathbf{f}} = \mathbf{y}$, since

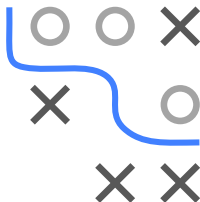
$$\tilde{f}^{(i)} = f(\mathbf{x}^{(i)} | \theta^{[t]}) - \frac{f(\mathbf{x}^{(i)} | \theta^{[t]}) - y^{(i)}}{1} = y^{(i)}$$

- Plugging into the Newton step:

$$\theta^{[t+1]} = (\mathbf{X}^\top \mathbf{D} \mathbf{X} + \lambda \mathbf{I})^{-1} \mathbf{X}^\top \mathbf{D} \tilde{\mathbf{f}} = (\mathbf{X}^\top \mathbf{X} + \lambda \mathbf{I})^{-1} \mathbf{X}^\top \mathbf{y}$$

recovering the well-known (regularized) LS estimator

- Since, in this specific case, neither $\mathbf{D}$ nor $\tilde{\mathbf{f}}$ depends on $\theta^{[t]}$, the problem is solved in **one** Newton step (from any starting point)



EXAMPLE: LOGISTIC REGRESSION

- Using $s'(f) = s(f)(1 - s(f))$, for the Bernoulli loss it follows:

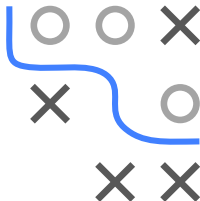
$$r_i = - \left(\frac{y^{(i)}}{\pi^{(i)}} - \frac{1 - y^{(i)}}{1 - \pi^{(i)}} \right) s' \left(f(\mathbf{x}^{(i)} \mid \boldsymbol{\theta}) \right) = \pi^{(i)} - y^{(i)},$$

$$d_i = \frac{\partial r_i}{\partial f(\mathbf{x}^{(i)} \mid \boldsymbol{\theta})} = s' \left(f(\mathbf{x}^{(i)} \mid \boldsymbol{\theta}) \right) = \pi^{(i)} (1 - \pi^{(i)}),$$

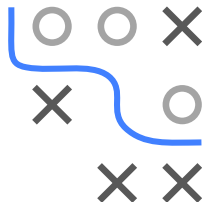
$$\tilde{f}^{(i)} = f(\mathbf{x}^{(i)} \mid \boldsymbol{\theta}^{[t]}) - \frac{\pi^{(i)} - y^{(i)}}{\pi^{(i)}(1 - \pi^{(i)})}$$

(for a detailed derivation see [► Aggarwal 2020](#))

- Interpretation:** Weights $d_i = \pi^{(i)}(1 - \pi^{(i)}) \in (0, \frac{1}{4}]$ are small for confidently classified obs. – they barely influence the update



EXAMPLE: L2-SVM



- For $i \in A$, the working response collapses to the label (using $(y^{(i)})^2 = 1$):

$$\tilde{f}^{(i)} = f(\mathbf{x}^{(i)} \mid \boldsymbol{\theta}^{[t]}) + y^{(i)} \left(1 - y^{(i)} f(\mathbf{x}^{(i)} \mid \boldsymbol{\theta}^{[t]}) \right) = y^{(i)}$$

- Consequently, the Newton step is a ridge regression on the current margin violators with their true labels as targets ($\mathbf{X}_A, \mathbf{y}_A$: rows/entries with $i \in A$):

$$\boldsymbol{\theta}^{[t+1]} = (2\mathbf{X}_A^\top \mathbf{X}_A + \lambda \mathbf{I})^{-1} 2\mathbf{X}_A^\top \mathbf{y}_A$$