

goodness-of-fit measures

(a) Root-mean-square error:

$$RMSE(\vec{w}) = \sqrt{\frac{1}{N} \sum_{n=1}^N (y_n - \underbrace{\vec{w}' \cdot \vec{x}_n}_{d+1})^2}$$

(b) Coeff. of determination:

self. of determination:

$$R^2 = 1 - \frac{\sum_n (y_n - \hat{y}_n)^2 \leftarrow \text{var in our model}}{\sum_n (y_n - \bar{y})^2 \leftarrow \text{var in 'null' model } \bar{y}}$$

$$\hat{y}_n = \underline{w}' \cdot \underline{x}_n, \quad \bar{y} = \frac{1}{N} \sum_n y_n.$$

$0 \leq R^2 \leq 1$

↑
 model does just as well as the $\bar{y}$ model

↑
 model reproduces y_n exactly

Larger $R^2 \Rightarrow$ better fit

Ridge regression

Introduce prior $p(\vec{w}) = \mathcal{N}(\vec{w} | \vec{0}, \tau^2 \mathbb{I})$ to avoid overfitting.

$$\vec{w}'_{\text{MAP}} = \underset{\vec{w}}{\text{argmin}} \left\{ \frac{1}{2\sigma^2} (\vec{y} - X\vec{w}')^T (\vec{y} - X\vec{w}') + \frac{1}{2\tau^2} \underbrace{\vec{w}^T \vec{w}}_{\text{"}|\vec{w}|^2\text{"}} \right\}$$

b is found as before $\approx D\text{-dim}$
and is not constrained

Consider $J(\vec{w}) = (\vec{y} - X\vec{w})^T (\vec{y} - X\vec{w}) + \lambda \underbrace{|\vec{w}|^2}_{\frac{\sigma^2}{\tau^2}}$

$\uparrow$
D-dim

Then $\frac{\partial}{\partial \vec{w}} J(\vec{w}) = 2 (X^T X \vec{w} - X^T \vec{y} + \lambda \vec{w}) = 0$, so that

$\left[\begin{array}{l} X_{N \times D} \text{ is the centered} \\ \text{data matrix,} \\ \vec{y} \text{ is a vector of centered observations} \end{array} \right]$

$$(X^T X + \lambda \mathbb{I}_D) \vec{w}_{\text{MAP}} = X^T \vec{y},$$

$$\vec{w}_{\text{MAP}} = (X^T X + \lambda \mathbb{I}_D)^{-1} X^T \vec{y} =$$

$$= \left(\underbrace{\sum_n (\vec{x}_n \vec{x}_n^T)}_{D \times D} + \lambda \mathbb{I}_D \right)^{-1} \times \sum_n y_n \vec{x}_n$$

$X^T X + \lambda \mathbb{I}_D$ should be more stable to invert than $X^T X$.

SVD solution

Consider $X = \underbrace{U}_{N \times D} \underbrace{S}_{N \times N} \underbrace{V^T}_{N \times D \ D \times D}$, then

$$\vec{w}_{\text{MAP}} = ((V S^T U^T)(U S V^T) + \lambda \mathbb{I}_D)^{-1} (V S^T U^T) \vec{y} \quad \ominus$$

Note that in general,

$$(+)$$

$$\underbrace{\left(\mathbb{I}_M + \underbrace{A B}_{M \times M} \right)^{-1}}_{M \times M} A_{M \times N} = A_{M \times N} \underbrace{\left(\mathbb{I}_N + \underbrace{B_{N \times M} A_{M \times N}}_{N \times N} \right)^{-1}}_{N \times N}$$

Using Eq. (+) with $\begin{cases} A \rightarrow X^T \\ B \rightarrow X \end{cases}$, we obtain:

$$\Leftrightarrow X^T \lambda^{-1} [\mathbb{I}_N + \underbrace{\lambda X X^T}_{N \times N}]^{-1} \bar{y} \Leftrightarrow$$

$$X X^T = (U S V^T)(V S^T U^T) = (\underbrace{U S}_{N \times D})(\underbrace{S^T U^T}_{D \times N}) = R R^T$$

$$\Leftrightarrow (V R^T) \lambda^{-1} [\mathbb{I}_N + \lambda^{-1} R R^T]^{-1} \bar{y} =$$

$$= (V R^T) [\lambda \mathbb{I}_N + R R^T]^{-1} \bar{y} =$$

$$= V [\lambda \mathbb{I}_D + \underbrace{R^T R}_{D \times D}] R^T \bar{y}$$

It may be more numerically stable to compute $\bar{w}_{MAP}$ in N -space.

—○—

Next, consider

$$\hat{\bar{y}} = X \bar{w}_{MAP} = (U S V^T) V [\lambda \mathbb{I}_D + \underbrace{R^T R}_{D \times D}]^{-1} (S^T U^T) \bar{y} \Leftrightarrow$$

$$\underbrace{(S^T U^T)(U S)}_{S^T S_{D \times N} \quad N \times D} \equiv \underbrace{S^2}_{D \times D \text{ diag.}}$$

$$\Leftrightarrow U S [\lambda \mathbb{I}_D + S^2]^{-1} S^T U^T \bar{y}$$

$N \times N$ matrix with D non-zero elements on the diagonal: $\frac{\sigma_j^2}{\sigma_j^2 + \lambda}$

$j = 1, \dots, D$

$\frac{\sigma_j^2}{\sigma_j^2 + \lambda}$

Finally,

$$\underbrace{\hat{\vec{y}}}_{N\text{-dim}} = \sum_{j=1}^D \underbrace{\vec{u}_j}_{\substack{\text{N-dim} \\ \text{N-dim}}} \frac{\sigma_j^2}{\sigma_j^2 + \lambda} \underbrace{\vec{u}_j^T \vec{y}}_{\substack{\text{N-dim} \\ \text{N-dim}}}.$$

If $\sigma_j^2 \ll \lambda$, $\frac{\sigma_j^2}{\sigma_j^2 + \lambda} \approx \frac{\sigma_j^2}{\lambda} \ll 1$, and the corresponding $\vec{u}_j$ does not contribute to $\hat{\vec{y}}$.

One can define the effective # of degrees of freedom (DoF):

$$D(\lambda) = \sum_{j=1}^D \frac{\sigma_j^2}{\sigma_j^2 + \lambda}.$$

If $\lambda \rightarrow 0$, $D(\lambda) \rightarrow D$;

If $\lambda \rightarrow \infty$, $D(\lambda) \rightarrow 0$.

Now recall that

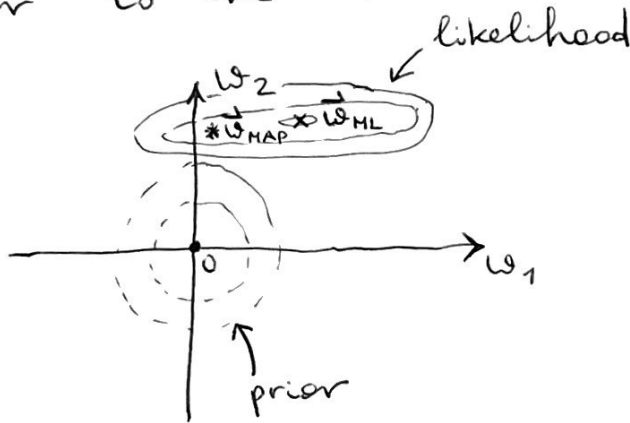
$$\begin{aligned} \hat{\vec{y}}_{MLE} &= X \vec{w}_{MLE} = X (X^T X)^{-1} X^T \vec{y} = \\ &= (U S V^T) \underbrace{(V S^T U^T U S V^T)^{-1}}_{\text{II}} (V S^T U^T) \vec{y} = \\ &= (U S V^T) \underbrace{(V S^2 V^T)^{-1}}_{D \times D} (V S^T U^T) \vec{y} = \\ &= \underbrace{(U S V^T) (V^T)^{-1}}_{\text{II}} (S^2)^{-1} \underbrace{V^{-1} V S^T U^T}_{\text{II}} \vec{y} \quad \text{②} \end{aligned}$$

$$\textcircled{=}\quad \underbrace{U S (S^2)^{-1} S^T U^T}_{N \times N \text{ matrix}} \bar{y} = \sum_{j=1}^D \bar{u}_j \bar{u}_j^T \bar{y}$$

with D 1's on the diagonal

all $\bar{u}_j$'s are equally important,
there's no regularization.

Qualitatively, the ^{'soft'} directions that are not strongly constrained by the data will be pulled towards 0 by the prior, whereas the 'hard' directions will stay much closer to the ML solution:



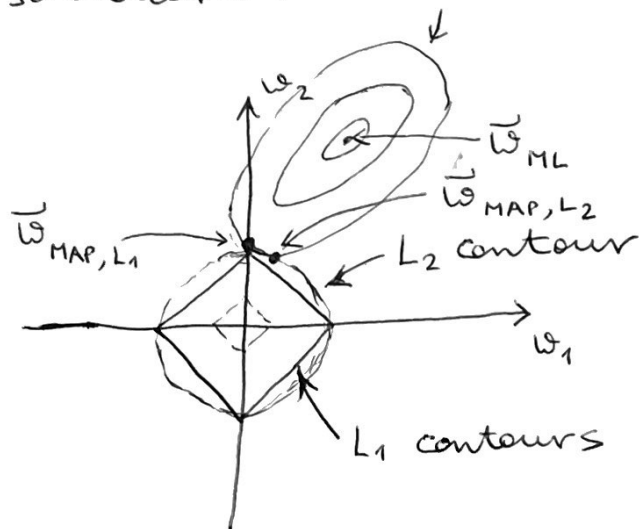
$$\bar{w} = (w_1, w_2)$$

λ in the MAP procedure can be chosen using standard methods such as 5-fold cross-validation.

Lasso regression

L_1 rather than L_2 norm is employed:
 $\lambda |\vec{w}|^2$ is replaced by $\lambda \sum_{j=1}^D |w_j|$.

This turns out to yield much sparser solutions: feature selection



Ex.

$\vec{w} = (1, 0)$ sparse

$\vec{w}' = (\frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}})$ non-sparse

$$|\vec{w}|^2 = |\vec{w}'|^2 = 1, \text{ but } \sum_{i=1}^2 |w_i| = 1 < \sum_{i=1}^2 |w'_i| = \sqrt{2}$$

↓
 sparse solution preferred under the L_1 norm

However, using L_1 penalty requires numerical optimization $\Rightarrow$ there are ~~no~~ adapted specifically to this problem. methods

The problem is that the L_1 penalty is non-smooth $\Rightarrow$ use proximal grad descent

Consider minimizing $J(\vec{\theta}) = \underbrace{J_1(\vec{\theta})}_{\text{smooth (e.g., -likelihood)}} + \underbrace{J_2(\vec{\theta})}_{\text{non-smooth (e.g., a lasso constraint)}}$

Proximal grad. method:

$$\vec{\theta}_{t+1} = \text{Pr} \left[\underbrace{\vec{\theta}_t - \eta_t \nabla J_1(\vec{\theta}_t)}_{\vec{\theta}'_{t+1}} \right],$$

where $\Pr[\bar{\theta}] = \arg\min_{\bar{z}} [\mathcal{I}_2(\bar{z}) + \frac{1}{2\lambda} \|\bar{z} - \bar{\theta}\|^2]$.
 $\uparrow$
 proximal operator
 (*)

So, the grad step produces $\bar{\theta}'_{t+1}$ & the projection step produces $\bar{\theta}_{t+1}$ using $\bar{\theta}'_{t+1}$ as an input prn into Eq. (*).

For ex., consider a 1D constraint:

$$\mathcal{I}_2(\theta) = \begin{cases} 0 & \text{if } -a \leq \theta \leq a \\ \infty & \text{otherwise} \end{cases}$$

Then $\Pr[\theta] = \begin{cases} +a, & \theta > a \\ \theta, & -a \leq \theta \leq a \\ -a, & \theta < -a \end{cases}$

What about the $\mathcal{I}_1$ penalty?

$\mathcal{I}_2(\bar{z}) = \sum_{i=1}^D |\theta_i|$, each component can be projected separately.

For a single component θ , we need

to find $\Pr[\theta] = \arg\min_z \left\{ |z| + \frac{1}{2\lambda} (z - \theta)^2 \right\} =$
 $= \arg\min_z \left\{ \underbrace{\lambda |z| + \frac{1}{2} (z - \theta)^2}_{f(z)} \right\}.$

It turns out that

$$\lambda > 0$$

$$Pr[\theta] = \begin{cases} \theta - \lambda & \theta \geq \lambda \\ 0 & |\theta| \leq \lambda \\ \theta + \lambda & \theta < -\lambda \end{cases} = \begin{matrix} \text{soft-thresholding} \\ \text{operator:} \\ \text{Soft}(\theta, \lambda) \end{matrix}$$

$$\text{Indeed, } \frac{df(z)}{dz} = \begin{cases} \lambda + z - \theta & , z \geq 0 \\ -\lambda + z - \theta & , z < 0 \end{cases}$$

$$\text{If } \theta > \lambda \Rightarrow \begin{matrix} z^* = \theta - \lambda > 0 \\ \text{expect} \\ z^* > 0 @ \min \\ \text{unique} \end{matrix}$$

$$\text{If } \theta < -\lambda \Rightarrow \begin{matrix} z^* = \theta + \lambda < 0 \\ \text{expect } z^* < 0 \end{matrix}$$

$$\text{If } -\lambda \leq \theta \leq \lambda \Rightarrow \begin{matrix} f(z) \text{ has a minimum} \\ \text{at } z = 0 \Rightarrow f(z) = \frac{\theta^2}{2} \text{ in} \\ \text{this case.} \end{matrix}$$

Then Lasso can be implemented using

$$\vec{w}_{t+1} = \text{Soft}(\vec{w}_t - \eta_t \underbrace{\vec{g}(\vec{w})}_{\text{grad of negative log-likelihood}}, \eta_t \lambda) \quad (+)$$

penalty depends on step-size

Soft(...) acts on $\vec{w}$
component - by - component.

Eq. (+) implements iterative soft thresholding algorithm (ISTA)