

Linear Regression

$$P(y|x, \theta) = \mathcal{N}(y|w^T x, \sigma^2)$$

multivariate linear regression

Multiple data points

$$y_i = \text{~~W~~ } w^T x_i + \epsilon_i \quad \epsilon_i \sim \mathcal{N}(0, \sigma^2)$$

$$y_i = \sum_a w_a x_{ia} + \epsilon_i$$

Sometimes we have nonlinear relation expressed through a basis expansion

$$P(y|x, \theta) = \mathcal{N}(y|w^T \phi(x), \sigma^2)$$

For example $\phi(x) = [1 \ x \ \dots \ x^d]$

The structure of the problems ~~is~~ ^{are} the same

$$\begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{pmatrix} = \begin{pmatrix} x_{11} & \dots & x_{1D} \\ x_{21} & \dots & x_{2D} \\ \vdots & & \vdots \\ x_{N1} & \dots & x_{ND} \end{pmatrix} \begin{pmatrix} w_1 \\ \vdots \\ w_D \end{pmatrix} + \begin{pmatrix} \epsilon_1 \\ \epsilon_2 \\ \vdots \\ \epsilon_N \end{pmatrix}$$

or

$$\begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{pmatrix} = \begin{pmatrix} 1 & x_1 & x_1^2 & \dots & x_1^d \\ 1 & x_2 & x_2^2 & \dots & x_2^d \\ \vdots & \vdots & \vdots & & \vdots \\ 1 & x_N & x_N^2 & \dots & x_N^d \end{pmatrix} \begin{pmatrix} w_0 \\ w_1 \\ \vdots \\ w_d \end{pmatrix} + \begin{pmatrix} \epsilon_1 \\ \epsilon_2 \\ \vdots \\ \epsilon_N \end{pmatrix}$$

If ϵ_i 's are independent,

$$\log \prod p(y_i | w) = \log e^{-\frac{\sum (y_i - w^T x_i)^2}{2\sigma^2}} \frac{1}{(\sqrt{2\pi}\sigma)^n}$$

$$= -\frac{1}{2\sigma^2} \sum (y_i - w^T x_i)^2 - \frac{n}{2} \log 2\pi\sigma^2$$

Max over $w \Rightarrow \min \sum (y_i - w^T x_i)^2$

Least Square fit.

$$RSS(w) = \sum (y_i - w^T x_i)^2$$

↑
Residual sum of squares

$$RSS(w) = \| \epsilon \|^2$$

In general we will write

$$y = Xw + \epsilon$$

$$\begin{pmatrix} x_{11} & \dots & x_{1d} \\ \vdots & & \vdots \\ x_{N1} & \dots & x_{Nd} \end{pmatrix} \text{ or } \begin{pmatrix} 1 & x_1^1 & \dots & x_1^d \\ \vdots & \vdots & & \vdots \\ 1 & x_N^1 & \dots & x_N^d \end{pmatrix} \text{ or } \begin{pmatrix} \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot \end{pmatrix}$$

X is called the design matrix.

$$P(y/w) \propto e^{-\frac{(y-Xw)^T(y-Xw)}{2\sigma^2}}$$

So, the object to minimize is negative log likelihood (NLL) upto some linear transformation.

$$\begin{aligned} NLL(w) &= \frac{1}{2}(y-Xw)^T(y-Xw) \\ &= \frac{1}{2}w^T X^T X w - w^T X^T y + \frac{1}{2}y^T y \end{aligned}$$

Now take w derivative and set to zero

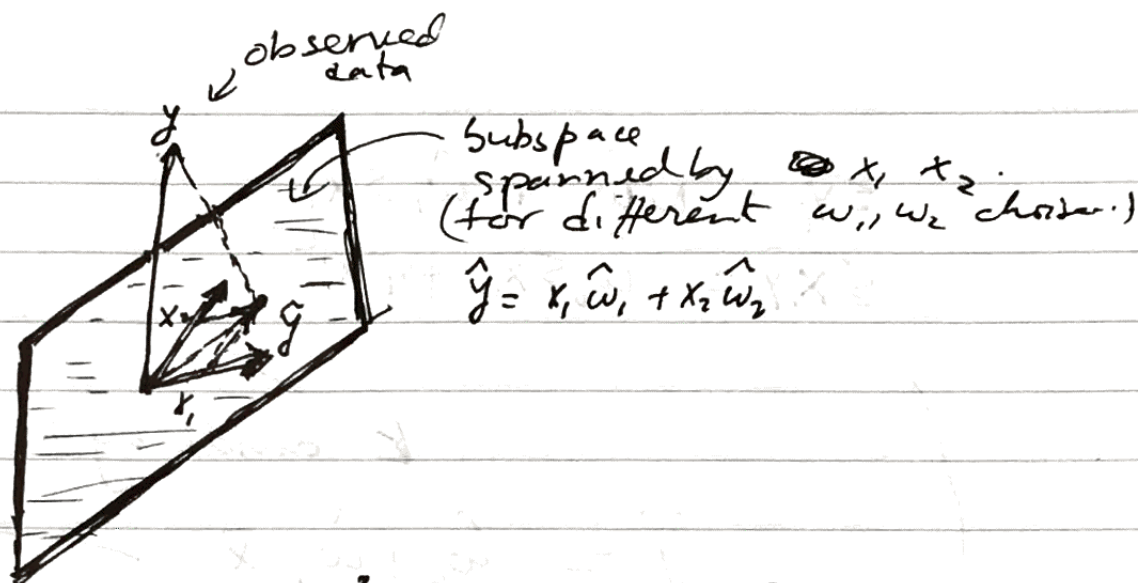
$$X^T X w = X^T y$$

$$\hat{w}_{OLS} = (X^T X)^{-1} X^T y$$

↑
ordinary
least square

$$y = x_1 w_1 + x_2 w_2 + \dots + \epsilon$$

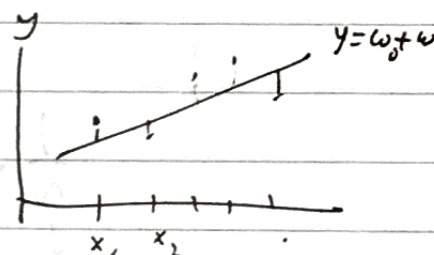
↑ ↑
Columns of the design matrix



$(y - \hat{y})^2$ is the ~~the~~ RSS

Let us take ~~the~~ the matrix formulation out for a test drive on the familiar road of linear regression

$$y_i = w_0 + w_1 x_i + \epsilon_i$$



$$\begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{pmatrix} = \underbrace{\begin{pmatrix} 1 & \cdots & x_1 \\ \vdots & & \vdots \\ 1 & \cdots & x_N \end{pmatrix}}_X \begin{pmatrix} w_0 \\ w_1 \end{pmatrix} + \begin{pmatrix} \epsilon_1 \\ \vdots \\ \epsilon_N \end{pmatrix}$$

$$X^T X = \begin{pmatrix} N & \sum x_i \\ \sum x_i & \sum x_i^2 \end{pmatrix} \quad X^T y = \begin{pmatrix} \sum y_i \\ \sum x_i y_i \end{pmatrix}$$

$$X^T X \hat{w} = X^T y \Rightarrow \begin{pmatrix} N & \sum x_i \\ \sum x_i & \sum x_i^2 \end{pmatrix} = \begin{pmatrix} \sum y_i \\ \sum x_i y_i \end{pmatrix}$$

$$\sum y_i = N\hat{\omega}_0 + \hat{\omega}_1 \sum x_i$$

$$\sum x_i y_i = \hat{\omega}_0 \sum x_i + \hat{\omega}_1 \sum x_i^2$$

← divide by N

$$\bar{y} = \hat{\omega}_0 + \hat{\omega}_1 \bar{x}$$

The line passes through $(\bar{x}, \bar{y})$

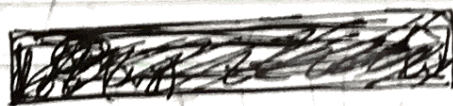
divide
by N

$$\frac{1}{N} \sum x_i y_i = \hat{\omega}_0 \bar{x} + \hat{\omega}_1 \frac{1}{N} \sum x_i^2$$

multiplies
by $\bar{x}$
and subtract

$$\frac{1}{N} \sum x_i y_i - \bar{x} \bar{y} = \hat{\omega}_1 \left(\frac{1}{N} \sum x_i^2 - \bar{x}^2 \right)$$

Empirical
covariance



empirical
variance of
x

$$\left[\frac{1}{N} \sum (x_i - \bar{x})(y_i - \bar{y}) \right] = \hat{\omega}_1 \left[\frac{1}{N} \sum (x_i - \bar{x})^2 \right]$$

Correlation coefficient

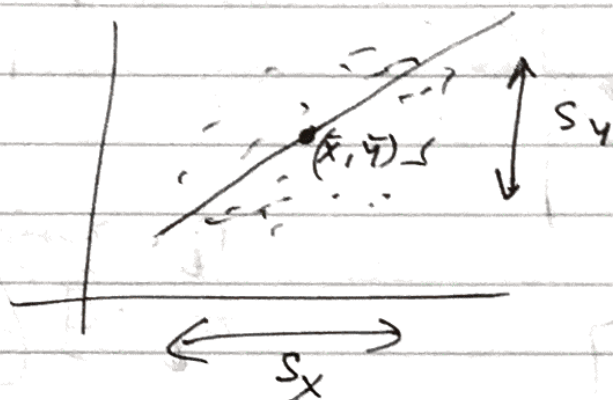
$$r_{x,y} = \frac{\frac{1}{N} \sum (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\frac{1}{N} \sum (x_i - \bar{x})^2} \sqrt{\frac{1}{N} \sum (y_i - \bar{y})^2}}$$

s_x

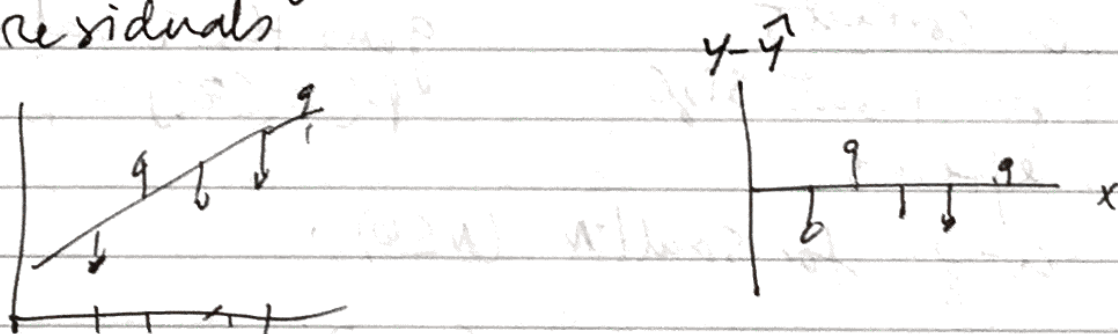
s_y

$$r_{x,y} S_x S_y = \hat{\omega}_1 S_x^2$$

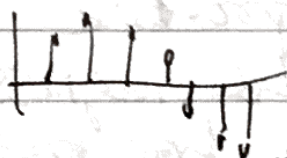
$$\hat{\omega}_1 = \frac{r_{x,y} S_y}{S_x}$$



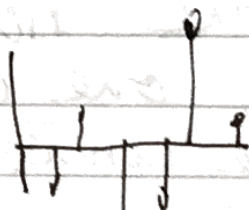
It is a good practice to check what the residuals



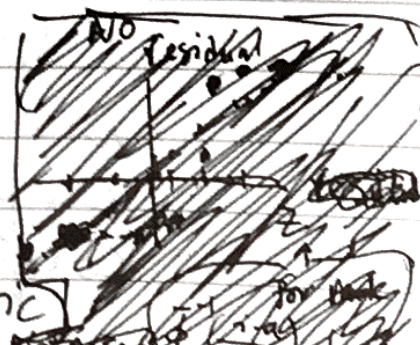
and check whether the residuals look like independent, homoskedastic and normal



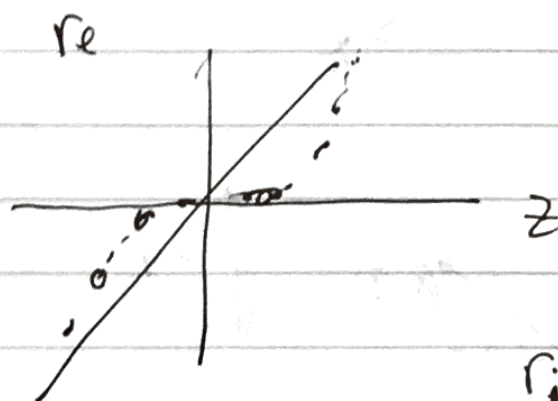
Not uncorrelated



heteroskedastic



Checking normality
Imagine the residual distrib.



normal
probability
plot

Matlab function
norm plot

$$r_i \text{ residual} \\ z_i = \Phi^{-1}\left(\frac{i-a}{N+1-2a}\right)$$

approx $\frac{i}{N}$ $\Phi^{-1}\left(\frac{i}{N}\right)$
gives the z value with
 $P(Z \leq z_i) = \frac{i}{N}$

a corrects
for finite size
effects

$a = \frac{3}{8}$ for small N ($N \leq 10$)

$a = .5$ for large N ($N > 10$)

When the normal error model
is right, one could provide
confidence limits for the regression
coefficients w_a 's as well for predicted

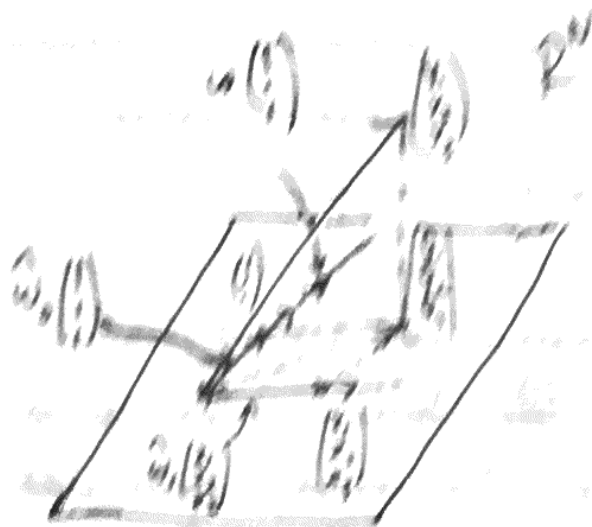
values of y . One could also do hypothesis tests.

We will discuss an example that it related to the ~~issue~~ issue of complexity, but is often posed as a test for ~~an~~ explanatory relevance ~~of~~ of feature.

Consider these nested models

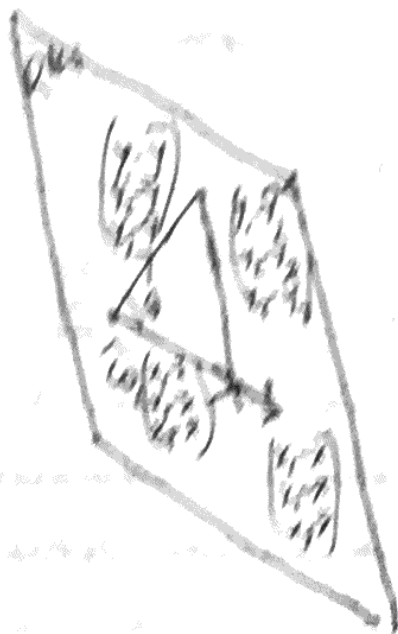
$$\begin{aligned} y_i &= \omega_0 + \varepsilon_i \\ y_i &= \omega_0 + \omega_1 x_i + \varepsilon_i \end{aligned} \quad \left. \begin{array}{l} \nearrow \\ \nearrow \end{array} \right\} \omega_1 = 0$$

If we do two parameter ~~least~~ least square, we will ~~usually~~ usually find ~~a~~ a nonzero $\hat{\omega}_1$. One could ask whether one should just stick to the simpler model with $\omega_1 = 0$.



How do I test
anything about
 w_1

Project to the subspace perpendicular
to $\begin{pmatrix} 1 \\ 1 \end{pmatrix}$



$$wt0 = \frac{\hat{w}_1 \sqrt{\sum_i (x_i - \bar{x})^2}}{\sqrt{\sum_i (y_i - \bar{y})^2}}$$

$$t = \frac{\hat{w}_1 \sqrt{\sum_i (x_i - \bar{x})^2}}{\sqrt{\sum_i \frac{(y_i - \bar{y})^2}{N-2}}} = \frac{\hat{w}_1 \sqrt{\sum_i \frac{(x_i - \bar{x})^2}{N}}}{\sqrt{\sum_i \frac{(y_i - \bar{y})^2}{N}}} \sqrt{N-2}$$

$$= \frac{\hat{w}_1 s_x}{s_{y.x}} \sqrt{N-2} \quad \text{d.f. } N-2$$

Note that $\sum_i (y_i - \bar{y})^2 = \sum_i (y_i - \hat{y}_i + \hat{y}_i - \bar{y})^2$
 $= \sum_i (y_i - \hat{y}_i)^2 + 2 \sum_i (y_i - \hat{y}_i)(\hat{y}_i - \bar{y}) + \sum_i (\hat{y}_i - \bar{y})^2$

$$\begin{aligned} \hat{y}_i - \bar{y} &= \hat{\omega}_0 + \hat{\omega}_1 x_i - (\hat{\omega}_0 + \hat{\omega}_1 \bar{x}) \\ &= \hat{\omega}_1 (x_i - \bar{x}) \end{aligned}$$

$$\sum_i (y_i - \hat{y}_i)(x_i - \bar{x}) = 0$$

So $\downarrow$ Total ~~variance~~ $\sum_i (y_i - \bar{y})^2$ = Unexplained var $\sum_i (y_i - \hat{y}_i)^2$ + Explained variance $\sum_i (\hat{y}_i - \bar{y})^2$

~~RSS~~ = ~~(RSS)~~ + ESS $\hat{\omega}_1^2 \sum_i (x_i - \bar{x})^2$

$$So \quad t^2 = \frac{\sum_i (\hat{y}_i - \bar{y})^2 / 1}{\sum_i (y_i - \hat{y}_i)^2 / (N-2)} = F \text{ statistic with dfs } 1, N-2$$

In general, if we went from a p_0 parameters to p_1 parameters and ~~RSS~~ reduced to ~~RSS~~, then

$$F = \frac{(\text{RSS}_0 - \text{RSS}_1) / (p_1 - p_0)}{\text{RSS}_1 / (N - p_1 - 1)}$$

would be distributed as $F_{p-p, N-p-1}$ if the smaller model is sufficient and the larger model is only fitting noise.

For polynomial fits, there is a natural structure: ~~the~~ degree of the polynomial. Many multivariate linear regression problems, it is not obvious what to fit first and what next.

Many different methods has been proposed how to choose a subset of variables to fit. In general they run into the trouble that there are many such subsets.

Stepwise Forward method chooses the variable that explains the most variance significantly. It then considers adding one more etc etc. They all have essentially some multitesting issues

Alternative $\rightarrow$ ~~Shrinkage~~ Shrinkage based methods.

Ridge regression

Prior $p(w) = \prod_j \mathcal{N}(w_j | 0, \tau^2)$

Now the MAP estimation of w

$$\arg \max_w \sum_{i=1}^N \log \mathcal{N}(y_i | w_0 + w^T x_i, \sigma^2) + \sum_{j=1}^D \log \mathcal{N}(w_j | 0, \tau^2)$$

Equivalent to minimizing

$$J(w) = \sum_{i=1}^N (y_i - (w_0 + w^T x_i))^2 + \lambda \|w\|_2^2$$

where $\lambda = \frac{\sigma^2}{\tau^2}$ and $\|w\|_2^2 = \sum_{j=1}^D w_j^2 = w^T w$

The parameter

w_0 is there to handle $\bar{y}$ and is not regularized

Ridge regression or penalized least squares

The intuition is that w_i would be small if it does not improve the ~~loss~~ square error significantly
 w_0 derivative set to zero

$$\sum_i [y_i - (w_0 + w^T x_i)] = 0$$

$$w_0 = \bar{y} - w^T \bar{x}$$

Let us center both y_i and x_i

$$y_i - \bar{y} \rightarrow y_i$$

$$x_i - \bar{x} \rightarrow x_i$$

$$\sum_i (y_i - w^T x_i)^2 + \lambda w^T w$$

$$= (y - Xw)^T (y - Xw) + w^T w$$

(w deriv)

$$2X^T(X\hat{w} - y) + 2\lambda \hat{w} = 0$$

$$(\lambda I_D + X^T X) \hat{w} = X^T y$$

$$\hat{w}_{ridge} = (\lambda I_D + X^T X)^{-1} X^T y$$

Relationship to PCA, SVD etc.

A quick introduction to singular value decomposition

$$X = U S V^T$$

$N \times D$ $N \times N$ $N \times D$ $D \times D$

U, V unitary (orthogonal for real X) $\Rightarrow U^T U = I_N$ $V^T V = V V^T = I_D$

$$S = \begin{pmatrix} \sigma_1 & & & \\ & \ddots & & \\ & & \sigma_D & \\ & & & 0 \end{pmatrix}$$

$\begin{matrix} \uparrow D \\ \downarrow N-D \end{matrix}$ $\begin{matrix} \leftarrow D \\ \rightarrow N-D \end{matrix}$

(We assume $N > D$)

$$\begin{pmatrix} \begin{matrix} \text{D columns} \\ \text{N-D columns} \end{matrix} \end{pmatrix} \begin{pmatrix} \sigma_1 & & \\ & \ddots & \\ & & \sigma_D \\ & & & 0 \end{pmatrix} \begin{pmatrix} \text{D rows} \\ \text{N-D rows} \end{pmatrix}$$

u_1, u_2 S v_1^T, v_2^T

$$X = \sum_{i=1}^D \sigma_i u_i v_i^T$$

You can think of it this way

~~The~~ The Gramians $X^T X$ and $X X^T$ are symmetric matrices. They can be diagonalized $X^T X = V \Lambda V^T$ and $X X^T = U \tilde{\Lambda} U^T$.
Since $\text{Tr}(X^T X) = \text{Tr}(X X^T)$

~~They share~~ $\sum_{i=1}^D \lambda_i = \sum_{j=1}^N \tilde{\lambda}_j$

Only way to satisfy this $\{\tilde{\lambda}_j\}_{j=1}^N$

$= \{\lambda_1, \dots, \lambda_D, 0, 0, \dots, 0\}$
 $\leftarrow N-D \rightarrow$

~~$\lambda_i \geq 0$~~ $\lambda_i \geq 0$ since $X^T X$ is positive semidefinite. Choose $\sigma_i = \sqrt{\lambda_i}$

Only ambiguity in the choice of U, V

is $V \Rightarrow \tilde{V} P_R$ where $P_R = \begin{pmatrix} \pm 1 & & \\ & \pm 1 & \\ & & \dots \end{pmatrix}$

$U \Rightarrow \tilde{U} P_L$ $P_L = \begin{pmatrix} \pm 1 & & \\ & \pm 1 & \\ & & \dots \end{pmatrix}$

Consider the matrix $\tilde{V}$ ~~$P_R V$~~ and its rows. $\begin{pmatrix} \tilde{v}_1^T \\ \tilde{v}_2^T \\ \vdots \end{pmatrix}$

and the matrix $\tilde{U}$ and its columns $(\tilde{u}_1, \tilde{u}_2, \dots)$

Note that $(X^T X) \tilde{v}_j = \lambda_j \tilde{v}_j$

Then $X X^T (X \tilde{v}_j) = \lambda_j X \tilde{v}_j$

But $X X^T \tilde{u}_j = \lambda_j \tilde{u}_j$

So $X \tilde{v}_j = \tilde{u}_j \Rightarrow \tilde{v}_j^T X^T X \tilde{v}_j = \tilde{u}_j^T \tilde{u}_j$
But then $\tilde{u}_j^T \tilde{u}_j = \lambda_j \Rightarrow \tilde{u}_j = \pm \tilde{v}_j$

Using the ± 1 in P_L, P_R we ~~can~~ could
define u_j, v_j st.

$X v_j = u_j$ etc...

$$X = U S V^T$$

Now, for centered Y, X in ridge regression

$$\hat{w}_{ridge} = (\lambda I_D + X^T X)^{-1} X^T Y$$

using SVD representation of X .

$$\hat{w}_{ridge} = V (\lambda I_D + \underbrace{S^T S}_{I_D})^{-1} V^T V S U^T Y = V (S^2 + \lambda I)^{-1} S U^T Y$$

Ridge regression predictions are

$$\hat{y} = X \hat{w}_{\text{ridge}} = U S U^T v (S^2 + \lambda I)^{-1} S U^T y \\ = \sum_{j=1}^D u_j \tilde{s}_{jj} u_j^T y$$

$$\text{Where } \tilde{s}_{jj} = \frac{\sigma_j^2}{\sigma_j^2 + \lambda}$$

$$\text{So } \hat{y} = \sum_{j=1}^D u_j \frac{\sigma_j^2}{\sigma_j^2 + \lambda} u_j^T y$$

The corresponding least square answer is

$$\hat{y} = \sum_{j=1}^D u_j u_j^T y$$

For directions u_j where $\sigma_j^2 \ll \lambda$, u_j would not have much of an effect on prediction in ridge regression as opposed to pure least square.

One might consider $\sum \frac{\sigma_i^2}{\sigma_i^2 + \lambda}$ to be an effective degree of freedom here.

Why is this a good thing?

Well $\log \text{Likelihood} \sim -\frac{1}{2\sigma^2} (Xw - y)^T (Xw - y)$

$\Rightarrow$ Fisher Information $= \frac{1}{\sigma^2} X^T X$

$\Rightarrow \text{Cov}(\hat{w}) \sim \sigma^2 (X^T X)^{-1}$

That means $\hat{w}$ is most ambiguous in directions ~~co~~ which are given by eigenvectors of $X^T X$ corresponding to lowest eigenvalues $\lambda_i = \sigma_i^2$. ~~The~~ The factor $\frac{\sigma_i^2}{\sigma_i^2 + \lambda}$

suppresses these directions, preventing effect of noise from influencing $\hat{w}$ and subsequently the prediction.

Applied to the familiar territory of single variable linear regression,

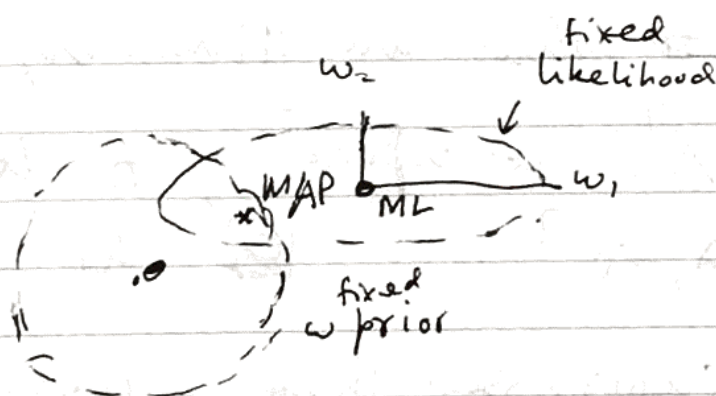
$$\underset{w_1}{\operatorname{argmin}} \left[\sum_i \{y_i - \bar{y} - w_1(x_i - \bar{x})\}^2 + \lambda w_1^2 \right]$$

$$(\hat{w}_1)_{\text{ridge}} = \frac{\sum_i (x_i - \bar{x})(y_i - \bar{y})}{\sum_i (x_i - \bar{x})^2 + \lambda} = \frac{r_{xy} s_x s_y}{s_x^2 + \lambda/N}$$

$$= (\hat{w}_1)_{\text{OLS}} \frac{NS_x^2}{NS_x^2 + \lambda}$$

$\hat{w}_1$ is 'shrunk'

Two variables

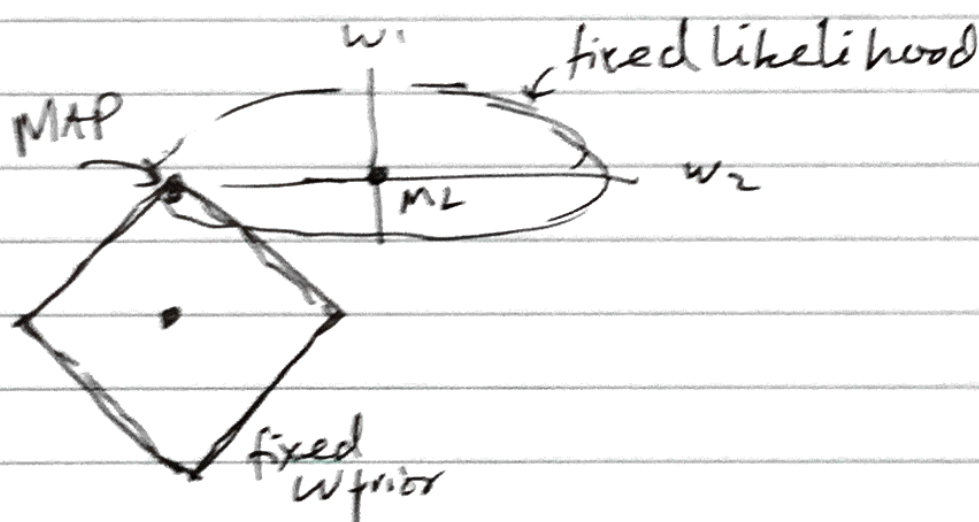


$\hat{w}_1$, which is more error-prone, is shrunk more than, $\hat{w}_2$.

If we really want some w_j 's set to zero, one could use the Laplace prior

$$P(w) \propto e^{-\lambda \|w\|_1} = e^{-\lambda \sum_{j=1}^D |w_j|}$$

$\downarrow$ l_1 norm



The l_1 penalty could lead to sparse w 's : many $w_j = 0$.

$\Rightarrow$ LASSO (least absolute shrinkage and selection operator).

Both ~~ridge~~ ridge regression & LASSO are convex optimization problems.