

Solution 1

(a) Writing $\hat{\mu} - \mu = H_\lambda y - \mu = H_\lambda(y - \mu) + (H_\lambda - I)\mu$ gives

$$(\hat{\mu} - \mu)^T(\hat{\mu} - \mu) = (y - \mu)^T H_\lambda^T H_\lambda (y - \mu) + 2(y - \mu)^T H_\lambda^T (H_\lambda - I)\mu + \mu^T (H_\lambda - I)^T (H_\lambda - I)\mu.$$

The expected value of the second term is zero, because $E(y - \mu) = 0$, and the final term is the constant $\|(I - H_\lambda)\mu\|_2^2$, while the first term has expectation

$$E[\text{tr}\{(y - \mu)^T H_\lambda^T H_\lambda (y - \mu)\}] = \text{tr}[E\{(y - \mu)(y - \mu)^T H_\lambda^T H_\lambda\}] = \text{tr}(\sigma^2 I_n H_\lambda^T H_\lambda),$$

which gives the result.

(b) As $E(yy^T) = \text{var}(y) + E(y)E(y)^T$ we have

$$\begin{aligned} E[(y - \hat{\mu})^T(y - \hat{\mu})] &= E[y^T(I - H_\lambda)^T(I - H_\lambda)y] \\ &= E[\text{tr}\{(I - H_\lambda)^T(I - H_\lambda)yy^T\}] \\ &= \text{tr}\{(I - H_\lambda)^T(I - H_\lambda)(\mu\mu^T + \sigma^2 I_n)\} \\ &= \mu^T(I - H_\lambda)^T(I - H_\lambda)\mu + \sigma^2(n - 2\nu_1 + \nu_2) \\ &= \|(I - H_\lambda)\mu\|_2^2 + \sigma^2(n - 2\nu_1 + \nu_2), \end{aligned}$$

so $\hat{\sigma}_\lambda^2$ is biased upwards unless $\|(I - H_\lambda)\mu\|_2 = 0$, i.e., unless μ lies in the kernel of $I - H_\lambda$.

In a standard setting H_λ is a projection matrix, so $H_\lambda^T H_\lambda = H_\lambda H_\lambda = H_\lambda$, and thus $\nu_1 = \nu_2 = \text{rank}(H_\lambda)$, so $\hat{\sigma}_\lambda^2$ becomes the usual unbiased variance estimator for a linear model.

Solution 2

(a) We have $X^T X = I_p$ and hence $\hat{\beta} = (X^T X)^{-1} X^T y = X^T y$. The sum of squares is

$$(y - X\beta)^T(y - X\beta) = y^T y - 2y^T X\beta + \beta^T X^T X\beta = y^T y - 2\hat{\beta}^T \beta + \beta^T \beta,$$

so the function to be minimised is

$$L = \frac{1}{2} (y^T y - 2\hat{\beta}^T \beta + \beta^T \beta) + \lambda \sum_{r=1}^p |\beta_r| \equiv \sum_{r=1}^p (\beta_r^2/2 - \hat{\beta}_r \beta_r + \lambda |\beta_r|).$$

This is a sum of p separate functions, each of which is a sum of the two convex functions of the forms $x^2 - ax$ and $b|x|$ for $b > 0$, so each is convex. Each of the p individual summands can be minimised individually.

(b) The function L is differentiable in β_r except at $\beta_r = 0$, and the minimum is either at $\beta_r = 0$ or elsewhere. Now

$$\partial L / \partial \beta_r = \beta_r - \hat{\beta}_r + \lambda \text{sign}(\beta_r),$$

so

$$\lim_{\beta_r \rightarrow 0+} \partial L / \partial \beta_r = \lambda - \hat{\beta}_r, \quad \lim_{\beta_r \rightarrow 0-} \partial L / \partial \beta_r = -\lambda - \hat{\beta}_r.$$

- For a minimum at $\beta_r = 0$ we must have $\lambda - \hat{\beta}_r > 0$ and $-\lambda - \hat{\beta}_r < 0$, or equivalently $|\hat{\beta}_r| < \lambda$. In this case $\tilde{\beta}_r = 0$.
- For a minimum not at $\beta_r = 0$, setting $\partial L / \partial \beta_r = 0$ gives

$$\tilde{\beta}_r = \hat{\beta}_r - \lambda \operatorname{sign}(\tilde{\beta}_r),$$

so if $\tilde{\beta}_r > 0$, then $\tilde{\beta}_r = \hat{\beta}_r - \lambda$, whereas if $\tilde{\beta}_r < 0$, then $\tilde{\beta}_r = \hat{\beta}_r + \lambda$.

Putting these two cases together gives

$$\tilde{\beta}_r = \operatorname{sign}(\hat{\beta}_r)(|\hat{\beta}_r| - \lambda)I(|\hat{\beta}_r| > \lambda),$$

as required.