

Regularization methods in multiple regression

Malgorzata Bogdan

University of Wrocław

May 12, 2020

High dimensional regression

Ridge regression (1)

$$Y_{nx1} = X_{n \times p} \beta_{p \times 1} + z_{nx1}, \quad z \sim N(0, \sigma^2 I)$$

$Y = (Y_1, \dots, Y_n)^T$ - wektor of trait values for n individuals

$X_{n \times p}$ - matrix of regressors

When $n > p$ but p is large (say $n/2$) the variance of LS estimates may be very large

When $p > n$ the matrix $X'X$ is singular and the LS estimate of β does not exist

Ridge regression:

$$\hat{\beta} = \operatorname{argmin}_{\beta \in \mathbb{R}^p} L(\beta), \text{ where } L(\beta) = \|Y - X\beta\|^2 + \gamma \| \beta \|^2$$

$$\frac{\partial L(\beta)}{\partial \beta} = -2X'(Y - X\beta) + 2\gamma\beta = 0$$

$$-X'Y + (X'X + \gamma I)\beta = 0 \Leftrightarrow \beta = (X'X + \gamma I)^{-1}X'Y$$

$$\hat{\beta} = (X'X + \gamma I)^{-1} X'Y, \text{ where } \gamma > 0$$

$$X'Xu = \lambda u$$

$$\hat{Y} = X\hat{\beta} = MY, \text{ with } M = X(X'X + \gamma I)^{-1}X'$$

$$(X'X + \gamma I)u = (\lambda + \gamma)u, \quad (X'X + \gamma I)^{-1}u = \frac{1}{\lambda + \gamma}u$$

$$Tr[M] = Tr[(X'X + \gamma I)^{-1}X'X]$$

$$(X'X + \gamma I)^{-1}X'Xu = \frac{\lambda}{\lambda + \gamma}u, \quad Tr(M) = \sum_{i=1}^n \frac{\lambda_i(X'X)}{\lambda_i(X'X) + \gamma}$$

$$Tr[M] = \sum_{i=1}^p \lambda_i(M), \text{ where } \lambda_1(M), \dots, \lambda_n(M) \text{ are eigenvalues of } M$$

$$\hat{P}E = RSS + 2\sigma^2 \sum_{i=1}^p \frac{\lambda_i(X'X)}{\lambda_i(X'X) + \gamma}$$

$$X'X = I, \quad \hat{\beta} = \frac{1}{1 + \gamma} X'Y = \frac{1}{1 + \gamma} (\beta + X'\epsilon)$$

When ridge is better than LS ?

$$Z = X'\epsilon \sim N(0, \sigma^2 I)$$

$$\frac{\gamma^2 \|\beta\|^2 + p\sigma^2}{(1 + \gamma)^2} < p\sigma^2$$

$$E(\hat{\beta}_i - \beta_i)^2 = E\left(\frac{1}{1 + \gamma}\beta_i - \beta_i + \frac{1}{1 + \gamma}Z_i\right)^2$$

Ridge is always better than LS when $\|\beta\|^2 < p\sigma^2$
Otherwise, when

$$= \frac{\gamma^2}{(1 + \gamma)^2} \beta_i^2 + \frac{\sigma^2}{(1 + \gamma)^2}$$

$$\|\beta\|^2 < \frac{\gamma + 2}{\gamma} p\sigma^2$$

$$E\|\hat{\beta} - \beta\|^2 = \frac{\gamma^2}{(1 + \gamma)^2} \|\beta\|^2 + \frac{p\sigma^2}{(1 + \gamma)^2}$$

$$\gamma < \frac{2p\sigma^2}{\|\beta\|^2 - p\sigma^2}$$

$$Y = X\beta$$

Basis Pursuit (Chen and Donoho, 1994): when $p > n$ recover β by minimizing $\|b\|_1 = \sum_{i=1}^n |b_i|$ subject to $Y = Xb$.

BP can recover β if it is *identifiable* with respect to L_1 norm, i.e.

$$\text{If } X\gamma = X\beta \text{ and } \gamma \neq \beta \text{ then } \|\gamma\|_1 > \|\beta\|_1.$$

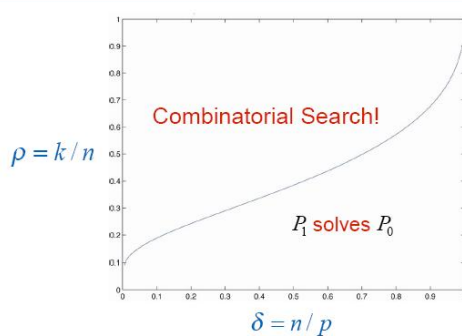
$$k = \|\beta\|_0 = \#\{i : \beta_i \neq 0\}$$

Basis Pursuit can recover β if k is small enough.

Let's assume that $p \rightarrow \infty$, $n/p \rightarrow \delta$ and $k/n \rightarrow \epsilon$.

If X_{ij} are iid $N(0, \tau^2)$ then the probability that BP recovers β converges to 1 if $\epsilon < \rho(\delta)$ and to 0 if $\epsilon > \rho(\delta)$, where $\rho(\delta)$ is the *transition curve*.

Phase Transition: (l_1, l_0) equivalence



$$Y_{n \times 1} = X_{n \times p} \beta_{p \times 1} + z_{n \times 1}, \quad z \sim N(0, \sigma I)$$

Convex program: Minimize $\|b\|_1$ subject to $\|Y - Xb\|_2^2 \leq \epsilon$

Or alternatively: $\min_{b \in \mathbb{R}^p} \|y - Xb\|_2^2 + \lambda \|b\|_1$

BPDN (Chen and Donoho, 1994) or LASSO (Tibshirani, 1996)

- General rule: the reduction of λ_L results in identification of more elements from the true support (true discoveries) but at the same time it produces more falsely identified variables (false discoveries)
- The choice of λ_L is challenging- e.g. crossvalidation typically leads to many false discoveries
- When $X^T X = I$ Lasso selects X_j iff $|\hat{\beta}_j^{LS}| > \lambda$
- Selection $\lambda = \sigma \Phi^{-1}(1 - \alpha/(2p)) \approx \sigma \sqrt{2 \log p}$ corresponds to Bonferroni correction and controls FWER.

The sign vector of β is defined as

$$S(\beta) = (S(\beta_1), \dots, S(\beta_p)) \in \{-1, 0, 1\}^p,$$

where for $x \in \mathbb{R}$, $S(x) = \mathbf{1}_{x>0} - \mathbf{1}_{x<0}$

Let $I := \{i \in \{1, \dots, p\} \mid \beta_i \neq 0\}$, and let X_I, X_{I^c} be matrices whose columns are respectively $(X_i)_{i \in I}$ and $(X_i)_{i \notin I}$.

Irrepresentable condition:

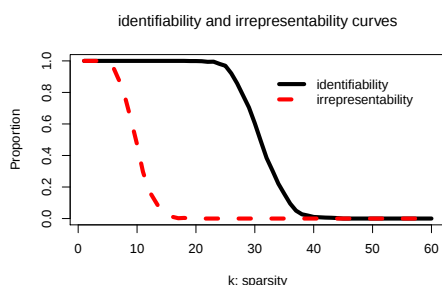
$$\|X_{I^c}^T X_I (X_I^T X_I)^{-1} S(\beta_I)\|_\infty \leq 1$$

When

$$\|X_{I^c}^T X_I (X_I^T X_I)^{-1} S(\beta_I)\|_\infty > 1$$

then probability of the support recovery by LASSO is smaller than 0.5 (Wainwright, 2009).

$n=100$, $p=300$, elements of X were generated as iid $N(0,1)$



Definition (Identifiability)

Let X be a $n \times p$ matrix. The vector $\beta \in R^p$ is said to be identifiable with respect to the l^1 norm if the following implication holds

$$X\gamma = X\beta \text{ and } \gamma \neq \beta \Rightarrow \|\gamma\|_1 > \|\beta\|_1. \quad (1)$$

Theorem (Tardivel, Bogdan, 2019)

For any $\lambda > 0$ LASSO can separate well the causal and null features if and only if vector β is identifiable with respect to l_1 norm and $\min_{i \in I} |\beta_i|$ is sufficiently large.

Corollary

Appropriately thresholded LASSO can properly identify the sign of sufficiently large β if and only if β is identifiable with respect to l_1 norm.

Conjecture

Adaptive (reweighted) LASSO can properly identify the sign of sufficiently large β if and only if β is identifiable with respect to l_1 norm.

Intuitive explanation:

$$\hat{\beta} = \eta_{\lambda}(\beta_i + X_i'z + v_i)$$

$$v_i = \langle X_i, \sum_{j \neq i} X_j(\beta_j - \hat{\beta}_j) \rangle$$

$$\eta_{\lambda}(t) = \text{sign}(t)(|t| - \lambda)_+, \quad \text{applied componentwise}$$

If $X^T X = I$ then $X_i'z = Z_i \sim N(0, 1)$, $v_i = 0$ and H_{0i} is rejected if $\beta_i + Z_i > \lambda$

When the design is not orthogonal: $v_i \neq 0$ - additional noise, dependent on λ (level of shrinkage), the level of sparsity and magnitude of true signals

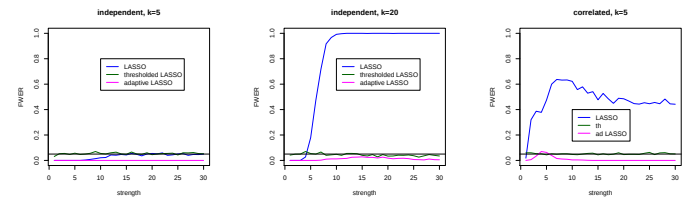
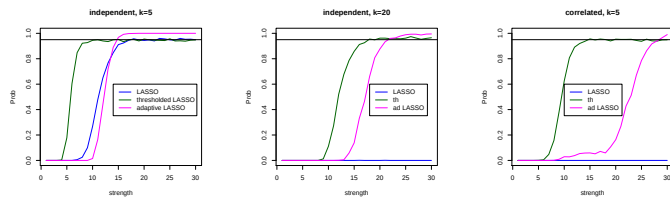
Adaptive LASSO [Zou, JASA 2006], [Candès, Wakin and Boyd, J. Fourier Anal. Appl. 2008]

$$\beta_{aL} = \underset{b}{\operatorname{argmin}} \left\{ \frac{1}{2} \|y - Xb\|_2^2 + \lambda \sum_{i=1}^P w_i |b|_i \right\}, \quad (2)$$

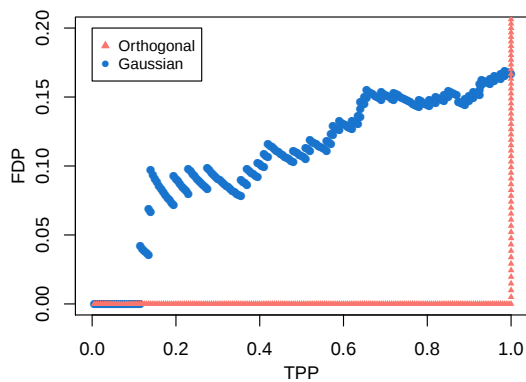
where $w_i = \frac{1}{\hat{\beta}_i}$, and $\hat{\beta}_i$ is some consistent estimator of β_i .

Reduces bias and improves model selection properties

1. λ for LASSO selected as to control FWER at the level 0.05 for $k = 5$ (theoretical result in (Tardivel and Bogdan, 2019))
2. λ_{AMP} for thresholded LASSO and independent gaussian design selected according to AMP theory for LASSO (see e.g. (Wang, Weng, Maleki, 2018))
3. For correlated design (off diagonal covariance 0.9) we used 0.5 λ_{AMP}
4. For adaptive LASSO - weights based on LASSO estimator with λ as in 2 and 3, selection based on LASSO with λ as in 1
5. Threshold selected by using knockoff control variables (Foygel-Barber and Candès, 2015; Candès, Fan, Janson, Lv, 2016)



Su, Bogdan and Candes, (2017), $\delta = 1$, $\epsilon = 0.2$



LASSO solution

$$\hat{\beta} = \eta_{\lambda}(\hat{\beta} - X'(X\hat{\beta} - y)) = \eta_{\lambda}(\hat{\beta} - X'X(\hat{\beta} - \beta) + X'z),$$

where $\eta_{\lambda}(t) = \text{sgn}(t)(|t| - \lambda)_+$, applied componentwise

$$\hat{\beta}_i = \eta_{\lambda}(\beta_i + Z_i + v_i),$$

$$\text{where } v_i = \langle X_i, \sum_{j \neq i} X_j(\beta_j - \hat{\beta}_j) \rangle \text{ and } Z_i \sim N(0, \sigma_i^2)$$

$$X_{ij} \sim \mathcal{N}(0, 1/n), \quad z_i \sim \mathcal{N}(0, \sigma^2)$$

$\beta_1, \dots, \beta_p$: iid, distributed as the random variable Π , such that $\mathbb{E} \Pi < \infty$, $\mathbb{P}(\Pi \neq 0) = \epsilon \in (0, 1)$.

$$\tau^2 = \sigma^2 + \frac{1}{\delta} \mathbb{E} \left(\eta_{\alpha\tau}(\Pi + \tau Z) - \Pi \right)^2,$$

$$\lambda = \left(1 - \frac{1}{\delta} \mathbb{P}(|\Pi + \tau Z| > \alpha\tau) \right) \alpha\tau.$$

Theorem

For any pseudo-Lipschitz function φ , the lasso solution $\hat{\beta}$ with fixed λ obeys

$$\frac{1}{p} \sum_{i=1}^p \varphi(\hat{\beta}_i, \beta_i) \rightarrow \mathbb{E} \varphi(\eta_{\alpha\tau}(\Pi + \tau Z), \Pi)$$

$\hat{\mathcal{S}}$ - set of variables selected by LASSO

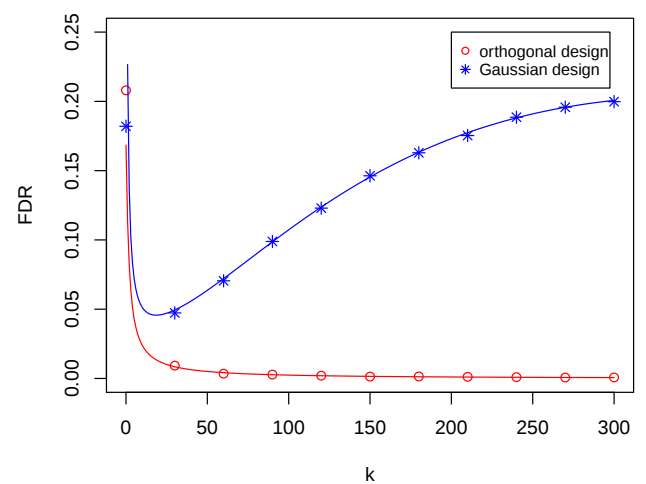
$$FDP \equiv \frac{|\hat{\mathcal{S}} \cap \mathcal{H}_0|}{|\hat{\mathcal{S}}|}$$

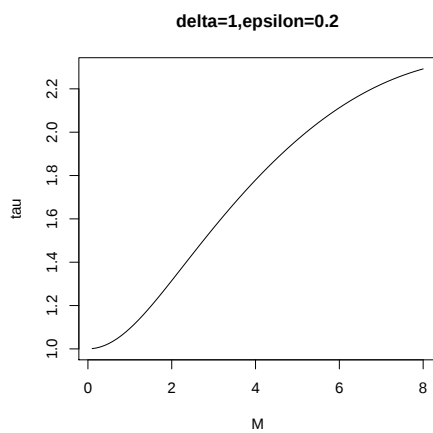
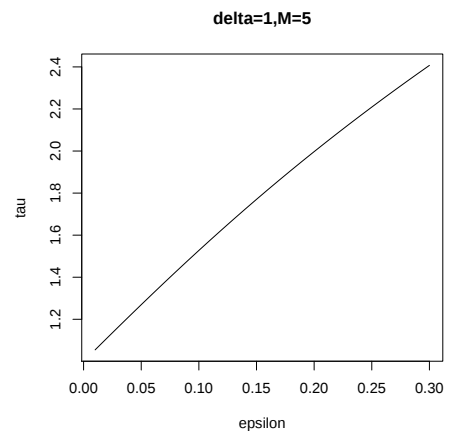
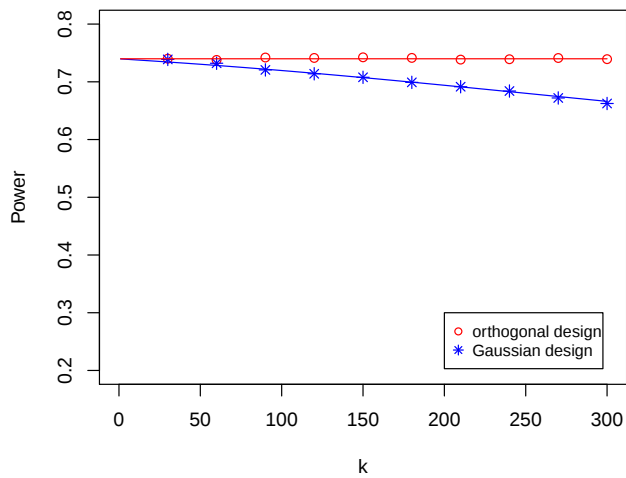
$$FDR = E(FDP)$$

Bogdan, van den Berg, Su and Candés, 2013

$$FDR \rightarrow \frac{2\mathbb{P}(\Pi = 0)\Phi(-\alpha)}{\mathbb{P}(|\Pi + \tau Z| > \alpha\tau)},$$

$$\text{Power} \rightarrow \mathbb{P}(|\Pi + \tau Z| > \alpha\tau | \Pi \neq 0).$$





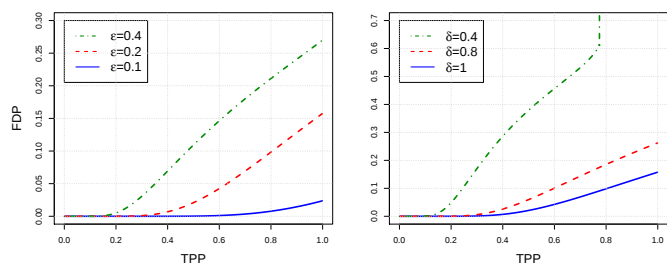
Theorem (Su, Bogdan, Candes, 2017)

Fix $\delta \in (0, \infty)$ and $\epsilon \in (0, 1)$. Then the event

$$\bigcap_{\lambda \geq 0.01} \left\{ \text{FDP}(\lambda) \geq q^*(\text{TPP}(\lambda)) - 0.001 \right\} \quad (3)$$

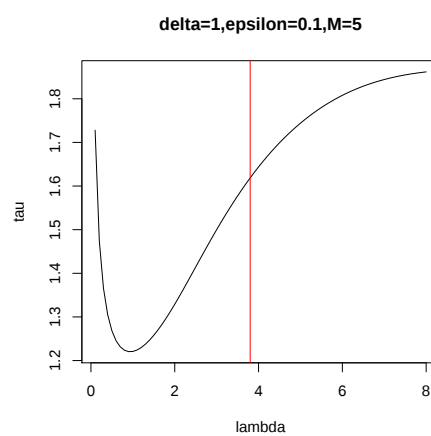
holds with probability tending to one.

FDR-Power trade-off (2)



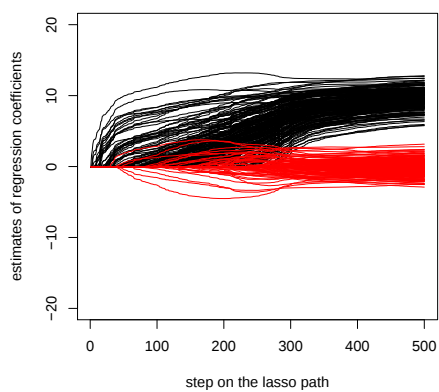
Malgorzata Bogdan Regularization

Magnitude of noise



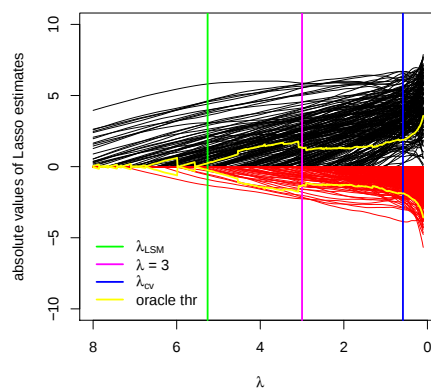
Malgorzata Bogdan Regularization

Thresholded LASSO (1)

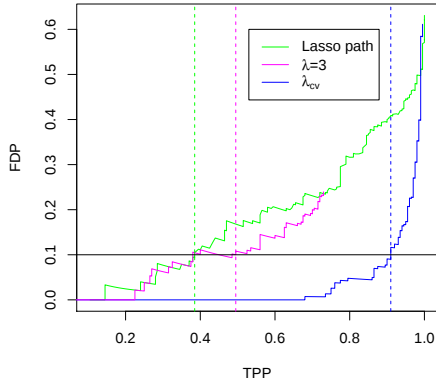


Malgorzata Bogdan Regularization

Thresholded LASSO (2)



Malgorzata Bogdan Regularization



Candès, Fan, Janson and Lv (2017) - model free knockoffs

Consider n i.i.d random vectors $(Y_i, X_{1i}, \dots, X_{pi})$

Variable X_j is a null variable, if $Y \perp X_j \mid X_{-j}$, where X_{-j} denotes the remaining $p - 1$ variables excluding X_j

Construct a set of “fake” covariates $\tilde{X} = (\tilde{X}_1, \tilde{X}_2, \dots, \tilde{X}_p)$ which satisfy:

- 1 **Exchangeability:** for any subset $S \subset \{1, \dots, p\}$,

$$(X, \tilde{X})_{\text{swap}(S)} \stackrel{d}{=} (X, \tilde{X}), \quad (4)$$

where $\text{swap}(S)$ is obtained by swapping the entries X_j and $\tilde{X}_j$ for each $j \in S$.

- 2 **Unimportant variables:** $\tilde{X} \perp Y \mid X$, which can be guaranteed if $\tilde{X}$ is constructed without looking at Y .

If $X \sim \mathcal{N}(0, \Sigma)$, then a joint distribution of $(X, \tilde{X})$ can be:

$$(X, \tilde{X}) \sim \mathcal{N}(0, G), \quad \text{where } G = \begin{bmatrix} \Sigma & \Sigma - \text{diag}(s) \\ \Sigma - \text{diag}(s) & \Sigma \end{bmatrix}, \quad (5)$$

with any choice of the diagonal matrix $\text{diag}(s)$ s.t. G is positive semidefinite. Possible choice of s :

$$s_j = 2\lambda_{\min}(\Sigma) \wedge 1, \quad \forall j, \quad (6)$$

Sample $\tilde{X}$ from its conditional distribution:

$$\tilde{X} | X \stackrel{d}{=} \mathcal{N}(\mu_c, \Sigma_c),$$

where

$$\begin{aligned} \mu_c &= X - X\Sigma^{-1} \text{diag}(s) \\ \Sigma_c &= 2\text{diag}(s) - \text{diag}(s)\Sigma^{-1}\text{diag}(s). \end{aligned} \quad (7)$$

$$W = (W_1, \dots, W_p), \quad W_j = w_j([X, \tilde{X}], Y)$$

flip-sign property:

$$w_j([X, \tilde{X}]_{\text{swap}(S)}, Y) = \begin{cases} w_j([X, \tilde{X}], Y), & j \notin S \\ -w_j([X, \tilde{X}], Y), & j \in S \end{cases}$$

Example:

$$T = (Z, \tilde{Z}) = t([X, \tilde{X}], Y)$$

$$(Z, \tilde{Z})_{\text{swap}(S)} = t([X, \tilde{X}]_{\text{swap}(S)}, Y)$$

$$w_j = f_j(Z_j, \tilde{Z}_j), \quad f(v, u) = -f(u, v)$$

Lasso Coefficient Difference (LCD) statistic:

$$W_j = |\hat{\beta}_j(\lambda)| - |\hat{\beta}_{j+p}(\lambda)|$$

Define a random threshold as

$$\hat{t}(\lambda) = \min \left\{ t > 0 : \frac{1 + \#\{j : W_j(\lambda) \leq -t\}}{\#\{j : W_j(\lambda) \geq t\}} \leq q \right\}$$

and select

$$\widehat{S(\lambda)} = \{j : W_j(\lambda) \geq \hat{t}(\lambda)\},$$

Candès, Fan, Janson and Lv (2017) - The above knockoff procedure $KN(\lambda, q)$ controls FDR at the level q .

Su, Weinstein, Bogdan, Candès (2019)

Definition

A random variable Π is said to be ϵ -sparse if $\mathbb{E} \Pi^2 < \infty$ and $\mathbb{P}(\Pi \neq 0) = \epsilon$.

Theorem

Assume that $\epsilon/2 < \epsilon_{DT}(\delta/2)$, where $\epsilon_{DT}(\delta)$ is a point on the Donoho-Tanner transition curve. Then for any fixed $0 < \lambda_1 < \lambda_2, 0 < q < 1$, and any $\nu > 0$, there exists an ϵ -sparse prior Π and n' such that

$$\mathbb{P} \left(\inf_{\lambda_1 \leq \lambda \leq \lambda_2} \text{TPP}(\lambda, \Pi, q, n, p) > 1 - \nu \right) \geq 1 - \nu$$

if $n \geq n'$.

We show that the triples $(\beta_j, \hat{\beta}_j, \hat{\beta}_{p+j})$ are independent, and each is distributed as $(\Pi, \eta_{\alpha\tau}(\Pi + \tau W), \eta_{\alpha\tau}(\tau \tilde{W}))$, where W and $\tilde{W}$ are independent $\mathcal{N}(0, 1)$ random variables that are furthermore independent of Π ; and (α, τ) are determined by λ as the solution to

$$\begin{aligned} \tau^2 &= \sigma^2 + \frac{1}{\delta} \mathbb{E} [\eta_{\alpha\tau}(\Pi + \tau W) - \Pi]^2 + \frac{1}{\delta} \mathbb{E} \eta_{\alpha\tau}(\tau W)^2 \\ \lambda &= \left[1 - \frac{1}{\delta} \mathbb{P}(|\Pi + \tau W| > \alpha\tau) - \frac{1}{\delta} \mathbb{P}(|\tau W| > \alpha\tau) \right] \alpha\tau. \end{aligned} \quad (8)$$

For fixed $\lambda > 0$, let $t^\infty = t^\infty(q) > 0$ be such that

$$\frac{\mathbb{P}(\omega(\eta_{\alpha\tau}(\Pi + \tau W), \tau \eta_\alpha(\tilde{W})) \leq -t^\infty)}{\mathbb{P}(\omega(\eta_{\alpha\tau}(\Pi + \tau W), \tau \eta_\alpha(\tilde{W})) \geq t^\infty)} = q, \quad (9)$$

where (α, τ) is the solution to (8). Then

- 1 The quantity $t^\infty(q)$ exists and is unique for any $q \in (0, 1)$. Furthermore, it has a limit as $q \rightarrow 0$ (and, for fixed λ , this limit depends on Π only).
- 2 Knockoff random threshold $\hat{t}$ satisfies $\hat{t} \rightarrow t^\infty(q)$ in probability.

The asymptotic power is given by

$$\text{TPP} \rightarrow \mathbb{P}(\omega(\eta_{\alpha\tau}(\Pi + \tau W), \tau \eta_\alpha(\tilde{W})) \geq t^\infty | \Pi \neq 0)$$

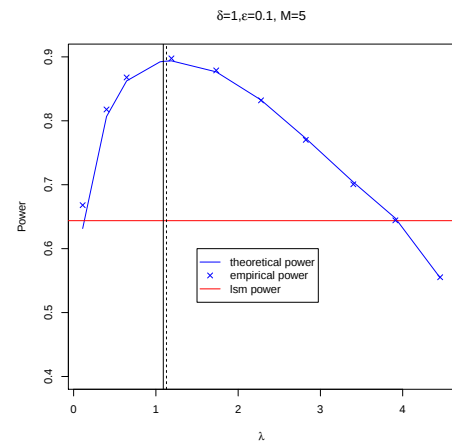
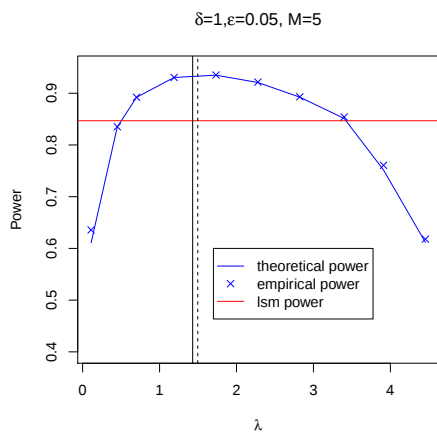
Given that

$$\hat{\beta}_i \sim \tau \eta_\alpha \left(\frac{\Pi}{\tau} + Z \right)$$

the "best" ordering of $\hat{\beta}_i$ occurs when τ is minimal. Bayati and Montanari (2012):

$$\frac{1}{p} \|\hat{\beta} - \beta\|^2 \rightarrow \delta(\tau^2 - \sigma^2)$$

Thus minimizing τ corresponds to minimizing the prediction error. Optimal τ can be identified through crossvalidation



Other examples of applications of AMP theory

G. Reeves, 2017, neural networks

P.Sur and E.J.Candès, 2018, maximum likelihood estimators in logistic regression