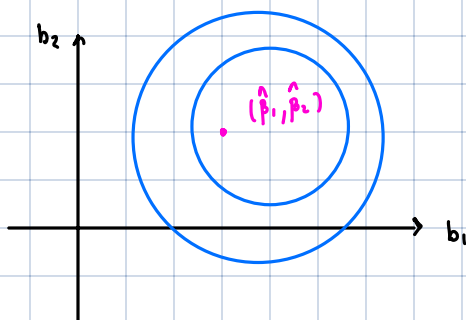


Recordemos nuestro modelo lineal

$$y_i = x_i' \beta + u_i, \quad E[u_i | x_i] = 0.$$

Tenemos que $\hat{\beta}_n$ es la solución a

$$\hat{\beta}_n = \underset{b}{\operatorname{argmin}} \sum_{i=1}^n \frac{(y_i - x_i' b)^2}{2n}$$



Las curvas de nivel
se minimizan en
 $(\hat{\beta}_1, \hat{\beta}_2)$

3. Ridge

La regresión Ridge propone estimar

$$\hat{\beta}_n^R = \underset{b}{\operatorname{argmin}} \sum_{i=1}^n \frac{(y_i - x_i' b)^2}{2n} + \underbrace{\lambda b' b / 2},$$

Penalizamos
 $\|b\|_2^2 := \sum_{j=1}^p b_j^2$

donde λ es un parámetro externo que elegimos (no estimamos directamente). A esto se le conoce como un **hiperparámetro**.

Derivando

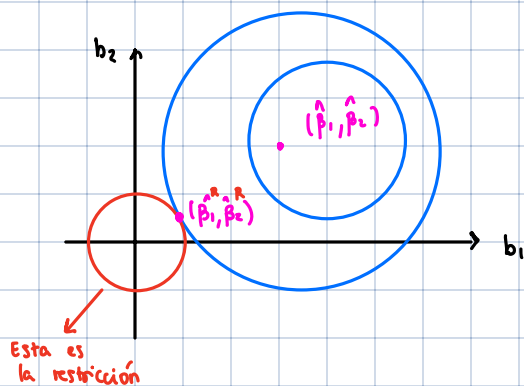
$$Q_n(b; \lambda) = \frac{1}{2n} (y - Xb)'(y - Xb) + \frac{\lambda}{2} b' b$$

$$- X' (y - X \hat{\beta}_n^R) + \lambda \hat{\beta}_n^R = 0$$

$$\Rightarrow X' X \hat{\beta}_n^R + \lambda \hat{\beta}_n^R = X' y$$

$$\Rightarrow \hat{\beta}_n^R = \underbrace{(X' X + \lambda I_k)^{-1}} X' y$$

Incluso cuando $X' X$ es cercano a ser singular, al sumar este último término hacemos que una solución siempre exista.



* la regresión Ridge acerca los coeficientes hacia cero (shrinkage)

Sin embargo, noten que hemos introducido sesgo incluso en muestras finitas. Supongamos $E[u_i | x_i] = 0$.

$$\hat{\beta}_n^R = (X^T X + \lambda I_k)^{-1} X^T X \beta + (X^T X + \lambda I_k)^{-1} X^T u$$

$$\Rightarrow E[\hat{\beta}_n^R | X] = (X^T X + \lambda I_k)^{-1} X^T X \beta \neq \beta.$$

Es decir, Ridge introduce sesgo a costa de reducir la varianza de las predicciones. Noten como estabiliza la varianza en el caso $k=1$ y homocedástico:

$$\hat{\beta}_{1n} = \frac{\sum_{i=1}^n x_i y_i}{\sum_{i=1}^n (x_i^2 + \lambda)}, \quad \text{Var}(\hat{\beta}_1 | X) = \frac{\sigma^2}{\sum_{i=1}^n (x_i^2 + \lambda)}$$

Este denominador está más lejos de 0.

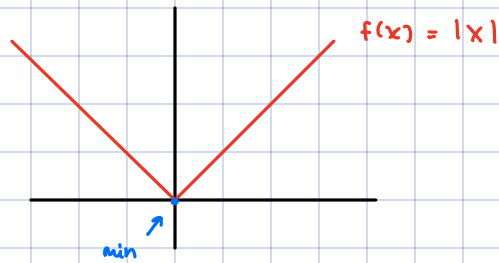
Finalmente, ya que Ridge NO impone una solución de esquina (i.e. que algún $\hat{\beta}_j$ sea 0), no sirve para hacer selección de regresores. Por ello, estudiaremos un siguiente estimador.

2. Minimización Convexa

Existen casos donde queremos minimizar una función objetivo que es convexa. Por tanto, definitivamente podemos encontrar un óptimo. Sin embargo, no siempre son objetos diferenciables.



Se halla como $\partial f(x) = 2x = 0$



¿cómo se halla?

2.1. Subgradiente

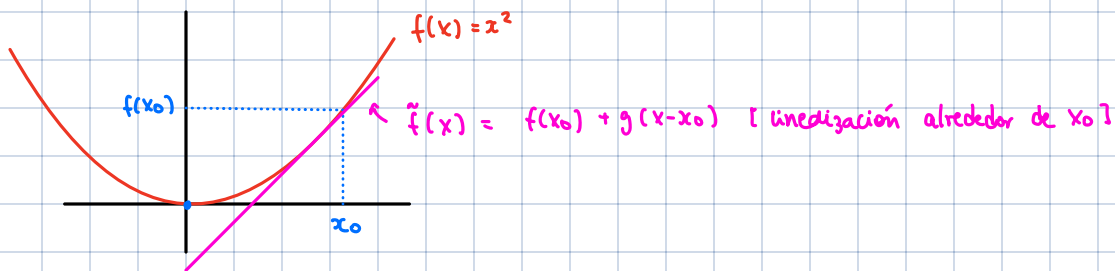
Def:- g es subgradiente de f en x_0 si $\forall x \in \mathbb{R}^k$ tenemos

$$f(x) - f(x_0) \geq g^T(x - x_0)$$

Para ganar intuición consideremos $k=1$. La definición implica que

$$\begin{aligned} f(x) &\geq f(x_0) - g(x_0) + g x \\ &= f(x_0) + g(x - x_0) \end{aligned}$$

Ejemplo:



Es decir, la subgradiente es una de las posibles pendientes de las líneas que pasan por el punto $(x_0, f(x_0))$, y que siempre estén por debajo de $f(x)$. Vemos que en el caso de $f(x) = x^2$ solo hay una posible línea.

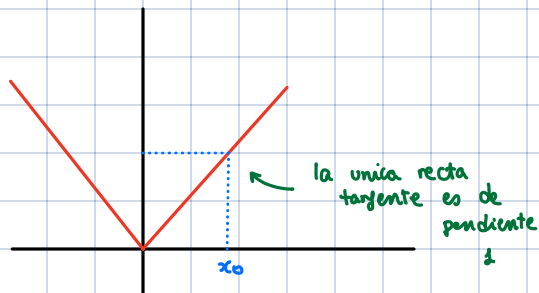
Def:- El set de todos los subgradientes en x_0 se conoce como **subdiferencial** en x_0 .

$$\partial f(x_0) = \{ g \in \mathbb{R}^k : f(x) - f(x_0) \geq g^T(x - x_0) \quad \forall x \in \mathbb{R}^k \}$$

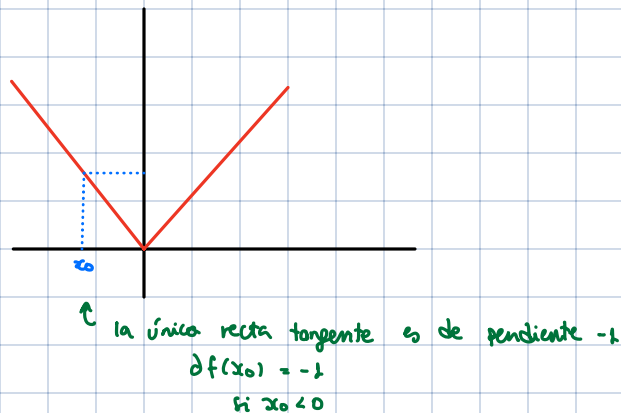
Decimos que una función es diferenciable en x_0 si su subdiferencial en x_0 solo incluye un único elemento.

$$\partial f(x_0) = \{g\} \iff g = \nabla f(x).$$

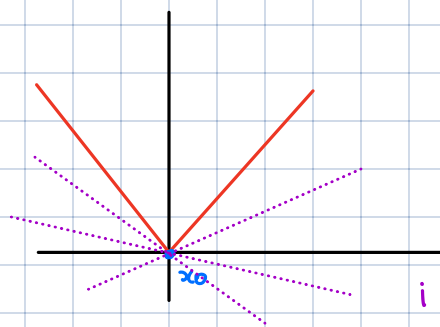
Veamos ahora el valor absoluto:



$$\partial f(x_0) = 1 \quad \text{si } x_0 > 0$$



$$\partial f(x_0) = -1 \quad \text{si } x_0 < 0$$



¡ Hay varias posibles tangentes!

$$\partial f(x_0) = [-1, 1] \quad \text{si } x_0 = 0$$

Es decir, cualquier pendiente entre -1 y 1 es un subgradiente.

Por ello, decimos que:

$$\partial |x| = \begin{cases} -1 & , \quad x < 0 \\ [-1, 1] & , \quad x = 0 \\ 1 & , \quad x > 0 \end{cases}$$

Proposición.- Sea M el set de puntos mínimos de una función convexa f :

$$M = \{x^* \in \mathbb{R}^k : f(x^*) \leq f(x) \quad \forall x \in \mathbb{R}^k\}$$

Entonces

$$x^* \in M \iff 0 \in \partial f(x^*).$$

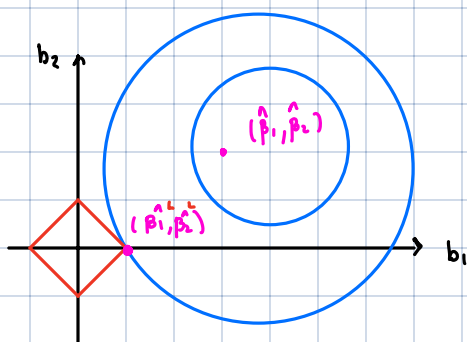
Por ello, podemos decir que $x_0 = 0$ minimiza $f(x) = |x|$:

$$0 \in \partial |x| \text{ cuando } x = x_0 \Rightarrow x^* \in M.$$

3. LASSO

La regresión Ridge propone estimar

$$\hat{\beta}_n^L = \underset{b}{\operatorname{argmin}} \sum_{i=1}^n \frac{(y_i - x_i' b)^2}{2n} + \lambda \sum_{j=1}^p |b_j| = \frac{\|Y - Xb\|_2^2}{2n} + \lambda \|b\|_1$$



Está diseñado para hallar soluciones de esquina donde algún $\hat{\beta}_j$ será igual a 0. Es decir, selecciona regresores.

Importantemente, este es un problema de minimización convexa.

3.1. Solución Analítica (caso especial)

Consideremos

$$\sum_{i=1}^n \frac{x_{ij} x_{i\ell}}{n} = \begin{cases} 0 & j \neq \ell \\ 1 & j = \ell \end{cases} \quad \left. \vphantom{\sum_{i=1}^n} \right\} \text{ regresores ortogonales y segundo momento igual a 1.}$$

En ese caso

$$\frac{X'X}{n} = I_K$$

$$\Rightarrow \hat{\beta}^{OLS} = \frac{X'Y}{n},$$

por lo que

$$\hat{\beta}_j^{OLS} = \sum_{i=1}^n \frac{x_i y_i}{n}.$$

Vamos a definir $(x)^+ = \max\{x, 0\}$

$$\operatorname{sign}(x) = \begin{cases} -1 & , x < 0 \\ 0 & , x = 0 \\ 1 & , x > 0 \end{cases}$$

Proposición:- Suponer $X'X/n = I_K$. Entonces el estimador LASSO resuelve la minimización de

$$Q_{n,\lambda}(b) = \left\{ \frac{1}{2n} \|y - Xb\|_2^2 + \lambda \|b\|_1 \right\}$$

y satisface

$$\hat{\beta}_j = \text{sign}(\hat{\beta}_j^{\text{OLS}}) (|\hat{\beta}_j^{\text{OLS}}| - \lambda)^+ \quad \forall j = 1, 2, \dots, K.$$

Achica los coeficientes en una magnitud λ .

Si $\lambda > |\hat{\beta}_j^{\text{OLS}}|$ lo reemplaza

$\hat{\beta}_j^{\text{OLS}}$ por $\hat{\beta}_j = 0$.

Shrinkage + selection

proof:-

Definimos

$$\underbrace{v}_{K \times 1} + \underbrace{c}_{1 \times 1} \underbrace{S}_{\text{set de vectores } K \times 1} = \{v + cg : g \in S\}.$$

Usando esa notación definimos

$$\begin{aligned} \partial Q_{n,\lambda}(b) &= - \frac{X'(y - Xb)}{n} + \lambda \partial \|b\|_1 \\ &= - \frac{X'y}{n} + \frac{X'X}{n} b + \lambda \partial \|b\|_1 \\ &= - \hat{\beta}^{\text{OLS}} + I_K b + \lambda \partial \|b\|_1 \\ &= - (\hat{\beta}^{\text{OLS}} - b) + \lambda \partial \|b\|_1 \\ &= - \begin{pmatrix} \hat{\beta}_1^{\text{OLS}} - b_1 & - \lambda \partial |b_1| \\ \vdots \\ \hat{\beta}_K^{\text{OLS}} - b_K & - \lambda \partial |b_K| \end{pmatrix}. \end{aligned}$$

Noten que podemos resolver elemento por elemento. En concreto, el estimador LASSO satisface

$$0 \in \hat{\beta}_j^{\text{OLS}} - \hat{\beta}_j - \lambda \partial |\hat{\beta}_j|$$

Entonces $\hat{\beta}_j = 0$ si y solo si

$$0 \in \hat{\beta}_j^{\text{OLS}} - 0 - \lambda \partial |0| = \hat{\beta}_j^{\text{OLS}} - \lambda [-1, 1] = [\hat{\beta}_j^{\text{OLS}} - \lambda, \hat{\beta}_j^{\text{OLS}} + \lambda]$$

Esto es equivalente a decir

$$\hat{\beta}^{OLS} - \lambda \leq 0 \leq \hat{\beta}^{OLS} + \lambda$$

$$\Rightarrow -\lambda \leq \hat{\beta}^{OLS} \leq \lambda$$

$$\Rightarrow |\hat{\beta}^{OLS}| \leq \lambda$$

$$\Rightarrow (|\hat{\beta}^{OLS}| - \lambda)^+ = 0.$$

El estimador LASSO $\hat{\beta}_j$ sera $\neq 0$ si y solo si $|\hat{\beta}^{OLS}| > \lambda$, que ocurre si

$$\hat{\beta}^{OLS} > \lambda \text{ o } \hat{\beta}^{OLS} < -\lambda.$$

En esos casos tenemos:

$$\underbrace{\hat{\beta}_j > \lambda}_{\hat{\beta}_j > 0} \Rightarrow \hat{\beta}_j^{OLS} - \hat{\beta}_j - \lambda \partial |\hat{\beta}_j| = \hat{\beta}_j^{OLS} - \hat{\beta}_j - \lambda \cdot 1 = 0$$

$$\Rightarrow \hat{\beta}_j = \hat{\beta}_j^{OLS} - \lambda = - (|\hat{\beta}_j^{OLS}| - \lambda)^+.$$

$$\underbrace{\hat{\beta}_j < -\lambda}_{\hat{\beta}_j < 0} \Rightarrow \hat{\beta}_j^{OLS} - \hat{\beta}_j - \lambda \partial |\hat{\beta}_j| = \hat{\beta}_j^{OLS} - \hat{\beta}_j - \lambda \cdot (-1) = 0$$

$$\Rightarrow \hat{\beta}_j = \hat{\beta}_j^{OLS} + \lambda = - (|\hat{\beta}_j^{OLS}| - \lambda)^+.$$

Ahora, recuerden que $\hat{\beta}_j^{OLS} = \beta_j + O_p(n^{-1/2})$, que para $\beta_j = 0$ implica $\hat{\beta}_j^{OLS} = O_p(n^{-1/2})$.

Para que LASSO seleccione correctamente, λ_n debe dominar ese componente de ruido.

Sin embargo λ_n tampoco puede ser muy grande o enviara todo a 0. Esto empezara a ocurrir si $\lambda > \min_{j \in A_0} |\beta_j|$.

Supongamos que $\min_j |\beta_j| \geq \delta > 0$, $\lambda_n = \sqrt{\frac{\log n}{n}}$.

$$\begin{aligned} \Pr(\hat{\beta}_j = 0 \mid \beta_j = 0) &= \Pr(|\hat{\beta}_j^{OLS}| < \lambda_n \mid \beta_j = 0) \\ &= \Pr(O_p(n^{-1/2}) < \sqrt{\frac{\log n}{n}}) \rightarrow 1. \end{aligned}$$

$$\begin{aligned} \Pr(\hat{\beta}_j \neq 0 \mid \beta_j \neq 0) &= \Pr(|\hat{\beta}_j^{OLS}| > \lambda_n \mid \beta_j \neq 0) \\ &= \Pr(|\beta_j + O_p(n^{-1/2})| > \sqrt{\frac{\log n}{n}} \mid \beta_j \neq 0) \\ &= \Pr(\delta + O_p(n^{-1/2}) > \sqrt{\frac{\log n}{n}}) \rightarrow 1. \end{aligned}$$

3.2. Caso General

NO EXISTE una fórmula cerrada de la solución, pero la vamos a caracterizar. En concreto LASSO resuelve

$$-\frac{1}{n} \sum_{i=1}^n x_i (y_i - x_i' \hat{\beta}) + \lambda_n \hat{g} = 0,$$

$$\text{donde } \hat{g} \in \|\hat{\beta}\|_1 : \quad \hat{g}_j = \begin{cases} \text{sign}(\hat{\beta}_j) & \hat{\beta}_j \neq 0 \\ \in [-1, 1] & \hat{\beta}_j = 0 \end{cases}$$

Denotamos A_n^* como los regresores seleccionados por Lasso

$$A_n^* = \{j : \hat{\beta}_j \neq 0\}$$

Podemos escribir la condición de primer orden así

$$\frac{1}{n} \sum_{i=1}^n x_{i, A_n^*} (y_i - x_{i, A_n^*}' \hat{\beta}_{A_n^*}) - \lambda_n \underbrace{\text{sign}(\hat{\beta}_{A_n^*})}_{\text{sign}(\cdot) \text{ de cada uno de sus elementos}} = 0$$

Tenemos que Lasso excluye regresores ($j \in A_n^{*c}$ o $\hat{\beta}_j = 0$) cuando

$$0 = -\frac{1}{n} \sum_{i=1}^n x_{ij} (y_i - x_{i, A_n^*}' \hat{\beta}_{A_n^*}) + \lambda_n \cdot \hat{g}_j, \quad \text{donde } |\hat{g}_j| \leq 1$$

o equivalentemente

$$\left| \frac{1}{n} \sum_{i=1}^n x_{ij} (y_i - x_{i, A_n^*}' \hat{\beta}_{A_n^*}) \right| \leq \lambda_n$$

Denotamos $A_0 = \{j : \beta_j \neq 0\}$ como el set de regresores relevantes, tal que el verdadero modelo es

$$y_i = x_{i, A_0}' \beta_{A_0} + u_i$$

Finalmente, fíjense que $\text{sign}(\hat{\beta}^*) = \text{sign}(\beta)$ implica que seleccionamos los regresores correctos

$$\begin{aligned} \beta_j = 0 &\Rightarrow \hat{\beta}_j^* = 0 & (\text{ocurre si } j \in A_0^c) &\Rightarrow A_n^* = A_0 \\ \beta_j \neq 0 &\Rightarrow \hat{\beta}_j^* \neq 0 & (\text{ocurre si } j \in A_0) & \end{aligned}$$

Vamos a demostrar cuando ocurre esto.

Proposición - Supongamos que $E X_{i,A_0} X'_{i,A_0'}$ es p.d. Entonces $\text{sign}(\hat{\beta}) = \text{sign}(\beta)$
si y solo si

$$\textcircled{1} \quad \text{sign}(\beta_{A_0}) = \text{sign} \left(\beta_{A_0} + \left(\frac{\sum X_{i,A_0} X'_{i,A_0'}}{n} \right)^{-1} \left(\frac{1}{n} \sum X_{i,A_0} u_i - \lambda_n \text{sign}(\beta_{A_0}) \right) \right)$$

y para todo $j \in A_0^c$

$$\textcircled{2} \quad \left| \frac{1}{n} \sum_{i=1}^n x_{ij} x'_{i,A_0} \left(\frac{\sum x_{i,A_0} x'_{i,A_0}}{n} \right)^{-1} \left(\frac{\sum x_{i,A_0} u_i}{n} - \lambda_n \text{sign}(\beta_{A_0}) \right) - \frac{\sum x_{ij} u_i}{n} \right| \leq \lambda_n$$

proof: $(\Rightarrow)$ Supongamos $\text{sign}(\hat{\beta}) = \text{sign}(\beta)$. Entonces a simple

$$0 = \frac{1}{n} \sum x_{i,A_0} (y_i - x'_{i,A_0'} \hat{\beta}_{A_0}) - \lambda_n \text{sign}(\beta_{A_0})$$

$$\Rightarrow 0 = \frac{1}{n} \sum x_{i,A_0} (u_i - x'_{i,A_0} (\hat{\beta}_{A_0} - \beta_{A_0})) - \lambda_n \text{sign}(\beta_{A_0})$$

$$\Rightarrow \hat{\beta}_{A_0} = \beta_{A_0} + \left(\frac{\sum x_{i,A_0} x'_{i,A_0}}{n} \right)^{-1} \left(\frac{1}{n} \sum x_{i,A_0} u_i - \lambda_n \text{sign}(\beta_{A_0}) \right)$$

lo que significa que

$$\text{sign}(\beta_{A_0}) = \text{sign}(\hat{\beta}_{A_0})$$

$$= \text{sign} \left(\beta_{A_0} + \left(\frac{\sum x_{i,A_0} x'_{i,A_0}}{n} \right)^{-1} \left(\frac{1}{n} \sum x_{i,A_0} u_i - \lambda_n \text{sign}(\beta_{A_0}) \right) \right).$$

(obtenemos $\textcircled{1}$)

la ecuación $\textcircled{2}$ viene de la cond. de primer orden, ya que $A_n^* = A_0$
excluye a j si:

$$\left| \frac{1}{n} \sum_{i=1}^n x_{ij} (y_i - x'_{i,A_0} \hat{\beta}_{A_0}) \right| \leq \lambda_n$$

reemplazamos la solución de arriba y nos da $\textcircled{2}$

$(\Leftarrow)$ revisenla en las notas, pero sigue técnicas que ya usamos.

Podemos re-escribir la proposición

$$① \quad \text{sign}(\beta_{A_0}) = \text{sign} \left(\beta_{A_0} + \left(\sum \frac{x_{iA_0} x_{iA_0}'}{n} \right)^{-1} \left(\frac{1}{n} \sum x_{iA_0} u_i - \lambda_n \text{sign}(\beta_{A_0}) \right) \right)$$

para todo $j \in A_0^c$

$$② \quad \left| \frac{1}{n} \sum_{i=1}^n x_{ij} x_{iA_0}' \left(\sum \frac{x_{iA_0} x_{iA_0}'}{n} \right)^{-1} \left(\frac{\sum x_{iA_0} u_i}{n} - \lambda_n \text{sign}(\beta_{A_0}) \right) - \frac{\sum x_{ij} u_i}{n} \right| \leq \lambda_n$$

como

$$\left| \beta_j + \left[\left(\sum \frac{x_{iA_0} x_{iA_0}'}{n} \right)^{-1} \left(\frac{1}{n} \sum x_{iA_0} u_i - \lambda_n \text{sign}(\beta_{A_0}) \right) \right]_j \right| > 0, \quad j \in A_0$$

$$\left| \frac{1}{n} \sum_{i=1}^n x_{ij} x_{iA_0}' \left(\sum \frac{x_{iA_0} x_{iA_0}'}{n} \right)^{-1} \left(\frac{\sum x_{iA_0} u_i}{n} - \lambda_n \text{sign}(\beta_{A_0}) \right) - \frac{\sum x_{ij} u_i}{n} \right| \leq \lambda_n, \quad j \in A_0^c$$

Supongamos que

$$\lambda_n = \sqrt{\frac{\log n}{n}}.$$

Entonces, para todo $j \in A_0$

$$\left| \beta_j + \left[\left(\sum \frac{x_{iA_0} x_{iA_0}'}{n} \right)^{-1} \left(\frac{1}{n} \sum x_{iA_0} u_i - \lambda_n \text{sign}(\beta_{A_0}) \right) \right]_j \right|$$

$$= \left| \beta_j + O_p(1) \left(O_p\left(\frac{1}{\sqrt{n}}\right) - O\left(\sqrt{\frac{\log n}{n}}\right) \right) \right|$$

$$= |\beta_j + o_p(1)| > 0$$

con probabilidad aprox. a 1 siempre
que $\min_{j \in A_0} |\beta_j| \geq \delta > 0$.

Ahora, para todo $j \in A_0^c$

$$\left| \frac{1}{n} \sum_{i=1}^n x_{ij} x_{iA_0}' \left(\sum \frac{x_{iA_0} x_{iA_0}'}{n} \right)^{-1} \underbrace{\left(\frac{\sum x_{iA_0} u_i}{n} - \lambda_n \text{sign}(\beta_{A_0}) \right)}_{O_p(1/\sqrt{n})} - \underbrace{\frac{\sum x_{ij} u_i}{n}}_{O_p(1/\sqrt{n})} \right| \leq \lambda_n$$

$$\Rightarrow \left| O_p(1/\sqrt{n}) + \lambda_n \sum_{i=1}^n x_{ij} x_{iA_0}' \left(\sum \frac{x_{iA_0} x_{iA_0}'}{n} \right)^{-1} \text{sign}(\beta_{A_0}) \right| \leq \lambda_n$$

Fijense que debe cumplirse que en muestras grandes

$$\left| E X_{ij} X'_{iA_0} (E X_{i,A_0} X'_{i,A_0})^{-1} \text{sign}(\beta_{A_0}) \right| \leq 1$$

para $j \in A_0^c$. Esta condición se llama **irrepresentabilidad**. Lo que dice es que los regresores irrelevantes no deben estar fuertemente correlacionados con los relevantes. De lo contrario, Lasso mantendrá regresores irrelevantes.