

# Datos direccionales. Densidad.

Leyenda Rodríguez, María

4 de abril de 2011

- 1 Estrategias de modelado para datos circulares bivariantes
  - Introducción
  - Propiedades de simetría
  - Modelo híbrido
  - Simulación
  - Estudio
- 2 Distribución von Mises Multivariante
  - Introducción
  - Distribución von Mises Multivariante
  - Enfoques para la inferencia
- 3 Estimación no paramétrica de la densidad circular
  - Estimación de la densidad circular tipo núcleo
  - Selección del parámetro de suavizado. Plug-in
  - Selección del parámetro de suavizado. Validación cruzada
  - Análisis mensual
- 4 Estimación de modelo de series de tiempo

- En el toro hay dos comunes aproximaciones para construir la distribución la cual es análoga a la distribución normal bivalente en el plano.
- Estas aproximaciones son a menudo obtenidas por la aproximación en términos del modelo del seno y coseno, respectivamente; el modelo del coseno en sus dos versiones.
- Cada aproximación tiene sus ventajas y sus inconvenientes.

- Un modelo bidimensional circular general, introducido por Mardia(1975a)

### Distribución von Mises bidimensional completa"

$$f(\theta_1, \theta_2) \propto e^{\kappa_1 \cos(\theta_1 - \mu_1) + \kappa_2 \cos(\theta_2 - \mu_2) + [\cos(\theta_1 - \mu_1) \sin(\theta_1 - \mu_1)] A \begin{bmatrix} \cos(\theta_2 - \mu_2) \\ \sin(\theta_2 - \mu_2) \end{bmatrix}}$$

- Donde  $\theta_1, \theta_2 \in (-\pi, \pi)$  se encuentran en el toro, es decir, un cuadrado con lados opuestos identificados.
- A es una matriz  $2 \times 2$ .
- Este modelo tiene 8 parámetros y permite dependencia entre los 2 ángulos.

- Con el fin de simplificar la presentación, ponemos  $\mu_1 = \mu_2 = 0$ , así que los tres parámetros restantes describen la concentración de cada ángulo y su concentración. Estos modelos son

- 1 el modelo coseno con iteración positiva

$$f(\theta_1, \theta_2) \propto \exp \{ \kappa_1 \cos(\theta_1) + \kappa_2 \cos(\theta_2) + \gamma_1 \cos(\theta_1 - \theta_2) \}$$

- 2 El modelo seno con iteración negativa

$$f(\theta_1, \theta_2) \propto \exp \{ \kappa_1 \cos(\theta_1) + \kappa_2 \cos(\theta_2) + \gamma_2 \cos(\theta_1 + \theta_2) \}$$

- 3 el modelo sel seno (Singh et al., 2002)

$$f(\theta_1, \theta_2) \propto \exp \{ \kappa_1 \cos(\theta_1) + \kappa_2 \cos(\theta_2) + \delta \operatorname{sen}(\theta_1) \operatorname{sen}(\theta_2) \}$$

- Mardia et al.(2007) dan un estudio sistemático de los modelos (i)-(iii) y una comparación entre ellos.

- Bajo alta concentración sobre  $(\theta_1, \theta_2) = (0, 0)$ , cada modelo se comporta siguiendo una distribución normal bidimensional y la inversa de la matriz de covarianzas es de la forma

$$\sum_1^{-1} = \begin{pmatrix} \kappa_1 + \gamma_1 & -\gamma_1 \\ -\gamma_1 & \kappa_2 + \gamma_1 \end{pmatrix}$$

$$\sum_2^{-1} = \begin{pmatrix} \kappa_1 + \gamma_2 & -\gamma_2 \\ -\gamma_2 & \kappa_2 + \gamma_2 \end{pmatrix}$$

$$\sum_3^{-1} = \begin{pmatrix} \kappa_1 & -\delta \\ -\delta & \kappa_2 \end{pmatrix}$$

- Bajo alta concentración suponemos que los parámetros relevantes  $\kappa_1, \kappa_2, \delta, \gamma_1, \gamma_2$  se hacen grandes sin dejar de ser una proporción constante entre sí, bajo la restricción de que la matriz inversa de covarianzas es definida positiva.
- Para los tres modelos esta restricción se reduce a
  - 1  $\kappa_1 + \gamma_1 > 0, \kappa_2 + \gamma_1 > 0, \gamma_1^2 < (\kappa_1 + \gamma_1)(\kappa_2 + \gamma_1),$
  - 2  $\kappa_1 + \gamma_2 > 0, \kappa_2 + \gamma_2 > 0, \gamma_2^2 < (\kappa_1 + \gamma_2)(\kappa_2 + \gamma_2),$
  - 3  $\kappa_1 > 0, \kappa_2 > 0, \delta^2 < \kappa_1 \kappa_2,$
- Además, para el modelo coseno también se necesita  $\kappa_1 > 0, \kappa_2 > 0$  para garantizar el modo global de  $f$  en  $(\theta_1, \theta_2) = (0, 0)$ .

- Para cada uno de estos modelos, es posible escoger los parámetros de modo que coincida con cualquier matriz inversa de covaianzas definida positiva.
- Sin embargo, el modelo del coseno puede mostrar algún comportamiento multimodal poco atractivo si  $\kappa_1 < 0$  o  $\kappa_2 < 0$ . Por lo tanto, restringimos la atención a los modelos coseno en los que  $\kappa_1 > 0$  o  $\kappa_2 > 0$ .
- En este caso, la correspondiente matriz de covarianza inversa  $\Sigma^{-1}$  posee, lo que puede llamarse la propiedad "covarianza dominada", es decir,  $|\sigma^{12}| < \sigma^{11}$  y  $|\sigma^{12}| < \sigma^{22}$

- Por construcción, cada modelo del seno y del coseno es simétrico si,  $f(\theta_1, \theta_2) = f(-\theta_1, -\theta_2)$ . Esta simetría proporciona un patrón elíptico en los contornos de masa de probabilidad constante de  $f$  sobre la moda  $(\theta_1, \theta_2) = (0, 0)$
- Para cada uno de los tres modelos,  $f$  tiene una simetría más ya que es una función continua en el toro,

$$f(\theta_1, \pi) = f(\theta_1, -\pi), \quad f(\pi, \theta_2) = f(-\pi, \theta_2)$$

- Un patrón elíptico en un contorno de masa de probabilidad constante para  $f$  será generalmente distorsionado cuando  $(\theta_1, \theta_2)$  se aproxima al límite del cuadrado en el que  $f$  está definida.
- Esta distorsión complica el desarrollo de los algoritmos de simulación eficientes sobre 2-dimensiones ya que la densidad no será necesariamente monótona decreciente en los rayos desde el origen hasta el borde del cuadrado.

- Por simplicidad nos restringiremos a la situación de covarianza dominada.
  - Resulta que el modelo coseno con interacción positiva es el que menos distorsiona bajo correlación positiva ( $\gamma_1 > 0$ ) y que el modelo coseno con interacción negativa es el que menos distorsiona bajo correlación negativa ( $\gamma_2 < 0$ )
  - Parece oportuno emplear el modelo coseno con interacción positiva bajo correlación positiva entre  $\text{sen}\theta_1$  y  $\text{sen}\theta_2$  y el modelo coseno con interacción negativa bajo correlación negativa.
- Desafortunadamente el cruce entre los dos modelos no es continuo en el modelo de independencia.
- Luego consideraremos un modelo híbrido para proporcionar una suave transición.

- Para poca concentración, el carácter exacto de cualquier desviación de la distribución uniforme no es demasiado importante.
- Por lo tanto, sugerimos el siguiente modelo híbrido

$$f(\theta_1, \theta_2) \propto e^{\kappa_1 \cos \theta_1 \cos \theta_2} + \beta[(\cosh \lambda - 1) \cos \theta_1 \cos \theta_2 + \sinh \lambda \sin \theta_1 \sin \theta_2]$$

- $\beta$  es un parámetro de ajuste, el cual fijaremos 1 por simplicidad.
- Para  $\lambda$  próximo a cero el modelo se comporta como el modelo del seno con  $\beta\lambda \approx \delta$ .
- Para valores grandes de  $\lambda > 0$  (o grandes  $lambda > 0$ ) el modelo se comporta como el modelo coseno con iteracción positiva con  $\beta \exp(\lambda)/2 \approx \gamma_1$  (o  $\beta \exp(-\lambda)/2 \approx \gamma_2$ )

## Conclusión

Para poca correlación el modelo se comporta como el modelo seno y para mucha correlación el modelo se comporta como el modelo coseno, con positiva o negativa iteracción, según el caso.

- Consideramos el problema de simular la Distribución von Mises bidimensional completa.<sup>o</sup> una de sus subfamilias.

Una posibilidad es usar un algoritmo MCMC basado en el hecho de que las distribuciones condicionales de  $\theta_1|\theta_2$  y  $\theta_2|\theta_1$  son von Mises (Mardia et al., 2008a). Sin embargo esta estrategia puede ser demasiado engorrosa

Al menos para los modelos seno y coseno hay un enfoque más sencillo.

## Algoritmo

- 1 Simulamos  $\theta_1$  de su marginal distribución.
  - 2 Simulamos  $\theta_2|\theta_1$  de la distribución von-Mises usando, por ejemplo el algoritmo Best-Fisher (Best y Fisher, 1979).
- 
- Mardia et al. (2007) y Boomsma et al. (2008) estudian la selección empírica de una distribución von Mises en el caso unimodal (o una mixtura de dos distribuciones en el caso bimodal) para utilizarla en un algoritmo de aceptación-rechazo.
  - Más recientemente, en un trabajo inédito una justificación teórica ha sido encontrada para confirmar lo apropiado y eficiente que es la von Mises sobre el caso unimodal.

- La geometría del toro implica que no es posible obtener una única plenamente satisfactoria análoga a la distribución normal bidimensional.
- Una completa comparación entre las virtudes del modelo coseno y seno no es todavía posible, es posible extraer algunas conclusiones.
  - En muchas situaciones no hay mucha diferencia entre el modelo seno y coseno. Además, bajo alta concentración, utilizando uno u otro modelo es equivalente al ajuste de una distribución normal bidimensional en el plano tangente.
  - Para aplicaciones rutinarias el modelo seno es un poco más fácil de usar, ya que puede ser adaptado a cualquier matriz definida positiva  $\Sigma^{-1}$ , mientras que el modelo coseno son limitados para el caso de covarianza dominada.

- Sin embargo, si es necesario un modelo más refinado, el modelo coseno o híbrido pueden proporcionar un mejor ajuste.
- Para cualquiera de los modelos, la inferencia estadística es intratable usando máxima verosimilitud, pero comienza a ser sencillo usando una verosimilitud compuesta (psudo-verosimilitud) obtenida por el producto de densidades condicionales (Mardia et al. 2008a). Las pruebas limitadas en la actualidad sugieren que la distribución marginal angular sea próxima a la distribución von Mises tanto para el modelo coseno como para el modelo seno, y que la estimación mediante verosimilitud compuesta será más eficiente en esta situación.
- El modelo seno y coseno en el toro bidimensional pueden ser fácilmente extendidos a dimensiones más altas en el toro.

- Singh, Hnizdo, Demchuck (2002) propusieron una distribución circular bidimensional en el toro que es análoga a la distribución normal bidimensional.

### Función de masa de probabilidad

$$f(\theta_1, \theta_2) = \frac{e^{\kappa_1 \cos(\theta_1 - \mu_1) + \kappa_2 \cos(\theta_2 - \mu_2) + \lambda_{12} \sin(\theta_1 - \mu_1) \sin(\theta_2 - \mu_2)}}{T(\kappa_1, \kappa_2, \lambda_{12})}$$

$$-\pi \leq \theta_1, \theta_2 \leq \pi \quad \kappa_1, \kappa_2 \geq 0 \quad -\pi \leq \mu_1, \mu_2 \leq \pi$$

$$T(\kappa_1, \kappa_2, \lambda_{12}) = 4\pi^2 \sum_{m=0}^{\infty} \binom{2m}{m} \left(\frac{\lambda}{2}\right)^{2m} \kappa_1^{-m} I_m(\kappa_1) \kappa_2^{-m} I_m(\kappa_2)$$

- Función de masa de probabilidad  $\Theta^T = (\Theta_1, \Theta_2, \dots, \Theta_p)$

A set of small navigation icons typically found in Beamer presentations, including symbols for back, forward, search, and other slide controls.

- A set of small navigation icons typically found in Beamer presentations, including symbols for back, forward, search, and other slide controls.

## Propiedades

- Para alta concentración en variables circulares, tenemos aproximadamente

$$\Theta = (\Theta_1, \Theta_2, \dots, \Theta_p)^T \sim N_p(0, \Sigma)$$

$$\left(\sum_{i=1}^{-1}\right)_{ii} = \kappa_i, \quad \left(\sum_{i=1}^{-1}\right)_{ij} = -\lambda_{ij}, \quad i \neq j$$

- El modelo bivalente de Singh, Hnizdo, Demchuck (2002) se obtiene que las distribuciones marginales son simétricas entorno a la media y son unimodales o bimodales.
- Cuando  $\Theta$  sigue una distribución von Mises Multivariante, es claro que  $\Theta$  y  $-\Theta$  sigue la misma distribución. Luego todas las distribuciones marginales son simétricas.

- Sea  $\theta_i = (\theta_{1i}, \dots, \theta_{pi})$ ,  $i = 1, \dots, n$  siguiendo la distribución von Mises multivariante.
- Necesitamos estimar los parámetros del modelo.
- Como no tenemos una expresión para la constante normalizadora, el estimador de máxima verosimilitud tiene que trabajar a través de un método de optimización numérica en la constante normalizadora la cual es calculada en cada etapa de la solución iterativa de los valores del parámetro

## Método de los momentos

- Consideremos todas las distribuciones condicionales univariantes y multivariantes.
- Tenemos  $\hat{\mu}_i = \bar{x}_{0i}$ ,  $i = 1, \dots, p$ , dónde  $\bar{x}_{0i}$  son las medias circulares de las distribuciones marginales de  $\Theta_i$ ,  $i = 1, \dots, p$ .
- Para alta correlación, la matriz de covarianzas puede ser calculada para las correspondientes variables en la distribución von Mises multivariante.
- Dado el estimador de los momentos de  $\Sigma$  como  $\hat{\Sigma} = (\bar{S}_{ij})$  con

$$\bar{S}_{ij} = \frac{1}{n} \sum_{r=1}^n \text{sen}(\theta_{ir} - \bar{x}_{0i}) \text{sen}(\theta_{jr} - \bar{x}_{0j})$$

## Método de los momentos

- $\bar{S}$  es la matriz estandard de covarianzas, y debe ser interpretada como tal.
- Tomando la inversa de  $\hat{\Sigma}$  obtemos los estimadores de  $\kappa_i$ ,  $i = 1, \dots, p$  y  $\lambda_{ij}$ ,  $i \neq j$

## Método de máxima pseudo-verosimilitud

- Definimos la pseudo-verosimilitud basada en una muestra aleatoria de  $n$  observaciones  $\Theta = (\Theta_1, \dots, \Theta_p)^T$  por

$$PL = \prod_{j=1}^p \prod_{i=1}^n g_j(\Theta_{ji} | \Theta_{1i}, \dots, \Theta_{j-1,i}, \Theta_{j+1,i}, \dots, \Theta_{pi}; q)$$

- dónde  $g_j(\cdot | \dots, q)$  es la distribución condicional cuyos parámetros dependen de  $j$  y  $q$  es un parámetro desconocido de longitud  $r$ .
- Para la von Mises  $p$ -dimensional tenemos además

$$PL = (2\pi)^{-pn} \prod_{j=1}^p \prod_{i=1}^n g_j[l_0(\kappa_{j-rest}^{(i)})] e^{\kappa_{j-rest}^{(i)} \cos(\theta_{ji} - \mu_{j-rest}^{(i)})}$$

## Método de máxima pseudo-verosimilitud

- dónde para la  $i$ -ésima observación, tenemos los coeficientes de las distribuciones marginales las cuales son dadas por los parámetros en el modelo

$$\mu_{j-rest}^{(i)} = \mu_j + \tan^{-1} \{ \lambda_{jl} \sin(\theta_{li} - \mu_l) / \kappa \}$$

$$\kappa_{j-rest}^{(i)} = \left\{ \kappa_j^2 + \left[ \sum_{l \neq j} \lambda_{jl} (\theta_{li} - \mu_l) \right] \right\}$$

- Estimación paramétrica basada en máxima verosimilitud maximizando PL respecto a los  $p + p(p+1)/2$  parámetros desconocidos

- Dada una muestra aleatoria de ángulos  $\theta_1, \dots, \theta_n \in [0, 2\pi]$ .
- Consideraremos que el estimador no paramétrico de la densidad con función núcleo la distribución von Mises, viene dado por

$$\hat{f}(\theta, \nu) = \frac{c_0(\nu)}{n} \sum_{i=1}^n L(\nu \cos(\theta - \theta_i))$$

$$c_0(\nu) = \frac{1}{2\pi I_0(\nu)} \quad L(x) = e^x$$

- Donde  $I_r(\nu)$  es la función de Bessel modificada de orden  $r$ .
- El parámetro de concentración  $\nu$  ahora ha asumido el papel de (la inversa del) parámetro de suavizado.

- El error asintótico cuadrático medio integrado es

$$AMISE(\nu) = 3\kappa^2 I_2(2\kappa) / \{32\pi\nu^2 I_0(\kappa)^2\}$$

- El error asintótico cuadrático medio integrado es de la forma  $a\nu^{-2} + b\nu^{1/2}$  la cual puede ser minimizada diferenciando respecto de  $\nu$  e igualando a cero.
- Esto permite una Regla plug-in von Mises-escala" para el parámetro de suavizado  $\nu$  basado en la estimación de  $\kappa$ .

### Regla plug-in von Mises-escala

$$\nu = [3n\hat{\kappa}^2 I_2(2\hat{\kappa}) \{4\pi^{1/2} I_0(\hat{\kappa})^2\}^{-1}]^{2/5}$$

$$\nu = Cn^{2/5}$$

$$C = [3\hat{\kappa}^2 I_2(2\hat{\kappa}) \{4\pi^{1/2} I_0(\hat{\kappa})^2\}^{-1}]^{2/5}$$

- Introducimos la validación cruzada para minimizar la función pérdida dada por error cuadrático medio y la función de pérdida dada por Kullback-Leiber.
- Consideremos la estimación de la densidad construida dejando fuera el valor  $\theta_j$  de la muestra.

$$\hat{f}_j(\theta, \nu) = \frac{c_0(\nu)}{n} \sum_{i \neq j} L(\nu \cos(\theta - \theta_i))$$

$$c_0(\nu) = \frac{1}{2\pi I_0(\nu)} \quad L(x) = e^x$$

- Sea

$$cv_2(\nu) = 2n^{-1} \sum_{i=1}^n \hat{f}_i(\theta_i, \nu) - \int \hat{f}^2(\theta, \nu) d\theta$$

$$cv_{KL} = n^{-1} \sum_{i=1}^n \hat{f}_i(\theta_i, \nu)$$

- Luego  $-cv_2(\nu) + \int f^2$  y  $-cv_{KL}(\nu) + \int f \log(f)$  son estimaciones insesgadas de la pérdida del error cuadrático  $L_2(\nu)$  y la pérdida dada por Kullback-Leiber, respectivamente.

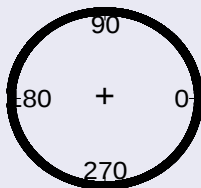
$$\nu_2 = \operatorname{argmin}_{\nu \geq 0} -cv_2(\nu)$$

$$\nu_{KL} = \operatorname{argmin}_{\nu \geq 0} -cv_{KL}(\nu)$$

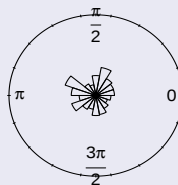
- De estas dos expresiones obtenemos dos posibles valores para el parámetro de suavizado.

## Enero

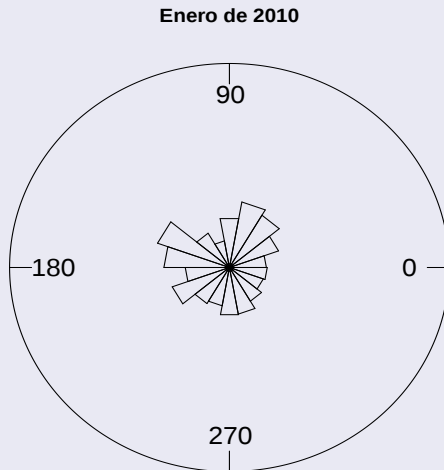
Enero de 2010 (min)



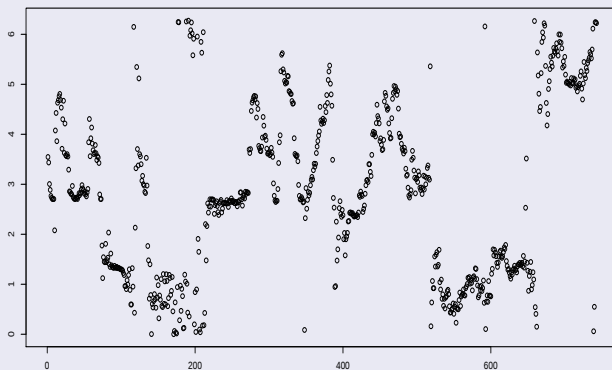
Enero de 2010 (min)



## Enero

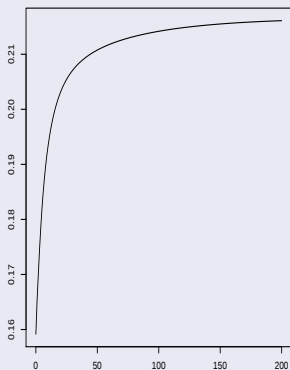


## Enero

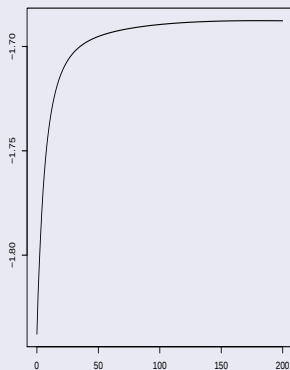


## Enero

Cross validation. Squared-error loss

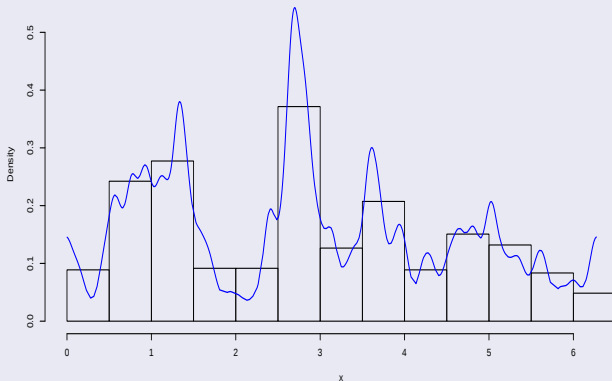


Cross validation. Kullback-Leibler

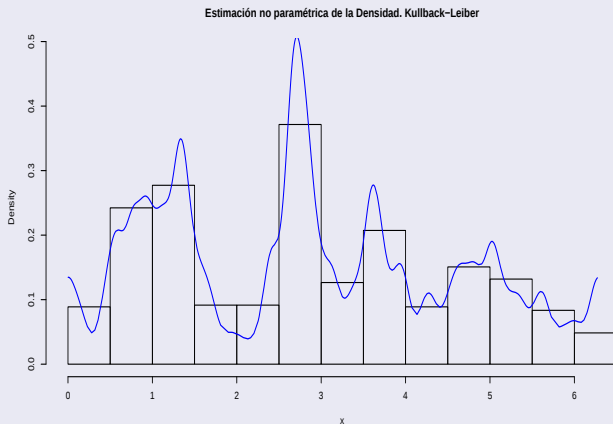


## Enero

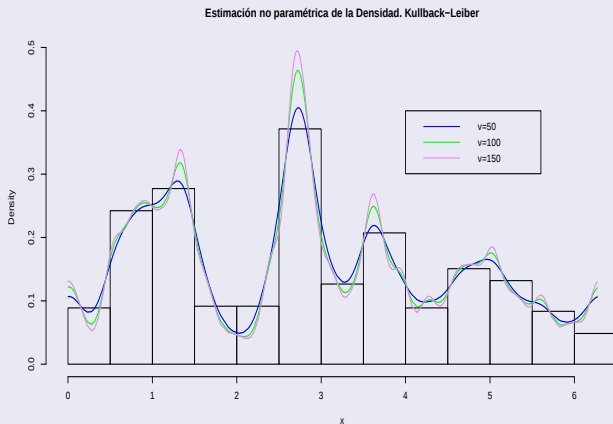
Estimación no paramétrica de la Densidad. L2



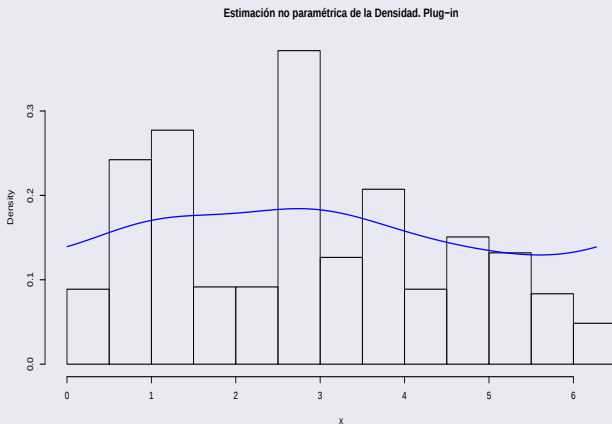
## Enero



## Enero

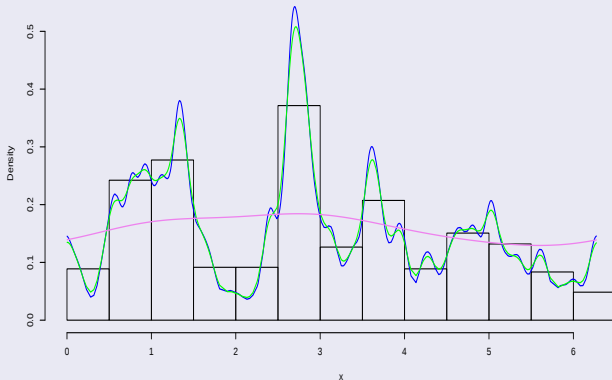


## Enero



## Enero

Estimación no paramétrica de la Densidad.



## Autocorrelación k-lag circular

$$\begin{aligned}\phi_i &= \theta_{i+k}, \quad i = 1, \dots, n-k \\ &(\phi_1, \theta_1), (\phi_2, \theta_2), \dots, (\phi_{n-k}, \theta_{n-k}) \\ \hat{\rho}_{kT} &= \frac{\sum_{1 \leq i < j \leq n-k} \text{sen}(\theta_i - \theta_j) \text{sen}(\phi_i - \phi_j)}{\sum_{1 \leq i < j \leq n-k} \text{sen}^2(\theta_i - \theta_j) \sum_{1 \leq i < j \leq n-k} \text{sen}^2(\phi_i - \phi_j)}\end{aligned}$$

## El modelo Möbius de series de tiempo

- La componente determinística del modelo de regresión estudiado por Downs y Mardia(2002) une a la variable angular dependiente  $v$  con la variable angular independiente  $u$  mediante

$$\tan \frac{1}{2}(v - \beta) = \omega \tan \frac{1}{2}(u - \alpha)$$
$$v = \beta + 2 \tan^{-1} \left\{ \omega \tan \frac{1}{2}(u - \alpha) \right\}$$

- $\omega \in [-1, 1]$  es el parámetro pendiente.
- $-\pi \leq \alpha, \beta < \pi$  son los parámetros angulares de localización.

## El modelo Möbius de series de tiempo

- Reemplazamos el ángulo  $v$  por  $\theta_t$  y el ángulo  $u$  por  $\theta_{t-1}$ ,  
 $t = 2, \dots, n$ .  $\alpha = \beta$

$$\tan \frac{1}{2}(\theta_t - \alpha) = \omega \tan \frac{1}{2}(\theta_{t-1} - \alpha)$$

$$\theta_t = \alpha + 2 \tan^{-1} \left\{ \omega \tan \frac{1}{2}(\theta_{t-1} - \alpha) \right\}$$

$$\theta_t | \theta_{t-1} \sim M(\alpha + 2 \tan^{-1} \left\{ \omega \tan \frac{1}{2}(\theta_{t-1} - \alpha) \right\}, k)$$

## El modelo Möbius de series de tiempo

- Por tanto el modelo de serie de tiempo se convierte en

$$\theta_t = \alpha + 2 \tan^{-1} \left\{ \omega \tan \frac{1}{2} (\theta_{t-1} - \alpha) \right\} + \epsilon_t$$

$$\epsilon_t \sim M(0, k)$$

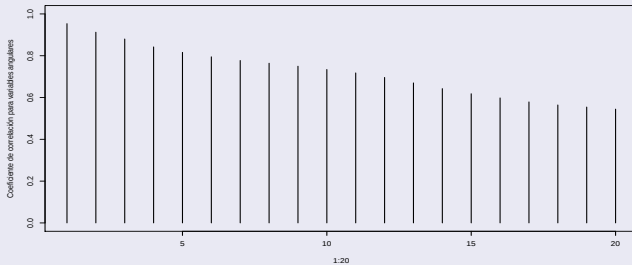
$$\mu_t = \alpha + 2 \tan^{-1} \left\{ \omega \tan \frac{1}{2} (\theta_{t-1} - \alpha) \right\}$$

## El modelo Möbius de series de tiempo

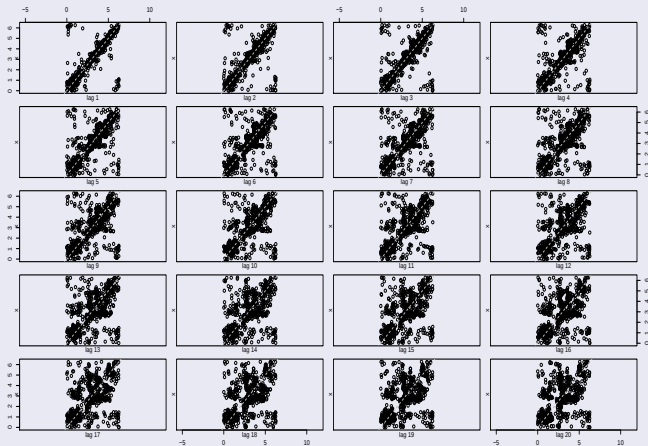
- Estimamos los parámetros  $\alpha$  y  $\omega$  mediante máxima verosimilitud; es decir, maximizando la función

$$l(\alpha, \beta) = \sum_{i=2}^n \cos \left[ \theta_t - \alpha - 2 \tan^{-1} \left\{ \omega \tan \frac{1}{2} (\theta_{t-1} - \alpha) \right\} \right]$$

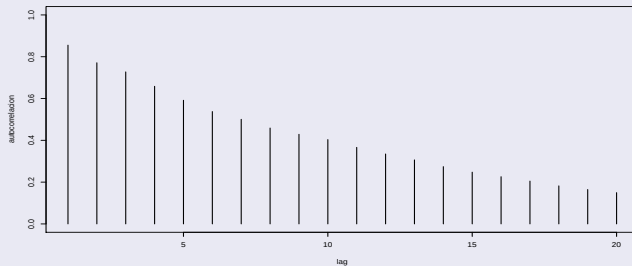
## Enero



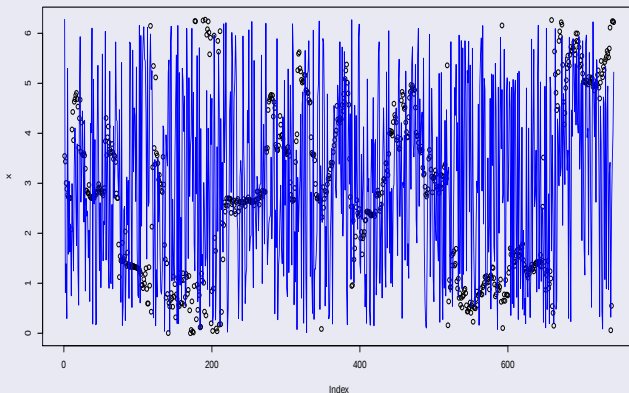
# Enero



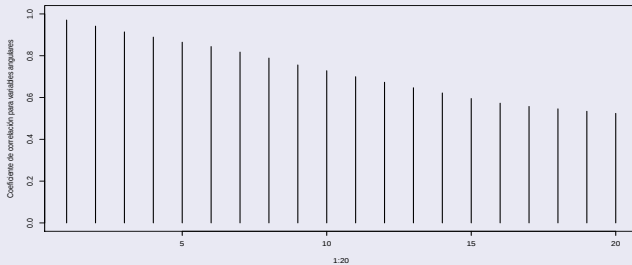
## Enero



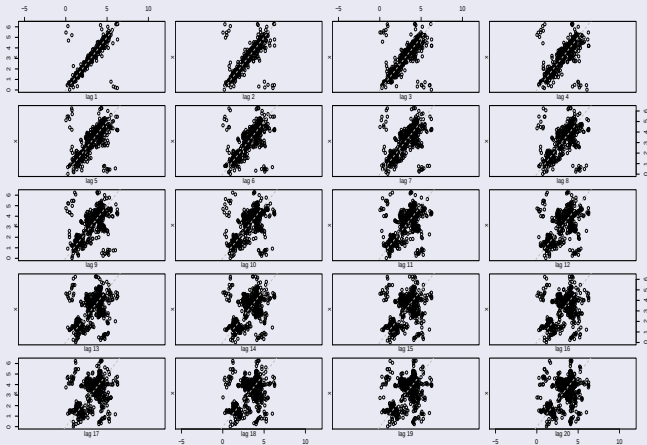
## Enero



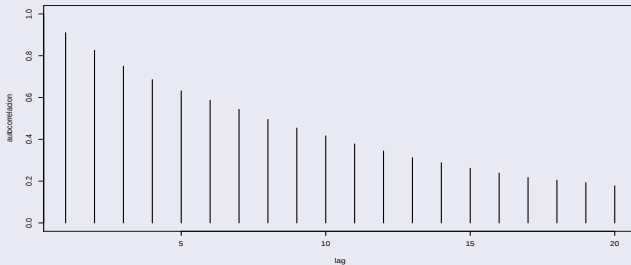
## Febrero



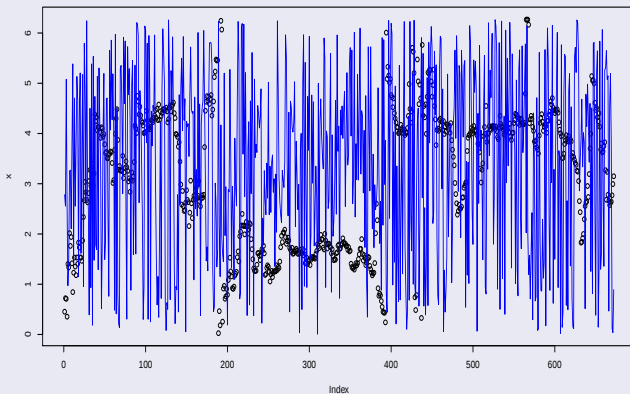
## Febrero



## Febrero



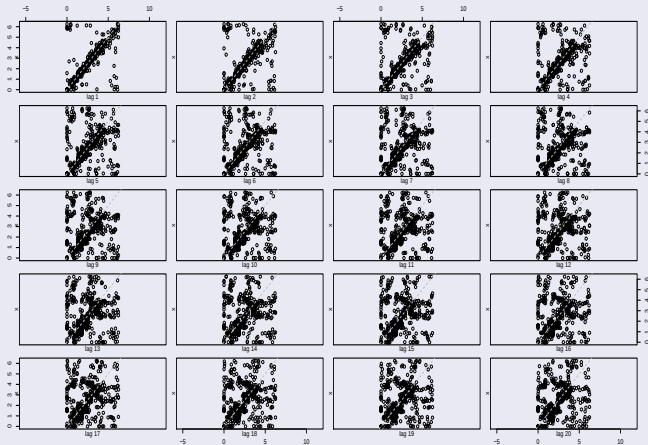
## Febrero



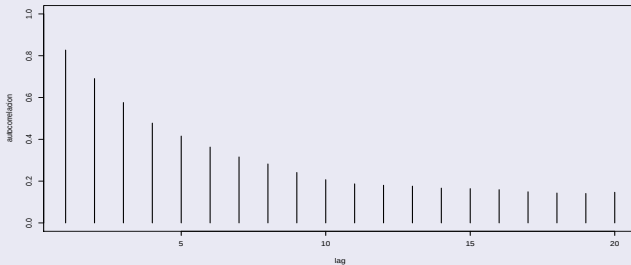
## Marzo



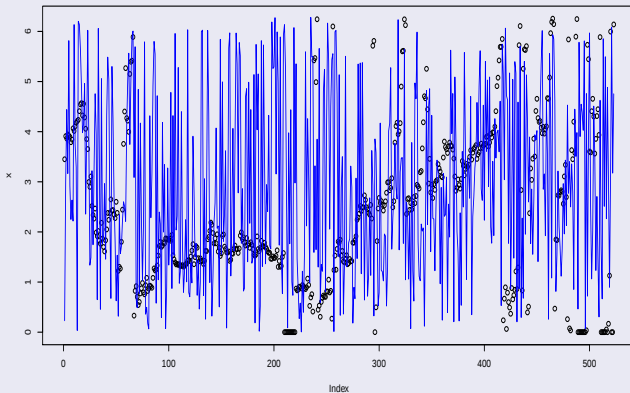
## Marzo



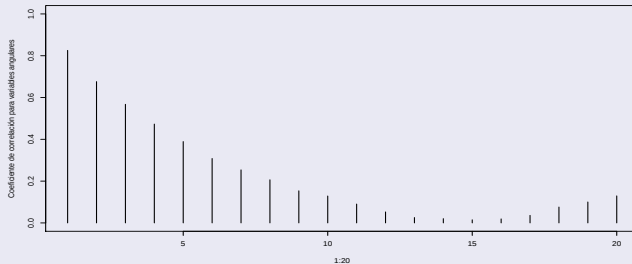
## Marzo



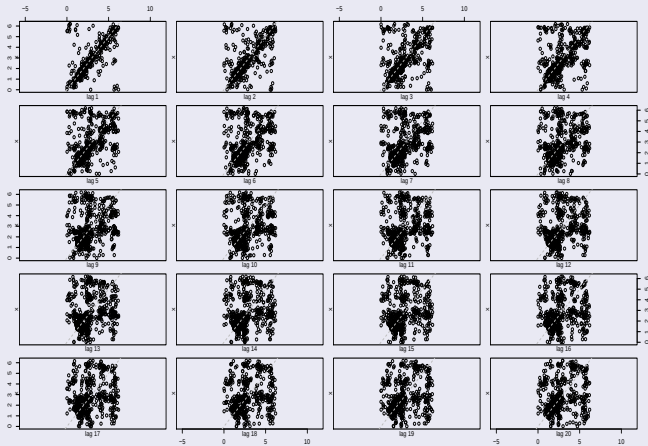
## Marzo



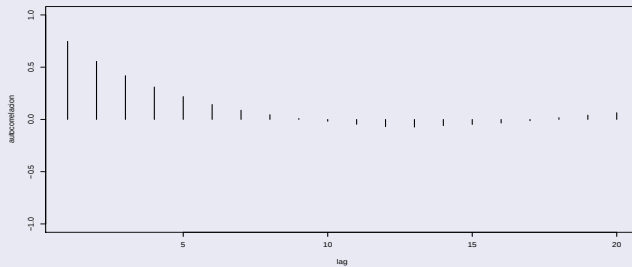
Abril



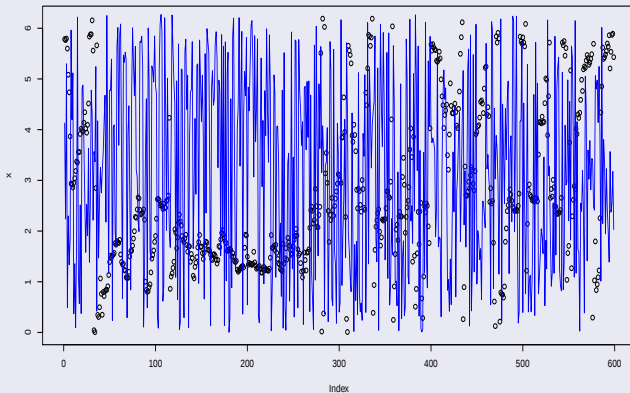
## Abril



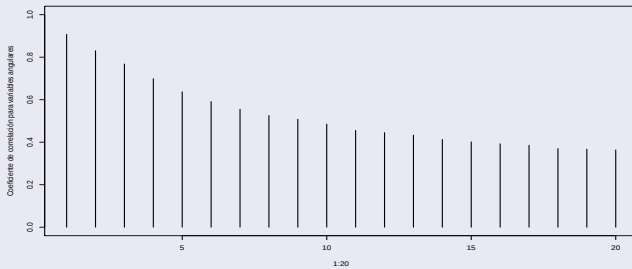
Abril



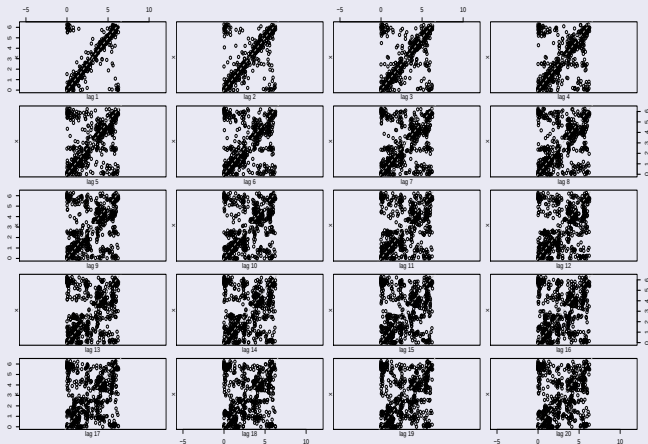
Abril



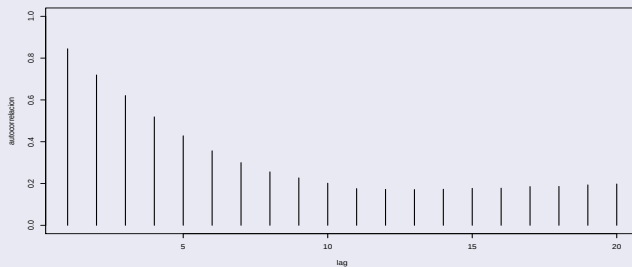
Mayo



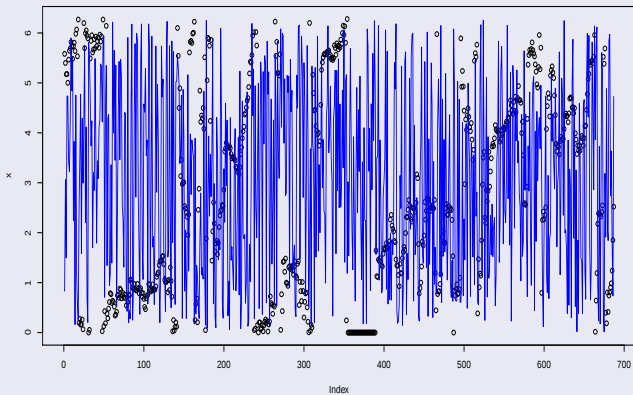
# Mayo



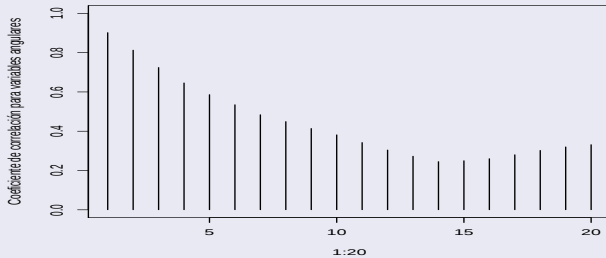
Mayo



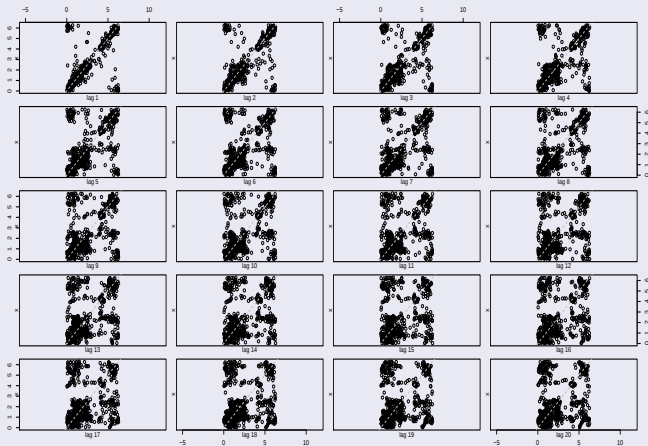
Mayo



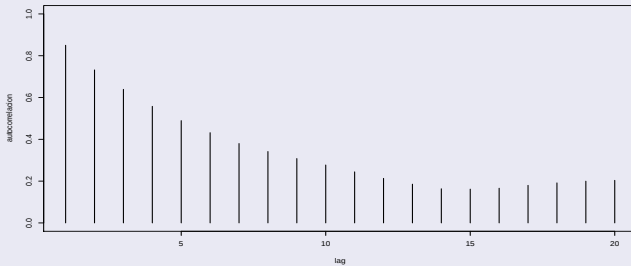
## Junio



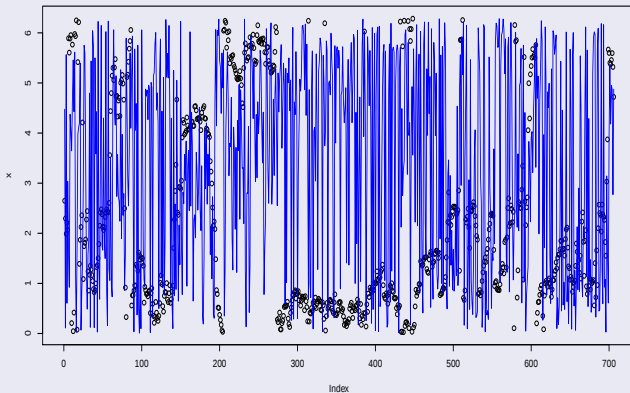
## Junio



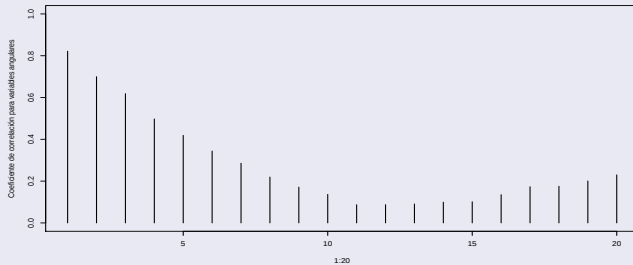
Junio



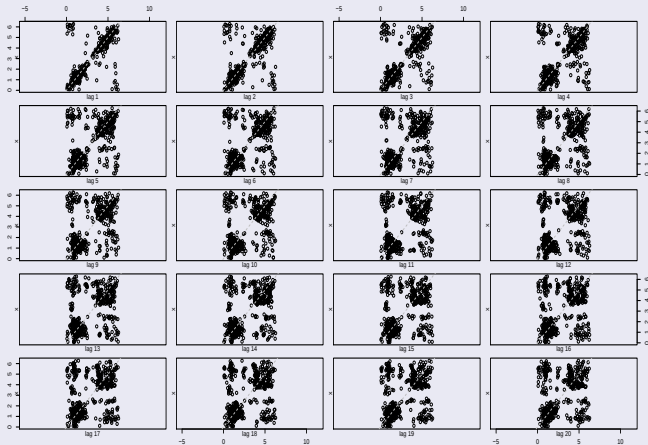
## Junio



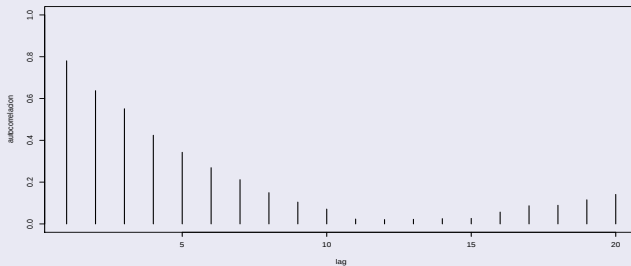
## Julio



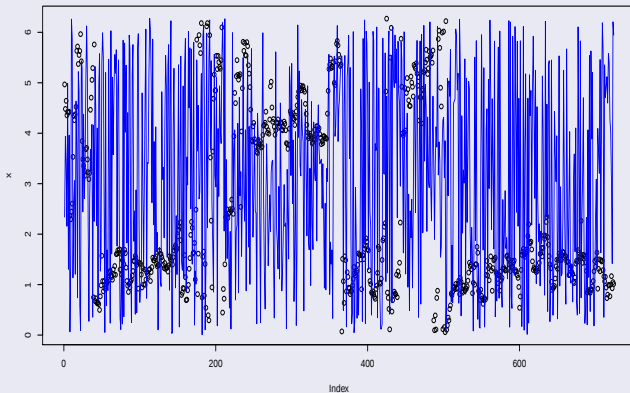
## Julio



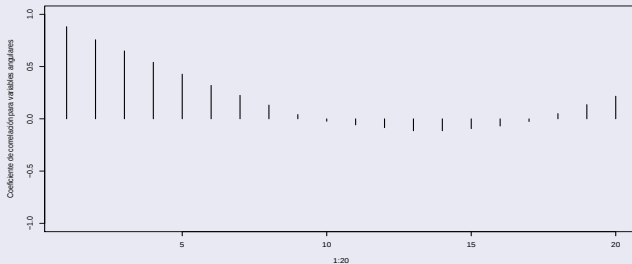
Julio



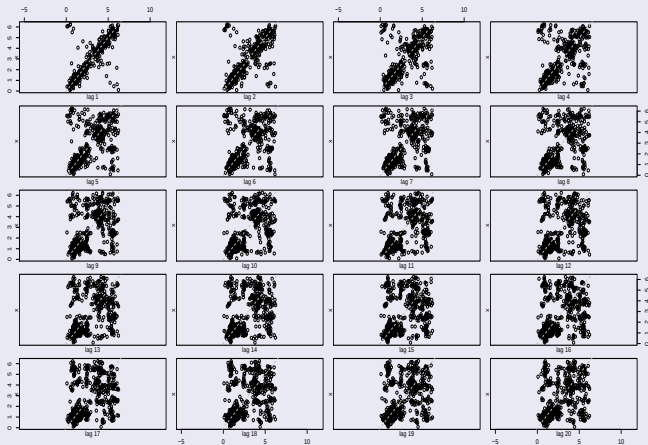
## Julio



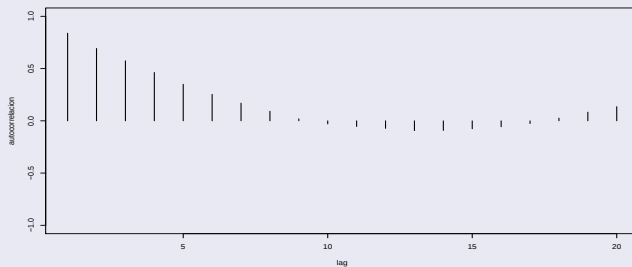
## Agosto



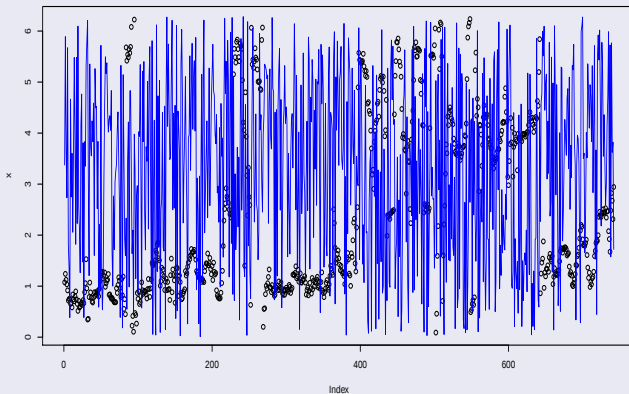
# Agosto



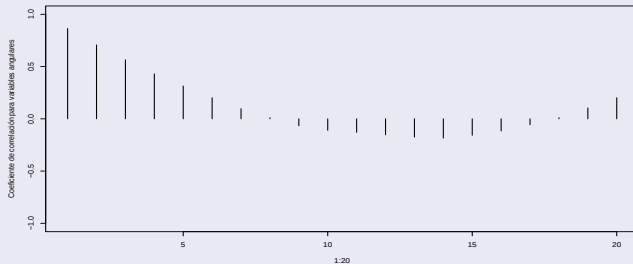
## Agosto



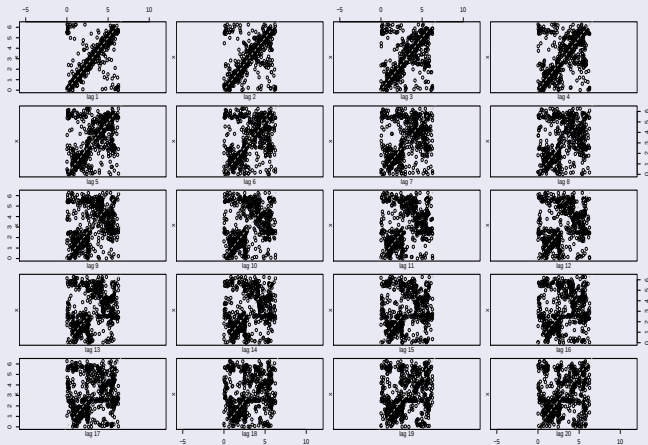
## Agosto



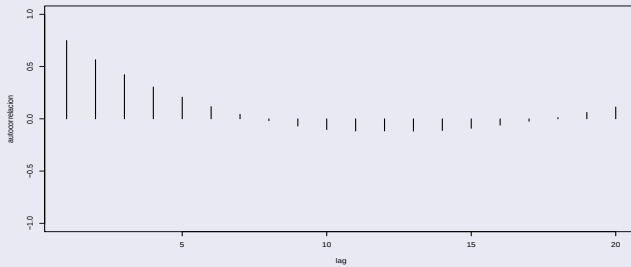
## Septiembre



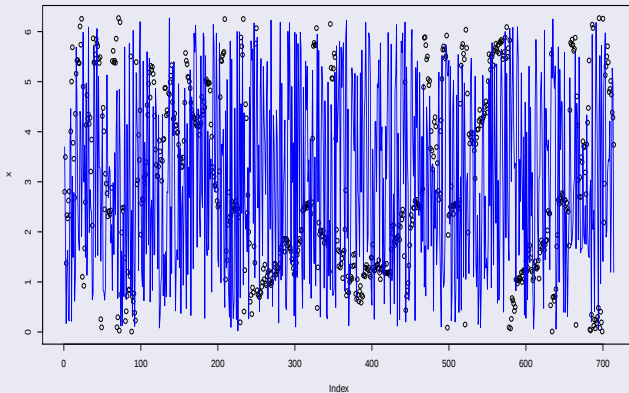
## Septiembre



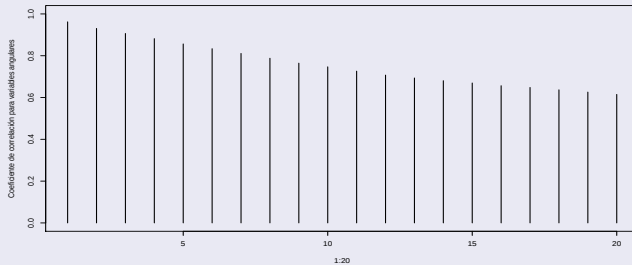
## Septiembre



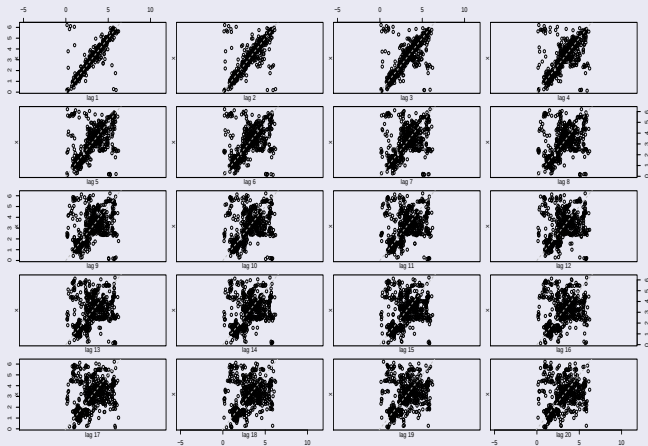
## Septiembre



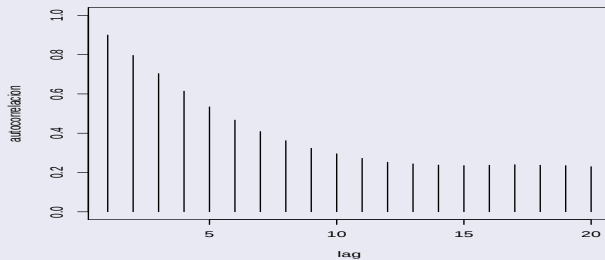
## Octubre



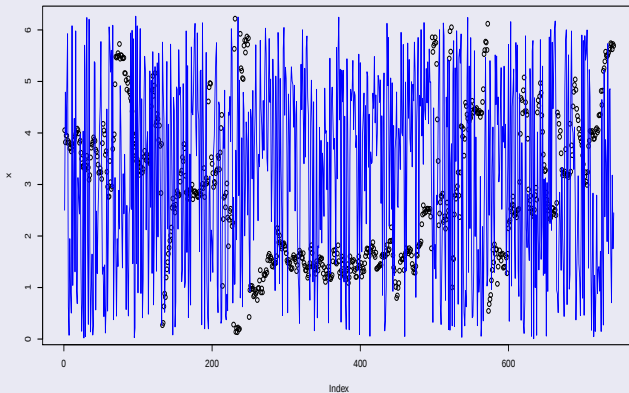
# Octubre



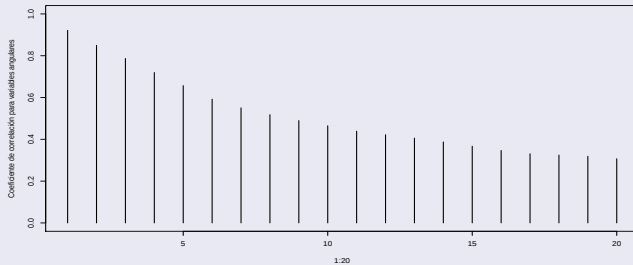
## Octubre



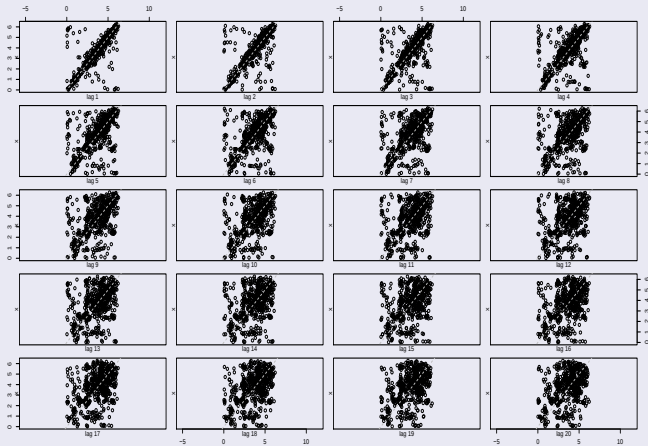
## Octubre



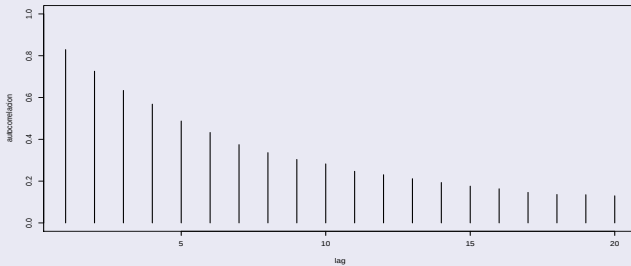
## Noviembre



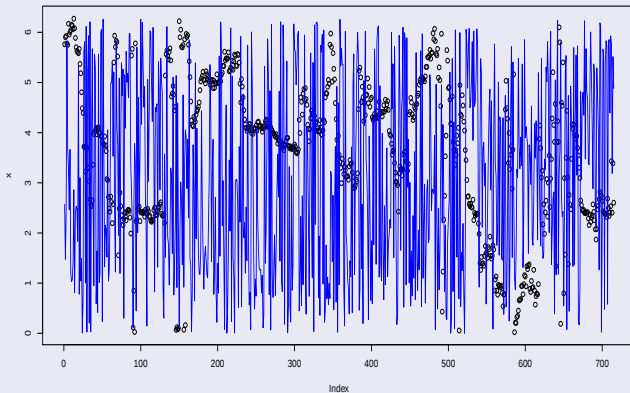
## Noviembre



## Noviembre



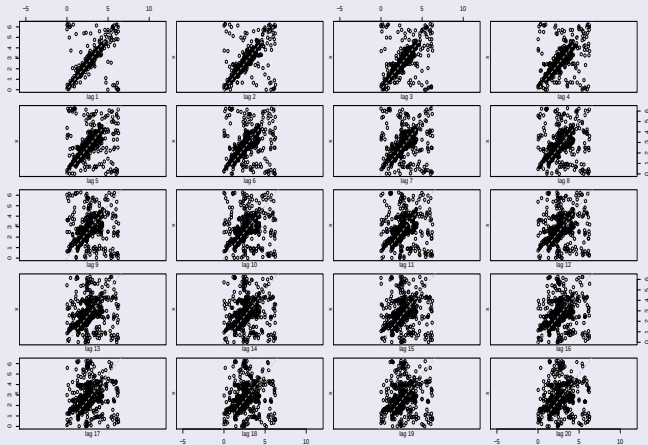
## Noviembre



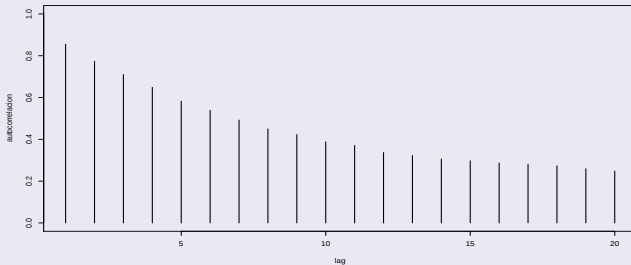
## Diciembre



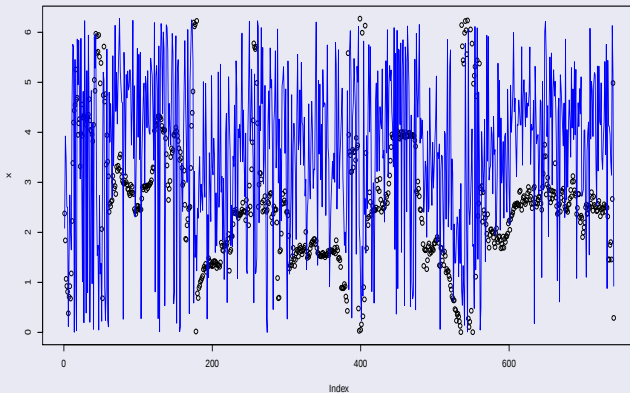
## Diciembre



## Diciembre



## Diciembre



## El proceso Linked

- Fisher y Lee (1994), usan la idea de funciones *link* entre los modelos de series de tiempo lineales a un contexto direccional.
- Proceso estacionario direccional  $\{\theta_t\}$  con media direccional con media  $\mu$  es un proceso *link* autoregresivo de medias móviles (LARMA); si y sólo si el proceso  $\{x_t\} = \{g^{-1}(\theta_t - \mu)\}$  es un proceso autoregresivo de medias móviles (ARMA)
- Dónde  $g$  es par, es una función monótona  $(-\pi, \pi]$  y  $g(0)=0$
- k-lag función circular de autocorrelación de  $\{\theta_t\}$  es dado como

$$\rho_T(k) = \rho_T \{g(X_t), g(X_{t+k})\}$$

$$\rho_T = \frac{E[\text{sen}(\Theta_1 - \Theta_2)\text{sen}(\Phi_1 - \Phi_2)]}{\{E[\text{sen}^2(\Theta_1 - \Theta_2)]E[\text{sen}^2(\Phi_1 - \Phi_2)]\}^{1/2}}$$

# Datos direccionales. Densidad.

Leyenda Rodríguez, María

4 de abril de 2011