

# Generalised Additive Models for Location Scale and Shape

Mikis Stasinopoulos<sup>1</sup>   Bob Rigby<sup>1</sup>

<sup>1</sup>STORM, London Metropolitan University

GAMLSS: Short course Cordoba, Argentina, Oct 2012



# Summary of the GAMLSS course

## Day 1:

- Introduction to GAMLSS
- introduction to the GAMLSS software

## Day 2:

- Continuous distributions in GAMLSS
- Discrete, mixed and finite mixture distributions in GAMLSS

## Day 3:

- Smoothing and random effects within GAMLSS
- Centile estimation and further work

- 1 Motivating Examples
  - Centile estimation
  - The fish species data
  - A stylometric application
  - The film data
- 2 The statistical modelling philosophy
  - Introduction
  - What is GAMLSS?
- 3 Regression: recent history
  - Regression type models
  - The linear model
  - The generalised linear model
  - The generalised additive model
- 4 generalised additive models for location scale and shape
  - The model
- 5 GAMLSS: the current state
  - Distributions
  - The additive terms
  - Model estimation

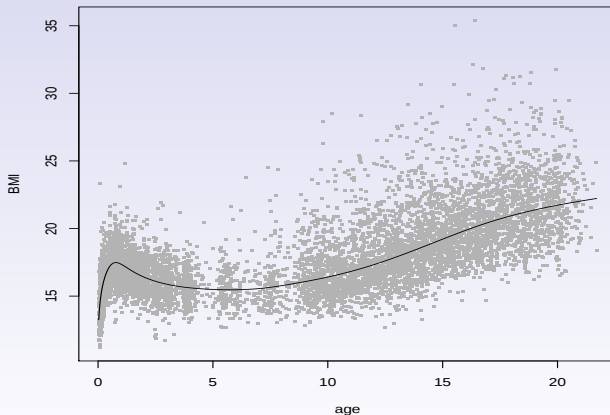
# The Dutch boys data

**BMI** : the BMI of 7294 boys

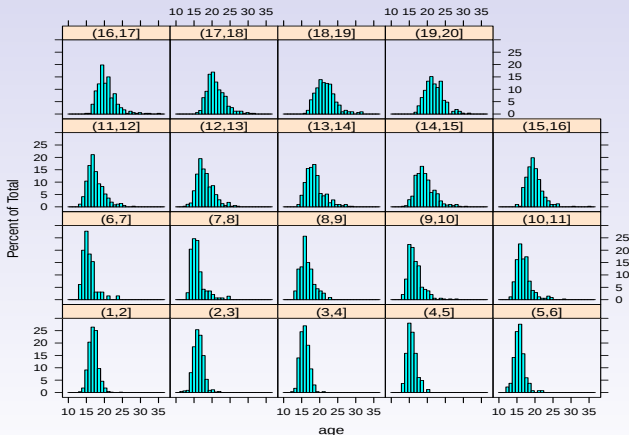
**age** : the age in years

Source: Buuren and Fredriks (2001)

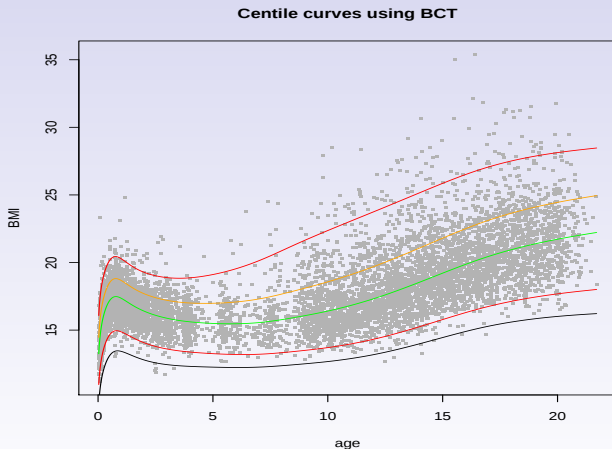
# The Dutch boys data: statistical challenges



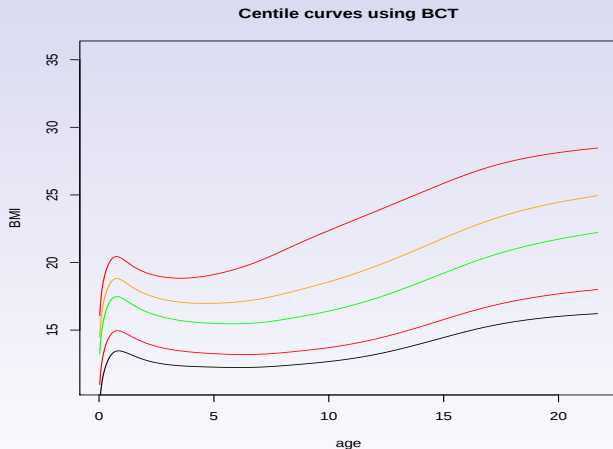
# The Dutch boys data: Histograms by age



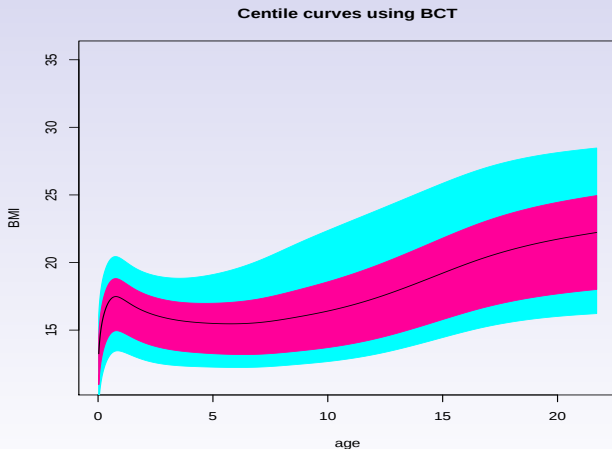
# The Dutch boys data: centile estimation



# The Dutch boys data: centiles



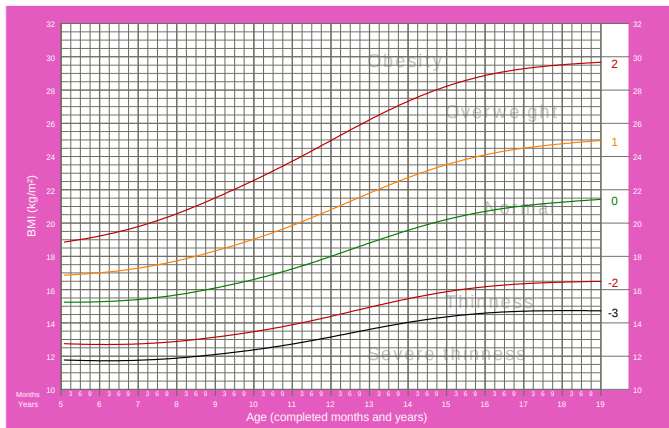
# The Dutch boys data: fan plot



# World Health Organisation Child Growth Standards: Girls

## BMI-for-age GIRLS

5 to 19 years (z-scores)

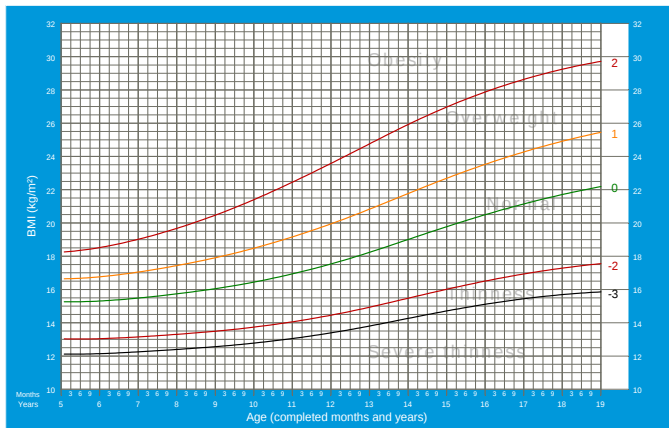


gamlss

# World Health Organisation Child Growth Standards: Boys

## BMI-for-age BOYS

5 to 19 years (z-scores)



gamlss

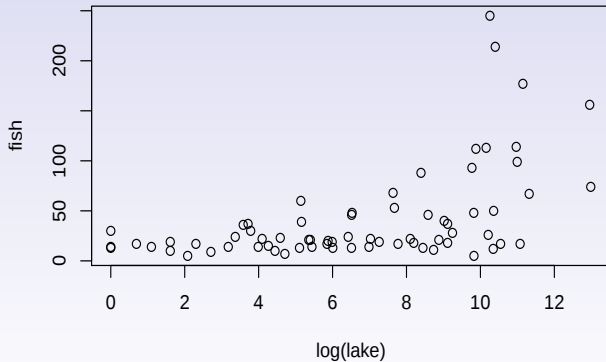
## The fish species data

`fish` : the number of different species in 70  
lakes in the world

`lake` : the lake area

Source: Stein and Juritz (1988)

# The fish species data



## A stylometric application

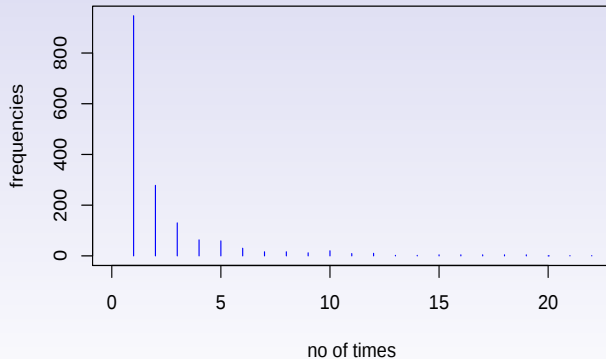
64 observations

**word** : is the number of times a word appears in  
a single text

**freq** : the number of different words which occur exactly  
word times in the text

Source: Dr Mario Cortina-Borja

# The stylometric data



# The film data

4031 film openings in US 1988-1999:

`lborev1` : the response variable

log of the (total revenues - the first week revenue),  
calculated in 1987 prices

`lbopen` : the log the first week revenue, calculated in 1987  
prices

`lnosc` : the log of the number of screens

`dist` : the type of distributor i.e. whether independent or  
major distributor

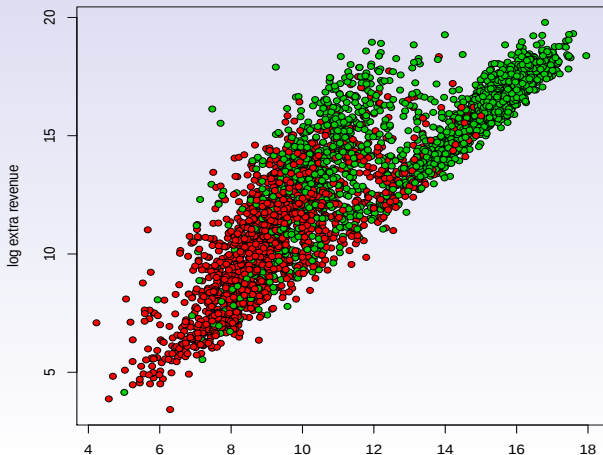
Source: Voudouris *et al.* (2010)

[Go to Quiz](#)

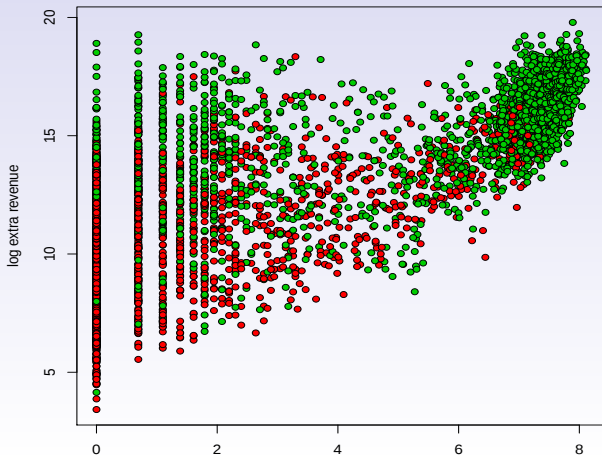
## The film data: Question of interest

- Checking whether the “nobody knows anything” principle is valid
- Can we predict the extra revenue given the revenue at the first week?
- Which other explanatory variables are important?
- What is the conditional distribution of the response variable given the explanatory variables?

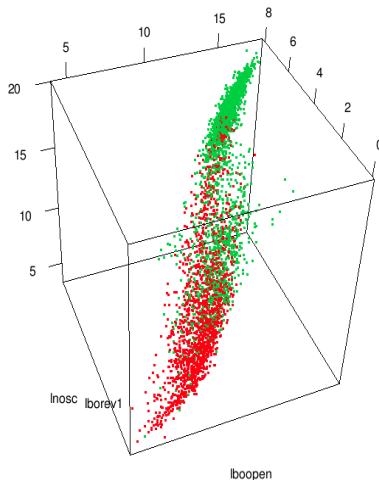
# The film data: log extra revenue against log first week revenue



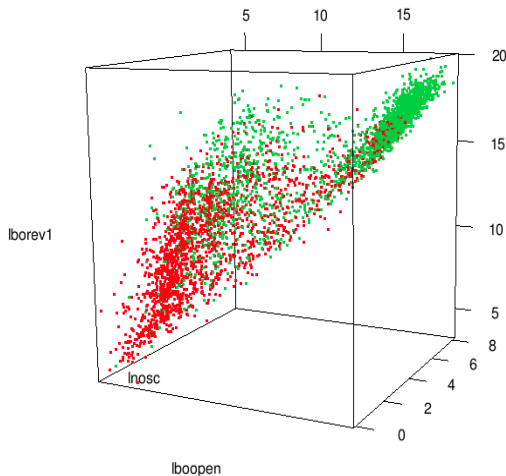
# The film data: log extra revenue against log number of screens



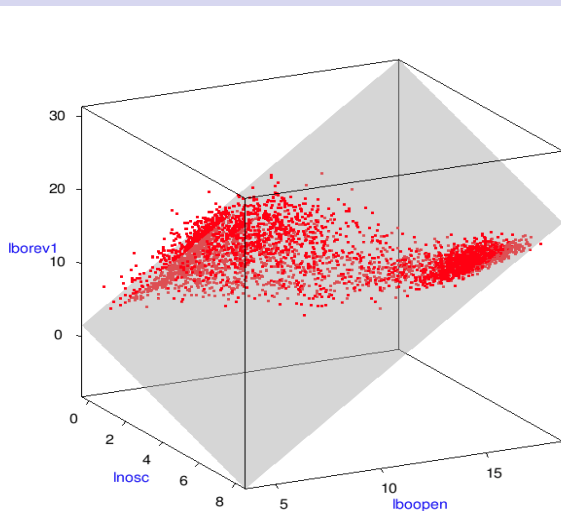
# The film data: 3d plots



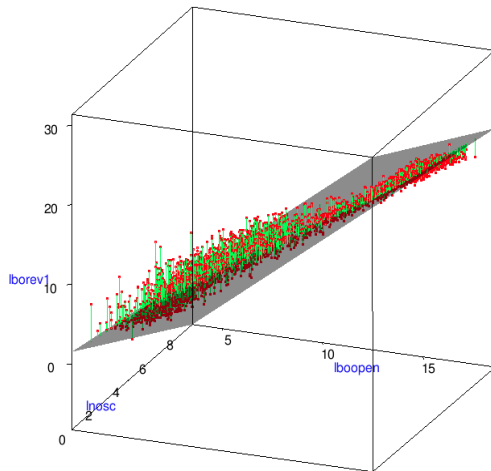
# The film data: 3d plots



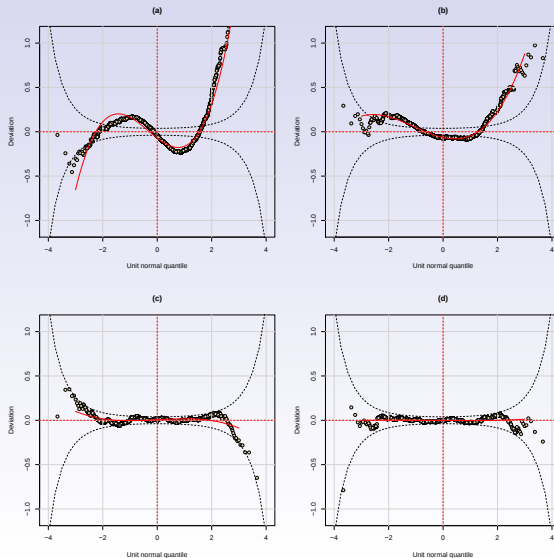
# The film data: Least Square fit



# The film data: the residuals from least squares fit



# The film data: worm plot of the residuals



# What we need for modelling the above data?

We need

- flexible distributions for the response variable
- to be able to deal with heterogeneity in the data
- to be able to model skewness and kurtosis
- to be able to model overdispersion in count data
- We need modelling all the parameters of the distributions
- flexible functions to model the relationship between the parameter of the distribution and the explanatory variables

# The statistical modelling philosophy

*Statistical modelling* is the art of building **parsimonious** statistical models for a better understanding of the phenomena of interest.

- Any **model** is a simplification of reality therefore **no model is correct** but some of them are useful
- **Occam's Razor** which states '*entities should not be multiplied beyond necessity*' or **KISS** (Keep It Simple Stupid)

# The statistical modelling philosophy

- Far better an **approximate** answer to the **right** question, which is often vague, than an exact answer to the wrong question, which can always be made precise. – John W. Tukey (1962), "The future of data analysis", Annals of Mathematical Statistics 33, 1-67 (page 13)
- *"no matter how beautiful your theory, no matter how clever you are or what your name is, if it disagrees with experiment, it's wrong"* (Richard Feynman)
- As statisticians we should try **different** models and choose the one appropriate for the data

# What is GAMLSS?

**GAMLSS**: are semi-parametric regression type models.

- **regression type**: we have many explanatory variables  $\mathbf{X}$  and one response variable  $\mathbf{y}$  and we believe that  $\mathbf{X} \rightarrow \mathbf{y}$
- **parametric**: a parametric distribution assumption for the response variable,
- **semi**: the parameters of the distribution, as functions of explanatory variables, may involve non-parametric smoothing functions
- GAMLSS philosophy: try different models

GAMLSS is a generalization of GLM and GAM models.

# Regression type models

$$\text{input} \longrightarrow \text{output} \quad (1)$$

$$X \longrightarrow Y \quad (2)$$

$X$  : *input variable(s), explanatory, independent variable(s)*

$Y$  : *response variable, dependent variable, target, y-variable, output*

one response variable is *univariate* analysis otherwise *multivariate*

# The linear model

## Linear regression model

$$\mathbf{y} = \mathbf{X}\beta + \epsilon \text{ where } \epsilon \sim N(\mathbf{0}, \sigma^2)$$

$$\mu = \mathbf{X}\beta \text{ where } \mathbf{y} \sim N(\mu, \sigma^2)$$

## Assumptions

- 1 The  $X$ 's effect  $Y$  linearly
- 2 The error terms are independently distributed random variables
- 3 The variance of  $\epsilon$  and therefor  $Y$  is constant

# The linear model assumptions

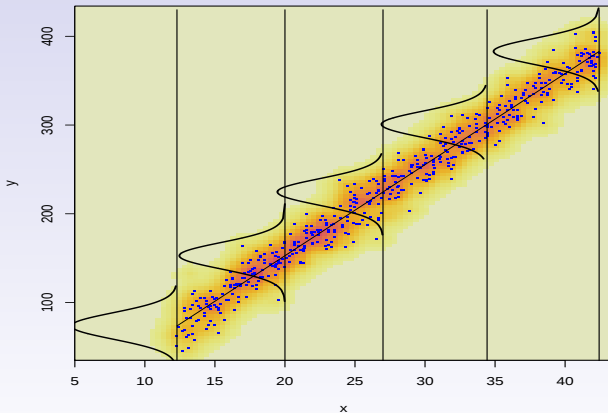


Figure: The simple regression assumptions

# The estimation of the linear model

Least Squares:

$$\boldsymbol{\epsilon}^\top \boldsymbol{\epsilon} = (\mathbf{Y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}) = \sum_{i=1}^n \epsilon_i^2.$$

$$\hat{\boldsymbol{\beta}} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{Y}$$

Maximum Likelihood

$$\ell(\boldsymbol{\beta}, \sigma^2) = \frac{n}{2} \log(2\pi\sigma^2) - \frac{1}{2\sigma^2} (\mathbf{Y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{Y} - \mathbf{X}\boldsymbol{\beta})$$

$$\hat{\sigma}^2 = \frac{(\mathbf{Y} - \mathbf{X}\hat{\boldsymbol{\beta}})^\top (\mathbf{Y} - \mathbf{X}\hat{\boldsymbol{\beta}})}{n}.$$

# The weighted linear model

## Weighted Linear regression model

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon} \text{ where } \boldsymbol{\epsilon} \sim N(\mathbf{0}, \boldsymbol{\Sigma})$$

$$\boldsymbol{\mu} = \mathbf{X}\boldsymbol{\beta} \text{ where } \mathbf{y} \sim N(\boldsymbol{\mu}, \boldsymbol{\Sigma})$$

$$\text{often } \boldsymbol{\Sigma} = \sigma^2 \mathbf{W}^{-1}$$

where  $\mathbf{W}$  is a diagonal matrix

$$\hat{\boldsymbol{\beta}} = \left( \mathbf{X}^\top \mathbf{W} \mathbf{X} \right)^{-1} \mathbf{X}^\top \mathbf{W} \mathbf{y}. \quad (3)$$

# The weighted linear model assumptions

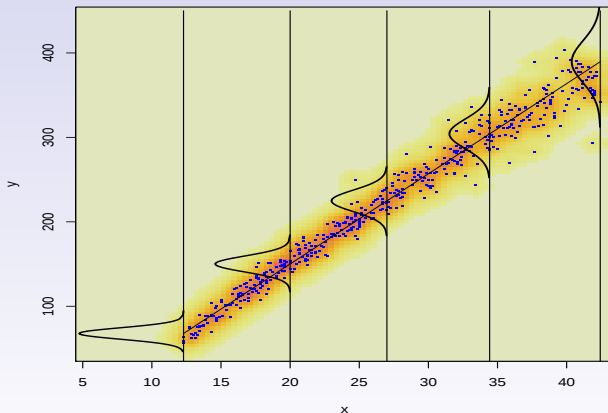


Figure: The weighted least squares regression assumptions 

# The generalised linear model

generalised Linear Model, Wedderburn and Nelder (1972)

$$g(\mu) = \mathbf{X}\beta \text{ where } \mathbf{y} \sim \text{ExpFamily}(\mu, \phi)$$

where  $g()$  is the link function

The exponential family

- ① normal
- ② Gamma
- ③ inverse Gaussian
- ④ Poisson
- ⑤ binomial

# The generalised linear model

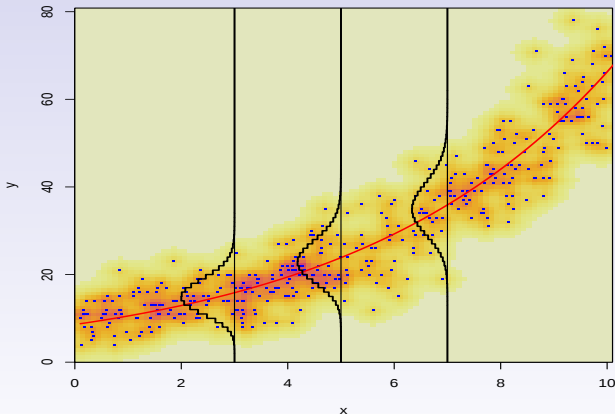


Figure: The Poisson regression assumptions

# The generalised linear model: parameter estimation

Maximize the log-Likelihood with respect to  $\beta$ .

$$\ell(\theta) = \log L(\theta) = \sum_{i=1}^n \log f(y_i | \theta)$$

This is done using Iterative Weighted Least Squares

$$\hat{\beta}^{(i)} = \left( \mathbf{X}^\top \mathbf{W}^{(i)} \mathbf{X} \right)^{-1} \mathbf{X}^\top \mathbf{W}^{(i)} \mathbf{z}^{(i)}. \quad (4)$$

where  $\mathbf{z}^{(i)}$  iterative response variable and  $\mathbf{W}^{(i)}$  are the iterative weights.

# The generalised additive model

generalised additive model Hastie and Tibshirani (1990)

$$g(\boldsymbol{\mu}) = \mathbf{X}\boldsymbol{\beta} + h_1(\mathbf{x}_1) + \dots + h_k(\mathbf{x}_k) \text{ where } \mathbf{y} \sim \text{ExpFamily}(\boldsymbol{\mu}, \phi)$$

where  $h_j(\mathbf{x}_j)$  are smooth functions of the  $X$ 's.

To estimate the parameters  $\boldsymbol{\beta}$  and the smoothing functions  
maximize the penalised log-Likelihood

$$\ell_p = \ell - \frac{1}{2} \sum_{j=1}^{J_k} \lambda_j \mathbf{Q}_j(\mathbf{h}_j)$$

where  $\mathbf{Q}_j(\mathbf{h}_j)$  is a quadratic function of the  $h_j$ 's

# The generalised additive model

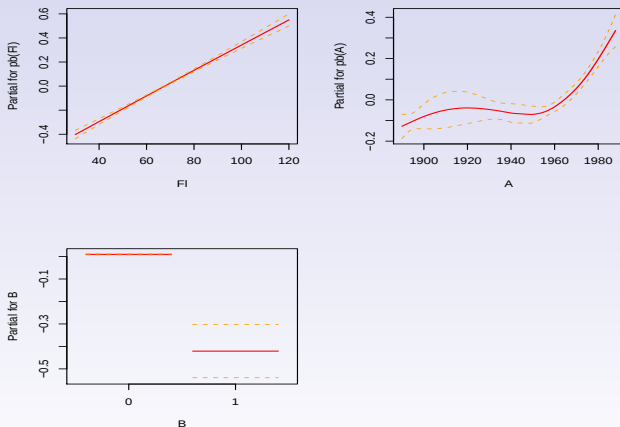
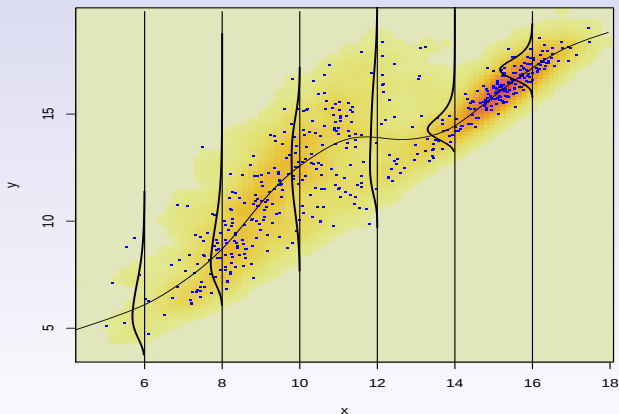


Figure: A generalised additive model fit

## GAMLSS: the model



## GAMLSS: the model

Let  $Y$  is the response variable

$$Y_i \sim f_Y(y_i | \theta_1, \theta_2, \theta_3, \theta_4) = f_Y(y_i | \mu_i, \sigma_i, \nu_i, \tau_i)$$

where  $f_Y()$  can be **any** distribution (non only the exponential family).

**any** parameter  $\theta_i$  can be a function of explanatory variables.

## GAMLSS: the model

$$g_1(\mu) = \eta_1 = \mathbf{X}_1 \beta_1 + \sum_{j=1}^{J_1} h_{j1}(x_{j1})$$

$$g_2(\sigma) = \eta_2 = \mathbf{X}_2 \beta_2 + \sum_{j=1}^{J_2} h_{j2}(x_{j2})$$

$$g_3(\nu) = \eta_3 = \mathbf{X}_3 \beta_3 + \sum_{j=1}^{J_3} h_{j3}(x_{j3})$$

$$g_4(\tau) = \eta_4 = \mathbf{X}_4 \beta_4 + \sum_{j=1}^{J_4} h_{j4}(x_{j4}).$$

# GAMLSS: the model (more general)

$$g_1(\mu) = \eta_1 = \mathbf{x}_1\beta_1 + \sum_{j=1}^{J_1} \mathbf{z}_{j1}\gamma_{j1}$$

$$g_2(\sigma) = \eta_2 = \mathbf{x}_2\beta_2 + \sum_{j=1}^{J_2} \mathbf{z}_{j2}\gamma_{j2}$$

$$g_3(\nu) = \eta_3 = \mathbf{x}_3\beta_3 + \sum_{j=1}^{J_3} \mathbf{z}_{j3}\gamma_{j3}$$

$$g_4(\tau) = \eta_4 = \mathbf{x}_4\beta_4 + \sum_{j=1}^{J_4} \mathbf{z}_{j4}\gamma_{j4}.$$

## GAMLSS: the random effects

where each  $\gamma_{jk}$  is a random variable

$$\gamma_{jk} \sim N_{q_{jk}}(\mathbf{0}, \mathbf{G}_{jk}^{-1})$$

where  $\mathbf{G}_{jk}^{-1}$  is the (generalised) inverse of a symmetric matrix

$$\mathbf{G}_{jk} = \mathbf{G}_{jk}(\lambda_{jk})$$

hyperparameters:  $\lambda_{jk}$

## GAMLSS: submodels

*semi-parametric additive*

$$g_k(\theta_k) = \eta_k = \mathbf{X}_k \beta_k + \sum_{j=1}^{J_k} h_{jk}(\mathbf{x}_{jk})$$

*parametric linear*

$$g_k(\theta_k) = \eta_k = \mathbf{X}_k \beta_k$$

*non-linear parametric + additive*

$$g_k(\theta_k) = \eta_k = h_k(\mathbf{X}_k, \beta_k) + \sum_{j=1}^{J_k} h_{jk}(\mathbf{x}_{jk})$$

# GAMLSS: Distributions

There are around 70 **discrete** `discrete`, **continuous** `continuous`, and **mixed** `mixed` distributions, implemented as `gamlss.family` in the R including highly skew and kurtotic distributions `Shapes`,

- creating a **new** distribution is relatively easy
- **truncating** `truncated` an existing distribution
- using a **censored** version of an existing distribution
- **mixing** `mixture` different distributions to create a new finite mixture distribution.
- **discretise** `discretise` continuous distributions
- **log** or **logic** any continuous distribution in  $(-\infty, \infty)$

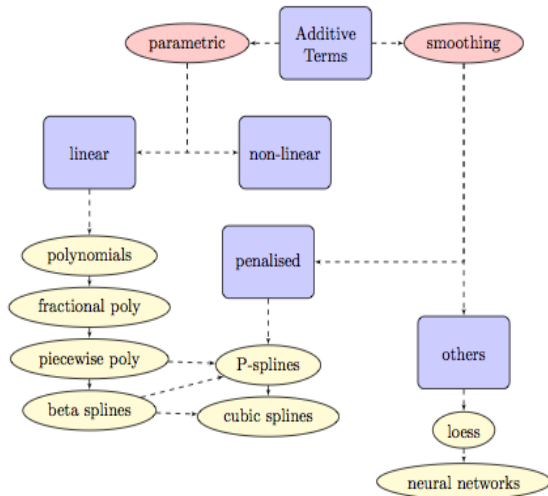
next page Additive

# GAMLSS: additive terms

| Additive terms                     | R function names                                                               |
|------------------------------------|--------------------------------------------------------------------------------|
| boosting                           | <code>boost()</code>                                                           |
| cubic splines based                | <code>cs()</code> , <code>scs()</code>                                         |
| decision trees                     | <code>tr()</code>                                                              |
| fractional and power polynomials   | <code>fp()</code> , <code>pp()</code>                                          |
| free knot smoothing (break points) | <code>fk()</code>                                                              |
| Penalised lgs                      | <code>la()</code>                                                              |
| loess                              | <code>lo()</code>                                                              |
| neural networks                    | <code>nn()</code>                                                              |
| non-linear fit                     | <code>nl()</code>                                                              |
| penalized Beta splines based       | <code>pb()</code> , <code>ps()</code> , <code>cy()</code> , <code>pvc()</code> |
| random effects                     | <code>random()</code>                                                          |
| ridge regression                   | <code>ri()</code> , <code>ridge()</code>                                       |
| Simon Wood's gam                   | <code>ga()</code>                                                              |

**Table:** Additive terms implemented within the **gamlss** packages

# GAMLSS: additive terms



# What is GAMLSS?: The quantities involved

- $y$  response variable
- $\mathbf{X} = (\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_p)$  fixed effect design matrices
- $\mathbf{Z} = (\mathbf{Z}_1, \mathbf{Z}_2, \dots, \mathbf{Z}_p)$  random effect design matrices
- $\boldsymbol{\beta}^\top = (\boldsymbol{\beta}_1^\top, \boldsymbol{\beta}_2^\top, \dots, \boldsymbol{\beta}_p^\top)$  fixed effect parameters
- $\boldsymbol{\gamma}^\top = (\boldsymbol{\gamma}_1^\top, \boldsymbol{\gamma}_2^\top, \dots, \boldsymbol{\gamma}_p^\top)$  random effect parameters
- $\boldsymbol{\lambda}^\top = (\lambda_{11}^\top, \lambda_{21}^\top, \dots, \lambda_{J_p p}^\top)$  hyperparameters

# Model estimation

## Model estimation in GAMLSS

- case 1 (Linear or non-linear parametric) GAMLSS model i.e  $\beta$ 's but no  $\gamma$ 's or  $\lambda$ 's
- case 2 Random effects or semi-parametric GAMLSS model i.e  $\beta$ 's,  $\gamma$ 's and  $\lambda$ 's

# Model estimation

In [case 1](#) the parametric case: GAMLSS estimates the  $\beta$ 's by

maximum likelihood estimation i.e. maximizing the log-likelihood of the data.

$$\ell = \sum_{i=1}^n \log f(y_i | \mu_i, \sigma_i, \nu_i, \tau_i)$$

with respect to  $\beta$ 's.

# Model estimation

In [case 2](#) the random effect or semi-parametric case:

The main GAMLSS algorithm estimates  $(\beta, \gamma)$ 's for given  $\lambda$ 's by maximizing the penalized log-likelihood (equivalent to posterior mode estimation (MAP) assuming a uniform prior for  $\beta$ 's):

$$\ell_p = \ell - \frac{1}{2} \sum_{k=1}^p \sum_{j=1}^{J_k} \lambda_{jk} \gamma'_{jk} \mathbf{G}_{jk} \gamma_{jk}$$

# Model estimation

MAP estimation of  $(\beta, \gamma)$  given  $\lambda$

MAP=maximum *a posteriori*= posterior mode

$$f(\beta, \gamma | \mathbf{y}, \lambda) \propto f(\mathbf{y} | \beta, \gamma) f(\gamma | \lambda)$$

assuming a uniform prior for  $\beta$

# Model estimation

$$\begin{aligned}\log f(\boldsymbol{\beta}, \boldsymbol{\gamma} | \mathbf{y}, \boldsymbol{\lambda}) &= \log f(\mathbf{y} | \boldsymbol{\beta}, \boldsymbol{\gamma}) + \log f(\boldsymbol{\gamma} | \boldsymbol{\lambda}) + \mathbf{c}_1(\mathbf{y}, \boldsymbol{\lambda}) \\ &= \ell_h + \mathbf{c}_1(\mathbf{y}, \boldsymbol{\lambda}) \\ &= \log f(\mathbf{y} | \boldsymbol{\beta}, \boldsymbol{\gamma}) - \frac{1}{2} \sum_{k=1}^p \sum_{j=1}^{J_k} \lambda_{jk} \boldsymbol{\gamma}'_{jk} \mathbf{G}_{jk} \boldsymbol{\gamma}_{jk} + \mathbf{c}_2(\mathbf{y}, \boldsymbol{\lambda}) \\ &= \ell_p + \mathbf{c}_2(\mathbf{y}, \boldsymbol{\lambda})\end{aligned}$$

since  $\boldsymbol{\gamma}_{jk} \sim N(\mathbf{0}, \mathbf{G}_{jk}^{-1})$  and assuming a uniform prior for  $\boldsymbol{\beta}$

# Model estimation

MAP estimation of  $(\beta, \gamma)$  given  $\lambda$

Hence given  $\lambda$

MAP estimation of  $(\beta, \gamma)$

$\equiv$  maximizing  $\ell_h$ , hierarchical log likelihood

$\equiv$  maximizing  $\ell_p$  penalized log likelihood

with respect to  $(\beta, \gamma)$

# Model estimation

Estimation of the hyper parameters  $\lambda$

- maximizing a profile GAIC over  $\lambda$
- maximizing the marginal likelihood of  $\lambda$

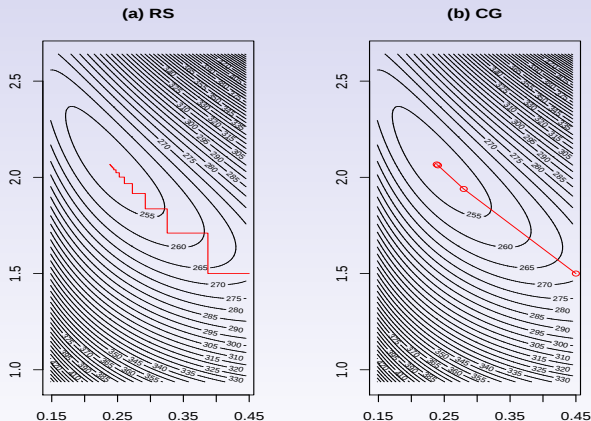
# Algorithms

**RS** a generalization of the MADAM algorithm, Rigby and Stasinopoulos (1996a)

**CG** a generalization of Cole and Green (1992) algorithm

**mixed** a mixture of RS+CG (i.e.  $j$  iterations of RS, followed by  $k$  iterations of CG)

# Algorithms



# Algorithms: properties

## Advantages of algorithms

- ① flexible modular fitting procedure
- ② easy implementation of new distributions
- ③ easy implementation of new additive terms
- ④ simple starting values for  $(\mu, \sigma, \nu, \tau)$  easily found
- ⑤ stable and reliable algorithms
- ⑥ fast fitting (for fixed hyperparameters)

## Algorithms: details

CG algorithm for MAP estimation of  $(\boldsymbol{\beta}, \boldsymbol{\gamma})$

The  $r^{th}$  iteration of the Newton Raphson algorithm gives:

$$\begin{bmatrix} \mathbf{D}_{11} & \mathbf{D}_{12} \\ \mathbf{D}_{21} & \mathbf{D}_{22} \end{bmatrix}^{(r)} \begin{bmatrix} \begin{pmatrix} \boldsymbol{\beta} \end{pmatrix}^{(r+1)} \\ \begin{pmatrix} \boldsymbol{\gamma} \end{pmatrix}^{(r+1)} \end{bmatrix} - \begin{bmatrix} \begin{pmatrix} \boldsymbol{\beta} \end{pmatrix}^{(r)} \\ \begin{pmatrix} \boldsymbol{\gamma} \end{pmatrix}^{(r)} \end{bmatrix} = \begin{bmatrix} \frac{\partial \ell_h}{\partial \boldsymbol{\beta}} \\ \frac{\partial \ell_h}{\partial \boldsymbol{\gamma}} \end{bmatrix}^{(r)}$$

where  $\boldsymbol{\beta}^\top = (\boldsymbol{\beta}_1^\top, \boldsymbol{\beta}_2^\top, \dots, \boldsymbol{\beta}_p^\top)$ ,  $\boldsymbol{\gamma}^\top = (\boldsymbol{\gamma}_1^\top, \boldsymbol{\gamma}_2^\top, \dots, \boldsymbol{\gamma}_p^\top)$  and  $\ell_h$  is the hierarchical likelihood.

The Fisher scoring algorithm uses the expected information matrix instead of the observed information matrix  $\mathbf{D}$

## Algorithms : Updating $\beta_k$

Rearranging the row of the Fisher scoring equation for  $\beta_k$  gives

$$\mathbf{X}_k \beta_k^{(r+1)} = \mathbf{H}_k^{(r)} \epsilon_{ok}^{(r)}$$

updates  $\beta_k$  where

$$\mathbf{H}_k^{(r)} = \mathbf{X}_k \left( \mathbf{X}_k^T \mathbf{W}_{kk}^{(r)} \mathbf{X}_k \right)^{-1} \mathbf{X}_k^T \mathbf{W}_{kk}^{(r)}$$

i.e. a weighted regression of partial residuals  $\epsilon_{ok}^{(r)}$  against  $\mathbf{X}_k$

## Algorithms : Updating $\gamma_k$

Rearranging the row of the Fisher scoring equation for  $\gamma_{jk}$  gives

$$\mathbf{z}_{jk}\gamma_{jk}^{(r+1)} = \mathbf{s}_{jk}^{(r)}\epsilon_{jk}^{(r)}$$

updates  $\gamma_{jk}$  where

$$\mathbf{s}_{jk}^{(r)} = \mathbf{z}_{jk} \left( \mathbf{z}_{jk}^T \mathbf{w}_{kk}^{(r)} \mathbf{z}_{jk} + \mathbf{G}_{jk} \right)^{-1} \mathbf{z}_{jk}^T \mathbf{w}_{kk}^{(r)}$$

i.e. a generalised ridge regression of partial residuals  $\epsilon_{jk}^{(r)}$  against  $\mathbf{z}_{jk}$

# Algorithms : convergence

Note that both the CG and the RS algorithms update each

$(\beta_1, \beta_2, \dots, \beta_p)$  and

$(\gamma_{11}, \gamma_{21}, \dots, \gamma_{J_p, p})$

until convergence of  $(\beta^{(r+1)}, \gamma^{(r+1)})$

## Normalised quantile residuals-continuous response

If  $y_i$  is an observation from a continuous response variable then

$$\hat{u}_i = F(y_i | \hat{\theta}^i)$$

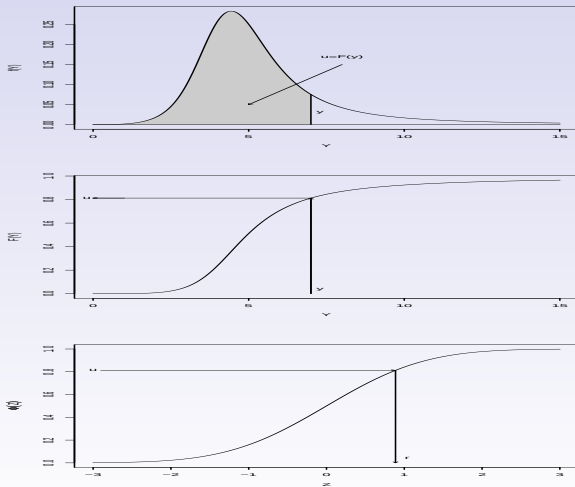
where  $u_i = F(y_i | \theta^i)$  is the assumed cumulative distribution function for case  $i$ .

$$r_i = \Phi^{-1}(u_i),$$

So  $r_i$  will have an approximately standard normal distribution

$r_i =$  z-scores

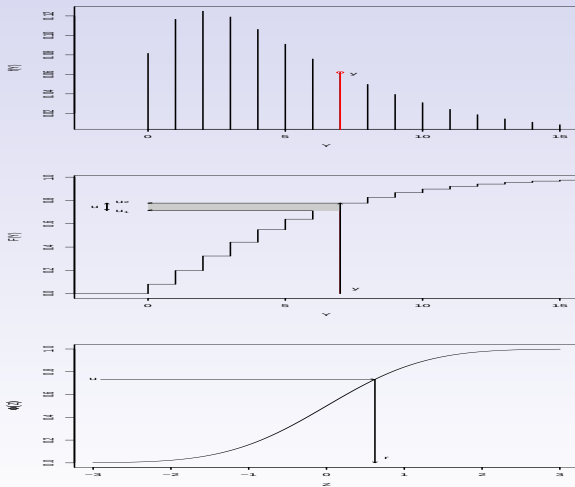
# Normalised quantile residuals-continuous response



## Normalised quantile residuals-discrete response

If  $y_i$  is an observation from a discrete integer response then  $\hat{u}_i$  is a random value from the uniform distribution on the interval  $[u_1, u_2] = [F(y_i - 1 | \hat{\theta}^i), F(y_i | \hat{\theta}^i)]$ .  
 $\hat{r}_i = \Phi^{-1}(\hat{u}_i)$ .

# Normalised quantile residuals-discrete response



# Conclusions

- Unified framework for univariate regression type of models
- Allows any distribution for the response variable  $Y$
- Models all the parameters of the distribution of  $Y$
- Allows a variety of additive terms in the models for the distribution parameters
- The fitted algorithm is modular, where different components can be added easily
- Models can be fitted easily and fast
- Explanatory tool to find appropriate set of models
- It deals with overdispersion, skewness and kurtosis

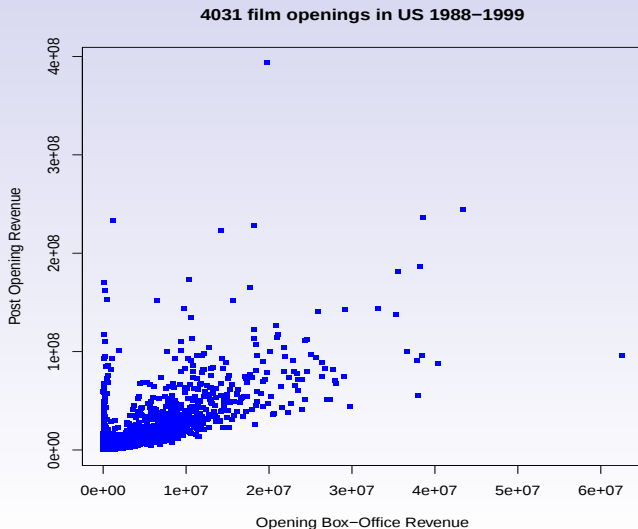
End

END

for more information see

[www.gamlss.org](http://www.gamlss.org)

# Film Data

[Go back](#)

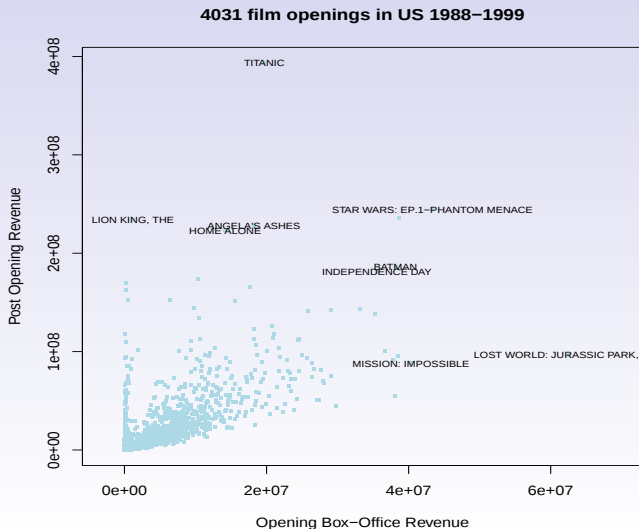
## Film Data: the quiz

### *A small Quiz*

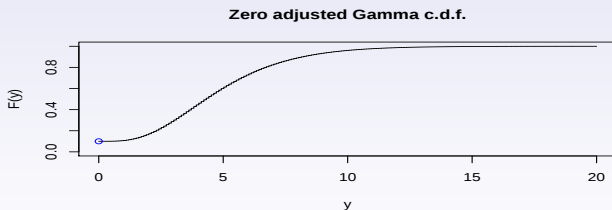
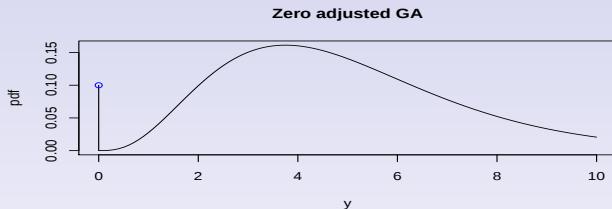
- Question 1 which film in the 90's has made the most money?
- Question 2 which film in the 90's did badly in the first week but subsequently did very well?
- Question 3 which film in the 90's did well in the first week but flopped after that?

Go back

# Film Data: Answers to the quiz

[Go back](#)

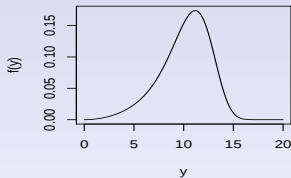
# Example of mixed distribution distributions

[Go back Distributions](#)

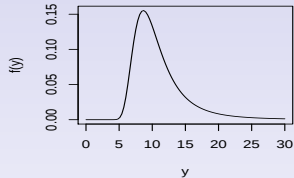
# Continuous distributions: different shapes

[Go back](#) [Distributions](#)

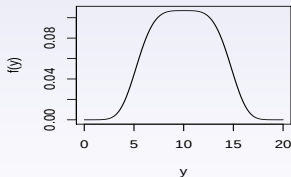
**negative skewness**



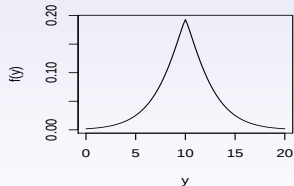
**positive skewness**



**platy-kurtosis**

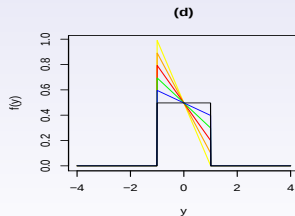
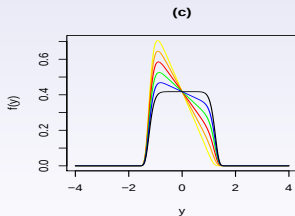
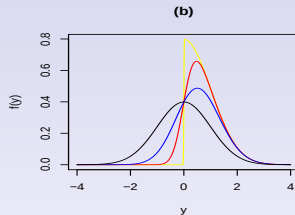
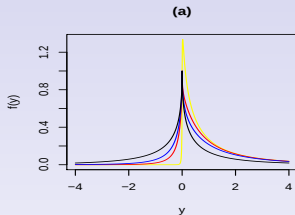


**lepto-kurtosis**

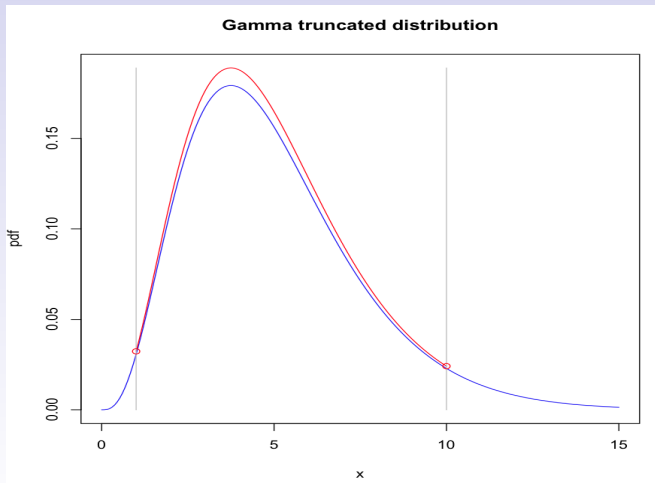


# Continuous distributions: different types

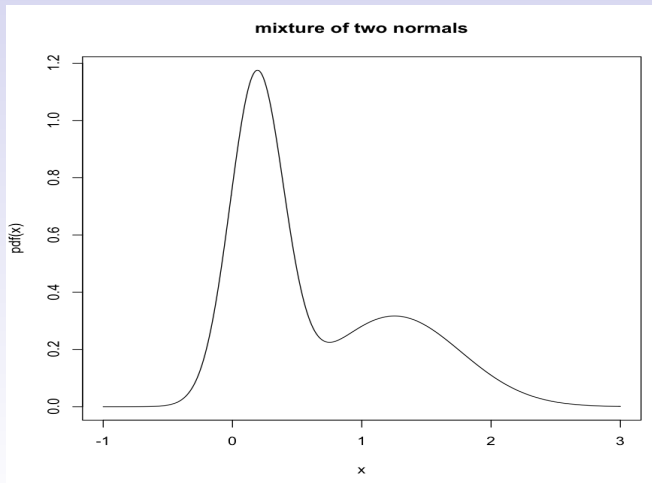
[Go back](#) [Distributions](#)



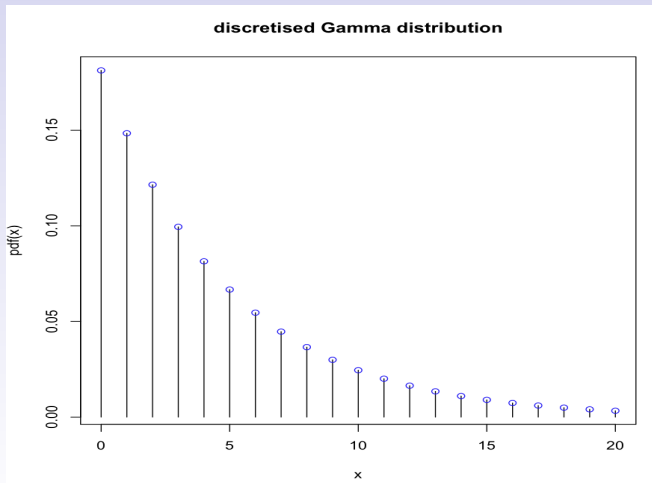
## Continuous distributions: different types [Go back Distributions](#)



# Continuous distributions: different types

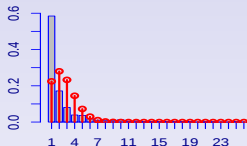
[Go back Distributions](#)

# Continuous distributions: different types

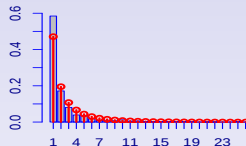
[Go back Distributions](#)

# The stylometric data, [Go back to distributions](#)

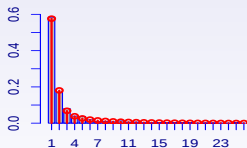
**(b) Poisson**



**(c) negative binomial II**



**(c) Delaporte**



**(d) Sichel**

