

A User's Guide to the SpatialExtremes Package

Mathieu Ribatet

Copyright ©2009

Chair of Statistics
École Polytechnique Fédérale de Lausanne
Switzerland

Contents

Introduction	1
1 An Introduction to Max-Stable Processes	3
1.1 The Smith Model	3
1.2 The Schlather Model	6
2 Unconditional Simulation of Random Fields	11
2.1 Simulation of Gaussian Random Fields	11
2.1.1 Direct Approach	11
2.1.2 Circulant Embedding Method	12
2.1.3 Turning Bands Method	12
2.2 Simulation of Max-Stable Random Fields	15
2.2.1 The Smith Model	16
2.2.2 The Schlather Model	17
3 Spatial Dependence of Max-Stable Random Fields	19
3.1 The Extremal Coefficient	19
3.2 Madogram-based approaches	21
3.2.1 Madogram	21
3.2.2 F -Madogram	23
3.2.3 λ -Madogram	24
4 Fitting a Unit Fréchet Max-Stable Process to Data	27
4.1 Least Squares	27
4.2 Pairwise Likelihood	28
4.2.1 Misspecification	28
4.2.2 Assessing Uncertainties	31
5 Model Selection	33
5.1 Takeuchi Information Criterion	33
5.2 Likelihood Ratio Statistic	34
5.2.1 Adjusting the χ^2 distribution	35
5.2.2 Adjusting the composite likelihood	36

6	Fitting a Max-Stable Process to Data	37
6.1	Response Surfaces	38
6.1.1	Linear Regression Models	39
6.1.2	Semiparametric Regression Models	40
6.2	Building Response Surfaces for the GEV Parameters	43
7	Conclusion	47
A	P-splines with radial basis functions	49
A.1	Model definition	49
A.2	Fast computation of p-splines	50

List of Figures

1.1	Two simulations of the Smith model with different Σ matrices. Left panel: $\sigma_{11} = \sigma_{22} = 9/8$ and $\sigma_{12} = 0$. Right panel: $\sigma_{11} = \sigma_{22} = 9/8$ and $\sigma_{12} = 1$. The max-stable processes are transformed to unit Gumbel margins for viewing purposes.	6
1.2	Plots of the Whittle–Matérn, the powered exponential, the Cauchy and the Bessel correlation functions - from left to right. The sill and the range parameters are $c_1 = c_2 = 1$ while the smooth parameters are given in the legends.	7
1.3	Two simulations of the Schlather model with different correlation functions having approximately the same practical range. Left panel: Whittle–Matérn with $c_1 = c_2 = \nu = 1$. Right panel: Powered exponential with $c_1 = \nu = 1$ and $c_2 = 1.5$. The max-stable processes are transformed to unit Gumbel margins for viewing purposes.	8
1.4	Contour plots of an isotropic (left panel) and anisotropic (right panel) correlation function. Powered exponential family with $c_1 = c_2 = \nu = 1$. The anisotropy parameters are: $\varphi = \pi/4$, $r = 0.5$	9
2.1	One realisation of a gaussian random field having a Whittle-Matern covariance function with sill, range and smooth parameters all equal to one using the turning bands methods. The continuous spectral method was used for simulating processes on the lines. Left panel: 1 line, middle panel: 25 lines, right panel: 625 lines.	14
2.2	Two simulations of the Smith model with different covariance matrices. Left panel: $\sigma_{11} = \sigma_{22} = 9/8$, $\sigma_{12} = 0$. Right panel: $\sigma_{11} = \sigma_{22} = 9/8$, $\sigma_{12} = 3/4$. The observations are transformed to unit Gumbel margins for viewing purposes.	17
3.1	Extremal coefficient functions for different max-stable models. Σ is the 2×2 identity matrix. Correlation function: Whittle–Matérn with $c_1 = c_2 = \nu = 1$	20
3.2	Pairwise extremal coefficient estimates from the Schlather and Tawn (left panel) and Smith (right panel) estimators from a simulated max-stable random field having a Whittle–Matérn correlation function - $c_1 = c_2 = \nu = 1$. The red lines are the theoretical extremal coefficient function.	21
3.3	Madogram (left panel) and binned madogram (right panel) with unit Gumbel margins for the Schlather model with the Whittle–Matérn correlation functions. The red lines are the theoretical madograms. Points are pairwise estimates.	22
3.4	Pairwise madograms (left panel) and extremal coefficients (right panel) estimates from a simulated max-stable random field having a Whittle–Matérn correlation function - $c_1 = c_2 = \nu = 1$. The red lines are the theoretical madogram and extremal coefficient function.	23

3.5	Pairwise F -madogram (left panel) and extremal coefficient (right panel) estimates for the Smith model with Σ equals to the identity matrix. The red lines are the theoretical F -madogram and extremal coefficient function.	23
3.6	Binned λ -madogram estimates for two Schlather models having a powered exponential (right panel) and a Cauchy covariance (left panel) functions. $c_1 = c_2 = \nu = 1$	26
6.1	Impact of the number of knots in the fitted p -spline. Left panel: $q = 2$, middle panel: $q = 10$, right panel: $q = 50$. The small vertical lines corresponds to the location of each knot.	41
6.2	Impact of the smoothing parameter λ on the fit. Left panel: $\lambda = 0$, middle panel: $\lambda = 0.1$ and right panel: $\lambda = 10$	41
6.3	Cross-validation and generalized cross validation curves and corresponding fitted curves.	43

Introduction

What is the SpatialExtremes package?

The **SpatialExtremes** package is an add-on package for the R [R Development Core Team, 2007] statistical computing system. It provides functions for the analysis of spatial extremes using (currently) max-stable processes.

All comments, criticisms and queries on the package or associated documentation are gratefully received.

Obtaining the package/guide

The package can be downloaded from CRAN (The Comprehensive R Archive Network) at <http://cran.r-project.org/>. This guide (in pdf) will be in the directory `SpatialExtremes/doc/` underneath wherever the package is installed. You can get it by invoking

```
> vignette("SpatialExtremesGuide")
```

To have a quick overview of what the package does, you might want to have a look at its own web page <http://spatialextremes.r-forge.r-project.org/>.

Contents

To help users to use properly the *SpatialExtremes* packages, this report introduces all the theory and references needed. Some practical examples are inserted directly within the text to show how it works in practice. Chapter 1 is an introduction to max-stable processes and introduces several models that might be useful in concrete applications. Chapter 2 describes simulation techniques for unconditional simulation of both Gaussian and max-stable random fields. Statistics and tools to analyze the spatial dependence of extremes are presented in Chapter 3. Chapter 4 tackles the problem of fitting max-stable process to data that are assumed to be unit Fréchet distributed. Chapter 5 is devoted to model selection. Chapter 6 presents models and procedures on how to fit max-stable processes to data that do not have unit Fréchet margins. Chapter ?? is devoted to model checking while Chapter ?? is devoted to inferential procedures and predictions. Lastly, Chapter 7 draws conclusions on spatial extremes. Note that several computations are reported in the Annex part.

Caveat

I have checked these functions as best I can but they may contain bugs. If you find a bug or suspected bug in the code or the documentation please report it to me at mathieu.ribatet@epfl.ch. Please include an appropriate subject line.

Legalese

This program is free software; you can redistribute it and/or modify it under the terms of the GNU General Public License as published by the Free Software Foundation; either version 3 of the License, or (at your option) any later version.

This program is distributed in the hope that it will be useful, but without any warranty; without even the implied warranty of merchantability or fitness for a particular purpose. See the GNU General Public License for more details.

A copy of the GNU General Public License can be obtained from <http://www.gnu.org/copyleft/gpl.html>.

Acknowledgements

This work has been supported by the Competence Center Environment and Sustainability <http://www.cces.ethz.ch/index> within the EXTREMES project <http://www.cces.ethz.ch/projects/hazri/EXTREMES>. I would like to thank S. A. Padoan for his support when we were both calculating the tedious partial derivatives required for the estimation of the asymptotic sandwich covariance matrix.

An Introduction to Max-Stable Processes

A max-stable process $Z(\cdot)$ is the limit process of maxima of independent identically distributed random fields $Y_i(x)$, $x \in \mathbb{R}^d$. Namely, for suitable $a_n(x) > 0$ and $b_n(x) \in \mathbb{R}$,

$$Z(x) = \lim_{n \rightarrow +\infty} \frac{\max_{i=1}^n Y_i(x) - b_n(x)}{a_n(x)}, \quad x \in \mathbb{R}^d \quad (1.1)$$

Note that (1.1) does not ensure that the limit exists. However, provided it does and from (1.1), we can see that max-stable processes might be appropriate models for modelling annual maxima of spatial data, for example.

Theoretically speaking, there is no loss of generality in transforming the margins to have a unit Fréchet scale i.e.

$$\Pr[Z(x) \leq z] = \exp\left(-\frac{1}{z}\right), \quad \forall x \in \mathbb{R}^d, \quad z > 0 \quad (1.2)$$

and we will first assume that the unit Fréchet assumption holds for which we have $a_n(x) = n$ and $b_n(x) = 0$.

Currently, there are two different characterisations of a max-stable process. The first one, often referred to as the *rainfall-storm* model, was first introduced by ?. More recently, Schlather [2002] introduced a new characterisation of a max-stable process allowing for a random shape. In this Chapter, we will present several max-stable models that might be relevant in studying spatial extremes.

1.1 The Smith Model

? was the first to proposed a characterisation of max-stable processes. Later, ? use this characterisation to provide a parametric model for spatial extremes. The construction was the following. Let $\{(\xi_i, y_i), i \geq 1\}$ denote the points of a Poisson process on $(0, +\infty) \times \mathbb{R}^d$ with intensity measure $\xi^{-2} d\xi \nu(dy)$, where $\nu(dy)$ is a positive measure on $\mathbb{R}^d$. Then one characterisation of a max-stable process with unit Fréchet margins is

$$Z(x) = \max_i \{\xi_i f(y_i, x)\}, \quad x \in \mathbb{R}^d \quad (1.3)$$

where $\{f(y, x), x, y \in \mathbb{R}^d\}$ is a non-negative function such that

$$\int_{\mathbb{R}^d} f(x, y) \nu(dy) = 1, \quad \forall x \in \mathbb{R}^d$$

To see that equation (1.3) defines a stationary max-stable process with unit Fréchet margins, we have to check that the margins are indeed unit Fréchet and $Z(x)$ has the max-stable property. Following Smith, consider the set defined by:

$$E = \left\{ (\xi, y) \in \mathbb{R}_*^+ \times \mathbb{R}^d : \xi f(y, x) > z \right\}$$

for a fixed location $x \in \mathbb{R}^d$ and $z > 0$. Then

$$\begin{aligned} \Pr[Z(x) \leq z] &= \Pr[\text{no points in } E] = \exp \left[- \int_{\mathbb{R}^d} \int_{z/f(y,x)}^{+\infty} \xi^{-2} d\xi \nu(dy) \right] \\ &= \exp \left[- \int_{\mathbb{R}^d} z^{-1} f(x, y) \nu(dy) \right] = \exp \left(-\frac{1}{z} \right) \end{aligned}$$

and the margins are unit Fréchet.

The max-stable property of $Z(\cdot)$ follows because the superposition of n independent, identical Poisson processes is a Poisson process with its intensity multiplied by n . More precisely, we have:

$$\left\{ \max_{i=1}^n Z_i(x_1), \dots, \max_{i=1}^n Z_i(x_k) \right\} \sim n \{Z(x_1), \dots, Z(x_k)\}, \quad k \in \mathbb{N}.$$

The process defined by (1.3) is often referred to as the rainfall-storm process, as one can have a more physical interpretation of the above construction. Think of y_i as realisations of rainfall storm centres in $\mathbb{R}^d$ and $\nu(dy)$ as the spatial distribution of these storm centres over $\mathbb{R}^d$ - usually $d = 2$. Each ξ_i represents the intensity of the i -th storm and therefore $\xi_i f(y_i, x)$ represents the amount of rainfall for this specific event at location x . In other words, $f(y_i, \cdot)$ drives how the i -th storm centred at y_i diffuses in space.

Definition (1.3) is rather general and Smith considered a particular setting where $\nu(dy)$ is the Lebesgue measure and $f(y, x) = f_0(y - x)$, where $f_0(y - x)$ is a multivariate normal density with zero mean and covariance matrix Σ ¹. With these additional assumptions, it can be shown that the bivariate CDF is

$$\Pr[Z(x_1) \leq z_1, Z(x_2) \leq z_2] = \exp \left[-\frac{1}{z_1} \Phi \left(\frac{a}{2} + \frac{1}{a} \log \frac{z_2}{z_1} \right) - \frac{1}{z_2} \Phi \left(\frac{a}{2} + \frac{1}{a} \log \frac{z_1}{z_2} \right) \right] \quad (1.4)$$

where Φ is the standard normal cumulative distribution function and, for two given locations 1 and 2

$$a^2 = \Delta x^T \Sigma^{-1} \Delta x, \quad \Sigma = \begin{bmatrix} \sigma_{11} & \sigma_{12} \\ \sigma_{12} & \sigma_{22} \end{bmatrix} \quad \text{or} \quad \Sigma = \begin{bmatrix} \sigma_{11} & \sigma_{12} & \sigma_{13} \\ \sigma_{12} & \sigma_{22} & \sigma_{23} \\ \sigma_{13} & \sigma_{23} & \sigma_{33} \end{bmatrix} \quad \text{and so forth}$$

where Δx is the distance vector between location 1 and location 2. Figure 1.1 plots two simulations of Smith's model with different covariance matrices.

¹Another form of Smith's model that uses a Student distribution instead of the normal one. However, it is not currently implemented.

Proof.

$$\begin{aligned}
-\log \Pr [Z(x_1) \leq z_1, Z(x_2) \leq z_2] &= \int_{\mathbb{R}^d} \int_{\min\{z_1/f_0(s-x_1), z_2/f_0(s-x_2)\}}^{+\infty} \xi^{-2} d\xi ds \\
&= \int \max \left\{ \frac{f_0(s-x_1)}{z_1}, \frac{f_0(s-x_2)}{z_2} \right\} ds \\
&= \int \frac{f_0(s-x_1)}{z_1} \mathbb{I} \left(\frac{f_0(s-x_1)}{z_1} > \frac{f_0(s-x_2)}{z_2} \right) ds \\
&+ \int \frac{f_0(s-x_2)}{z_2} \mathbb{I} \left(\frac{f_0(s-x_1)}{z_1} \leq \frac{f_0(s-x_2)}{z_2} \right) ds \\
&= \int \frac{f_0(s)}{z_1} \mathbb{I} \left(\frac{f_0(s)}{z_1} > \frac{f_0(s-x_2+x_1)}{z_2} \right) ds \\
&+ \int \frac{f_0(s)}{z_2} \mathbb{I} \left(\frac{f_0(s-x_1+x_2)}{z_1} \leq \frac{f_0(s)}{z_2} \right) ds \\
&= \frac{1}{z_1} \mathbb{E} \left[\mathbb{I} \left(\frac{f_0(X)}{z_1} > \frac{f_0(X-x_2+x_1)}{z_2} \right) \right] \\
&+ \frac{1}{z_2} \mathbb{E} \left[\mathbb{I} \left(\frac{f_0(X-x_1+x_2)}{z_1} \leq \frac{f_0(X)}{z_2} \right) \right]
\end{aligned}$$

where X is a r.v. having f_0 as density. To get the closed form of the bivariate distribution, it remains to compute the probabilities of the event $\{f_0(X)/z_1 > f_0(X-x_2+x_1)/z_2\}$.

$$\begin{aligned}
\frac{f_0(X)}{z_1} > \frac{f_0(X-x_2+x_1)}{z_2} &\iff 2\log z_1 + X^T \Sigma^{-1} X < 2\log z_2 + (X-x_2+x_1)^T \Sigma^{-1} (X-x_2+x_1) \\
&\iff X^T \Sigma^{-1} (x_1-x_2) > \log \frac{z_1}{z_2} - \frac{1}{2} (x_1-x_2)^T \Sigma^{-1} (x_1-x_2)
\end{aligned}$$

As X has density f_0 , $X^T \Sigma^{-1} (x_1-x_2)$ is normal with mean 0 and variance $a^2 = (x_1-x_2)^T \Sigma^{-1} (x_1-x_2)$. And finally, we get

$$\begin{aligned}
\frac{1}{z_1} \mathbb{E} \left[\mathbb{I} \left(\frac{f_0(X)}{z_1} > \frac{f_0(X-x_2+x_1)}{z_2} \right) \right] &= \frac{1}{z_1} \left\{ 1 - \Phi \left(\frac{\log z_1/z_2}{a} - \frac{a}{2} \right) \right\} \\
&= \frac{1}{z_1} \Phi \left(\frac{a}{2} + \frac{\log z_2/z_1}{a} \right)
\end{aligned}$$

and

$$\frac{1}{z_2} \mathbb{E} \left[\mathbb{I} \left(\frac{f_0(X-x_1+x_2)}{z_1} \leq \frac{f_0(X)}{z_2} \right) \right] = \frac{1}{z_2} \Phi \left(\frac{a}{2} + \frac{\log z_1/z_2}{a} \right)$$

This proves equation (1.4). □

In equation (1.4), a is the Mahanalobis distance and is similar to the Euclidean distance except that it gives different weights to each component of Δx . It is positive and the limiting cases $a \rightarrow 0^+$ and $a \rightarrow +\infty$ correspond respectively to perfect dependence and independence. Therefore, for Σ fixed, the dependence decreases monotonically and continuously as $\|\Delta x\|$ increases. On the contrary, if Δx is fixed, the dependence decreases monotonically as a increases.

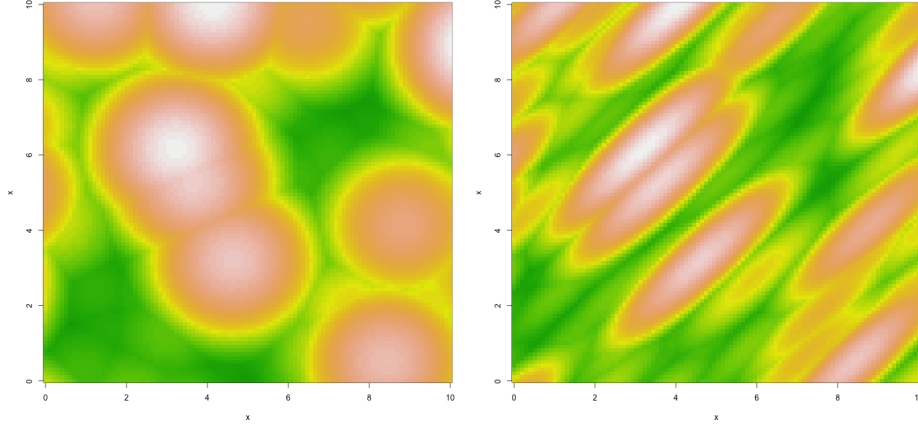


Figure 1.1: Two simulations of the Smith model with different Σ matrices. Left panel: $\sigma_{11} = \sigma_{22} = 9/8$ and $\sigma_{12} = 0$. Right panel: $\sigma_{11} = \sigma_{22} = 9/8$ and $\sigma_{12} = 1$. The max-stable processes are transformed to unit Gumbel margins for viewing purposes.

The covariance matrix Σ plays a major role in equation (1.4) as it determines the shape of the storm events. Indeed, due to the use of a multivariate normal distribution, the storm events have an elliptical shape. Considering the eigen-decomposition of Σ , we can write

$$\Sigma = U\Lambda U^T, \quad (1.5)$$

where U is a rotation matrix and Λ is a diagonal matrix of the eigenvalues. Thus, U controls the direction of the principal axes and Λ controls their lengths.

If Σ is diagonal and all the diagonal terms are identical, then $\Sigma = \Lambda$, so that the ellipsoids change to circles and model (1.4) becomes isotropic. Figure 1.1 is a nice illustration of this. The left panel corresponds to an isotropic random field while the right one depicts a clear anisotropy for which we have

$$\Sigma = \begin{bmatrix} 9/8 & 1 \\ 1 & 9/8 \end{bmatrix} = \begin{bmatrix} \cos(-3\pi/4) & \sin(-3\pi/4) \\ -\sin(-3\pi/4) & \cos(-3\pi/4) \end{bmatrix} \begin{bmatrix} 1/8 & 0 \\ 0 & 17/8 \end{bmatrix} \begin{bmatrix} \cos(-3\pi/4) & -\sin(-3\pi/4) \\ \sin(-3\pi/4) & \cos(-3\pi/4) \end{bmatrix},$$

so that the main direction of the major principal axis is $\pi/4$ and a one unit move along the direction $-\pi/4$ yields the same decrease in dependence as 17 unit moves along the direction $\pi/4$.

1.2 The Schlather Model

More recently, Schlather [2002] introduced a second characterisation of max-stable processes. Let $Y(\cdot)$ be a stationary process on $\mathbb{R}^d$ such that $\mathbb{E}[\max\{0, Y(x)\}] = 1$ and $\{\xi_i, i \geq 1\}$ be the points of a Poisson process on $\mathbb{R}_*^+$ with intensity measure $\xi^{-2}d\xi$. Then Schlather shows that a stationary max-stable process with unit Fréchet margins can be defined by:

$$Z(x) = \max_i \xi_i \max\{0, Y_i(x)\} \quad (1.6)$$

where the $Y_i(\cdot)$ are i.i.d copies of $Y(\cdot)$.

Figure 1.2: Plots of the Whittle–Matérn, the powered exponential, the Cauchy and the Bessel correlation functions - from left to right. The sill and the range parameters are $c_1 = c_2 = 1$ while the smooth parameters are given in the legends.

As before, the max-stable property of $Z(\cdot)$ stems from the superposition of n independent, identical Poisson processes, while the unit Fréchet margins holds by the same argument as for the Smith model. Indeed, let consider the following set:

$$E = \left\{ (\xi, y(x)) \in \mathbb{R}_*^+ \times \mathbb{R}^d : \xi \max(0, y(x)) > z \right\}$$

for a fixed location $x \in \mathbb{R}^d$ and $z > 0$. Then

$$\begin{aligned} \Pr[Z(x) \leq z] &= \Pr[\text{no points in } E] = \exp \left[- \int_{\mathbb{R}^d} \int_{z/\max(0, y(x))}^{+\infty} \xi^{-2} d\xi \nu(dy(x)) \right] \\ &= \exp \left[- \int_{\mathbb{R}^d} z^{-1} \max\{0, y(x)\} \nu(dy(x)) \right] = \exp \left(-\frac{1}{z} \right) \end{aligned}$$

As with the Smith model, the process defined in equation (1.6) has a practical interpretation. Think of $\xi_i Y_i(\cdot)$ as the daily spatial rainfall events so that all these events have the same spatial dependence structure but differ only in their magnitude ξ_i . This model differs slightly from Smith's one as we now have no deterministic shape such as a multivariate normal density for the storms but a random shape driven by the process $Y(\cdot)$.

The Schlather and Smith characterisations have strong connections. To see this, let consider the case for which $Y_i(x) = f_0(x - X_i)$ where f_0 is a probability density function and $\{X_i\}$ is a homogeneous Poisson process both in $\mathbb{R}^d$. With this particular setting, model (1.6) is identical to model (1.3).

Equation (1.6) is very general and we need additional assumptions to get practical models. Schlather proposed to take $Y_i(\cdot)$ to be a stationary standard Gaussian process with correlation function $\rho(h)$, scaled so that $\mathbb{E}[\max\{0, Y_i(x)\}] = 1$. With these new assumptions, it can be shown that the bivariate CDF of process (1.6) is

$$\Pr[Z(x_1) \leq z_1, Z(x_2) \leq z_2] = \exp \left[-\frac{1}{2} \left(\frac{1}{z_1} + \frac{1}{z_2} \right) \left(1 + \sqrt{1 - 2(\rho(h) + 1) \frac{z_1 z_2}{(z_1 + z_2)^2}} \right) \right] \quad (1.7)$$

where $h \in \mathbb{R}^+$ is the Euclidean distance between location 1 and location 2. Usually, $\rho(h)$, is chosen from one of the valid parametric families, such as

Whittle–Matérn	$\rho(h) = \frac{2^{1-\nu}}{\Gamma(\nu)} \left(\frac{h}{c_2} \right)^\nu K_\nu \left(\frac{h}{c_2} \right),$	$c_2 > 0, \quad \nu > 0$
Cauchy	$\rho(h) = \left[1 + \left(\frac{h}{c_2} \right)^2 \right]^{-\nu},$	$c_2 > 0, \quad \nu > 0$
Powered Exponential	$\rho(h) = \exp \left[- \left(\frac{h}{c_2} \right)^\nu \right],$	$c_2 > 0, \quad 0 < \nu \leq 2$
Bessel	$\rho(h) = \left(\frac{2c_2}{h} \right)^\nu \Gamma(\nu + 1) J_\nu \left(\frac{h}{c_2} \right),$	$c_2 > 0, \quad \nu \geq \frac{d-2}{2}$

where c_2 and ν are the range and the smooth parameters of the correlation function, Γ is the gamma function and J_ν and K_ν are the Bessel and the modified Bessel function of the third kind with order ν and d is the dimension of the random fields.

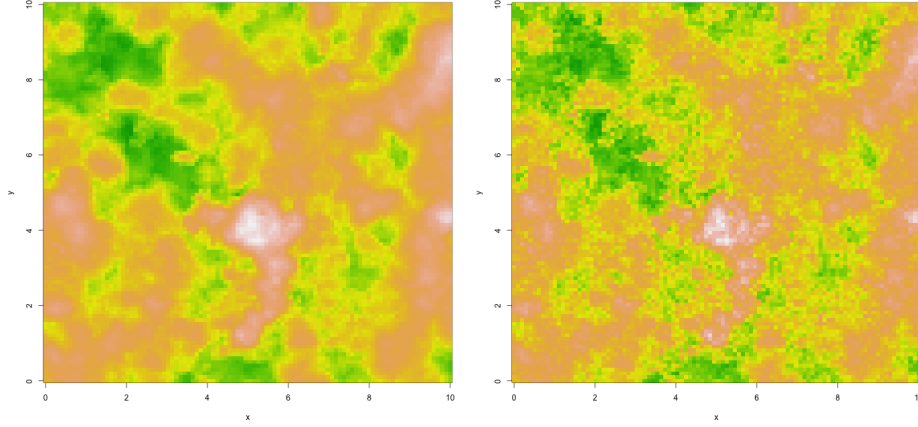


Figure 1.3: Two simulations of the Schlather model with different correlation functions having approximately the same practical range. Left panel: Whittle–Matérn with $c_1 = c_2 = \nu = 1$. Right panel: Powered exponential with $c_1 = \nu = 1$ and $c_2 = 1.5$. The max-stable processes are transformed to unit Gumbel margins for viewing purposes.

Accordingly to Gaussian processes, it is possible to add a sill c_1 and a nugget effect ν to these correlation functions i.e.

$$\rho_*(h) = \begin{cases} \nu + c_1, & h = 0 \\ c_1 \rho(h), & h > 0 \end{cases}$$

where ρ is one of the correlation functions introduced above. However, as Schlather [2002] consider stationary standard Gaussian processes, the sill and nugget parameters satisfy $\nu = 1 - c_1$ because the correlation function has to be equal to 1 at the origin.

Figure 1.2 plots the correlation functions for the parametric families introduced above. The left panel was generated with the following lines

```
> covariance(nugget = 0, sill = 1, range = 1, smooth = 4, cov.mod = "whitmat",
+           xlim = c(0,9), ylim = c(0, 1))
> covariance(nugget = 0, sill = 1, range = 1, smooth = 2, cov.mod = "whitmat",
+           add = TRUE, col = 2, xlim = c(0,9))
> covariance(nugget = 0, sill = 1, range = 1, smooth = 1, cov.mod = "whitmat",
+           add = TRUE, col = 3, xlim = c(0,9))
> covariance(nugget = 0, sill = 1, range = 1, smooth = 0.5, cov.mod = "whitmat",
+           col = 4, add = TRUE, xlim = c(0,9))
> legend("topright", c(expression(nu == 4), expression(nu == 2),
+           expression(nu == 1), expression(nu == 0.5)),
+       col = 1:4, lty = 1, inset = 0.05)
```

Figure 1.3 plots two realisations of the Schlather model with the powered exponential and Whittle–Matérn correlation functions. It can be seen that the powered exponential model leads to more rough random fields as, with this particular setting for the covariance parameters, the slope of the powered exponential correlation function near the origin is steeper than the Whittle–Matérn.

The correlation functions introduced above are all isotropic, but model (1.7) doesn't require this assumption. From a valid correlation function ρ it is always possible to get an elliptical correlation

Figure 1.4: Contour plots of an isotropic (left panel) and anisotropic (right panel) correlation function. Powered exponential family with $c_1 = c_2 = \nu = 1$. The anisotropy parameters are: $\varphi = \pi/4$, $r = 0.5$.

function ρ_e by using the following transformation:

$$\rho_e(\Delta x) = \rho\left(\sqrt{\Delta x^T A \Delta x}\right) \quad (1.8)$$

where Δx is the distance vector between two stations, A is any positive semi-definite matrix that may involve additional parameters. For example, if the spatial domain belongs to $\mathbb{R}^2$, a convenient parametrization for A is given by

$$A = \begin{bmatrix} \cos \varphi & r^2 \sin \varphi \\ r^2 \sin \varphi & \cos \varphi \end{bmatrix}$$

where $\varphi \in [0, \pi)$ is the rotation angle and $0 < r < 1$ is the ratio between the minor and major principal axes of the ellipse. Figure 1.4 plots the contour of an isotropic correlation function and an anisotropic one derived from equation (1.8).

The correlation coefficient $\rho(h)$ can take any value in $[-1, 1]$. Complete dependence is reached when $\rho(h) = 1$ while independence occurs when $\rho(h) = -1$. However, most parametric correlation families don't allow negative values so that independence is never reached.

2

Unconditional Simulation of Random Fields

In this chapter, we present some results useful for the simulation of stationary max-stable random fields. This chapter is organized as follows. As the max-stable process proposed by Schlather [2002] involves the simulation of gaussian processes, techniques for simulating such processes are first introduced. Next algorithms for simulating max-stable processes are presented.

2.1 Simulation of Gaussian Random Fields

The simulation of gaussian random fields has been of great interest since the pioneer work of Matheron [1973]. Various methods, more or less accurate and CPU demanding, were introduced. In this section we will focus especially on three different techniques: the direct approach, the circulant embedding method [Chan and Wood, 1997] and the turning bands [Matheron, 1973]. These simulation techniques are introduced in turn.

2.1.1 Direct Approach

The direct approach is certainly the simplest approach for simulating gaussian processes. Although this method is exact in principle and may be used in $\mathbb{R}^d$, $d \geq 1$, its use for a large number of locations is prohibited as we will see later. The direct simulation relies on the following property [Davis, 1987].

Proposition 2.1. *If X is a n vector of independent standard normal random variables, then*

$$Y = D^{1/2}X, \quad D = D^{1/2}(D^{1/2})^T \quad (2.1)$$

is a multivariate gaussian random vector with zero mean and covariance matrix D .

Proof. Y has a zero mean as

$$\mathbb{E}[Y] = \mathbb{E}[D^{1/2}X] = D^{1/2}\mathbb{E}[X] = 0.$$

The covariance of Y is therefore

$$\mathbb{E}[YY^T] = D^{1/2}\mathbb{E}[XX^T](D^{1/2})^T = D^{1/2}I_n(D^{1/2})^T = D^{1/2}(D^{1/2})^T = D$$

where I_n is the $n \times n$ identity matrix. □

The matrix $D^{1/2}$ can be obtained using either the Cholesky, the eigen or the singular value decompositions. These decompositions have in general a computational cost of order $O(n^3)$ while the product $D^{1/2}X$ has a cost of order $O(n^2)$ [Dietrich, 1995]. This prevents the use of the direct simulation for large n , typically for n larger than 1000.

To generate a gaussian random field with zero mean and covariance function γ observed at locations $x_1, \dots, x_n$, one can use the following scheme

1. Build the covariance matrix $D = [\gamma(x_i - x_j)]_{i,j}$, $1 \leq i, j \leq n$;
2. Compute $D^{1/2}$ by using the Cholesky or the eigen decompositions of D ;
3. Generate a n vector of independent standard normal random variables;
4. And use equation (2.1) to compute Y .

2.1.2 Circulant Embedding Method

2.1.3 Turning Bands Method

The turning bands method introduced by Matheron [1973] simplifies the problem of simulating a d -dimensional random field to the simulation of one-dimensional random fields. Consequently, this technique allows one to perform a simulation of a large dimensional random field at the cost of one-dimensional simulations.

Let $X(\cdot)$ be a zero mean isotropic random process in $\mathbb{R}$ with covariance function γ_X . The idea consists in considering a random field $Y(\cdot)$ in $\mathbb{R}^d$ such that Y equals X on a random line in $\mathbb{R}^d$ and is constant on all hyperplanes orthogonal to this line - see Figure 2.1 left panel. As the random line has to be centred at the origin, it is uniquely defined by a random vector U uniformly distributed over the half unit sphere S_d of $\mathbb{R}^d$. More formally, we have

$$Y(x) = X(\langle x, U \rangle), \quad \forall x \in \mathbb{R}^d \quad (2.2)$$

where $\langle \cdot, \cdot \rangle$ is the standard inner product in $\mathbb{R}^d$. The covariance of the process $Z(\cdot)$ is therefore

$$\begin{aligned} \gamma_Y(h) &= \mathbb{E}[Y(x), Y(x+h)] \\ &= \mathbb{E}\{\mathbb{E}[X(\langle x, U \rangle), X(\langle x+h, U \rangle)]\} \\ &= \mathbb{E}[\gamma_X(\langle h, U \rangle)] \\ &= \int_{S_d} \gamma_X(\langle h, u \rangle) \omega_d(du) \end{aligned}$$

where $\omega_d(du)$ is the uniform density over the half unit sphere in $\mathbb{R}^d$ and h is a vector in $\mathbb{R}^d$.

Matheron [1973] shows that the previous equation is a one-to-one mapping between the set of continuous isotropic covariance functions in $\mathbb{R}^d$ and that of continuous covariances in $\mathbb{R}$. Therefore, it allows one to simulate random fields having covariance γ_Y by the simulation of one-dimensional random fields with covariance γ_X through the turning bands operator $T(\gamma_X)$

$$T\{\gamma_X(h)\} = \int_{S_d} \gamma_Y(\langle h, u \rangle) \omega_d(du) \quad (2.3)$$

The turning band operator is often complicated when $d = 2$ and is defined by

$$\gamma_Y(h) = \frac{2}{\pi} \int_0^h \frac{\gamma_X(u)}{\sqrt{h^2 - u^2}} du \quad \text{or equivalently} \quad \gamma_X(h) = \frac{d}{dh} \int_0^h \frac{u \gamma_Y(u)}{\sqrt{h^2 - u^2}} du \quad (2.4)$$

while the case $d = 3$ leads to more tractable expressions

$$\gamma_Y(h) = \frac{1}{h} \int_0^h \gamma_X(u) du \quad \text{or equivalently} \quad \gamma_X(h) = \frac{d}{dh} \{h \gamma_Y(h)\} \quad (2.5)$$

As stated by [Chilès and Delfiner \[1999\]](#), equation (2.2) is not satisfactory as the random field $Y(\cdot)$ exhibits zonal anisotropy oriented along the random direction given by U so that its empirical covariance does not match the theoretical one. To bypass this hurdle, N i.i.d. random fields are superimposed i.e.

$$Y(x) = \frac{1}{\sqrt{N}} \sum_{i=1}^N Y_i(\langle x, U_i \rangle) \quad (2.6)$$

where U_i are i.i.d random variable uniformly distribution over S_d and $X_i(\cdot)$ are i.i.d. one-dimensional random fields having covariance γ_X .

For practical simulations, the number of lines generated is of order 50 and 500 for the two and three dimensional cases respectively. However, it has to be noticed that the simulation of the process on the line is often faster for the three dimensional case, or the only possible one, if the turning band operator in the two dimensional case leads to a complex covariance function in $\mathbb{R}$. For such cases, it might be preferable to simulate a three dimensional random field and take its values on an arbitrary random field.

Although the directions U_i were supposed so far to be uniformly distributed over S_d , the same result holds if the random directions are generated from an equidistributed sequence of points in S_d . For the case $d = 2$, several authors recommend the use of deterministic equal angles between the lines [[Chilès, 1997](#); [Schlather, 1999](#)]. For the case $d = 3$, [Freulon and de Fouquet \[1991\]](#) show that using a Van der Corput sequence instead of uniform directions leads to better ergodic properties for the simulated random fields.

The Van der Corput sequence computes the binary and ternary decomposition of any integer $n \in \mathbb{N}^*$ i.e.

$$n = a_0 + 2a_1 + \dots + 2^p a_p, \quad a_i = 0, 1 \quad (2.7)$$

$$n = b_0 + 3b_1 + \dots + 3^q b_q, \quad b_i = 0, 1, 2 \quad (2.8)$$

and leads to two numbers between 0 and 1

$$\begin{aligned} u_n &= \frac{a_0}{2} + \frac{a_1}{2^2} + \dots + \frac{a_p}{2^{p+1}} \\ v_n &= \frac{b_0}{3} + \frac{b_1}{3^2} + \dots + \frac{b_q}{3^{q+1}} \end{aligned}$$

from which we get a point lying in S_3

$$\left\{ \sqrt{1 - v_n^2} \cos(2\pi u_n), \sqrt{1 - v_n^2} \sin(2\pi u_n), v_n \right\} \quad (2.9)$$

For practical purposes, this sequence is defined up to a random rotation to avoid simulations based on the same set of directions.

Until now, we have not describe how to simulate a process on a line. The circulant embedding method could be used but this is not fully satisfactory. Although the locations in $\mathbb{R}^d$ defines a grid, their projections on the lines will usually not be regularly spaced. Fortunately, it is possible to perform a continuous simulation along the lines. An accurate method for this is the continuous spectral method [[Mantoglou, 1987](#)].

If γ_X is continuous, Bochner's theorem states that it is the Fourier transform of a positive measure χ , the spectral measure,

$$\gamma_X(h) = \int_{\mathbb{R}^d} \exp(i\langle u, h \rangle) d\chi(u) \quad (2.10)$$

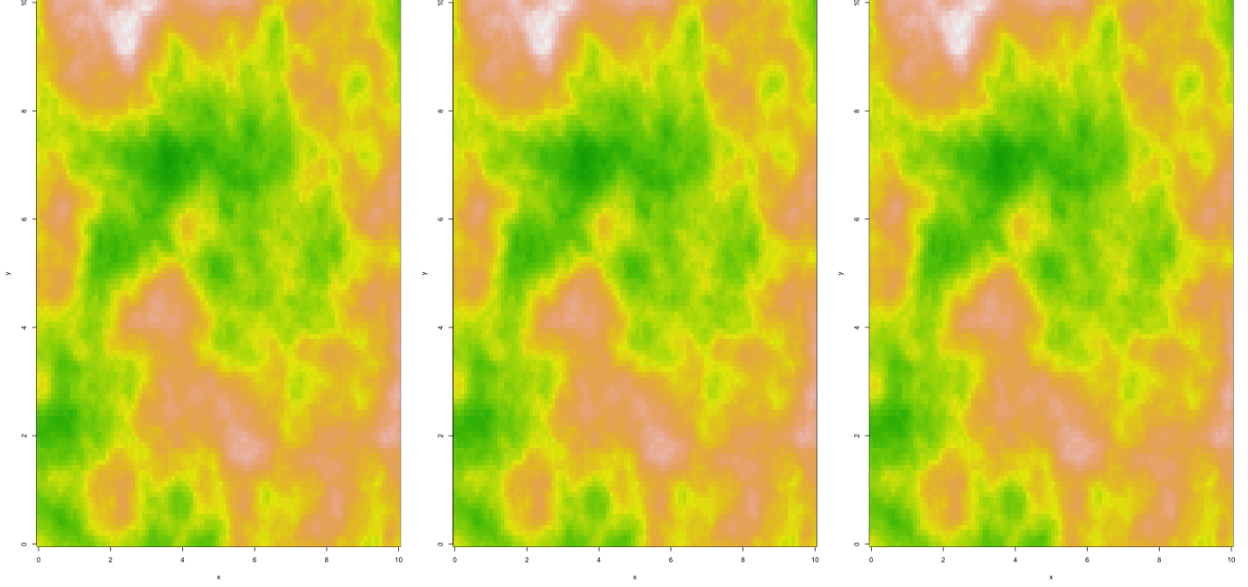


Figure 2.1: One realisation of a gaussian random field having a Whittle-Matern covariance function with sill, range and smooth parameters all equal to one using the turning bands methods. The continuous spectral method was used for simulating processes on the lines. Left panel: 1 line, middle panel: 25 lines, right panel: 625 lines.

Since $\gamma_X(0) = \sigma^2$, $\sigma^{-2}\chi$ is a probability distribution function. Hence, if Ω is a random vector with distribution $\sigma^{-2}\chi$ and $\Phi \sim U(0, 2\pi)$, Ω and Φ being mutually independent, then the random function

$$X(x) = \sqrt{2\sigma^2} \cos(\langle \Omega, x \rangle + \Phi) \quad (2.11)$$

has γ_X as its covariance function.

Proof.

$$\begin{aligned} \text{Cov}\{X(o), X(h)\} &= \mathbb{E}[X(o)X(h)] \\ &= 2\sigma^2 \mathbb{E}\{\mathbb{E}[\cos(\langle \Omega, h \rangle + \Phi) \cos \Phi]\} \\ &= \sigma^2 \mathbb{E}[\cos \langle \Omega, h \rangle] \\ &= \int_{\mathbb{R}^d} \exp(i\langle u, h \rangle) d\chi(u) \\ &= \gamma_X(h) \end{aligned}$$

□

A comprehensive overview of methods for simulating random processes on a line from various covariance families is given by [Emery and Lantu  joul \[2006\]](#).

The function `rgp` allows simulations of stationary isotropic gaussian fields. Figure 2.1 plots three realisations of a gaussian random field having a Whittle-Matern covariance function by using the turning bands method. The number of lines used for these simulations, from left to right, is 1, 25 and 625 respectively. The code used for generating this figure is

```

> x <- y <- seq(0, 10, length = 100)
> coord <- cbind(x, y)
> seed <- 3
> set.seed(seed)
> data1 <- rgp(1, coord, cov.mod = "whitmat", sill = 1, range = 1, smooth = 1,
+             grid = TRUE, control = list(nlines = 1))
> set.seed(seed)
> data2 <- rgp(1, coord, cov.mod = "whitmat", sill = 1, range = 1, smooth = 1,
+             grid = TRUE, control = list(nlines = 25))
> set.seed(seed)
> data3 <- rgp(1, coord, cov.mod = "whitmat", sill = 1, range = 1, smooth = 1,
+             grid = TRUE, control = list(nlines = 625))
> png("Figures/tbm.png", width = 1400, height = 700)
> par(mfrow=c(1,3))
> image(x, y, data1, col = terrain.colors(64))
> image(x, y, data2, col = terrain.colors(64))
> image(x, y, data3, col = terrain.colors(64))
> dev.off()

```

2.2 Simulation of Max-Stable Random Fields

The definition of a max-stable process involves the maximum over an infinite number of replication of a given random process. For simulation purposes, the number of replications is necessarily finite. However, Schlather [2002] shows that it is possible to get exact simulations on a finite sampling region.

Theorem 2.2. *Let Y be a measurable random function such that $\mathbb{E}[\int_{\mathbb{R}^d} \max\{0, Y(x)\} dx] = 1$, Π be a Poisson process on $\mathbb{R}^d \times (0, +\infty)$ with intensity measure $d\Lambda(y, \xi) = \xi^{-2} dy d\xi$ and $Z(\cdot)$ as defined by equation (1.6). Assume that Y is uniformly bounded by $C \in (0, +\infty)$ and has support in the ball $b(0, r)$ for some $r \in (0, +\infty)$. Let B be a compact set, Y_i be i.i.d. replications of Y , U_i be i.i.d. uniformly distributed on $B_r = \cup_{x \in B} b(x, r)$, ξ_i be i.i.d. standard exponential r.v. and Π , Y_i , ξ_i , U_i be mutually independent. Then, on B ,*

$$Z_*(x) = |B_r| \sup_i \frac{Y_i(x - U_i)}{\sum_{k=1}^i \xi_k}, \quad x \in B, \quad i = 1, 2, \dots \quad (2.12)$$

equals the max-stable random fields Z in distribution, and

$$Z_*(x) = |B_r| \sup \left\{ \frac{Y_i(x - U_i)}{\sum_{k=1}^i \xi_k} : i = 1, \dots, m, m \text{ is such that } \frac{C}{\sum_{k=1}^m \xi_k} \leq \max_{1 \leq i \leq m} \frac{Y_i(x - U_i)}{\sum_{k=1}^i \xi_k} \right\} \quad (2.13)$$

almost surely.

Proof. It is well known that a Poisson process on $\mathbb{R}^+$ with intensity 1 can be thought as a sum of i.i.d standard exponential random variables i.e. $\Pi = \{\sum_{i=1}^n \xi_i : n = 1, 2, \dots\}$. The application of the mapping $\xi \mapsto |B_r| \xi^{-1}$ to the points of Π yields to a new Poisson process on $\mathbb{R}^+$ with intensity measure $|B_r| \xi^{-2} d\xi$. As the U_i are i.i.d. uniformly distributed on B_r , we have that the random set

$$\left\{ \left(U_n, \frac{|B_r|}{\sum_{i=1}^n \xi_i} \right) : n = 1, 2, \dots \right\} \quad (2.14)$$

is a simple Poisson process on $B_r \times (0, +\infty)$ with intensity measure $d\Lambda(y, \xi) = \xi^{-2} dy d\xi$ and the process Z_* equals Z in distribution.

To show the second assertion, we first note that m is necessarily finite as the sequence $\sum_{k=1}^m \xi_k$ is non decreasing and Y is uniformly bounded by C so that

$$\lim_{m \rightarrow +\infty} \frac{C}{\sum_{k=1}^m \xi_k} = 0$$

Hence, there exists m finite such that

$$\Pr \left[\max_{1 \leq i \leq m} \frac{Y_i(x - U_i)}{\sum_{k=1}^i \xi_k} \geq \frac{C}{\sum_{k=1}^m \xi_k} \right] = 1$$

and

$$\Pr \left[\left| \max_{1 \leq i \leq m} \frac{Y_i(x - U_i)}{\sum_{k=1}^i \xi_k} - \sup_{i \geq 1} \frac{Y_i(x - U_i)}{\sum_{k=1}^i \xi_k} \right| \neq 0 \right] = 0$$

□

If the conditions of Theorem 2.2 are not met i.e. for random functions Y whose support is not included in a ball $b(o, r)$ or which are not uniformly bounded by a constant C , approximations for r and C should be used. Although the simulations will not be exact, it seems that approximations for r and C lead to accurate simulations.

2.2.1 The Smith Model

Recall that the Smith model is defined by

$$Z(x) = \max_i \xi_i Y_i(x), \quad Y_i(x) = f(x - U_i) \quad (2.15)$$

where f is the zero mean multivariate normal density with covariance matrix Σ and $\{(U_i, \xi_i)\}_{i \geq 1}$ are the points of a Poisson process with intensity measure $d\Lambda(y, \xi) = \xi^{-2} dy d\xi$.

As we just mentioned, the use of Theorem 2.2 to generate realisations from the Smith model is not theoretically justified as the support of the multivariate normal distribution is not included in a finite ball $b(o, r)$, $r < +\infty$. However, if r is large enough, the multivariate normal density should be close to 0 and the points in $x \in \mathbb{R}^d \setminus b(o, r)$ are unlikely to contribute to $Z(x)$.

Schlather [2002] in his seminal paper suggests to use r such that $\varphi(r) = 0.001$ i.e. $r \approx 3.46$, where φ denotes the standard normal density. By taking into account that the variance in the covariance matrix Σ might be large, a reasonable choice is to let

$$r = 3.46 \sqrt{\max \{\sigma_{ii} : i = 1, \dots, d\}} \quad (2.16)$$

where σ_{ii} are the diagonal elements of the covariance matrix Σ .

The function `rmaxstab` allows the simulation from the Smith model. Figure 2.2 plots two realisations on a 512×512 grid with two different covariance matrices Σ . The generation of these two random fields takes around 2.5 seconds each. The figures was generated by invoking the following lines

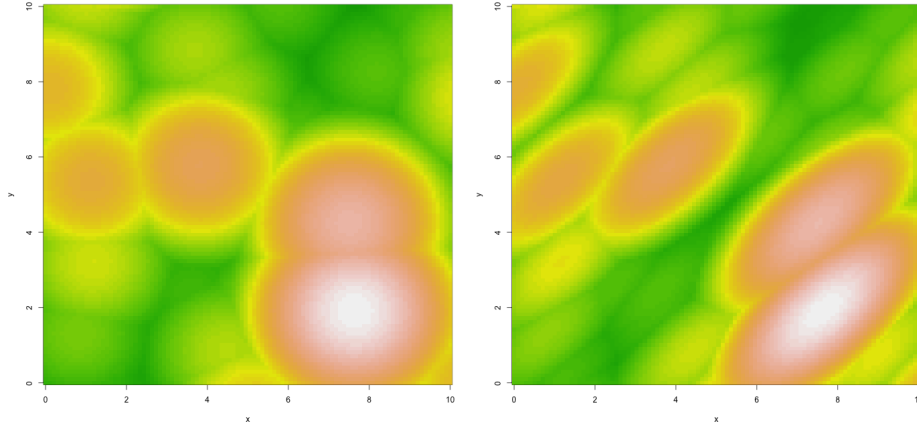


Figure 2.2: Two simulations of the Smith model with different covariance matrices. Left panel: $\sigma_{11} = \sigma_{22} = 9/8$, $\sigma_{12} = 0$. Right panel: $\sigma_{11} = \sigma_{22} = 9/8$, $\sigma_{12} = 3/4$. The observations are transformed to unit Gumbel margins for viewing purposes.

```
> x <- y <- seq(0, 10, length = 100)
> coord <- cbind(x, y)
> set.seed(8)
> M0 <- rmaxstab(1, coord, "gauss", cov11 = 9/8, cov12 = 0, cov22 = 9/8,
+               grid = TRUE)
> set.seed(8)
> M1 <- rmaxstab(1, coord, "gauss", cov11 = 9/8, cov12 = 3/4, cov22 = 9/8,
+               grid = TRUE)
> png("Figures/rmaxstabSmith.png", width = 1400, height = 700)
> par(mfrow = c(1,2))
> image(x, y, log(M0), col = terrain.colors(64))
> image(x, y, log(M1), col = terrain.colors(64))
> dev.off()
```

2.2.2 The Schlather Model

Recall that the Schlather model is defined by

$$Z(x) = \max_i \xi_i \max \{0, Y_i(x)\}, \quad (2.17)$$

where $Y_i(\cdot)$ are i.i.d. stationary standard gaussian processes with correlation function $\rho(h)$, scaled so that $\mathbb{E}[\max\{0, Y_i(x)\}] = 1$ and $\{\xi_i\}_{i \geq 1}$ are the points of a Poisson process on $\mathbb{R}_*^+$ with intensity measure $d\Lambda(\xi) = \xi^{-2}d\xi$.

Again, the conditions required by theorem 2.2 are not satisfied as the gaussian processes are not uniformly bounded. However, Schlather [2002] claims that by choosing a constant C such that $\Pr[\max\{0, Y_i(x)\} > C]$ is small, the simulation procedure is still accurate and recommends the use of $C = 3$.

The package allows for simulation of max-stable processes by using the `rmaxstab` function. A typical use for scattered location is

```
> coord <- matrix(runif(100, 0, 10), ncol = 2)
> data1 <- rmaxstab(100, coord, "whitmat", nugget = 0, range = 1, smooth = 1)
```

while for locations located on a grid, users should invoke

```
> x <- seq(0, 10, length = 100)
> coord <- cbind(x, x)
> data2 <- rmaxstab(1, coord, "powexp", nugget = 0, range = 1, smooth = 2,
+               grid = TRUE)
```

Spatial Dependence of Max-Stable Random Fields

Sooner or later, statistical modellers will be interested in knowing how dependence evolves in space. When dealing with non-extremal data, a common tool is the (semi-)variogram γ [Cressie, 1993]. Let $Y(\cdot)$ be a stationary gaussian process with correlation function ρ and variance σ^2 . It is well known that $Y(\cdot)$ is fully characterized by its mean and its covariance. Consequently, the variogram defined as

$$\gamma(x_1 - x_2) = \frac{1}{2} \text{Var} [Y(x_1) - Y(x_2)] = \sigma^2 \{1 - \rho(x_1 - x_2)\} \quad (3.1)$$

determines the degree of spatial dependence of $Y(\cdot)$.

When extreme values are of interest, the variogram is no longer a useful tool, as it may not even exist. Therefore, there is a need to develop more suitable tools for analyzing the spatial dependence of max-stable fields. In this chapter, we will present the extremal coefficient as a measure of the degree of dependence for extreme values, and variogram-based approaches that are especially well adapted to extremes.

3.1 The Extremal Coefficient

Let $Z(\cdot)$ be a stationary max-stable random field with unit Fréchet margins. The extremal dependence among N fixed locations in $\mathbb{R}^d$ can be summarised by the extremal coefficient, which is defined as:

$$\Pr [Z(x_1) \leq z, \dots, Z(x_N) \leq z] = \exp \left(-\frac{\theta_N}{z} \right) \quad (3.2)$$

where $1 \leq \theta_N \leq N$ with the lower and upper bounds corresponding to complete dependence and independence and thus provides a measure of the degree of spatial dependence between stations. Following this idea, θ_N can be thought as the effective number of independent stations.

Given the properties of the max-stable process with unit Fréchet margins, the finite-dimensional CDF belongs to the class of multivariate extreme value distributions

$$\Pr [Z(x_1) \leq z_1, \dots, Z(x_N) \leq z_N] = \exp \{-V(z_1, \dots, z_N)\} \quad (3.3)$$

where V is a homogeneous function of order -1 called the exponent measure [Pickands, 1981; Coles, 2001]. As a consequence, the homogeneity property of V implies a strong relationship between the exponent measure and the extremal coefficient

$$\theta_N = V(1, \dots, 1) \quad (3.4)$$

Figure 3.1: Extremal coefficient functions for different max-stable models. Σ is the 2×2 identity matrix. Correlation function: Whittle–Matérn with $c_1 = c_2 = \nu = 1$.

An important special case of equation (3.2) is to consider pairwise extremal coefficients, that is

$$\Pr[Z(x_1) \leq z, Z(x_2) \leq z] = \exp \left\{ -\frac{\theta(x_1 - x_2)}{z} \right\} \quad (3.5)$$

Following Schlather and Tawn [2003], $\theta(\cdot)$ is called the extremal coefficient function; it provides sufficient information about extremal dependence for many problems although it does not characterise the full distribution.

The extremal coefficient functions for max-stable models presented in Chapter 1 can be derived directly from their bivariate distribution by letting $z_1 = z_2 = z$. More precisely, we have:

$$\begin{aligned} \text{Smith} \quad \theta(x_1 - x_2) &= 2\Phi \left(\frac{\sqrt{(x_1 - x_2)^T \Sigma^{-1} (x_1 - x_2)}}{2} \right) \\ \text{Schlather} \quad \theta(x_1 - x_2) &= 1 + \sqrt{\frac{1 - \rho(x_1 - x_2)}{2}} \end{aligned}$$

Figure 3.1 plots the extremal coefficient function for the different max-stable models introduced. The Smith model covers the whole range of dependence while Schlather’s model has an upper bound of $1 + \sqrt{1/2}$ if the covariance function only takes positive values. More generally, the properties of isotropic positive definite functions [Matérn, 1986] imply that $\theta(\cdot) \leq 1.838$ in $\mathbb{R}^2$ and $\theta(\cdot) \leq 1.781$ in $\mathbb{R}^3$ for the Schlather model.

Schlather and Tawn [2003] prove several interesting properties for $\theta(\cdot)$:

1. $2 - \theta(\cdot)$ is a semi-definite positive function;
2. $\theta(\cdot)$ isn’t differentiable at 0 unless $\theta \equiv 1$;
3. if $d \geq 2$ and $Z(\cdot)$ is isotropic, $\theta(\cdot)$ has at most one jump at 0 and is continuous elsewhere.

These properties have strong consequences. The first indicates that the spatial dependence structure of a stationary max-stable process can be characterised by a correlation function. The second states that a valid correlation functions does not always lead to a valid extremal coefficient function. For instance, the Gaussian correlation model, $\rho(h) = \exp(-h^2)$, $h \geq 0$, is not allowed since it is differentiable at $h = 0$.

Equation (3.2) forms the basis for a simple maximum likelihood estimator. Suppose we have $Z_1(\cdot), \dots, Z_n(\cdot)$, independent replications of $Z(\cdot)$ observed at a set $A = \{x_1, \dots, x_{|A|}\}$ of locations. The log-likelihood based on equation (3.2) is given by:

$$\ell_A(\theta_A) = n \log \theta_A - \theta_A \sum_{i=1}^n \frac{1}{\max_{j \in A} \{Z_i(x_j) \overline{Z(x_j)}\}} \quad (3.6)$$

where the terms of the form $\log \max_{j \in A} \{Z_i(x_j)\}$ were omitted and $\overline{Z(x_j)} = n^{-1} \sum_{i=1}^n 1/Z_i(x_j)$. The scalings by $\overline{Z(x_j)}$ are here to ensure that $\hat{\theta}_A = 1$ when $|A| = 1$.

The problem with the MLE based on equation (3.6) is that the extremal coefficient estimates may not have the properties on the extremal coefficient function stated above. To solve this, Schlather

Figure 3.2: Pairwise extremal coefficient estimates from the Schlather and Tawn (left panel) and Smith (right panel) estimators from a simulated max-stable random field having a Whittle–Matérn correlation function - $c_1 = c_2 = \nu = 1$. The red lines are the theoretical extremal coefficient function.

and Tawn [2003] propose self consistent estimators for θ_A by either sequentially correcting the estimates obtained by equation (3.6) or by modifying the log-likelihood to ensure these properties.

? proposed another estimator for the pairwise extremal coefficients. As $Z(x)$ is unit Fréchet for all $x \in \mathbb{R}^d$, $1/Z(x)$ has a unit exponential distribution and, according to equation (3.5), $1/\max\{Z(x_1), Z(x_2)\}$ has an exponential distribution with rate $\theta(x_1 - x_2)$. By the law of large numbers $\sum_{i=1}^n 1/Z_i(x_1) = \sum_{i=1}^n 1/Z_i(x_2) \approx n^1$, this suggests a simple estimator for the extremal coefficient between locations x_1 and x_2 :

$$\hat{\theta}(x_1 - x_2) = \frac{n}{\sum_{i=1}^n \min\{Z_i(x_1)^{-1}, Z_i(x_2)^{-1}\}} \quad (3.7)$$

Figure 3.2 plots the pairwise extremal coefficient estimates from a simulated Schlather model having a Whittle–Matérn correlation function using equations (3.6) and (3.7). This figure was generated using the following code:

```
> n.site <- 40
> n.obs <- 100
> coord <- matrix(runif(2 * n.site, 0, 10), ncol = 2)
> data <- rmaxstab(n.obs, coord, "whitmat", nugget = 0, range = 1,
+               smooth = 1)
> par(mfrow=c(1,2))
> fitextcoeff(data, coord, loess = FALSE)
> fitextcoeff(data, coord, estim = "Smith", loess = FALSE)
```

3.2 Madogram-based approaches

As we already stated, variograms are useful quantities to assess the degree of spatial dependence for Gaussian processes. However their use for extreme observations is limited as variograms may not exist. To see this, consider a stationary max-stable process with unit Fréchet margins. For such random processes, both mean and variance are not finite and the variogram does not exist theoretically, so we need extra work to get reliable variogram-based tools.

3.2.1 Madogram

A standard tool in geostatistics, similar to the variogram, is the madogram [Matheron, 1987]. The madogram is

$$\nu(x_1 - x_2) = \frac{1}{2} \mathbb{E} [|Z(x_1) - Z(x_2)|], \quad (3.8)$$

where $Z(\cdot)$ is a stationary random field with mean assumed finite.

The problem with the madogram is almost identical to the one we emphasized with the variogram as unit Fréchet random variables have no finite mean. This led Cooley et al. [2006] to consider

¹In fact, these relations are exact if the margins were transformed to unit Fréchet by using the maximum likelihood estimates.

Figure 3.3: Madogram (left panel) and binned madogram (right panel) with unit Gumbel margins for the Schlather model with the Whittle–Matérn correlation functions. The red lines are the theoretical madograms. Points are pairwise estimates.

identical GEV margins with shape parameter $\xi < 1$ to ensure that the mean, and even the variance, are finite. It is possible to use the same strategy to ensure that the variogram exists theoretically but, as we will show later, we will see that working with the madogram is particularly suited for extreme values and has strong links with the extremal coefficient.

By using simple arguments and some results obtained by Hosking et al. [1985] on probability weighted moments, Cooley et al. [2006] show that

$$\theta(x_1 - x_2) = \begin{cases} u_\beta \left(\mu + \frac{\nu(x_1 - x_2)}{\Gamma(1-\xi)} \right), & \xi < 1, \xi \neq 0, \\ \exp \left(\frac{\nu(x_1 - x_2)}{\sigma} \right), & \xi = 0, \end{cases} \quad (3.9)$$

where μ , σ and ξ are the location, scale and shape parameters for the GEV distribution, $\Gamma(\cdot)$ is the Gamma function and

$$u_\beta(x) = \left(1 + \xi \frac{x - \mu}{\sigma} \right)_+^{1/\xi}$$

Equation (3.8) suggests a simple estimator

$$\hat{\nu}(x_1 - x_2) = \frac{1}{2n} \sum_{i=1}^n |z_i(x_1) - z_i(x_2)| \quad (3.10)$$

where $z_i(x_1)$ and $z_i(x_2)$ are the i -th observations of the random field at location x_1 and x_2 and n is the total number of observations. If isotropy is assumed, then it might be preferable to use a “binned” version of estimator (3.10)

$$\hat{\nu}(h) = \frac{1}{2n|B_h|} \sum_{(x_1, x_2) \in B_h} \sum_{i=1}^n |z_i(x_1) - z_i(x_2)| \quad (3.11)$$

where B_h is the set of pair of locations whose pairwise distances belong to $[h - \epsilon, h + \epsilon]$, for $\epsilon > 0$.

Figure 3.3 plots the theoretical madograms for the Schlather’s model having a Whittle–Matérn correlation function. Pairwise and binned pairwise estimates as defined by equations (3.10) and (3.11) are also reported. The code used to generate these madogram estimates was

```
> n.site <- 50
> n.obs <- 100
> coord <- matrix(runif(2 * n.site, 0, 10), ncol = 2)
> data <- rmaxstab(n.obs, coord, "whitmat", nugget = 0, range = 1, smooth = 1)
> par(mfrow=c(1,2))
> madogram(data, coord, which = "mado")
> madogram(data, coord, which = "mado", n.bins = 100)
```

Using a plugin estimate in equation (3.9) leads to a simple estimator for $\theta(\cdot)$:

$$\hat{\theta}(x_1 - x_2) = \begin{cases} u_\beta \left(\mu + \frac{\hat{\nu}(x_1 - x_2)}{\Gamma(1-\xi)} \right), & \xi < 1, \xi \neq 0, \\ \exp \left(\frac{\hat{\nu}(x_1 - x_2)}{\sigma} \right), & \xi = 0, \end{cases} \quad (3.12)$$

Figure 3.4: Pairwise madograms (left panel) and extremal coefficients (right panel) estimates from a simulated max-stable random field having a Whittle–Matérn correlation function - $c_1 = c_2 = \nu = 1$. The red lines are the theoretical madogram and extremal coefficient function.

Figure 3.5: Pairwise F -madogram (left panel) and extremal coefficient (right panel) estimates for the Smith model with Σ equals to the identity matrix. The red lines are the theoretical F -madogram and extremal coefficient function.

Figure 3.4 plots the madogram and extremal coefficient functions estimated from a simulated max-stable process with unit Fréchet margins and having a Whittle–Matérn correlation function. These estimates were obtained by using equations (3.10) and (3.12) respectively. The figure was generated using the code below.

```
> n.site <- 40
> n.obs <- 100
> coord <- matrix(runif(2 * n.site, 0, 10), ncol = 2)
> data <- rmaxstab(n.obs, coord, "whitmat", nugget = 0, range = 1, smooth = 1)
> madogram(data, coord)
```

3.2.2 F -Madogram

In the previous subsection, we introduced the madogram as a summary statistic for the spatial dependence structure. However, the choice of the GEV parameters to compute this madogram is somewhat arbitrary. Instead, Cooley et al. [2006] propose a modified madogram called the F -madogram

$$\nu_F(x_1 - x_2) = \frac{1}{2} \mathbb{E} [|F(Z(x_1)) - F(Z(x_2))|] \quad (3.13)$$

where $Z(\cdot)$ is a stationary max-stable random field with unit Fréchet margins and $F(z) = \exp(-1/z)$.

The F -madogram is well defined even in the presence of unit Fréchet margins as we work with $F(Z(x_1))$ instead of $Z(x_1)$. Obviously, equation (3.13) suggests a simple estimator:

$$\hat{\nu}_F(x_1 - x_2) = \frac{1}{2n} \sum_{i=1}^n |\hat{F}(z_i(x_1)) - \hat{F}(z_i(x_2))| \quad (3.14)$$

where $z_i(x_1)$ and $z_i(x_2)$ are the i -th observations of the random field at location x_1 and x_2 and n is the total number of observations.

The F -madogram has strong connections with the extremal coefficient introduced in Section 3.1. Indeed, we have

$$2\nu_F(x_1 - x_2) = \frac{\theta(x_1 - x_2) - 1}{\theta(x_1 - x_2) + 1} \quad (3.15)$$

or conversely

$$\theta(x_1 - x_2) = \frac{1 + 2\nu_F(x_1 - x_2)}{1 - 2\nu_F(x_1 - x_2)} \quad (3.16)$$

Proof. Let first note that

$$|x - y| = 2 \max \{x, y\} - (x + y) \quad (3.17)$$

Using equation (3.17) in equation (3.13), we have:

$$\begin{aligned}\nu_F(x_1 - x_2) &= \frac{1}{2} \mathbb{E} [|F(Z(x_1)) - F(Z(x_2))|] \\ &= \mathbb{E} [\max\{F(Z(x_1)), F(Z(x_2))\}] - \mathbb{E}[F(Z(x_1))] \\ &= \mathbb{E}[F(\max\{Z(x_1), Z(x_2)\})] - \frac{1}{2}\end{aligned}$$

where we used the fact that $F(Z(x_1))$ is uniformly distributed on $[0, 1]$ and F is monotone increasing. Now, from Section 3.1, we know that

$$\Pr[\max\{Z(x_1), Z(x_2)\} \leq z] = \exp\left(-\frac{\theta(x_1 - x_2)}{z}\right)$$

so that

$$\begin{aligned}\mathbb{E}[F(\max\{Z(x_1), Z(x_2)\})] &= \theta(x_1 - x_2) \int_0^{+\infty} \exp\left(-\frac{1}{z}\right) \exp\left(-\frac{\theta(x_1 - x_2)}{z}\right) z^{-2} dz \\ &= \frac{\theta(x_1 - x_2)}{\theta(x_1 - x_2) + 1}\end{aligned}$$

This proves equations (3.15) and (3.16). $\square$

As we can see from equation (3.16), there is a one-to-one relationship between the extremal coefficient and the F -madogram. This suggests a simple estimator for $\theta(x_1 - x_2)$

$$\hat{\theta}(x_1 - x_2) = \frac{1 + 2\hat{\nu}_F(x_1 - x_2)}{1 - 2\hat{\nu}_F(x_1 - x_2)} \quad (3.18)$$

Figure 3.5 plots the pairwise F -madogram and extremal coefficient estimates from 100 replications of the isotropic Smith model. The code used to produce this figure was:

```
> n.site <- 40
> n.obs <- 100
> coord <- matrix(runif(2 * n.site, 0, 10), ncol = 2)
> data <- rmaxstab(n.obs, coord, "gauss", cov11 = 1, cov12 = 0, cov22 = 1)
> par(mfrow=c(1,2))
> fmadogram(data, coord)
```

As with the madogram presented in the previous Section, it is also possible to have binned estimates of the F -madogram by passing the argument `n.bins` into the `fmadogram` function.

3.2.3 λ -Madogram

The extremal coefficient, and thus the (F -)madogram, does not fully characterize the spatial dependence of a random field. Indeed, from equation (3.2), it only considers $\Pr[Z(x_1) \leq z_1, Z(x_2) \leq z_2]$ where $z_1 = z_2 = z$. To solve this issue, Naveau et al. [2009] introduce the λ -madogram defined as

$$\nu_\lambda(x_1 - x_2) = \frac{1}{2} \mathbb{E} [|F^\lambda\{Z(x_1)\} - F^{1-\lambda}\{Z(x_2)\}|] \quad (3.19)$$

for any $\lambda \in [0, 1]$.

The idea beyond this is that by varying λ , we will focus on $\Pr[Z(x_1) \leq z_1, Z(x_2) \leq z_2]$ where $z_1 = \lambda z$ and $z_2 = (1 - \lambda)z$ and thus explore the whole space. The λ -madogram is related to the exponent measure V , namely

$$\nu_\lambda(x_1 - x_2) = \frac{V_{x_1, x_2}(\lambda, 1 - \lambda)}{1 + V_{x_1, x_2}(\lambda, 1 - \lambda)} - c(\lambda) \quad (3.20)$$

where $c(\lambda) = 3/\{2(1 + \lambda)(2 - \lambda)\}$.

Proof. Applying equation (3.17) with $x = F^\lambda\{Z(x_1)\}$ and $y = F^{1-\lambda}\{Z(x_2)\}$, we have

$$\begin{aligned} \nu_\lambda(x_1 - x_2) &= \mathbb{E} \left[\max\{F^\lambda\{Z(x_1)\}, F^{1-\lambda}\{Z(x_2)\}\} \right] - \frac{1}{2} \mathbb{E} \left[F^\lambda\{Z(x_1)\} \right] - \frac{1}{2} \mathbb{E} \left[F^{1-\lambda}\{Z(x_2)\} \right] \\ &= \mathbb{E} \left[\max\{F^\lambda\{Z(x_1)\}, F^{1-\lambda}\{Z(x_2)\}\} \right] - \frac{1}{2(1 + \lambda)} - \frac{1}{2(2 - \lambda)} \end{aligned}$$

where we used the fact that $\mathbb{E}[X^k] = 1/(1 + k)$, $X \sim U(0, 1)$, $k > 0$. From Section 3.1 we know that

$$\begin{aligned} \Pr \left[\max\{F^\lambda\{Z(x_1)\}, F^{1-\lambda}\{Z(x_2)\}\} \leq z \right] &= \Pr \left[Z(x_1) \leq -\frac{\lambda}{\log z}, Z(x_2) \leq -\frac{1 - \lambda}{\log z} \right] \\ &= \exp \{ -\log(z) V_{x_1, x_2}(\lambda, 1 - \lambda) \} \end{aligned}$$

where V_{x_1, x_2} is the homogeneous function of order -1 introduced in equation (3.3). Differentiating this distribution with respect to z gives the probability density function of the random variable $\max\{F^\lambda\{Z(x_1)\}, F^{1-\lambda}\{Z(x_2)\}\}$, so that we have

$$\begin{aligned} \mathbb{E} \left[\max\{F^\lambda\{Z(x_1)\}, F^{1-\lambda}\{Z(x_2)\}\} \right] &= \int_0^1 -V_{x_1, x_2}(\lambda, 1 - \lambda) \exp \{ -\log(z) V_{x_1, x_2}(\lambda, 1 - \lambda) \} dz \\ &= \frac{V_{x_1, x_2}(\lambda, 1 - \lambda)}{1 + V_{x_1, x_2}(\lambda, 1 - \lambda)} \end{aligned}$$

and finally

$$\nu_\lambda(x_1 - x_2) = \frac{V_{x_1, x_2}(\lambda, 1 - \lambda)}{1 + V_{x_1, x_2}(\lambda, 1 - \lambda)} - c(\lambda)$$

where $c(\lambda) = 3/\{2(1 + \lambda)(2 - \lambda)\}$. □

Again there is a one-to-one relationship between ν_λ and the dependence measure, so that we can express V_{x_1, x_2} as a function of ν_λ

$$V_{x_1, x_2}(\lambda, 1 - \lambda) = \frac{c(\lambda) + \nu_\lambda(x_1 - x_2)}{1 - c(\lambda) - \nu_\lambda(x_1 - x_2)} \quad (3.21)$$

As $V_{x_1, x_2}(0.5, 0.5) = 2\theta(x_1 - x_2)$, the previous equation induces that the madogram and the F -madogram are special cases of the λ -madogram when $\lambda = 0.5$. For instance, we have

$$\nu_{0.5}(x_1 - x_2) = \frac{8\nu_F(x_1 - x_2)}{3\{3 + 2\nu_F(x_1 - x_2)\}}$$

Equation (3.19) suggests a simple estimator

$$\hat{\nu}_\lambda(x_1 - x_2) = \frac{1}{2n} \sum_{i=1}^n |\hat{F}^\lambda\{z_i(x_1)\} - \hat{F}^{1-\lambda}\{z_i(x_2)\}| \quad (3.22)$$

Figure 3.6: Binned λ -madogram estimates for two Schlather models having a powered exponential (right panel) and a Cauchy covariance (left panel) functions. $c_1 = c_2 = \nu = 1$.

where $z_i(x_1)$ and $z_i(x_2)$ are the i -th observations of the random field at location x_1 and x_2 , n is the total number of observations and $\hat{F}$ is an estimate of the CDF at the specified location.

There is an issue with the previous estimator. Indeed, the boundary conditions for the λ -madogram when $\lambda = 0$ or $\lambda = 1$ are not always fulfilled. From equation (3.19), if $\lambda = 0$ or $\lambda = 1$, $\nu_\lambda(x_1 - x_2) = 1/4$. If F is estimated by MLE, then this condition fails while if $\hat{F}(x_{i:n}) = i/(n+1)$ we get the required conditions. To bypass this hurdle, Naveau et al. [2009] propose the following adjusted estimator

$$\begin{aligned} \hat{\nu}_\lambda(x_1 - x_2) &= \frac{1}{2n} \sum_{i=1}^n |\hat{F}^\lambda\{z_i(x_1)\} - \hat{F}^{1-\lambda}\{z_i(x_2)\}| - \frac{\lambda}{2n} \sum_{i=1}^n [1 - \hat{F}^\lambda\{z_i(x_1)\}] \\ &\quad - \frac{1-\lambda}{2n} \sum_{i=1}^n [1 - \hat{F}^{1-\lambda}\{z_i(x_2)\}] + \frac{1-\lambda+\lambda^2}{2(2-\lambda)(1+\lambda)} \end{aligned}$$

By plug in this estimator into equation (3.21) we get an estimator for the dependence measure

$$\hat{V}_{x_1, x_2}(\lambda, 1 - \lambda) = \frac{c(\lambda) + \hat{\nu}_\lambda(x_1 - x_2)}{1 - c(\lambda) - \hat{\nu}_\lambda(x_1 - x_2)} \quad (3.23)$$

Figure 3.6 plots the binned λ -madogram estimates from 100 replications of the Schlather model having a powered exponential and a Cauchy correlation functions. The code used to produce this figure was:

```
> n.site <- 40
> n.obs <- 100
> coord <- matrix(runif(2 * n.site, 0, 10), ncol = 2)
> data1 <- rmaxstab(n.obs, coord, "powexp", nugget = 0, range = 1, smooth = 1)
> data2 <- rmaxstab(n.obs, coord, "cauchy", nugget = 0, range = 1, smooth = 1)
> par(mfrow=c(1,2), pty = "s")
> lmadogram(data1, coord, n.bins = 60)
> lmadogram(data2, coord, n.bins = 60)
```

It might be useful to use the excellent *persp3d* function provided by the contributed *rgl* R package to explore dynamically the λ -madogram.

Fitting a Unit Fréchet Max-Stable Process to Data

In Chapter 1, we described max-stable processes and gave two different characterisations of them. However, we mentioned that only the bivariate distributions are analytically known so that the fitting of max-stable processes to data is not straightforward. In this Chapter, we will present two different approaches to fitting max-stable processes to data. The first one is based on least squares and was first introduced by ?. The second one uses the maximum composite likelihood estimator [Lindsay, 1988], more precisely the maximum pairwise likelihood estimator. We will consider these two different approaches separately.

4.1 Least Squares

As stated by equations (1.4) and (1.7), the density of a max-stable process is analytically known only for the bivariate case so that maximum likelihood estimators are not available. This observation led ? to propose an estimator based on least squares. This fitting procedure consists in minimizing the objective function

$$C(\psi) = \sum_{i < j} \left(\frac{\theta_{i,j} - \tilde{\theta}_{i,j}}{s(\tilde{\theta}_{i,j})} \right)^2 \quad (4.1)$$

where ψ is the vector parameter of the max-stable process, $\theta_{i,j}$ is the extremal coefficient predicted from the max-stable model for stations i and j , $\tilde{\theta}_{i,j}$ is a semi-parametric estimator defined by equation (3.7) for stations i and j and $s(\tilde{\theta}_{i,j})$ is the standard deviation related to the estimation of $\tilde{\theta}_{i,j}$, estimated by the jackknife estimator [Efron, 1982].

S.J. Neil, in his M.Phil. thesis, suggests the use of this weighted sum of squares criterion to avoid unsatisfactory fits in regions of high dependence i.e. when $\theta_{i,j}$ is close to 1.

Although ? proposed this estimator for his own max-stable model, there is no restriction in applying it to any of the max-stable models introduced in Chapter 1. An illustration of this fitting procedure is given by the following lines:

```
> n.site <- 40
> n.obs <- 80
> coord <- matrix(runif(2*n.site, 0, 10), ncol = 2)
> data <- rmaxstab(n.obs, coord, "gauss", cov11 = 9/8, cov12 = 1, cov22 = 9/8)
> fitcovmat(data, coord, marge = "emp")
```

```
Estimator: Least Squares
Model: Smith
Weighted: TRUE
Objective Value: 870.7347
```

Covariance Family: Gaussian

Estimates

Marginal Parameters:

Not estimated.

Dependence Parameters:

cov11 cov12 cov22

0.7598 0.6550 0.8091

Optimization Information

Convergence: successful

Function Evaluations: 144

This approach suffers from two major drawbacks. First, unless we use Monte-Carlo techniques, standard errors are not available and because the observations are far from being normal, the least squares estimator should be far from efficiency in the way given by the Cramér-Rao lower bound [Cramér, 1946] and the Gauss-Markov theorem. Secondly, for concrete analysis, it is hopeless that the observed (spatial) annual maxima have unit Fréchet margins so that we need first to transform the data to the unit Fréchet scale. This suggests the use of a more flexible estimator.

4.2 Pairwise Likelihood

As already stated, the “full” likelihood for the max-stable model presented in Chapter 1 is not analytically known if the number of stations under consideration is greater or equal to three. However, as the bivariate density is analytically known, this suggests the use of the pairwise likelihood in place of the full likelihood. The log pairwise-likelihood is given by

$$\ell_p(\mathbf{z}; \psi) = \sum_{i < j} \sum_{k=1}^{n_{i,j}} \log f(z_k^{(i)}, z_k^{(j)}; \psi) \quad (4.2)$$

where $\mathbf{z}$ is the data available on the whole region, $n_{i,j}$ is the number of common observations between sites i and j , $y_k^{(i)}$ is the k -th observation of the i -th site and $f(\cdot, \cdot)$ is the bivariate density of the unit Fréchet max-stable process.

4.2.1 Misspecification

Properties of the maximum composite likelihood estimator are well known [Lindsay, 1988; Cox and Reid, 2004] and belong to the class of maximum likelihood estimation under misspecification¹.

A statistical model $\{f(y; \psi), \psi \in \mathbb{R}^p\}$ is said *misspecified* if the observations y_i , $i = 1, \dots, n$ are drawn from a unknown true density g instead of f . We say that the model $\{f(y; \psi), \psi \in \mathbb{R}^p\}$ is *correct* if there exists $\psi_* \in \mathbb{R}^d$ such that $f(y; \psi_*) = g(y)$, for all y .

Let $\hat{\psi}$ be the maximum likelihood estimator. Because of the law of large numbers, we have

$$\frac{1}{n} \sum_{i=1}^n \log f(y_i; \hat{\psi}) \longrightarrow \int \log f(y; \psi_g) g(y) dy, \quad n \rightarrow +\infty \quad (4.3)$$

¹More rigorously we should say *partially specified*.

where ψ_g is the value that minimizes the Kullback–Leibler discrepancy defined as

$$D(f_\psi, g) = \int \log \left(\frac{g(y)}{f(y; \psi)} \right) g(y) dy \quad (4.4)$$

By definition of $\hat{\psi}$, we have:

$$\frac{1}{n} \sum_{i=1}^n \frac{\partial \log f(y_i; \hat{\psi})}{\partial \psi} = 0$$

so that, provided the log-likelihood is regular enough, a Taylor expansion about ψ_g yields

$$\begin{aligned} & \frac{1}{n} \sum_{i=1}^n \frac{\partial \log f(y_i; \psi_g)}{\partial \psi} + (\hat{\psi} - \psi_g)^T \frac{1}{n} \sum_{i=1}^n \frac{\partial^2 \log f(y_i; \psi_g)}{\partial \psi \partial \psi^T} \doteq 0 \\ \iff & \hat{\psi} \doteq \psi_g - \left\{ \frac{1}{n} \sum_{i=1}^n \frac{\partial^2 \log f(y_i; \psi_g)}{\partial \psi \partial \psi^T} \right\}^{-1} \left\{ \frac{1}{n} \sum_{i=1}^n \frac{\partial \log f(y_i; \psi_g)}{\partial \psi} \right\} \end{aligned}$$

It can be shown using the central limit theorem and the weak law of large numbers [Davison, 2003, p. 124] that the previous equation implies that

$$\hat{\psi} \dot{\sim} N(\psi_g, H(\psi_g)^{-1} J(\psi_g) H(\psi_g)^{-1}) \quad (4.5)$$

where

$$H(\psi_g) = n \int \frac{\partial^2 \log f(y; \psi)}{\partial \psi \partial \psi^T} g(y) dy \quad (4.6)$$

$$J(\psi_g) = n \int \frac{\partial \log f(y; \psi)}{\partial \psi} \frac{\partial \log f(y; \psi)}{\partial \psi^T} g(y) dy \quad (4.7)$$

Note that if by a lucky chance our candidate model $f(y; \psi)$ contains the true one, then $\psi_g = \psi$ and $H(\psi_g) = -J(\psi_g)$ so that equation (4.5) reduces to the usual asymptotic distribution for the MLE.

The use of pairwise likelihood, as a specific case of composite likelihood, leads to further simplifications. To see this, we consider the gradient of the log-pairwise likelihood

$$\nabla \ell_p(\mathbf{y}; \psi) = \sum_{i < j} \sum_{k=1}^{n_{i,j}} \nabla \log f(y_k^{(i)}, y_k^{(j)}; \psi) \quad (4.8)$$

For each fixed i and j ,

$$\sum_{k=1}^{n_{i,j}} \nabla \log f(y_k^{(i)}, y_k^{(j)}; \psi) = 0$$

is an unbiased estimating equation so that $\nabla \ell_p(\mathbf{y}; \psi) = 0$ is unbiased too as a linear combination of unbiased estimating equations. This leads to a modification of equation (4.5)

$$\hat{\psi}_p \dot{\sim} N(\psi, H(\psi)^{-1} J(\psi) H(\psi)^{-1}) \quad (4.9)$$

where $H(\psi)$ and $J(\psi)$ are given by equations (4.6) and (4.7).

Let consider a simple case study to see how it works in practice. Here we simulate independent replications of the Schlather model with a Whittle–Matérn correlation function having its sill, range and shape parameters equal to 0.8, 3 and 1.2 respectively.

```

> n.obs <- 80
> n.site <- 40
> set.seed(12)
> coord <- matrix(runif(2*n.site, 0, 10), ncol = 2)
> data <- rmaxstab(n.obs, coord, "whitmat", nugget = 0.8, range = 3, smooth = 1.2)
> fitmaxstab(data, coord, cov.mod = "whitmat")

```

```

      Estimator: MPLE
      Model: Schlather
      Weighted: FALSE
Pair. Deviance: 561442.9
      TIC: 561624
Covariance Family: Whittle-Matern

```

Estimates

```

Marginal Parameters:
Assuming unit Frechet.

```

```

Dependence Parameters:
nugget  range  smooth
0.7715  1.4972  4.3331

```

Standard Errors

```

nugget  range  smooth
0.07054  4.96117  24.43904

```

Asymptotic Variance Covariance

```

      nugget  range  smooth
nugget  4.976e-03 -2.797e-01  1.336e+00
range   -2.797e-01  2.461e+01 -1.208e+02
smooth   1.336e+00 -1.208e+02  5.973e+02

```

Optimization Information

```

Convergence: successful
Function Evaluations: 40

```

From this output, we can see that we indeed use the Schlather's representation with a Whittle–Matérn correlation function. The convergence was successful and the parameter estimates for the covariance function as well as the pairwise deviance are accessible. Large deviations from the theoretical values may be expected as the parameters of the Whittle–Matérn covariance function are far from orthogonal. Thus, the range and smooth estimates may be totally different while leading (approximately) to the same covariance function.

The `fitmaxstab` function provides a powerful option that can fix any parameter of the model under consideration. For instance, this could be especially useful when using the Whittle–Matérn correlation function as it is sometimes preferable to fix the smooth parameter using prior knowledge on the process smoothness [Diggle et al., 2007]. Obviously, this feature is not restricted to this specific case and one can consider more exotic models. Fixing parameters of the model is illustrated by the following lines

```
> fitmaxstab(data, coord, cov.mod = "whitmat", smooth = 1.2)
> fitmaxstab(data, coord, cov.mod = "whitmat", nugget = 0)
> fitmaxstab(data, coord, cov.mod = "whitmat", range = 3)
```

Although the Whittle–Matérn model is flexible, one may want to consider other families of covariance functions. This is achieved by invoking:

```
> fitmaxstab(data, coord, cov.mod = "cauchy")
> fitmaxstab(data, coord, cov.mod = "powexp")
```

One may also consider the Smith model, this could be done as follows

```
> fitmaxstab(data, coord, cov.mod = "gauss")
```

It is also possible to use different optimization routines to fit the model to data by passing the `method` argument. For instance, if one wants to use the BFGS method:

```
> fitmaxstab(data, coord, cov.mod = "gauss", cov12 = 0, method = "BFGS")
```

Instead of using the `optim` function, one may want to use the `nlm` or `nlminb` functions. This is done as before using the `method = "nlm"` or `method = "nlminb"` option.

Finally, it is important to note that maximizing the pairwise likelihood can be expensive. The choice of the numerical optimizer is important. In particular, optimizers that use the gradient of the pairwise likelihood might be clumsy. Indeed, if the Whittle–Matérn covariance function is considered and the smooth parameter has to be estimated, then the pairwise likelihood is not differentiable with respect to this specific parameter. In general, the Nelder–Mead [Nelder and Mead, 1965] approach seems to perform better even if the convergence is sometimes slow.

4.2.2 Assessing Uncertainties

Recall that the maximum pairwise likelihood estimator $\hat{\psi}_p$ satisfies

$$\hat{\psi}_p \sim \mathcal{N}(\psi, H(\psi)^{-1} J(\psi) H(\psi)^{-1}), \quad n \rightarrow +\infty,$$

where $H(\psi) = \mathbb{E}[\nabla^2 \ell_p(\psi; \mathbf{Y})]$ (the Hessian matrix) and $J(\psi) = \text{Var}[\nabla \ell_p(\psi; \mathbf{Y})]$, where the expectations are with respect to the “full” density.

In practice, to get the standard errors we need estimates of $H(\psi)$ and $J(\psi)$. The estimation of $H(\psi)$ is straightforward and is $\hat{H}(\hat{\psi}_p) = \nabla^2 \ell_p(\hat{\psi}_p; \mathbf{y})$; that is, the Hessian matrix evaluated at $\hat{\psi}_p$. Usually, standard optimizers are able to get finite-difference based estimates for $H(\hat{\psi}_p)$ so that no extra work is needed to get $\hat{H}(\hat{\psi}_p)$.

The estimation of $J(\hat{\psi}_p)$ is a little bit more difficult and can be done in two different ways [Varin and Vidoni, 2005]. First note that $J(\hat{\psi}_p)$ could be estimated using the “naive” estimator $\hat{J}(\hat{\psi}_p) = \nabla \ell_p(\hat{\psi}_p; \mathbf{y}) \nabla \ell_p(\hat{\psi}_p; \mathbf{y})^T$. This is unsatisfactory as most often the gradient of the log composite likelihood vanishes. To bypass this hurdle, $J(\hat{\psi}_p)$ can be estimated by

$$\hat{J}(\hat{\psi}_p) = \sum_{i=1}^n \nabla \ell_p(\hat{\psi}_p; \mathbf{y}_i) \nabla \ell_p(\hat{\psi}_p; \mathbf{y}_i)^T \quad (4.10)$$

Another estimator is given by noticing that $J(\psi)$ corresponds to the variance of the pairwise score equation $\ell_p(\psi; \mathbf{Y}) = 0$. The latter estimator is used to get standard errors. These two estimators are only accessible if independent replications of $\mathbf{Y}$ are available².

²which will mostly be the case for spatial extremes.

Model Selection

Model selection plays an important role in statistical modelling. According to Ockham's razor, given several models that fit the data equally well, we should focus on simple models rather than more complex ones. Depending on the models to be compared, several approaches exist for model selection. In this section, we will present theory on information criteria such as the Akaike Information Criterion (AIC) [Akaike, 1974] and the likelihood ratio statistic [Davison, 2003, Sec. 4.5]. We present these two approaches in turn.

5.1 Takeuchi Information Criterion

Having two different models, we want to know which one we should prefer for modelling our data. If two models have exactly the same maximized log-likelihoods, we should prefer the one which has fewer parameters because it will have a smaller variance. However, if these two maximized log-likelihoods only differ by a small amount, does this small increase worth the price of having additional parameters? To answer this question, we resort to the Kullback–Leibler discrepancy.

Let consider a random sample $Y_1, \dots, Y_n$ drawn from an unknown density g . Ignoring g , we fit a statistical model $f(y; \psi)$ by maximizing the log-likelihood. The Kullback–Leibler discrepancy measures the discrepancy of our fitted model f from the true one g

$$D(f_\psi, g) = \int \log \left(\frac{g(y)}{f(y; \psi)} \right) g(y) dy \quad (5.1)$$

The Kullback–Leibler discrepancy is always positive. Indeed, as $x \mapsto -\log(x)$ is a convex function, Jensen's inequality states

$$D(f_\psi, g) = \int \log \left(\frac{g(y)}{f(y; \psi)} \right) g(y) dy \geq -\log \left(\int \frac{f(y; \psi)}{g(y)} g(y) dy \right) = 0$$

where we used the fact that $f(y; \psi)$ is a probability density function. Furthermore, the lower bound is reached if and only if the convex function is not strictly convex which only occurs iff $f(y; \psi) = g(y)$.

Consequently, for model selection, we aim to choose models that minimize $D(f_\psi, g)$. However, $D(f_\psi, g)$ is not enough discriminant as several models may satisfy $D(f_\psi, g) = 0$. To solve this issue, suppose we have an estimate $\hat{\psi}$, we need to average $D(f_{\hat{\psi}}, g)$ over the distribution of $\hat{\psi}$. Intuitively, because of their larger sampling variance, averaging over $\hat{\psi}$ will penalized much more models having a larger number of parameters than those with fewer parameters.

Let ψ_g be the value that minimizes $D(f_\psi, g)$. A Taylor expansion of $\log f(y; \hat{\psi})$ around ψ_g gives

$$\log f(y; \hat{\psi}) \approx \log f(y; \psi_g) + \frac{1}{2} \left(\hat{\psi} - \psi_g \right)^T \frac{\partial^2 \log f(y; \psi_g)}{\partial \psi \partial \psi^T} \left(\hat{\psi} - \psi_g \right)$$

where we use the fact that ψ_g minimise the Kullback–Leibler discrepancy so its partial derivative with respect to ψ vanishes. By putting this into equation (5.1) we get

$$D(f_{\hat{\psi}}, g) \doteq D(f_{\psi_g}, g) - \frac{1}{2n} \text{tr} \left\{ (\hat{\psi} - \psi_g)(\hat{\psi} - \psi_g)^T J(\psi_g) \right\}$$

where $J(\psi_g)$ is given by equation (4.7). Finally, we have

$$n\mathbb{E}_g \left[D(f_{\hat{\psi}}, g) \right] \doteq nD(f_{\psi_g}, g) - \frac{1}{2} \text{tr} \left\{ J(\psi_g)^{-1} H(\psi_g) \right\} \quad (5.2)$$

where $H(\psi_g)$ is given by equation (4.6).

The AIC is, up to constant, an estimator of equation (5.2) and is defined as

$$\text{AIC} = -2\ell(\hat{\psi}) + 2p \quad (5.3)$$

The simplification $\text{tr}\{J(\psi_g)^{-1}H(\psi_g)\} = -p$ arises as the AIC supposes that there is no misspecification.

Because we see in Section 4.2.1 that by using the maximum pairwise likelihood estimator we work under misspecification, the AIC is not appropriate. Another estimator of equation (5.2), which allows for misspecification, is the Takeuchi information criterion (TIC)

$$\text{TIC} = -2\ell(\hat{\psi}) - 2\text{tr} \left\{ \hat{J}\hat{H}^{-1} \right\} \quad (5.4)$$

In accordance with the AIC, the best model will correspond to the one that minimizes equation (5.4). Recently, [Varin and Vidoni \[2005\]](#) rediscover this information criterion and demonstrate its use for model selection when composite likelihood is involved.

In practice, one can have a look at the output of the `fitmaxstab` function or use the `TIC` function.

```
> n.obs <- 80
> n.site <- 40
> coord <- matrix(runif(2*n.site, 0, 10), ncol = 2)
> data <- rmaxstab(n.obs, coord, "cauchy", nugget = 0.2, range = 3, smooth = 1.2)
> M0 <- fitmaxstab(data, coord, cov.mod = "powexp")
> M1 <- fitmaxstab(data, coord, cov.mod = "cauchy")
> TIC(M0, M1)

      M0      M1
570677.5 570698.0
```

The TIC for M1 is lower than the one for M0 so that we might prefer using M1.

5.2 Likelihood Ratio Statistic

TIC is useful when comparing different models. When dealing with nested models, one may prefer using the likelihood ratio statistic because of the Neyman–Pearson lemma [[Neyman and Pearson, 1933](#)].

Suppose we are interested in a statistical model $\{f(x; \psi)\}$ where $\psi^T = (\kappa^T, \phi^T)$ and more especially whether a particular value κ_0 of κ is relevant with our data. Let $(\hat{\kappa}^T, \hat{\phi}^T)$ be the maximum likelihood estimate for ψ and $\hat{\phi}_{\kappa_0}$ the maximum likelihood estimate under the restriction $\kappa = \kappa_0$.

A common statistic to check if κ_0 is relevant or not is the likelihood ratio statistic $W(\kappa_0)$ which satisfies, under mild regularity conditions,

$$W(\kappa_0) = 2 \left\{ \ell(\hat{\kappa}, \hat{\phi}) - \ell(\kappa_0, \hat{\phi}_{\kappa_0}) \right\} \longrightarrow \chi_p^2, \quad n \rightarrow +\infty \quad (5.5)$$

where p is the dimension of κ_0 .

Unfortunately, when working under misspecification, the usual asymptotic χ_p^2 distribution does not hold anymore. There's two ways to solve this issue: (a) adjusting the χ^2 distribution [Rotnitzki and Jewell, 1990] or (b) adjusting the composite likelihood so that the usual χ_p^2 holds [Chandler and Bate, 2007]. We will consider these two strategies in turn.

5.2.1 Adjusting the χ^2 distribution

If the model is misspecified, equation (5.5) has to be adjusted. More precisely, as stated by [Kent, 1982], we have

$$W(\kappa_0) = 2 \left\{ \ell(\hat{\kappa}, \hat{\phi}) - \ell(\kappa_0, \hat{\phi}_{\kappa_0}) \right\} \longrightarrow \sum_{i=1}^p \lambda_i X_i, \quad n \rightarrow +\infty \quad (5.6)$$

where λ_i are the eigenvalues of $((H^{-1}JH^{-1})_{\kappa}\{-(H^{-1})_{\kappa}\}^{-1})$, the X_i are independent χ_1^2 random variables and $(H^{-1}JH^{-1})_{\kappa}$ and $(H^{-1})_{\kappa}$ are the submatrices of $H^{-1}JH^{-1}$ and H^{-1} corresponding to the elements of κ and where the matrices H and J are given by equations (4.6) and (4.7). For practical purposes, the matrices H and J are substituted for their respective estimates as described in Section 4.2.2.

The problem with equation (5.6) is that generally the distribution of $\sum_{i=1}^p \lambda_i X_i$ is not known exactly. This led Rotnitzki and Jewell [1990] to consider $pW(\kappa_0)/\sum_{i=1}^p \lambda_i$ as a χ_p^2 random variable. However, a better approximation uses results on quadratic forms of normal random distribution.

An application of this approach is given by the following lines:

```
> n.obs <- 50
> n.site <- 30
> coord <- matrix(rnorm(2*n.site, sd = sqrt(.2)), ncol = 2)
> data <- rmaxstab(n.obs, coord, "gauss", cov11 = 100, cov12 = 25, cov22 = 220)
> M0 <- fitmaxstab(data, coord, "gauss", cov11 = 100)
> M1 <- fitmaxstab(data, coord, "gauss")
> anova(M0, M1)
```

Eigenvalue(s): 220.32

Analysis of Variance Table

	MDf	Deviance	Df	Chisq	Pr(> sum lambda Chisq)
M0	2	97942			
M1	3	97901	1	41.626	0.6638

From this output, we can see that the p -value is approximately 0.875 which turns out to be in favour of H_0 i.e. $\sigma_{11} = 100$ in Σ . Note that the eigenvalue estimate was also reported.

5.2.2 Adjusting the composite likelihood

Contrary to the approach of [Rotnitzki and Jewell \[1990\]](#), [Chandler and Bate \[2007\]](#) propose to adjust the composite likelihood instead of adjusting the asymptotic likelihood ratio statistic distribution. The starting point is that, under misspecification, the log-composite likelihood evaluated at its maximum likelihood estimate $\hat{\psi}$ has curvature $-\hat{H}^{-1}$ while it should be $\hat{H}^{-1}\hat{J}\hat{H}^{-1}$. This forms the basis for adjusting the log-composite likelihood is the following way,

$$\ell_A(\psi) = \ell(\psi^*), \quad \psi^* = \hat{\psi} + C(\psi - \hat{\psi}) \quad (5.7)$$

for some $p \times p$ matrix C .

It is straightforward to see that $\hat{\psi}$ maximizes ℓ_A with zero derivative. The key point is that its curvature at $\hat{\theta}$ is $C^T \hat{H}^{-1} C$. By choosing an appropriate C matrix, it is possible to ensure that ℓ_A has the right curvature for applying (5.5) directly. More precisely, by letting

$$C = M^{-1} M_A \quad (5.8)$$

where $M^T M = \hat{H}$ and $M_A^T M_A = \hat{H}^{-1} \hat{J} \hat{H}^{-1}$, we ensure that ℓ_A has curvature $\hat{H}^{-1} \hat{J} \hat{H}^{-1}$. If $p > 1$, the matrix square roots M and M_A are not unique and one may use the Cholesky or the singular value decompositions.

With this setting, we can apply (5.5) directly i.e.

$$W_A(\kappa_0) = 2 \left\{ \ell_A(\hat{\kappa}, \hat{\phi}) - \ell_A(\kappa_0, \hat{\phi}_A) \right\} \longrightarrow \chi_p^2$$

where $\hat{\phi}_A$ is the maximum adjusted likelihood estimated for the restricted model.

The problem with the above equation is that it requires the estimation of $\hat{\phi}_A$ which could be CPU prohibitive. Unfortunately, substituting $\hat{\phi}$ for $\hat{\phi}_A$ doesn't solve the problem as $\ell_A(\hat{\phi}) \leq \ell_A(\hat{\phi}_A)$ so that this substitution will inflate $W_A(\kappa_0)$ and thus $\Pr_{H_0}[W_a(\kappa_0) > w_a(\alpha)] > 1 - \alpha$, where $w_a(\alpha)$ is the $1 - \alpha$ quantile for the χ_p^2 distribution and α the significance level of the hypothesis test.

To solve these problems, [Chandler and Bate \[2007\]](#) propose to compensate for the use of $\hat{\phi}$ instead of $\hat{\phi}_A$ by considering a modified likelihood ratio statistic

$$W_A^*(\kappa_0) = 2c \left\{ \ell_A(\hat{\kappa}, \hat{\phi}) - \ell_A(\kappa_0, \hat{\phi}_A) \right\} \longrightarrow \chi_p^2 \quad (5.9)$$

where

$$c = \frac{(\hat{\kappa} - \kappa_0)^T \{ -(\hat{H}^{-1} \hat{J} \hat{H}^{-1})_{\kappa_0} \}^{-1} (\hat{\kappa} - \kappa_0)}{\{ (\hat{\kappa}^T, \hat{\phi}^T) \}^T \{ -(\hat{H}^{-1} \hat{J} \hat{H}^{-1})_{\kappa_0} \} \{ (\hat{\kappa}^T, \hat{\phi}^T) \}}$$

The next lines performs the same hypothesis test as that in Section 5.2.1.

```
> anova(M0, M1, method = "CB")
```

Analysis of Variance Table

	MDf	Deviance	Df	Chisq	Pr(> sum lambda Chisq)
M0	2	60576			
M1	3	60576	1	0.2089	0.6476

By using the Chandler and Bate methodology, we draw the same conclusion as in the previous section, i.e. the observations are consistent with the null hypothesis $\sigma_{11} = 100$.

6

Fitting a Max-Stable Process to Data With Unknown GEV Margins

In practice, the observations will never be drawn from a unit Fréchet distribution so that Chapter 4 won't help much with concrete applications. One way to avoid this problem is to fit a GEV to each location and then transform all data to the unit Fréchet scale. Given a continuous random variable Y whose cumulative distribution function is F , one can define a new random variable Z such as Z is unit Fréchet distributed

$$Z = -\frac{1}{\log F(Y)} \quad (6.1)$$

More precisely, if Y is a random variable distributed as a GEV with location, scale and shape parameters equal to μ , σ and ξ respectively, it turns out that equation (6.2) becomes

$$Z = \left(1 + \xi \frac{Y - \mu}{\sigma}\right)^{1/\xi} \quad (6.2)$$

The above transformation can be done by using the `gev2frech` function

```
> x <- c(2.2975896, 1.6448808, 1.3323833, -0.4464904, 2.2737603, -0.2581876,  
+       9.5184398, -0.5899699, 0.4974283, -0.8152157)  
> z <- gev2frech(x, 1, 2, .2)
```

or conversely if Z is a unit Fréchet random variable, then the random variable Y defined as

$$Y = \mu + \sigma \frac{Z^\xi - 1}{\xi} \quad (6.3)$$

is GEV distributed with location, scale and shape parameters equal to $\mu \in \mathbb{R}$, $\sigma \in \mathbb{R}_*^+$ and $\xi \in \mathbb{R}$ respectively.

```
> frech2gev(z, 1, 2, .2)
```

```
[1] 2.2975896 1.6448808 1.3323833 -0.4464904 2.2737603 -0.2581876  
[7] 9.5184398 -0.5899699 0.4974283 -0.8152157
```

The drawback of this approach is that standard errors are incorrect as the margins are fitted separately from the spatial dependence structure. Consequently, the standard errors related to the spatial dependence parameters are underestimated as we suppose that data were originally unit Fréchet.

One can solve this problem by fitting in *one step* both GEV and spatial dependence parameters [Padoan, 2008; Padoan et al., 2008; Gholam-Rezaee, 2009]. As the bivariate distributions for the max-stable models introduced in Chapter 1 were imposing unit Fréchet margins, we need to rewrite them for unknown GEV margins. To this aim, let define the transformation t such that

$$t : Y(x) \mapsto \left(1 + \xi(x) \frac{Y(x) - \mu(x)}{\sigma(x)} \right)^{1/\xi(x)} \quad (6.4)$$

where $Y(\cdot)$ is supposed to be a max-stable random field having GEV margins such that $Y(x) \sim \text{GEV}(\mu(x), \sigma(x), \xi(x))$, $\sigma(x) > 0$ for all $x \in \mathbb{R}^d$. Consequently, the bivariate distribution of $(Y(x_1), Y(x_2))$ is

$$\Pr[Y(x_1) \leq y_1, Y(x_2) \leq y_2] = \Pr[Z(x_1) \leq z_2, Z(x_2) \leq z_2]$$

where $z_1 = t(y_1)$ and $z_2 = t(y_2)$. Thus, one can relate the bivariate density for $(Y(x_1), Y(x_2))$ to the one for $(Z(x_1), Z(x_2))$ that we introduced in Chapter 1 and the log pairwise likelihood becomes

$$\ell_p(\mathbf{y}; \psi) = \sum_{i < j} \sum_{k=1}^{n_{i,j}} \left\{ \log f(z_k^{(i)}, z_k^{(j)}; \psi) + \log |J(y_k^{(i)}) J(y_k^{(j)})| \right\} \quad (6.5)$$

where $n_{i,j}$ is the sample size of common observations between site i and j and

$$z_k^{(i)} = \left(1 + \xi_i \frac{y_k^{(i)} - \mu_i}{\sigma_i} \right)^{1/\xi_i}$$

where μ_i, σ_i, ξ_i are the GEV parameters for the i -th site and $y_k^{(i)}$ is the k -th observation available at site i and $|J(t(y_k^{(i)}))|$ is the Jacobian of the mapping t evaluated at the $y_k^{(i)}$ observation i.e.

$$|J(t(y_k^{(i)}))| = \frac{1}{\sigma_i} \left(1 + \xi_i \frac{y_k^{(i)} - \mu_i}{\sigma_i} \right)^{1/\xi_i - 1}$$

Maximizing the log-pairwise likelihood given by equation (6.5) is possible by passing the option `fit.marge = TRUE` in the `fitmaxstab` function i.e.

```
> fitmaxstab(data, coord, "gauss", fit.marge = TRUE)
```

However, this will be really time consuming as such models will have $3n_{\text{site}} + p$ parameters to estimate, where p is the number of parameters related to the extremal spatial dependence structure. Another drawback is that prediction at unobserved locations won't be possible. Indeed, if no model is assumed for the evolution of the GEV parameters in space, it is therefore impossible to predict them where no data is available.

Another way may be to fit *response surfaces* for the GEV parameters. The next section aims to give an introduction to the use of response surfaces.

6.1 Response Surfaces

Response surfaces is a generic term when the problem under concern is to describe how a *response variable* y depends on *explanatory variables* $x_1, \dots, x_k$. For instance, with our particular problem

of spatial extremes, one may wonder how is it possible to predict the GEV parameters at a fixed location given the knowledge of extra covariables such as longitude, latitude, $\dots$. The goal of response surfaces is to get efficient predictions for the response variable while keeping, so far as we can, simple models.

In this section, we will first introduce the linear regression models. Next, we will increase in complexity and flexibility by introducing semiparametric regression models.

6.1.1 Linear Regression Models

Suppose we observe a response y through the $y_1, \dots, y_n$ values. For each observed y_i , we also have p related explanatory variables denoted by $x_{1,i}, \dots, x_{p,i}$. To predict y given the $x_{\cdot,i}$ values, one might consider the following model:

$$y_i = \beta_0 + \beta_1 x_{1,i} + \dots + \beta_p x_{p,i} + \epsilon_i$$

where $\beta_0, \dots, \beta_p$ are the regression parameters to be estimated and ϵ_i is an unobserved error term.

It is possible to write the above equation in a more compact way by using matrix notation e.g.

$$\mathbf{y} = \mathbf{X}\beta + \epsilon \quad (6.6)$$

where $\mathbf{y}$ is a $n \times 1$ vector, $\mathbf{X}$ is a $n \times p$ matrix called the *design matrix* and ϵ is a $p \times 1$ vector.

Model (6.6) is called a *linear model* as it is linear in β but not necessarily in the covariates x . For example, the two following models are linear models

$$\begin{aligned} y &= \beta_0 + \beta_1 x_1 + \beta_2 x_1^2 + \epsilon \\ y &= \beta_0 + \beta_1 x_1 + \beta_2 \log x_2 + \epsilon \end{aligned}$$

or equivalently in a matrix notation

$$\begin{aligned} \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} &= \begin{bmatrix} 1 & x_{1,1} & x_{1,1}^2 \\ \vdots & \vdots & \vdots \\ 1 & x_{1,n} & x_{1,n}^2 \end{bmatrix} \begin{bmatrix} \beta_0 \\ \beta_1 \\ \beta_2 \end{bmatrix} + \begin{bmatrix} \epsilon_0 \\ \epsilon_1 \\ \epsilon_2 \end{bmatrix} \\ \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} &= \begin{bmatrix} 1 & x_{1,1} & \log x_{2,1} \\ \vdots & \vdots & \vdots \\ 1 & x_{1,n} & \log x_{2,n} \end{bmatrix} \begin{bmatrix} \beta_0 \\ \beta_1 \\ \beta_2 \end{bmatrix} + \begin{bmatrix} \epsilon_0 \\ \epsilon_1 \\ \epsilon_2 \end{bmatrix} \end{aligned}$$

Usually, β is estimated by minimizing least squares

$$SS(\beta) = \sum_{i=1}^n (y_i - x_i^T \beta)^2 = (\mathbf{y} - \mathbf{X}\beta)^T (\mathbf{y} - \mathbf{X}\beta) \quad (6.7)$$

which amounts to solve the equations

$$\sum_{i=1}^n x_{j,i} (y_i - \beta^T x_i) = 0, \quad j = 1, \dots, p$$

or equivalently in a matrix notation

$$\mathbf{X}^T (\mathbf{y} - \mathbf{X}\beta) = 0$$

so that, provided that $\mathbf{X}^T \mathbf{X}$ is invertible, the least squares estimate for β is given by

$$\hat{\beta} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} \quad (6.8)$$

and the fitted $\mathbf{y}$ values are given by

$$\hat{\mathbf{y}} = \mathbf{X} \hat{\beta} = \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} \quad (6.9)$$

The matrix $\mathbf{H} = \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$ is called the *hat matrix* as it puts “hats” on $\mathbf{y}$. $\mathbf{H}$ is a projection matrix that orthogonally projects $\mathbf{y}$ onto the plane spanned by the columns of the design matrix $\mathbf{X}$.

The hat matrix plays an important role in parametric regression as it provides useful informations on the influence of some observations to the fitted values. Indeed, from equation (6.9), we have

$$\hat{y}_i = \sum_{j=1}^n H_{i,j} y_j$$

so that $H_{i,i}$ is the contribution of y_i to the estimate $\hat{y}_i$. Furthermore, if we consider the total influence of all the observations, we have

$$\begin{aligned} \sum_{i=1}^n H_{i,i} &= \text{tr}(\mathbf{H}) = \text{tr}\{\mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T\} \\ &= \text{tr}\{\mathbf{X}^T \mathbf{X} (\mathbf{X}^T \mathbf{X})^{-1}\} = \text{tr}(\mathbf{I}_p) = p \end{aligned}$$

and the total influence of all observations is equal to the degrees of freedom of the model.

6.1.2 Semiparametric Regression Models

In the previous section, we talked about linear regression models for which the relationship between the explanatory variables and the response has a deterministic shape and is supposed to be known. However, it may happened applications for which the data have a complex behaviour. For such cases, we benefit from using semiparametric regression models defined as

$$y_i = f(x_i) + \epsilon_i \quad (6.10)$$

where f is a smooth function with unknown shape.

The idea of semiparametric regression models is to decompose f into an appropriate basis for which equation (6.10) simplifies to equation (6.6) e.g.

$$f(x) = \sum_{j=1}^q b_j(x) \beta_j \quad (6.11)$$

where $b_j(\cdot)$ is the j -th basis function and β_j is the j -th element of the regression parameter β .

Several basis functions exist such as the polynomial basis, the cubic spline basis, B-splines, ... It is beyond the scope of this document to introduce all of them in details but the interested reader should have a look at [Ruppert et al. \[2003\]](#). Some details about the basis implemented in the package are reported in Annex A

Usually, the basis functions $b_j(\cdot)$ depends on *knots* κ so that equation (6.11) becomes

$$f(x) = \beta_0 + \beta_1 x + \sum_{j=1}^q b_j(x - \kappa_j) \quad (6.12)$$

Figure 6.1: Impact of the number of knots in the fitted p -spline. Left panel: $q = 2$, middle panel: $q = 10$, right panel: $q = 50$. The small vertical lines corresponds to the location of each knot.

Figure 6.2: Impact of the smoothing parameter λ on the fit. Left panel: $\lambda = 0$, middle panel: $\lambda = 0.1$ and right panel: $\lambda = 10$.

The problem with model (6.12) is that it is strongly affected by the number of knots. Figure 6.1 depicts this problem by fitting the same dataset to model (6.12) with $q = 2, 10$ and 50 . Clearly, the first fit is not satisfactory and we need to increase the number of knots. The second one seems plausible while the last one clearly overfits. This figure was generated with the following lines

```
> set.seed(12)
> x <- runif(100)
> fun <- function(x) sin(3 * pi * x)
> y <- fun(x) + rnorm(100, 0, 0.15)
> knots1 <- quantile(x, prob = 1:2 / 3)
> knots2 <- quantile(x, prob = 1:10 / 11)
> knots3 <- quantile(x, prob = 1:50 / 51)
> M0 <- rbpspline(y, x, knots = knots1, degree = 3, penalty = 0)
> M1 <- rbpspline(y, x, knots = knots2, degree = 3, penalty = 0)
> M2 <- rbpspline(y, x, knots = knots3, degree = 3, penalty = 0)
> par(mfrow=c(1,3))
> plot(x, y, col = "lightgrey")
> rug(knots1)
> lines(M0)
> plot(x, y, col = "lightgrey")
> rug(knots2)
> lines(M1, col = 2)
> plot(x, y, col = "lightgrey")
> rug(knots3)
> lines(M2, col = 3)
```

Consequently, there is a pressing need for a kind of “automatic knot selection”. One common strategy to overcome this issue is to resort to *penalized splines* or *p-splines*. The idea beyond this is to consider a large number of knots but to constrain, in a sense to be defined, their influence.

To avoid overfitting, one wish to minimize the sum of square subject to some constraint on the β parameter i.e.

$$\text{minimize } \|\mathbf{y} - \mathbf{X}\beta\|^2 \quad \text{subject to } \beta^T \mathbf{K} \beta \leq C$$

for a judicious choice of C and a given matrix $\mathbf{K}$. Annex A.1 gives further details for one possible choice of $\mathbf{K}$. Using a Lagrange multiplier argument, this constraint optimization problem is equivalent to choosing β to minimize

$$\|\mathbf{y} - \mathbf{X}\beta\|^2 + \lambda \beta^T \mathbf{K} \beta \tag{6.13}$$

for some $\lambda \geq 0$ called the *smoothing parameter* as it controls the amount of smoothing. Indeed, if $\lambda = 0$, then problem (6.13) is left unconstrained and leads to wiggly fits, while λ being large implies smoother fits. Figure 6.2 is a nice illustration of the impact of the smoothing parameter on the smoothness of the fitted curve. It was generated using the following code

```

> M0 <- rbpspline(y, x, knots = knots3, degree = 3, penalty = 0)
> M1 <- rbpspline(y, x, knots = knots3, degree = 3, penalty = 0.1)
> M2 <- rbpspline(y, x, knots = knots3, degree = 3, penalty = 10)
> par(mfrow=c(1,3))
> plot(x, y, col = "lightgrey")
> lines(M0)
> plot(x, y, col = "lightgrey")
> lines(M1, col = 2)
> plot(x, y, col = "lightgrey")
> lines(M2, col = 3)

```

It can be shown that problem (6.13) has the solution

$$\hat{\beta}_\lambda = (\mathbf{X}^T \mathbf{X} + \lambda \mathbf{K})^{-1} \mathbf{X}^T \mathbf{y} \quad (6.14)$$

and the corresponding fitted values for a penalized spline are given by

$$\hat{\mathbf{y}} = \mathbf{X}(\mathbf{X}^T \mathbf{X} + \lambda \mathbf{K})^{-1} \mathbf{X}^T \mathbf{y} \quad (6.15)$$

In accordance with the hat matrix with linear models, one can define the *smoother matrix* $\mathbf{S}_\lambda$ such that

$$\hat{\mathbf{y}} = \mathbf{S}_\lambda \mathbf{y} \quad (6.16)$$

where $\mathbf{S}_\lambda = \mathbf{X}(\mathbf{X}^T \mathbf{X} + \lambda \mathbf{K})^{-1} \mathbf{X}^T$. Consequently, a kind of effective degrees of freedom is given by $\text{tr}(\mathbf{S}_\lambda)$.

If the problem of knot selection seems to be resolved by using these constrained least squares minimisation, there is still some open questions: given our data and knots, what is the best value for λ ? Would it be possible to get an “automatic selection” for λ ?

One common tool for answering these two questions is known as *cross-validation* (CV)

$$CV(\lambda) = \sum_{i=1}^n \{y_i - \hat{f}_{-i}(x_i; \lambda)\}^2 \quad (6.17)$$

where $\hat{f}_{-i}$ corresponds to the semiparametric estimator applied to the data but with (x_i, y_i) omitted. Intuitively, large values of $CV(\lambda)$ corresponds to models that are wiggly and/or have a large variance in the parameter estimates so that minimising $CV(\lambda)$ is a nice option for an “automatic selection” of λ .

Unfortunately, the computation of equation (6.17) directly is often too CPU demanding. However, it can be shown [Ruppert et al., 2003] that

$$CV(\lambda) = \sum_{i=1}^n \left(\frac{y_i - \hat{y}_i}{1 - S_{\lambda,ii}} \right)^2 \quad (6.18)$$

where $S_{\lambda,ii}$ is the (i, i) element of $\mathbf{S}_\lambda$. Clearly, equation (6.18) does a better job than equation (6.17) as it only requires one fit to compute $CV(\lambda)$.

Sometimes, the weights $1 - S_{\lambda,ii}$ are replaced by the mean weight, $\text{tr}(\mathbf{Id} - \mathbf{S}_\lambda)/n$, where $\mathbf{Id}$ is the identity matrix, leading to the *generalized cross-validation* (GCV) score

$$GCV(\lambda) = n^2 \sum_{i=1}^n \left(\frac{y_i - \hat{y}_i}{\text{tr}(\mathbf{Id} - \mathbf{S}_\lambda)} \right)^2 \quad (6.19)$$

Figure 6.3: Cross-validation and generalized cross validation curves and corresponding fitted curves.

GCV has computational advantages over CV, and it has also computational advantages in term of invariance [Wood, 2006].

Figure 6.3 plots the CV and GCV curves for the data plotted in Figure 6.2 and the corresponding fitted p-spline. The selection of λ using CV or GCV yield approximately to the same smoothing parameter value. These “best” λ values are in accordance with the values we held fixed in Figure 6.2. The fitted curves using either CV or GCV lead to indistinguishable curves. The code used to generate Figure 6.3 was

```
> par(mfrow=c(1,3))
> lambda.cv <- cv(y, x, knots = knots3, degree = 3)$penalty
> abline(v = lambda.cv, lty = 2)
> lambda.gcv <- gcv(y, x, knots = knots3, degree = 3)$penalty
> abline(v = lambda.gcv, lty = 2)
> cv.fit <- rbpspline(y, x, knots3, degree = 3, penalty = "cv")
> gcv.fit <- rbpspline(y, x, knots3, degree = 3, penalty = "gcv")
> plot(x, y, col = "lightgrey")
> lines(cv.fit, col = 2)
> lines(gcv.fit, col = 3)
```

6.2 Building Response Surfaces for the GEV Parameters

In the previous section, we introduced the notion of response surfaces and we show that they should be used if one is interested in simultaneously fitting the GEV and the spatial dependence parameters of a max-stable process. However, one may wonder how to build accurate response surfaces for the GEV parameters. This is the aim of this section.

A first attempt could be to fit several max-stable models and identify the most promising ones by using the techniques on model selection introduced in Chapter 5. Although it is a legitimate approach, its use in practice is limited because the fitting procedure, due to the pairwise likelihood estimator, is CPU prohibitive.

A more pragmatic strategy is to consider only these response surfaces while omitting temporally the spatial dependence parameters. Although this strategy doesn't take into account all the uncertainties on the max-stable parameters, it should lead to accurate model selection as one expects the spatial dependence parameters and the GEV response surface parameters to be nearly orthogonal. The main asset of the latter approach is that fitting a (kind of) spatial GEV model to data is less CPU consuming.

This spatial GEV model is defined as follows:

$$Z(x) \sim GEV(\mu(x), \sigma(x), \xi(x)) \quad (6.20)$$

where the GEV parameters are defined through the following equations

$$\mu = X_\mu \beta_\mu, \quad \sigma = X_\sigma \beta_\sigma, \quad \xi = X_\xi \beta_\xi$$

where X . are design matrices and β . are parameters to be estimated.

The log-likelihood of the spatial GEV model is

$$\ell(\beta) = \sum_{i=1}^{n.site} \sum_{j=1}^{n.obs} \left\{ -\log \sigma_i - \left(1 + \xi_i \frac{z_{i,j} - \mu_i}{\sigma_i} \right)^{-1/\xi_i} - \left(1 + \frac{1}{\xi_i} \right) \log \left(1 + \xi_i \frac{z_{i,j} - \mu_i}{\sigma_i} \right) \right\} \quad (6.21)$$

where $\beta = (\beta_\mu, \beta_\sigma, \beta_\xi)$, μ_i , σ_i and ξ_i are the GEV parameters for the i -th site and $z_{i,j}$ is the j -th observation for the i -th site.

From equation (6.21), we can see that independence between stations is assumed. For most applications, this assumption is clearly incorrect and we require the use of the MLE asymptotic distribution under misspecification to get standard error estimates:

$$(\beta_\mu, \beta_\sigma, \beta_\xi) \overset{\cdot}{\sim} \mathcal{N}(\psi, H(\beta)^{-1} J(\beta) H(\beta)^{-1}), \quad n \rightarrow +\infty \quad (6.22)$$

where $H(\beta) = \mathbb{E}[\nabla^2 \ell_p(\beta; \mathbf{Y})]$ (the Hessian matrix) and $J(\beta) = \text{Var}[\nabla \ell_p(\beta; \mathbf{Y})]$.

In practice, the spatial GEV model is fitted to data through the `fitspatgev` function. The use of this function is similar to `fitmaxstab`.

Lets start by simulating a max-stable process with unit Fréchet margins and transform it to have a spatially structured GEV margins.

```
> n.site <- 20
> n.obs <- 50
> coord <- matrix(runif(2*n.site, 0, 10), ncol = 2)
> colnames(coord) <- c("lon", "lat")
> data <- rmaxstab(n.obs, coord, "gauss", cov11 = 100, cov12 = 25, cov22 = 220)
> param.loc <- -10 + 2 * coord[,2]
> param.scale <- 5 + 2 * coord[,1] + coord[,2]^2
> param.shape <- rep(0.2, n.site)
> for (i in 1:n.site)
+   data[,i] <- frech2gev(data[,i], param.loc[i], param.scale[i], param.shape[i])
```

Now we define appropriate response surfaces for our spatial GEV model and fit two different models.

```
> loc.form <- y ~ lat
> scale.form <- y ~ lon + I(lat^2)
> shape.form <- y ~ 1
> shape.form2 <- y ~ lon
> M1 <- fitspatgev(data, coord, loc.form, scale.form, shape.form)
> M2 <- fitspatgev(data, coord, loc.form, scale.form, shape.form2)
> M1
```

```
Model: Spatial GEV model
Deviance: 10539.99
TIC: 10647.06
```

```
Location Parameters:
locCoeff1 locCoeff2
-10.92      1.87
```

```

      Scale Parameters:
scaleCoeff1 scaleCoeff2 scaleCoeff3
      5.5874      1.8485      0.8579
      Shape Parameters:
shapeCoeff1
      0.2579

Standard Errors
      locCoeff1      locCoeff2      scaleCoeff1      scaleCoeff2      scaleCoeff3      shapeCoeff1
      1.7216      0.8970      2.1847      0.2414      0.1291      0.1098

Asymptotic Variance Covariance
      locCoeff1      locCoeff2      scaleCoeff1      scaleCoeff2      scaleCoeff3
locCoeff1      2.963756      0.925637      2.625197      -0.076452      0.115644
locCoeff2      0.925637      0.804575      1.153040      -0.066338      0.062434
scaleCoeff1      2.625197      1.153040      4.772983      -0.272605      0.199482
scaleCoeff2      -0.076452      -0.066338      -0.272605      0.058289      -0.001247
scaleCoeff3      0.115644      0.062434      0.199482      -0.001247      0.016677
shapeCoeff1      -0.057426      -0.018256      -0.047083      0.013160      -0.002727
      shapeCoeff1
locCoeff1      -0.057426
locCoeff2      -0.018256
scaleCoeff1      -0.047083
scaleCoeff2      0.013160
scaleCoeff3      -0.002727
shapeCoeff1      0.012066

Optimization Information
Convergence: successful
Function Evaluations: 1673

```

The output of model M1 is very similar to the one of a fitted max-stable process except the spatial dependence parameters are not present. As explained in Chapter 5, it is easy to perform model selection by inspecting the following output:

```

> anova(M1, M2)

Eigenvalue(s): 0.86

Analysis of Variance Table
      MDf Deviance Df Chisq Pr(> sum lambda Chisq)
M1      6      10540
M2      7      10539  1 0.551      0.4226

> TIC(M1, M2)

      M1      M2
10647.06 10648.58

```

From these two outputs, we can see that the p -value for the likelihood ratio test is around 0.72 which advocates the use of model **M1**. The TIC corroborates this conclusion.

Conclusion

A

P-splines with radial basis functions

A.1 Model definition

Let us recall that a general definition of a p-spline is given by

$$y_i = \beta_0 + \beta_1 x_i + \sum_{j=1}^q b_j(x_i - \kappa_j) \quad (\text{A.1})$$

for some basis functions b_j and knots κ_j .

As the purpose of this document is the modelling of spatial extremes, we benefit from using *radial basis* functions. Radial basis functions depend only on the distance $|x_i - \kappa_j|$ so that a generalisation to higher dimension, i.e. $\|\mathbf{x}_i - \kappa_j\|$, $\mathbf{x}_i, \kappa_j \in \mathbb{R}^d$, is straightforward.

The model for p-spline with radial basis function of order p , p being odd, is

$$f(x) = \beta_0 + \beta_1 x + \dots + \beta_{m-1} x^{m-1} + \sum_{j=1}^q \beta_{m+j} |x - \kappa_j|^{2m-1} \quad (\text{A.2})$$

where $p = 2m - 1$.

The fitting criterion is

$$\text{minimize } \|\mathbf{y} - \mathbf{X}\beta\|^2 + \lambda^{2m-1} \beta^T \mathbf{K} \beta \quad (\text{A.3})$$

where

$$\mathbf{X} = \begin{bmatrix} 1 & x_1 & \dots & x_1^{m-1} & |x_1 - \kappa_1|^{2m-1} & \dots & |x_1 - \kappa_q|^{2m-1} \\ \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_n & \dots & x_n^{m-1} & |x_n - \kappa_1|^{2m-1} & \dots & |x_n - \kappa_q|^{2m-1} \end{bmatrix}$$

and $\mathbf{K} = \mathbf{K}_*^T \mathbf{K}_*$ with

$$\mathbf{K}_* = \begin{bmatrix} 0 & \dots & 0 & 0 & \dots & 0 \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & 0 & \dots & 0 \\ 0 & \dots & 0 & |\kappa_1 - \kappa_1|^{m-1/2} & \dots & |\kappa_1 - \kappa_q|^{m-1/2} \\ \vdots & \ddots & \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & |\kappa_q - \kappa_1|^{m-1/2} & \dots & |\kappa_q - \kappa_q|^{m-1/2} \end{bmatrix}$$

where the m first rows and columns of $\mathbf{K}_*$ have zeros as elements.

A.2 Fast computation of p-splines

Although this section is included in the p-splines with radial basis functions, the methodology introduced here can be successfully applied to any other basis functions [Ruppert et al., 2003].

As we stated in Section 6.1.2, for a fixed smoothing parameter λ , the fitted values are given by

$$\hat{\mathbf{y}} = \mathbf{X}(\mathbf{X}^T \mathbf{X} + \lambda \mathbf{K})^{-1} \mathbf{X}^T \mathbf{y}$$

for some symmetric matrix $\mathbf{K}$.

Consequently, to perform automatic selection for λ by minimising the CV or GCV criterion might be computationally demanding and numerically unstable. Fortunately, the Demmler–Reinsch orthogonalisation often overcomes these issues. The following lines describe how it works in practice.

1. Obtain the Cholesky decomposition of $\mathbf{X}^T \mathbf{X}$ i.e.

$$\mathbf{X}^T \mathbf{X} = \mathbf{R}^T \mathbf{R}$$

where $\mathbf{R}$ is a square matrix and invertible

2. Obtain the singular value decomposition of $\mathbf{R}^{-T} \mathbf{K} \mathbf{R}^{-1}$ i.e.

$$\mathbf{R}^{-T} \mathbf{K} \mathbf{R}^{-1} = \mathbf{U} \mathbf{\Lambda} \mathbf{U}^T$$

3. Define $\mathbf{A} \leftarrow \mathbf{X} \mathbf{R}^{-1} \mathbf{U}$ and $\mathbf{b} \leftarrow \mathbf{A}^T \mathbf{y}$

4. The fitted values are

$$\hat{\mathbf{y}} = \mathbf{A} \frac{\mathbf{b}}{\mathbf{1} + \lambda \mathbf{\Lambda}}$$

with corresponding degrees of freedom

$$df(\lambda) = \mathbf{1}^T \frac{\mathbf{1}}{\mathbf{1} + \lambda \mathbf{\Lambda}}$$

Once the matrices $\mathbf{A}$ and $\mathbf{\Lambda}$ and the vector $\mathbf{b}$ have been computed, the fitted values $\hat{\mathbf{y}}$ and $df(\lambda)$ are obtained through a simple matrix multiplication. This is appealing as now the automatic selection for λ will be cheaper.

Bibliography

- Akaike, H. (1974). A new look at the statistical model identification. In *Automatic Control*, volume 19 of *IEEE Transactions on*.
- Chan, G. and Wood, A. (1997). An algorithm for simulating stationary Gaussian random fields. *Journal of the Royal Statistical Society Series C*, 46(1):171–181.
- Chandler, R. E. and Bate, S. (2007). Inference for clustered data using the independence loglikelihood. *Biometrika*, 94(1):167–183.
- Chilès, J.-P. (1997). *Géostatistique des phénomènes non-stationnaires (dans le plan)*. PhD thesis, Université de Nancy I.
- Chilès, J.-P. and Delfiner, P. (1999). *Geostatistics: Modelling Spatial Uncertainty*. Wiley, New York.
- Coles, S. (2001). *An Introduction to Statistical Modelling of Extreme Values*. Springer Series in Statistics. Springer Series in Statistics, London.
- Cooley, D., Naveau, P., and Poncet, P. (2006). Variograms for spatial max-stable random fields. In Springer, editor, *Dependence in Probability and Statistics*, volume 187, pages 373–390. Springer, New York, lecture notes in statistics edition.
- Cox, D. R. and Reid, N. (2004). A note on pseudolikelihood constructed from marginal densities. *Biometrika*, 91(3):729–737.
- Cramér, H. (1946). *Mathematical Methods of Statistics*, volume 9 of *Princeton Mathematical Series*. Princeton University Press.
- Cressie, N. A. C. (1993). *Statistics for Spatial Data*. Wiley Series in Probability and Statistics. John Wiley & Sons inc., New York.
- Davis, M. (1987). Production of conditional simulations via the lu triangular decomposition of the covariance matrix. *Mathematical Geology*, 19(2):91–98.
- Davison, A. (2003). *Statistical Models*. Cambridge Series in Statistical and Probabilistic Mathematics. Cambridge University Press.
- Dietrich, C. (1995). A simple and efficient space domain implementation of the turning bands method. *Water Resources Research*, 31(1):147–156.

- Diggle, P., Ribeiro, P., and Justiniano, P. (2007). *Model-based Geostatistics*. Springer Series in Statistics. Springer.
- Efron, B. (1982). The Jackknife, the Bootstrap and Other Resampling Plans. In *CBMS-NSF Regional Conference Series in Applied Mathematics*, Philadelphia. SIAM.
- Emery, X. and Lantuéjoul, C. (2006). Tbsim: A computer program for conditional simulation of three-dimensional gaussian random fields via the turning bands method. *Computer and Geosciences*, 32:1615–1628.
- Freulon, X. and de Fouquet, C. (1991). Remarques sur la pratique des bandes tournantes à trois dimensions. Technical report, Ecole des Mines de Paris.
- Gholam-Rezaee, M. (2009). *Spatial extreme value: A composite likelihood*. PhD thesis, École Polytechnique Fédérale de Lausanne.
- Hosking, J., Wallis, J., and Wood, E. (1985). Estimation of the Generalized Extreme Value Distribution by the Method of Probability Weighted Moments. *Technometrics*, 27(3):251–261.
- Kent, J. (1982). Robust properties of likelihood ratio tests. *Biometrika*, 69:19–27.
- Lindsay, B. (1988). *Composite likelihood methods*. Statistical Inference from Stochastic Processes. American Mathematical Society, Providence.
- Mantoglou, A. (1987). Digital-simulation of multivariate two-dimensional and three-dimensional stochastic processes with a spectral turning bands method. *Mathematical Geology*, 19:129–149.
- Matérn, B. (1986). *Spatial Variations*. Springer, Berlin, 2nd edition edition.
- Matheron, G. (1973). The intrinsic random functions and their applications. *Advances in Applied Probability*, 5:439–468.
- Matheron, G. (1987). Suffit-il, pour une covariance, d’être de type positif ? *Sciences de la Terre, série informatique géologique*, 26:51–66.
- Naveau, P., Guillou, A., and Cooley, D. (2009). Modelling pairwise dependence of maxima in space. *Submitted to Biometrika*.
- Nelder, J. and Mead, R. (1965). A simplex algorithm for function minimization. *Computer Journal*, 7:308–313.
- Neyman, J. and Pearson, E. (1933). On the problem of the most efficient tests of statistical hypothesis. *Philosophical Transactions of the Royal Statistical Society of London. Series A*, 231:289–337.
- Padoan, S. (2008). *Computational Methods for Complex Problems in Extreme Value Theory*. PhD thesis, University of Padova.
- Padoan, S., Ribatet, M., and Sisson, S. (2008). Likelihood-based inference for max-stable processes. *Submitted to the Journal of the American Statistical Association (Theory & Methods)*.
- Pickands, J. (1981). Multivariate Extreme Value Distributions. In *Proceedings 43rd Session International Statistical Institute*.

- R Development Core Team (2007). *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria. ISBN 3-900051-07-0.
- Rotnitzki, A. and Jewell, N. (1990). Hypothesis testing of regression parameters in semiparametric generalized linear models for cluster correlated data. *Biometrika*, 77:495–497.
- Ruppert, D., Wand, M., and Carroll, R. (2003). *Semiparametric regression*. Cambridge University Press.
- Schlather, M. (1999). An introduction to positive definite functions and to unconditional simulation of random fields. Technical report st 99-10, Dept. of Mathematics and Statistics, Lancaster University.
- Schlather, M. (2002). Models for stationary max-stable random fields. *Extremes*, 5(1):33–44.
- Schlather, M. and Tawn, J. (2003). A dependence measure for multivariate and spatial extremes: Properties and inference. *Biometrika*, 90(1):139–156.
- Varin, C. and Vidoni, P. (2005). A note on composite likelihood inference and model selection. *Biometrika*, 92(3):519–528.
- Wood, S. (2006). *Generalized Additive Models*. Chapman & Hall.

Index

- covariance function
 - bessel, 7
 - cauchy, 7
 - elliptical, 9
 - powered exponential, 7
 - range, 7
 - smooth, 7
 - Whittle–Matérn, 7
- cross-validation, 42
 - generalized, 42
- degrees of freedom, 40, 42
- design matrix, 39
- distance
 - Euclidean, 7
 - Mahanalobis, 5
- eigen-decomposition, 6
- eigenvalues, 6
- extremal coefficient, 19
 - function, 20
- Fréchet
 - unit, 3
- hat matrix, 40
- information criterion
 - AIC, 34
 - TIC, 34
- intensity measure, 3, 6
- Jacobian, 38
- Kullback–Leibler discrepancy, 29
- least squares, 27, 39
- likelihood ratio statistic, 35
- linear model, 39
- madogram, 21
 - F -madogram, 23
 - λ -madogram, 24
- max-stable
 - process, 3
 - property, 4
- misspecification, 28
- p-splines, 41
- pairwise-likelihood, 28
- Poisson process, 3
- rainfall-storm process, 3, 4
- score equation, 31
- smoother matrix, 42
- smoothing parameter, 41
- standard errors, 31
- variogram, 19